

INFORMAČNÍ BULLETIN



České statistické společnosti

Ročník 23, číslo 4, prosinec 2012

COMPOUND QUANTIFICATION OF IMPLICATIONS, DOUBLE-IMPLICATIONS AND EQUIVALENCY IN FOUR-FOLD TABLES

SDRUŽENÁ KVANTIFIKACE IMPLIKACÍ, DVOJITÝCH IMPLIKACÍ A EKVIVALENCE VE ČTYŘPOLNÍCH TABULKÁCH

Jiří Ivánek

Address: Department of Information and Knowledge Engineering,
Faculty of Informatics and Statistics, University of Economics, Prague

E-mail: ivanek@vse.cz

Abstract: Relations between two Boolean attributes derived from data can be quantified by functions defined on four-fold tables. In the paper, the method of a construction of the affiliated (logically nearest) double-implicational and equivalence quantifiers to a given implicational quantifier is recalled and applied to subclass of ratio-implicational quantifiers. Possible truth-configurations of the compounded set of four implications, two double-implications and equivalency (given some threshold) are discussed in details and interpreted as types of patterns hidden in the four-fold table.

Keywords: Four-fold Table, Implicational Quantifiers, Association Rules.

Abstrakt: Vztahy mezi dvěma boolovskými atributy odvozenými z dat mohou být kvantifikovány pomocí funkcí definovaných na čtyřpolních tabulkách. V tomto předkládaném článku aplikujeme metodu konstrukce sdružených (logicky nejbližších) dvojité-implikačních a ekvivalenčních kvantifikátorů k danému implikačnímu kvantifikátoru na podtřídu podílově-implikačních kvantifikátorů. Detailně diskutovány jsou možné pravdivostní konfigurace složené množiny čtyř implikací, dvou dvojitých implikací a ekvivalence (při zadané prahové hodnotě), které lze interpretovat jako typy vzorů skrytých ve čtyřpolní tabulce.

Klíčová slova: Čtyřpolní tabulka, implikační kvantifikátory, asociační pravidla.

1. Introduction

Assume having a data file and consider two Boolean (binary, dichotomic) attributes φ and ψ . A four-fold table $\langle a, b, c, d \rangle$ corresponding to these attributes is composed from numbers of objects in data satisfying four different Boolean combinations of attributes:

	ψ	$\neg\psi$	a – number of objects satisfying both φ and ψ ,
φ	a	b	b – number of objects satisfying φ and not satisfying ψ ,
$\neg\varphi$	c	d	c – number of objects not satisfying φ and satisfying ψ ,
			d – number of objects not satisfying φ and not satisfying ψ .

Four-fold table quantifier $\sim$ is a function with values from the interval $[0, 1]$ defined on the set of all four-fold tables $\langle a, b, c, d \rangle$. Several classes of quantifiers (implicational, equivalency) have been studied in the theory of the GUHA method [2], [3].

A quantifier $\sim (a, b)$ is *implicational* if $\sim (a', b') \geq \sim (a, b)$ when $a' \geq a, b' \leq b$. The most common example of implicational quantifier is the quantifier of basic implication (corresponds to the notion of a confidence of an association rule used in data mining, see [1], [8]):

$$\Rightarrow_0 (a, b) = \frac{a}{a + b}.$$

In the paper, a subclass of ratio-implicational quantifiers is investigated. The method of construction of the affiliated (logically nearest) double implicational and equivalence quantifiers to a given implicational quantifier is recalled and applied to the subclass of ratio-implicational quantifiers. Possible truth-configurations of the compounded set of seven formulae (given data and some threshold) are discussed in details.

2. Ratio-implicational quantifiers

This is one of the main properties of the basic implicational quantifier: the greater the ratio a/b , the greater the value of the quantifier. This property is stronger than that used in the definition of implicational quantifiers. Therefore we introduced a subclass of implicational quantifiers with this property [6]:

A quantifier $\sim (a, b)$ is *ratio-implicational*, if $\sim (a', b') \geq \sim (a, b)$ when $a'b \geq ab'$. For any $\theta > 0$ the following quantifier is ratio-implicational:

$$\Rightarrow_\theta (a, b) = \frac{a}{a + \theta b}.$$

It is clear that each ratio-implicational quantifier is also implicational. There are some other properties of ratio-implicational quantifiers proved in [6]:

- (i) If $a'b = ab'$ then $\Rightarrow^* (a', b') = \Rightarrow^* (a, b)$.

(ii) There are numbers m^*, M^* from $[0, 1]$ such that

$$\begin{aligned} m^* &= \Rightarrow^* (0, b) \text{ for all } b > 0, \\ M^* &= \Rightarrow^* (a, 0) \text{ for all } a > 0, \\ m^* &\leq \Rightarrow^* (a, b) \leq M^* \text{ for all } a, b > 0. \end{aligned}$$

(iii) There is a non-decreasing function g^* defined on non-negative rationals and ∞ such that

$$\Rightarrow^* (a, b) = g^* \left(\frac{a}{b} \right).$$

The function g^* is defined as follows:

$$g^*(0) = m^*, \quad g^*(\infty) = M^*, \quad g^* \left(\frac{i}{j} \right) = \Rightarrow^* (i, j) \quad \text{for all integers } i, j > 0.$$

Correctness of this definition follows from (i), (ii); monotonicity follows from the definition of ratio-implicational quantifiers.

The class of ratio-implicational quantifiers is a proper subclass of the class of implicational quantifiers. For a counterexample, let us observe that statistically motivated quantifier $\Rightarrow_p^?$ (where p is a parameter, $0 < p < 1$)

$$\Rightarrow_p^? (a, b) = \sum_{i=0}^a \frac{(a+b)!}{i! (a+b-i)!} p^i (1-p)^{a+b-i}$$

is implicational (see [3]), but is not ratio-implicational because for instance

$$\Rightarrow_p^? (0, b) = (1-p)^b \neq \Rightarrow_p^? (0, b+1) = (1-p)^{b+1}.$$

3. Affiliated double-implication and equivalency quantifiers

In the paper [4], the method of construction of triads of quantifiers is described.

Starting from an implicational quantifier $\Rightarrow^*$, *affiliated double-implicational quantifier* $\Leftrightarrow^*$ is given by the formula

$$\Leftrightarrow^* (a, b, c) = \Rightarrow^* (a, b+c),$$

and *affiliated equivalency quantifier* $\equiv^*$ is given by the formula

$$\equiv^* (a, b, c, d) = \Rightarrow^* (a+d, b+c).$$

Double-implicational quantifier $\Leftrightarrow^*$ measures the validity of bi-implication $(\varphi \Rightarrow \psi) \wedge (\psi \Rightarrow \varphi)$ in data taking into account only cases where φ or ψ is satisfied. Equivalency quantifier $\equiv^*$ measures the validity of equivalency $\varphi \equiv \psi$ in the whole data. Both affiliated quantifiers $\Leftrightarrow^*, \equiv^*$ naturally extend quantification of implication (given by a definition of particular implicational quantifier $\Rightarrow^*$) for covering also two types of symmetric relations between φ and ψ in data.

It is proved that the above constructed double-implicational quantifier $\Leftrightarrow^*$ is in some sense the least strict one (out of the class of so-called Σ -double implication quantifiers, see [7], [4]) satisfying required inequality:

$$\Leftrightarrow^* (a, b, c) \leq \min(\Rightarrow^* (a, b), \Rightarrow^* (a, c)).$$

Analogically, the above constructed equivalency quantifier $\equiv^*$ is in some sense the most strict one (out of the class of so-called Σ -equivalency quantifiers, see [7], [4]) satisfying inequality:

$$\equiv^* (a, b, c, d) \geq \max(\Leftrightarrow^* (a, b, c), \Leftrightarrow^* (d, b, c)).$$

In the case when the starting quantifier $\Rightarrow^*$ is ratio-implicational, there are further useful connections between it and affiliated equivalency quantifier $\equiv^*$ (proved in [5]):

For all a, b, c, d the value $\equiv^* (a, b, c, d)$ lies both

- (i) between the values $\Rightarrow^* (a, b)$, and $\Rightarrow^* (d, c)$;
- (ii) between the values $\Rightarrow^* (a, c)$, and $\Rightarrow^* (d, b)$.

4. Discussion of possible truth configurations

Let $\Rightarrow^*$ be a ratio-implicational quantifier, $\Leftrightarrow^*$ and $\equiv^*$ be its affiliated double-implicational and equivalency quantifiers. Assume some truth threshold t from $[0, 1]$ is given. A formulae $\varphi \sim \psi$ is treated as true in data if the value of the quantifier $\sim$ in the four-fold table $\langle a, b, c, d \rangle$ corresponding to the attributes φ, ψ is greater or equal to t . Using inequalities presented in the previous paragraph, we shall discuss possible truth configurations of the set of formulae

$$\begin{array}{lll} \text{Implications:} & \varphi \Rightarrow^* \psi, & \psi \Rightarrow^* \varphi, \quad \neg\varphi \Rightarrow^* \neg\psi, \quad \neg\psi \Rightarrow^* \neg\varphi; \\ \text{Double-implications:} & \varphi \Leftrightarrow^* \psi, & \neg\varphi \Leftrightarrow^* \neg\psi; \\ \text{Equivalency:} & \varphi \equiv^* \psi. & \end{array}$$

Relations among these formulae are illustrated graphically in Figure 1. In this example, all above mentioned formulae quantified by the triad of the basic quantifiers

$$\Rightarrow_0^* (a, b) = \frac{a}{a+b}, \quad \Leftrightarrow_0^* (a, b, c) = \frac{a}{a+b+c}, \quad \equiv_0^* (a, b, c, d) = \frac{a+d}{a+b+c+d}$$

are not true given the threshold $t = 0.7$.

Which subset of these formulae (given some general ratio-implicational quantifier and some threshold) could be true (in some four fold table)? There are formally $2^7 = 128$ such truth configurations, but we shall show that most of them are not possible in any data.

4.1. Discussion according to truthfulness of equivalency

e0) $\varphi \equiv^* \psi$ is not true. Then

- at most two implications could be true (but excluding the pair $\varphi \Rightarrow^* \psi, \neg\varphi \Rightarrow^* \neg\psi$, and the pair $\psi \Rightarrow^* \varphi, \neg\psi \Rightarrow^* \neg\varphi$);
- no double implications is true.

e1) $\varphi \equiv^* \psi$ is true. Then

- at least two implications are true;
- 0, 1 or 2 double-implications could be true.

4.2. Discussion according to truthfulness of double-implications

d0) Both $\varphi \Leftrightarrow^* \psi, \neg\varphi \Leftrightarrow^* \neg\psi$ are not true. Then

- the equivalency $\varphi \equiv^* \psi$ could be true or not;
- 0, 1, 2, 3 or 4 implications could be true.

d1) One of double-implications is true. Then

- the equivalency is true;
- at least two implications are true.

d2) Both double-implications are true. Then

- the equivalency is true;
- all implications are true.

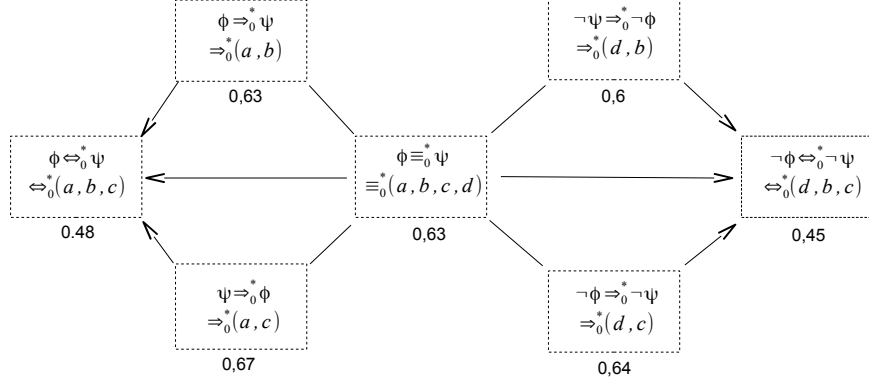


Figure 1: Truth configuration of four-fold table $\langle a = 10, b = 6, c = 5, d = 9 \rangle$, triad of basic quantifiers, threshold $t = 0,7$.

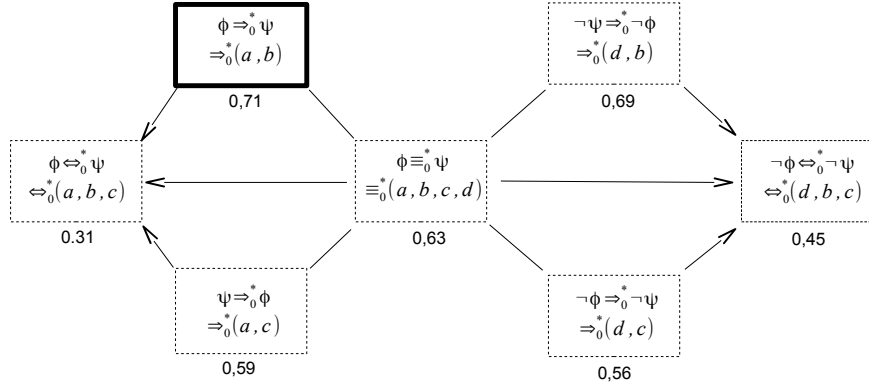


Figure 2: Truth configuration of four-fold table $\langle a = 10, b = 4, c = 7, d = 9 \rangle$, triad of basic quantifiers, threshold $t = 0,7$.

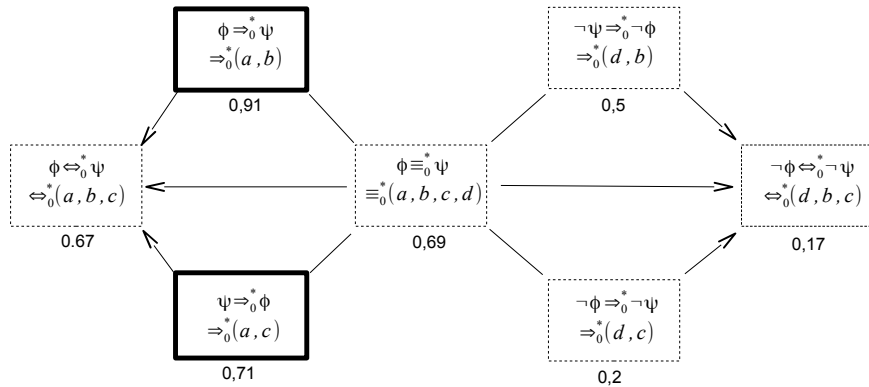


Figure 3: Truth configuration of four-fold table $\langle a = 10, b = 1, c = 4, d = 1 \rangle$, triad of basic quantifiers, threshold $t = 0,7$.

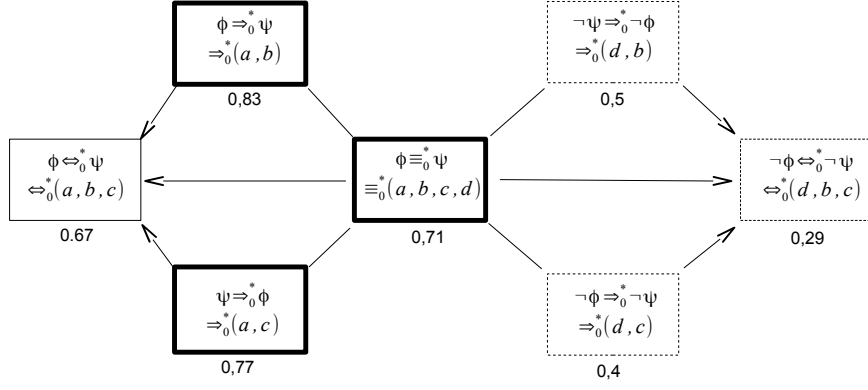


Figure 4: Truth configuration of four-fold table $\langle a = 10, b = 2, c = 3, d = 2 \rangle$, triad of basic quantifiers, threshold $t = 0,7$.

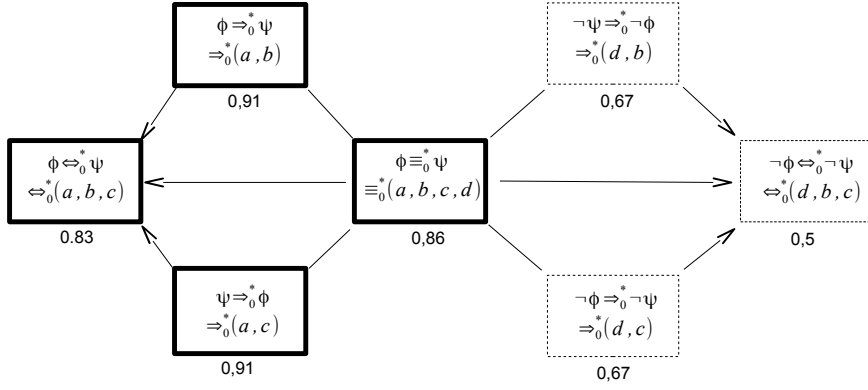


Figure 5: Truth configuration of four-fold table $\langle a = 10, b = 1, c = 1, d = 2 \rangle$, triad of basic quantifiers, threshold $t = 0,7$.

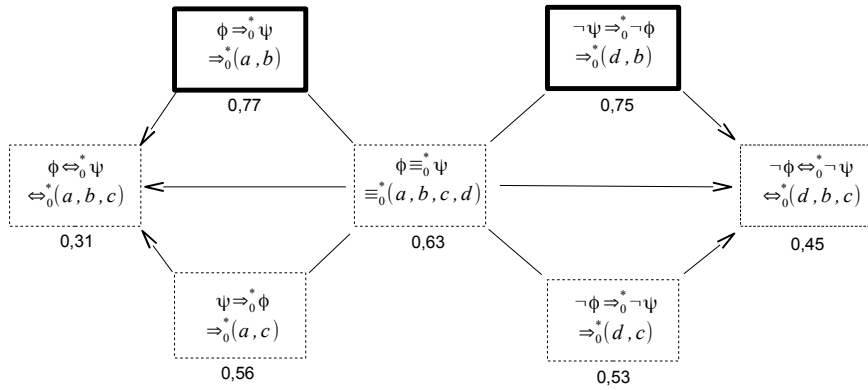


Figure 6: Truth configuration of four-fold table $\langle a = 10, b = 3, c = 8, d = 9 \rangle$, triad of basic quantifiers, threshold $t = 0,7$.

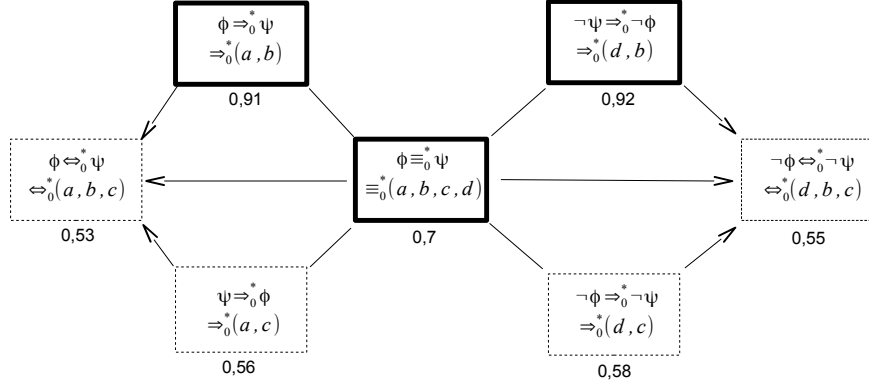


Figure 7: Truth configuration of four-fold table $\langle a = 10, b = 1, c = 8, d = 11 \rangle$, triad of basic quantifiers, threshold $t = 0,7$.

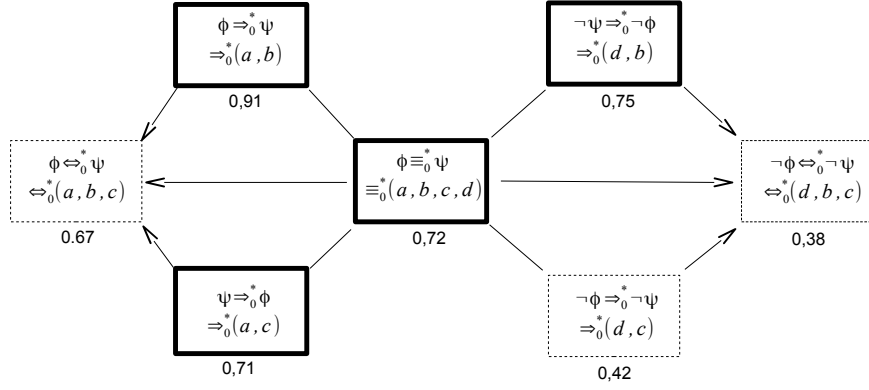


Figure 8: Truth configuration of four-fold table $\langle a = 10, b = 1, c = 4, d = 3 \rangle$, triad of basic quantifiers, threshold $t = 0,7$.

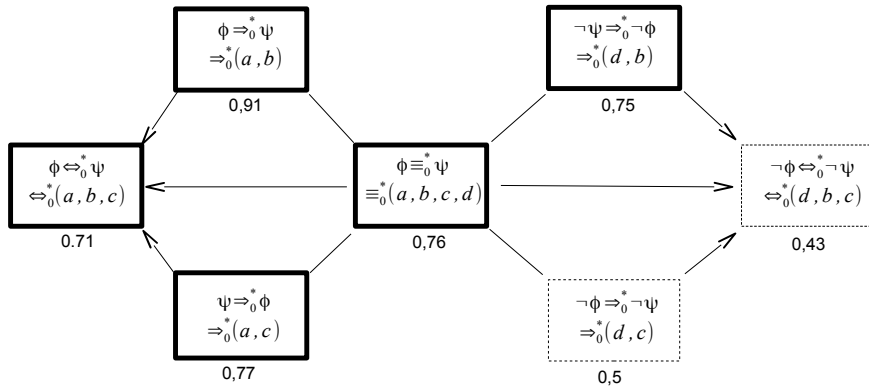


Figure 9: Truth configuration of four-fold table $\langle a = 10, b = 1, c = 3, d = 3 \rangle$, triad of basic quantifiers, threshold $t = 0,7$.

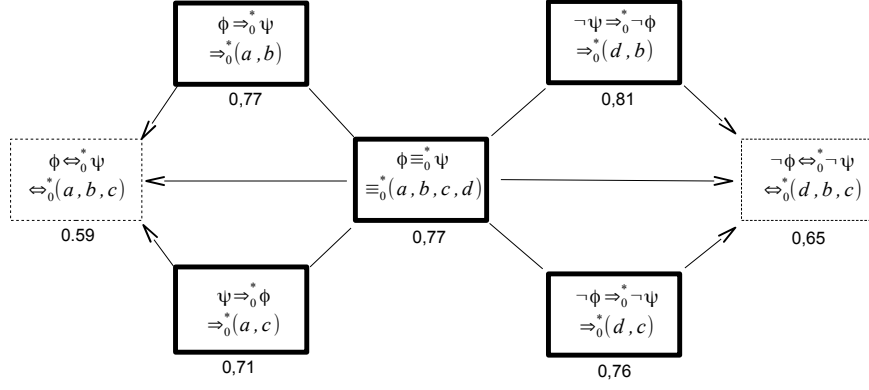


Figure 10: Truth configuration of four-fold table $\langle a = 10, b = 3, c = 4, d = 13 \rangle$, triad of basic quantifiers, threshold $t = 0,7$.

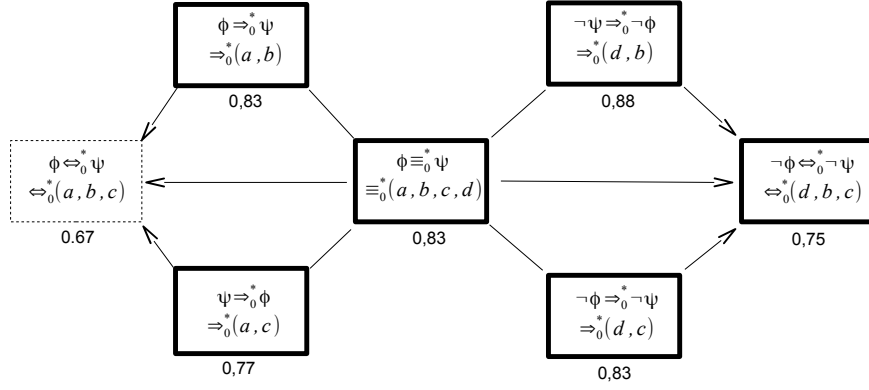


Figure 11: Truth configuration of four-fold table $\langle a = 10, b = 2, c = 3, d = 15 \rangle$, triad of basic quantifiers, threshold $t = 0,7$.

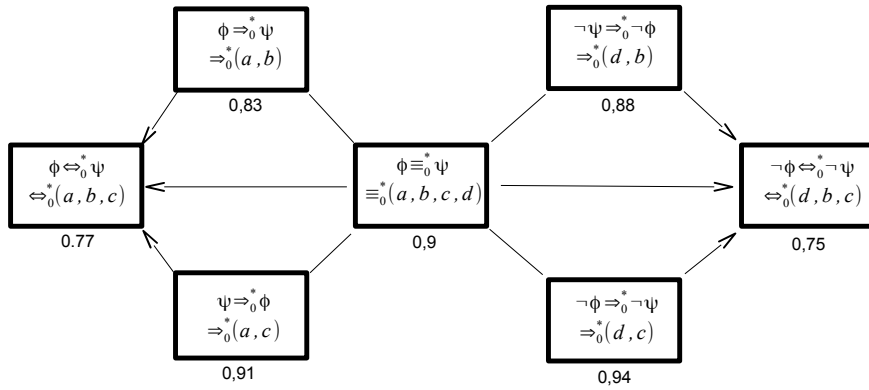


Figure 12: Truth configuration of four-fold table $\langle a = 10, b = 2, c = 1, d = 17 \rangle$, triad of basic quantifiers, threshold $t = 0,7$.

4.3. Discussion according to truthfulness of implications

This discussion will be provided in details with distinction of different situations. Types of configurations will be described by numbers $i/d/e$ of true implications, double-implications and equivalency, respectively. Each type of configuration will be provided with an appropriate example of four-fold table $\langle a, b, c, d \rangle$ quantified by the triad of the basic quantifiers

$$\Rightarrow_0^* (a, b) = \frac{a}{a+b}, \quad \Leftrightarrow_0^* (a, b, c) = \frac{a}{a+b+c}, \quad \equiv_0^* (a, b, c, d) = \frac{a+d}{a+b+c+d}$$

with the list of true formulae given the threshold $t = 0.7$.

- i0) No implication is true. Then no formulae of double-implication or equivalency could be true.

Type: 0/0/0

Example: Four-fold table $\langle a = 10, b = 6, c = 5, d = 9 \rangle$

Truth configuration: no true formula.

(see Figure 1)

- i1) Exactly one implication is true. Then no formulae of double-implication or equivalency could be true.

Type: 1/0/0

Example: Four-fold table $\langle a = 10, b = 4, c = 7, d = 9 \rangle$

Truth configuration: $\varphi \Rightarrow^* \psi$.

(see Figure 2)

- i2) Exactly two implications are true. Two pairs out from six pairs of implications are excluded (namely the pair $\varphi \Rightarrow^* \psi$, $\neg\varphi \Rightarrow^* \neg\psi$ and the pair $\psi \Rightarrow^* \varphi$, $\neg\psi \Rightarrow^* \neg\varphi$, because in these cases also the equivalency would be true, hence some third implication also would be true). Remaining four possible pairs of true implications lead to two different cases:

- i2a) Either $\varphi \Rightarrow^* \psi$, $\psi \Rightarrow^* \varphi$ are true or $\neg\varphi \Rightarrow^* \neg\psi$, $\neg\psi \Rightarrow^* \neg\varphi$ are true. Then corresponding double-implication (either $\varphi \Leftrightarrow^* \psi$ or $\neg\varphi \Leftrightarrow^* \neg\psi$) could be true and also the equivalency could be true. Possible types of configurations:

Type: 2/0/0

Example: Four-fold table $\langle a = 10, b = 1, c = 4, d = 1 \rangle$

Truth configuration: $\varphi \Rightarrow^* \psi$, $\psi \Rightarrow^* \varphi$.

(see Figure 3)

Type: 2/0/1

Example: Four-fold table $\langle a = 10, b = 2, c = 3, d = 2 \rangle$

Truth configuration: $\varphi \Rightarrow^* \psi, \psi \Rightarrow^* \varphi, \varphi \equiv^* \psi$.

(see Figure 4)

Type: 2/1/1

Example: Four-fold table $\langle a = 10, b = 1, c = 1, d = 2 \rangle$

Truth configuration: $\varphi \Rightarrow^* \psi, \psi \Rightarrow^* \varphi, \varphi \Leftrightarrow^* \psi, \varphi \equiv^* \psi$.

(see Figure 5)

- i2b) Either $\varphi \Rightarrow^* \psi, \neg\psi \Rightarrow^* \neg\varphi$ are true or $\psi \Rightarrow^* \varphi, \neg\varphi \Rightarrow^* \neg\psi$ are true. Then no double-implication could be true. Possible types of configurations:

Type: 2/0/0

Example: Four-fold table $\langle a = 10, b = 3, c = 8, d = 9 \rangle$

Truth configuration: $\varphi \Rightarrow^* \psi, \neg\psi \Rightarrow^* \neg\varphi$.

(see Figure 6)

Type: 2/0/1

Example: Four-fold table $\langle a = 10, b = 1, c = 8, d = 11 \rangle$

Truth configuration: $\varphi \Rightarrow^* \psi, \neg\psi \Rightarrow^* \neg\varphi, \varphi \equiv^* \psi$.

(see Figure 7)

- i3) Exactly three implications are true. Then the equivalency is true and at most one double-implication is true. Possible types of configurations:

Type: 3/0/1

Example: Four-fold table $\langle a = 10, b = 1, c = 4, d = 3 \rangle$

Truth configuration: $\varphi \Rightarrow^* \psi, \psi \Rightarrow^* \varphi, \neg\psi \Rightarrow^* \neg\varphi, \varphi \equiv^* \psi$.

(see Figure 8)

Type: 3/1/1

Example: Four-fold table $\langle a = 10, b = 1, c = 3, d = 3 \rangle$

Truth configuration: $\varphi \Rightarrow^* \psi, \psi \Rightarrow^* \varphi, \neg\psi \Rightarrow^* \neg\varphi, \varphi \Leftrightarrow^* \psi, \varphi \equiv^* \psi$.

(see Figure 9)

- i4) All four implications are true. Then the equivalency is true and 0, 1, or 2 double-implications could be true. Possible types of configurations:

Type: 4/0/1

Example: Four-fold table $\langle a = 10, b = 3, c = 4, d = 13 \rangle$

Truth configuration: $\varphi \Rightarrow^* \psi, \psi \Rightarrow^* \varphi, \neg\psi \Rightarrow^* \neg\varphi,$
 $\neg\varphi \Rightarrow^* \neg\psi, \varphi \equiv^* \psi.$

(see Figure 10)

Type: 4/1/1

Example: Four-fold table $\langle a = 10, b = 2, c = 3, d = 15 \rangle$

Truth configuration: $\varphi \Rightarrow^* \psi, \psi \Rightarrow^* \varphi, \neg\psi \Rightarrow^* \neg\varphi,$
 $\neg\varphi \Rightarrow^* \neg\psi, \neg\varphi \Leftrightarrow^* \neg\psi, \varphi \equiv^* \psi.$

(see Figure 11)

Type: 4/2/1

Example: Four-fold table $\langle a = 10, b = 2, c = 1, d = 17 \rangle$

Truth configuration: $\varphi \Rightarrow^* \psi, \psi \Rightarrow^* \varphi, \neg\psi \Rightarrow^* \neg\varphi,$
 $\neg\varphi \Rightarrow^* \neg\psi, \varphi \Leftrightarrow^* \psi, \neg\varphi \Leftrightarrow^* \neg\psi, \varphi \equiv^* \psi.$

(see Figure 12)

Summary of discussion: only 10 types of configurations are possible out from $5 \cdot 3 \cdot 2 = 30$ formally existing types. More detailed analysis would show that there are exactly 27 possible configurations out of 128 formally existing ones.

5. Conclusions

In the paper, the class of ratio-implicational quantifiers defined on four-fold tables was discussed and following properties were presented:

- each ratio-implicational quantifier can be represented by a non-decreasing function on rationals;
- there are double-implication and equivalency quantifiers affiliated to a given ratio-implicational quantifier which leads to the compounded set of seven formulae

$$\varphi \Rightarrow^* \psi, \psi \Rightarrow^* \varphi, \neg\psi \Rightarrow^* \neg\varphi, \neg\varphi \Rightarrow^* \neg\psi, \varphi \Leftrightarrow^* \psi, \neg\varphi \Leftrightarrow^* \neg\psi, \varphi \equiv^* \psi$$

connected by the set of inequalities among their values in data;

- number of possible truth-configurations of these formulae in data given some threshold is significantly reduced;

- resulting truth-configuration for the given four-fold table could be illustrated by a simple figure.

The approach described in the paper can serve to formulate new data-mining tasks seeking for both asymmetric and symmetric association rules in an unified way or to filter or visualize sets of association rules for more user-oriented outputs of data-mining procedures.

Acknowledgements This work has been supported by the Ministry of Education, Youths and Sports of the Czech Republic (MSM 6138439910).

References

- [1] Aggraval, R. et al.: Fast Discovery of Association Rules. In Fayyad, V. M. et al.: *Advances in Knowledge Discovery and Data Mining*. AAAI Press / MIT Press 1996, pp. 307–328.
- [2] Hájek, P., Havránek, T.: *Mechanising Hypothesis Formation – Mathematical Foundations for a General Theory*. Springer-Verlag, Berlin 1978, 396 p.
- [3] Hájek, P., Havránek, T., Chytil M.: *Metoda GUHA*. Academia, Praha 1983, 314 p. (in Czech).
- [4] Ivánek, J.: On the Correspondence between Classes of Implicational and Equivalence Quantifiers. In *Principles of Data Mining and Knowledge Discovery*. Proc. PKDD’99 Prague (Zytkow, J. and Rauch, J., eds.), Springer-Verlag, Berlin 1999, p. 116–124.
- [5] Ivánek, J.: *Affiliated ratio-implicational and equivalency data-mining quantifiers and their truth configurations*. In: WUPES’2012, Mariánské Lázně (submitted).
- [6] Ivánek, J.: Construction of Implicational Quantifiers from Fuzzy Implications. *Fuzzy Sets and Systems* 151, 2005, pp. 381–391.
- [7] Rauch, J.: Classes of Four-Fold Table Quantifiers. In *Principles of Data Mining and Knowledge Discovery* (Quafafou, M. and Zytkow, J., eds.), Springer Verlag, Berlin 1998, pp. 203–211.
- [8] Zembowicz, R.; Zytkow, J.: From Contingency Tables to Various Forms of Knowledge in Databases. In Fayyad, U.M. et al.: *Advances in Knowledge Discovery and Data Mining*. AAAI Press / The MIT Press 1996, pp. 329–349.

JSOU MEZE PRO UŽITÍ DIDAKTICKÝCH TESTŮ K ODHADU VĚDOMOSTI ŽÁKŮ A STUDENTŮ?

Zdeněk Půlpán

Adresa: Prof. RNDr. PhDr. Zdeněk Půlpán, CSc.,
Univerzita Hradec Králové, Přírodovědecká fakulta,
Katedra matematiky, Rokitanského 62, 500 03 Hradec Králové 3
E-mail: zdenek.pulpan@uhk.cz

Abstrakt: Stať upozorňuje na důležitost analýzy podmínek aplikací statistických metod a s tím související možné omezení interpretace výsledných numerických hodnot. V důsledku vysoké variability proměnných, definovaných v souborech živých jedinců, by se měla pedagogika také zaměřit na užívání informačních teorií. Nabádá k tomu v této souvislosti i moderní přístup kognitivní psychologie.

Klíčová slova: Vzdělávání, kognitivní psychologie, interpretace statistiky, výzkumné metody.

Keywords: Education, Cognitive Psychology, Interpretation of Statistical Methods, Research Methods.

*„Měřit je snadné – nesnadné je měřit
přesnost a spolehlivost měření.“*

Stanislav Komenda
Myšlenky, Olomouc 1997

*„Tradiční kvantitativně orientované
výzkumy představují v současné době
(ve světě i u nás) nejčastější
typ pedagogických výzkumů.“*

Jan Průcha
Pedagogická encyklopedie,
Portál, Praha 2009, str. 817

1. Rozumíme statistice?

Statistika se opírá o *matematické modely* náhodných pokusů. *Náhodný pokus* je děj, jehož výsledek není zcela jednoznačně určen jeho podmínkami. Teorie pravděpodobnosti, která právě studuje matematické modely náhodných pokusů, se nezabývá libovolnými pokusy, ale pouze těmi, které se vyznačují *statistickou stabilitou*. Ta je charakterizována *stabilitou relativních četností*.

Snadno je pozorovatelná u dějů, kde se náhoda uplatňuje ve své „čistě“ podobě, např. při hodech mincí nebo kostkou. Z uvedených dějů, které jsme

schopni snadno pochopit, vyplývá, že např. se zvyšováním počtu hodů se poměr počtu padnutí líců k počtu všech hodů blíží 0,5.

Pozn. Přesto významní matematici i v minulém století tento pokus opakovali, např. Karl Pearson házel 24 000krát mincí. Proč tak činil, nechť si laskavý čtenář uvědomí třeba na základě dalšího našeho výkladu.

Číslo, kolem něhož kolísají poměrné četnosti jevu „padnutí líce“ se nazývá pravděpodobností tohoto jevu. Můžeme ale takto realizovat odhad pravděpodobnosti jevu, že žák X. Y. vyřeší danou matematickou úlohu? Jakou *interpretaci* dáme „pravděpodobnosti vyřešení dané úlohy“ žákem X. Y.? Odhad této pravděpodobnosti můžeme za *určitých předpokladů* realizovat nepřímo, například konstrukcí sady „podobných“ úloh (v testu). Pokud žák například vyřešil sedm z deseti úloh, nabízí se odhad hledané pravděpodobnosti číslem 0,7 (to je ta relativní četnost). Můžeme toto číslo považovat za dobrý odhad pravděpodobnosti správného vyřešení některé úlohy? Vzhledem k velmi malému počtu a nesterpně obtížných úloh asi ne. Musíme tedy o postupu žáka při řešení úloh vyslovit několik předpokladů:

- a) žák úlohu buď vyřeší, nebo nevyřeší;
- b) žák řeší úlohy nezávisle na sobě (vyřešení některé z nich není ovlivněno řešením úloh předchozích, nezáleží na tom, v jakém pořadí žák úlohy řeší);
- c) žákova pravděpodobnost vyřešení kterékoliv úlohy je stejná a rovna p .

Zřejmě pak úroveň žáka (z hlediska schopnosti řešit úlohy daného testu) je možné posuzovat podle hodnoty pravděpodobnosti p . Za uvedených podmínek a), b), c) je počet m správně vyřešených úloh z daných $n = 10$ úloh náhodnou veličinou s pravděpodobnostní funkcí $P(m, p)$ ve tvaru

$$P(m; p) = \binom{10}{m} \cdot p^m \cdot (1 - p)^{10-m}. \quad (1)$$

Výsledek $m = 7$ je dosažen s pravděpodobností

$$P(7; p) = \binom{10}{7} \cdot p^7 \cdot (1 - p)^3 = 120p^7(1 - p)^3. \quad (2)$$

Ve vztahu (2) však p neznáme. Je přirozené předpokládat, že se výsledek $m = 7$ realizoval s maximální pravděpodobností (výsledky s menší pravděpodobností mají menší šanci se vyskytnout). Snadno zjistíme, že když $p = 0,7$, je hodnota výrazu (2) největší a rovna přibližně 0,267. Získali jsme tak jednu možnost interpretce pro $p = 0,7$.

Jak se může m měnit, když p bude rovno 0,7? Pravděpodobnosti $P(m; 0,7)$ pro $m = 0, 1, 2, \dots, 10$ jsou uvedeny v tabulce Tab. 1.

Z tabulky Tab. 1 vyplývá, že s pravděpodobností 0,95 za uvedených předpokladů při úrovni žáka $p = 0,7$ může žák vyřešit 4, 5, 6, 7, 8 nebo 9 úloh

Tabulka 1: Tabulka pravděpodobností $P(m; 0,7)$.

m	0	1	2	3	4	5	6	7	8	9	10
$P(m; 0,7)$	0	0	0	0,01	0,04	0,10	0,20	0,27	0,23	0,12	0,03

(nebo také 5, 6, 7, 8, 9 nebo 10 úloh). Jakou informaci o úrovni žáka nám pak dává počet $m = 7$ správně vyřešených úloh? Nebo, jinak řečeno, jaké může být p , když $m = 7$?

Testujeme proto hypotézu $H_0: p = 0,4$ proti alternativě $H: p \neq 0,4$ (případně alternativě $H^*: p < 0,4$).

Když by žák z 10 úloh vyřešil m úloh, pak za předpokladu H_0 je

$$P(m; p = 0,4) = \binom{10}{m} \cdot 0,4^m \cdot 0,6^{10-m}$$

Tabulka 2: Tabulka pravděpodobností $P(m; 0,4)$.

m	0	1	2	3	4	5	6	7	8	9	10
$P(m; 0,4)$	0	0,04	0,12	0,21	0,25	0,20	0,11	0,04	0,01	0,00	0,00

Z tabulky Tab. 2 vidíme, že s pravděpodobností 0,95 (případně 0,96) je možné, aby m bylo přirozeným číslem z intervalu $\langle 2; 7 \rangle$, (případně z intervalu $\langle 2; 10 \rangle$). Tedy při $m = 7$ se nezamítá hypotéza $H_0: p = 0,4$ na hladině významnosti 5 % (případně 4 %).

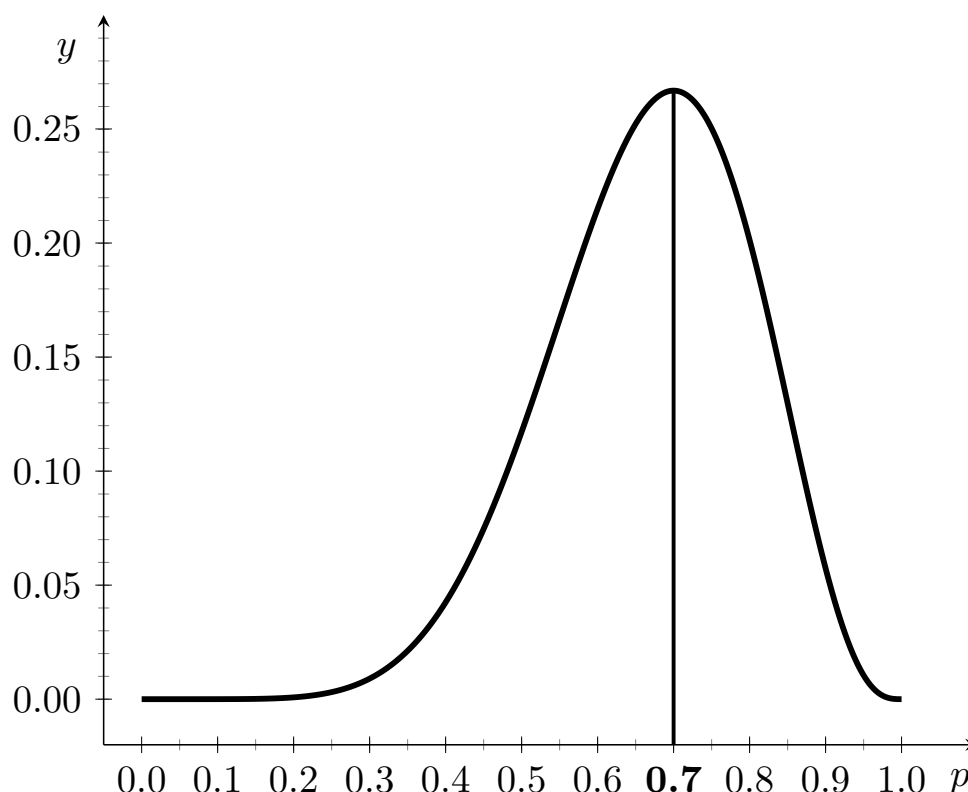
Podobně se ukáže, že se nezamítá na stejné hladině a za stejných podmínek ($m = 7$) ani hypotéza $H_0: p = 0,9$. Úroveň žáka, odhadovaná číslem p , tedy může být s pravděpodobností 0,95 v dosti širokém intervalu $\langle 0,4; 0,9 \rangle$. Uvedenou situaci popisuje také graf závislosti (2), viz Graf 1.

Informace o úrovni respondenta, odhadovaná z relativní četnosti 0,7 správně vyřešených úloh je v tomto případě velmi malá. Statistik by tu doporučil zvětšit počet řešených úloh aspoň na 30, což je z pedagogického hlediska nerealistické. Pro $n = 30$ a stejnou relativní četnost 0,7 je však 95% interval spolehlivosti (odhadovaný na základě hrubé normální aproximace) pro úroveň respondenta p jen o málo užší:

$$0,52 = 0,7 - 1,96 \cdot \frac{1}{2\sqrt{30}} \leq p \leq 0,7 + 1,96 \cdot \frac{1}{2\sqrt{30}} = 0,88.$$

Pozn. Odhad získaný přesnější aproximací ale nebude podstatně užší.

Zkusme proto dát počtu správně vyřešených úloh ještě jinou interpretaci.



Obrázek 1: Graf závislosti pravděpodobnosti $P(7; p)$ na p .

Nebudeme-li sledovat při řešení testu jen jediného žáka, ale např. 30 žáků, můžeme z jejich výsledků $x_1, x_2, \dots, x_{30}$ vytvořit uspořádanou posloupnost

$$x_{(1)} \leq x_{(2)} \leq \dots \leq x_{(30)}$$

a z té pak empirickou distribuční funkci $F_{30}(x)$:

$$F_{30}(x) = \begin{cases} 0; & x < x_{(1)} \\ \frac{i}{30}; & x \in \langle x_{(i)}, x_{(i+1)} \rangle \\ 1; & x > x_{(30)} \end{cases}$$

Z hodnoty $F_{30}(7)$ lze odhadnout úroveň našeho respondenta ve zkoumané populaci na škále $\langle 0; 1 \rangle$; čím je tato hodnota bližší 1, tím vyšší úroveň vykazuje. Úroveň respondenta je však zde poměřována nejen s úrovní testu (jako ovšem i v předchozích úvahách), ale i s úrovní vybrané skupiny 30 žáků. Používáme-li daný test vícekrát, pak jeho objektivní škálování bychom mohli odvodit z mnohonásobné aplikace aproximací empirických distribučních funkcí funkcí „limitní“ $F_{\infty}(x)$. To je ovšem jen teoretická úvaha.

Předchozí úvahy o odhadu úrovně respondenta pomocí pravděpodobnosti p a správné odpovědi byly podmíněny předpokladem nezávislosti tes-

tových položek. *Obsahovou analýzou* položek však většinou zjistíme, že testové položky jsou „*obsahově*“ závislé. Statistická závislost položek v případě malého počtu respondentů se obtížně prokazuje. Jak tedy máme postupovat, když chceme ohodnotit výkon respondenta a položky nebudou obsahově nebo statisticky nezávislé? Interpretace výkonu se může opírat o normu danou buď testem a určitou populací nebo danou jen testem (u tzv. výkonových $C - R$ testů).

V prvním případě pro daný test a určitý druh populace stanovíme teoretickou distribuční funkci $F_\infty(x) \in \langle 0; 1 \rangle$. Interval $\langle 0; 1 \rangle$ rozdělíme do k disjunktních intervalů $I_1, I_2, \dots, I_k$, kde

$$\bigcup_{i=1}^k I_i = \langle 0; 1 \rangle$$

a určíme intervaly $F^{-1}(I_i) = J_i, J_i \subset \langle 0; 10 \rangle, i = 1, 2, \dots, k$, kterými „pokryjeme“ množinu možných testových výsledků $\{0, 1, 2, \dots, 10\}$. Pak ovšem musíme ještě přistoupit k interpretacím žákovských výkonů dané populace z intervalů J_i . Interpretace musí být podloženy dalším studiem, jehož výsledkem je k tvrzení tvaru: „Jestliže žák dosáhl výkonu $V = v_i$ z intervalu J_i , pak by měl znát...“.

Ve druhém případě je apriorní normou test, konstruovaný na základě jistých požadavků, nezávisle na populaci. Takový test by ovšem měl projít i ověřením ve výběrovém souboru, který je hodně podobný tomu, ve kterém bude test používán. Od respondenta se vyžaduje splnit určitou (minimální) podmínku, tj. např. vyřešit aspoň 7 úloh z 10. Splnění této minimální podmínky by mělo být podmíněno určitým souborem znalostí. Interpretace „splnil podmínky testu“ zde je „žák má (minimálně následující znalosti...“.

Uvažujme ještě jednou test s celkem n položkami, který řeší respondent úrovně p . Pak při statistické nezávislosti a rovnocennosti testových položek můžeme předpokládat, že výsledek $X = m$ bude dosažen s pravděpodobností

$$P(m|p) = \binom{n}{m} p^m (1-p)^{n-m}. \quad (3)$$

Je přirozené předpokládat, že p není pevně dáno, ale je náhodnou veličinou. Zavedeme proto veličinu p , která bude mít beta rozdělení, konjugované s (3), s parametry α, β :

$$f(p) = K(\alpha, \beta) \cdot p^{\alpha-1} \cdot (1-p)^{\beta-1}, \quad p \in (0; 1), \quad (4)$$

kde $K(\alpha, \beta) = \frac{\Gamma(\alpha+\beta)}{\Gamma(\alpha) \cdot \Gamma(\beta)}$, Γ je gama funkce:

$$\Gamma(\lambda) = \int_0^{\infty} x^{\lambda-1} \cdot e^{-x} dx.$$

Aposteriorní rozdělení p , při kterém uvažujeme jak apriorní rozdělení (4), tak výsledek pozorování, získáme ze spojitě verze Bayesova vzorce ([12])

$$f(p|m) = \frac{P(m|p) \cdot f(p)}{\int_0^1 P(m|p) \cdot f(p) dp}. \quad (5)$$

Výsledná aposteriorní hustota $f(p|m)$ pak je

$$f(p|m) = K(m + \alpha; n + \beta - m) \cdot p^{m+\alpha-1} \cdot (1-p)^{n+\beta-m-1}. \quad (6)$$

Nyní máme dvě možnosti odhadu pravděpodobnosti p (úrovně respondent) v případě, že jeho testový výsledek X byl m :

a) pomocí *střední* hodnoty rozdělení (6):

$$\widehat{p}_1 = \frac{m + \alpha}{n + \alpha + \beta}, \quad (7)$$

b) pomocí *modu* rozdělení (6):

$$\widehat{p}_2 = \frac{m + \alpha - 1}{n + \alpha + \beta - 2}. \quad (8)$$

Hodnoty α a β určujeme na základě apriorní zkušenosti. Nemáme-li žádnou apriorní zkušenost, volíme $\alpha = \beta = 1$ (pak je apriorní rozdělení rovnoměrné na intervalu $(0; 1)$). Pak ze (7) pro náš diskutovaný příklad $n = 10$ a $m = 7$ dostáváme $\widehat{p}_1 = 0,67$ a z (8) za stejných předpokladů pro $\widehat{p}_2 = 0,70$; příslušný bayesovský 95% oboustranný interval spolehlivosti pro $\widehat{p}_1$ je $(0,40; 0,94)$. Pro testy s n položkami, u nichž máme k nabídnutých možností odpovědi, z nichž právě jedna je správná, můžeme za apriorní hodnoty α, β volit hodnoty, předpokládající jen hádání:

$$\alpha = \frac{n}{k} + 1; \quad \beta = n \left(1 - \frac{1}{k} \right) + 1.$$

Z této volby dostaneme při $n = 1$, $m = 7$, $k = 4$ odhady

$$\widehat{p}_1 = 0,43; \quad \widehat{p}_2 = 0,475,$$

a pak bayesovský 95% interval spolehlivosti pro $\widehat{p}_1$ je $(0,21; 0,65)$. Pozn. Ale i zde jsme vzhledem k malému rozsahu statistického souboru na hraně použitelnosti. Čtenáři dávám k úvaze interpretaci tohoto bayesovského intervalu spolehlivosti.

Bayesovské odhady umožňují využívat jak prvotní zkušenost, tak i zkušenost aktuální (experimentální) a hodí se pro odhady v souborech malého rozsahu. Pozn. Dávají obecně např. užší intervaly spolehlivosti než postupy klasické.

Pro volbu určitého matematického prostředku ke zpracování výsledků testu je důležitá informace o tom, zda testem chceme hodnotit např. výsledky výuky nebo také jednotlivce. Důležitá je také adekvátní představa pedagoga o mechanismu tvorby odpovědi na položky (tady statistik nemůže pedagoga zastoupit). Teorie odpovědi na položku (item response theory) z určitých matematických představ o charakteru rozložení odpovědí vychází. Ukážeme ještě, jaká je významnost celkových skóre v testu s k nabídnutými odpověďmi. Z našich úvah například vyplyne, že test s 10 položkami a 4 nabídnutými odpověďmi je pro spolehlivější diagnózu jedince bezcenný.

2. Významnost celkových skóre v testu s k nabídnutými odpověďmi

Uvažujme test s n položkami, z nichž každá má k možných odpovědí. Postavme se na nejextrémnější pozici respondenta, který „řeší“ test s nabídnutými odpověďmi jen hádáním. Respondent ví, že právě jedna z nabídnutých odpovědí je správná; pravděpodobnost uhodnutí správné odpovědi je pak $p = 1/k$ (a nesprávné $1 - 1/k$). Počet X všech takto uhodnutých odpovědí označených jako „správné“ má binomické rozložení:

$$P(X = x) = \binom{n}{x} \left(\frac{1}{k}\right)^x \left(1 - \frac{1}{k}\right)^{n-x}. \quad (9)$$

Porovnáme-li střední hodnotu tohoto rozložení $\mu_n = n \cdot 1/k$ se směrodatnou odchylkou $\sigma_n = \sqrt{n \cdot 1/k(1 - 1/k)}$, vidíme, že sice pro velké hodnoty n , tj. při velkém počtu položek, je směrodatná odchylka mnohem menší než střední hodnota:

$$\frac{\sigma_n}{\mu_n} \sim \frac{\sqrt{k}}{\sqrt{n}},$$

ale v reálném případě, kdy je například $k = 4$, $n = 25$, je $\mu_{25} = 6,25$ a $\sigma_{25} = 2,17$. Modus binomického rozložení x_0 by za uvedených okolností splňoval podmínku $\mu_n - (1 - 1/k) \leq x_0 \leq \mu_n + 1/k$, a tedy pro naše hodnoty $x_0 = 6$. K uvedeným charakteristikám by tedy konvergovaly při zvětšujícím se počtu respondentů jejich odhady z výběrových souborů jen „hádačích“ respondentů. Protože se zvětšujícím se n se zmenšuje poměr velikosti intervalu

možných hodnot okolo np vzhledem k np na stupnici výsledků testu pro „hádající“ respondenty, je pro diagnostikování hádání výhodnější test s větším počtem položek.

Upřesněme ještě trochu uvedená tvrzení a zeptejme se takto: Kolik položek musí mít náš test, aby u libovolného respondenta „řešícího“ test byla pravděpodobnost, že výsledek testu X je menší nebo roven x , rovna nejméně 0,9? Na tuto otázku snadněji odpovíme, když si napíšeme příslušnou nerovnost v následujícím tvaru:

$$P(X \leq \mu_n + \beta\sigma_n) \geq 0,9; \quad x = \mu_n + \beta\sigma_n. \quad (10)$$

Pro jednoduchost aproximujme binomické rozložení náhodné veličiny X s parametry n , $p = 1/k$ rozložením normálním s odhadovanou střední hodnotou $\mu = np$ a rozptylem $\sigma^2 = np(1 - p)$. Vzhledem k této aproximaci můžeme určit veličinu β z podmínky

$$P\left(\frac{X - \mu_n}{\sigma_n} \leq \beta\right) = 0,9. \quad (11)$$

Z tabulek hodnot distribuční funkce normálního rozložení dostaneme odhad $\beta = 1,28$. Dosazením tohoto odhadu do výrazu $X \leq \mu_n + \beta\sigma_n = x_n$ je možné pro každé n stanovit horní mez x_n výsledku testu s n otázkami při hádání z k nabídnutých možností, z nichž právě jedna je správná, viz tabulka 3 pro $k = 4$.

Tabulka 3: Tabulka horních mezí x_n výsledku testu s n položkami při hádání a za předpokladu (5).

n	1	2	3	4	5	6	7	8	9	10
x_n	0,8	1,2	1,7	2,2	2,5	2,9	3,2	3,6	3,9	4,2

n	11	12	13	14	15	16	17	18	19	20
x_n	4,5	4,9	5,2	5,5	5,9	6,2	6,5	6,8	7,2	7,4

n	21	22	23	24	25	26	27	28	29	30
x_n	7,7	8,0	8,4	8,7	9,0	9,3	9,6	9,9	10,2	10,5

Je-li tedy například $n = 25$, pak respondent, který v celém testu jen hádá, by v celém testu hodnoty větší než 9 „správně zodpovězených“ položek dosáhl nejvýše s pravděpodobností 0,1.

Nás však bude zajímat, kolik z X „správně zodpovězených“ položek mohl (s jistou pravděpodobností) respondent uhodnout a kolik jich skutečně nejméně musel vyřešit. Uvažujme takto. Jestliže zná správnou odpověď například na $m = 7$ položek, pak v případě zbývajících $25 - 7 = 18$ položek hádá (což je lepší strategie než neřešit, a tedy „nebodovat“). Potom ale (viz Tab. 3) hádáním může dosáhnout nejvýše $x_{18} = 7$ „správných“ položek. Tedy minimálně s pravděpodobností 0,9 může dosáhnout nejvýše $X = 13$ „správných“ položek. Podobně i v ostatních případech, jak ukazuje následující Tab. 4.

Uvědomme si, že tabulka 4 zatím udává, že za předpokladu m (nebo $m\%$) skutečně vyřešených položek může být s pravděpodobností nejméně 0,9 maximální zisk X . Je však možná i obrácená úvaha: jestliže byl dosažen zisk X , pak respondent musel vyřešit (nejméně s pravděpodobností 0,9) minimálně m ($m\%$) položek bez hádání. Obrácená úvaha je možná proto, že funkce $X = X(m)$, definovaná tabulkou 4, je monotónní. Můžeme tedy říci, že respondent, který získá 17 bodů, ovládá... (nejméně s pravděpodobností 0,9) minimálně 48–52 % učiva... zahrnutého v testu.

Tabulka 4: Tabulka minimálního počtu m položek, které musel vyřešit (bez hádání) respondent s testovým výsledkem X .

m	9	10	11	12	13	14	15	16	17	18	19	20	21	22	23	24	25
x_{25-m}	6	5	5	5	4	4	4	3	3	3	2	2	2	1	1	0	0
X	15	15	16	17	17	18	19	19	20	21	21	22	23	23	24	24	25
$m\%$	36	40	44	48	52	56	60	64	68	72	76	80	84	88	92	96	100

Naše úvahy vedou k tomu, že když budeme například v didaktických testech užívat položek s nabídnutými odpověďmi, tak je, vzhledem k požadavkům spolehlivosti, nutné, aby počet položek byl co největší ($n \geq 25$). To je však na hranici jejich praktické upotřebitelnosti. Konkrétněji, při $n = 5$ a $k = 4$, žák, který nezná správnou odpověď na určité dvě položky, s vysokou pravděpodobností (0,43) může mít užitím strategie hádání zisk větší než 3. Takový test není schopen odlišit výkon odpovídající skutečným znalostem od výkonu získaného náhodně.

Na první pohled se zdá, že by mohlo být jednodušší se soustředit (s ohledem na možnosti statistiky) místo na jedince, na popis určité výběrové populace. Uvažujme proto opět dotazník s 10 položkami. Hodnotíme-li odpovědi na dotazník jen dichotomicky (ano - ne, splnil – nesplnil apod.), což je nejjednodušší způsob, můžeme získat ale až $2^{10} = 1024$ různých odpovědí. Abychom tento velký počet možných odpovědí zmenšili, použijeme náhod-

nou veličinu m , která je počtem kladných odpovědí. Ta může nabývat jen 11 různých hodnot z množiny $\{0; 1; 2; \dots; 10\}$. Náhodná veličina m je diskrétní a z pokusu pro ni pak hledáme odpovídající typ teoretického rozdělení z empirického rozdělení relativních četností. Řešit tento úkol může mít smysl jen tehdy, je-li výběrový soubor dostatečně rozsáhlý a můžeme-li předpokládat určitou stabilitu odhadovaných relativních četností. To je ovšem velmi těžký úkol, zvláště v pedagogice, protože zkušenost říká, že variabilita řešení úloh respondenty je vysoká a má vliv na rozložení relativních četností. Přitom ale víme, že stejného výsledku (stejného m) může dosáhnout velké množství jedinců s rozdílnými rozloženími správných odpovědí. Jak pak můžeme interpretovat určitý výsledek, hodnocený číslem m ?

Proto například kognitivní psychologové jsou v používání statistických metod zdrženliví (Eysenck, Keane, 2008). Z hlediska strategií řešení úloh se doporučuje studovat informační procesy místo „nejistých“ statistických analýz testů. Informační přístup totiž musí respektovat i širší vazby experimentu na prostředí, historii příjmu informace i předpokládaný „model“ ukládání informace, např. s ohledem na informace „dříve“ uložené. Nová informace při učení nemůže být uložena jako nezávislý soubor svých prvků, ale jako doplněk již dříve uložené struktury, do které se „vnoří“ tak, že s ní vhodně kooperuje, pojmově i sémanticky je kompatibilní s informací původní.

3. Testy statistických hypotéz mají také svá omezení

Při testování statistických hypotéz by se měly uvažovat obě chyby, chyba prvního druhu s hodnotou pravděpodobnosti α i chyba druhého druhu s hodnotou pravděpodobnosti β . Praxe je ovšem taková, že se uvažuje „jen“ chyba prvního druhu v hodnotách pravděpodobnosti 0,05 nebo 0,01, případně ještě 0,1 (platí to pro humanitní obory). Přitom se předpokládá, že za vhodných podmínek, tj. při dostatečném rozsahu zkoumaného souboru a dostatečné „vzdálenosti“ hypotézy nulové od hypotézy alternativní, pravděpodobnost chyby druhého druhu β nepřekročí nepřijatelnou mez (nepřijatelná je např. $\beta > 0,2$). Porušíme-li však některý z uvedených dvou předpokladů, hodnota β se nám 15) může zvětšit nekontrolovaně nad únosnou mez a znehodnotit tak celou proceduru testování aniž to zjistíme. Ukážeme si to na jednoduchém příkladě.

Uvažujme znovu náš dotazník s $n = 10$ zcela rovnocennými a statisticky nezávislými položkami a testujme hypotézu $H_0: p = 0,7$ proti alternativě $H: p < 0,7$ (tato formulace alternativní hypotézy je rozumná zvláště u výkonových testů a usnadňuje interpretaci).

Kritickou oblastí W (viz Tab. 1) je za předpokladu $\alpha = 0,05$ (nebo $0,10$) množina $W_{0,05} = \{m; m \in \{0, 1, 2, 3, 4\}\}$, a tedy

$$\alpha = P(m \in W_{0,05} | H_0) \approx 0,047 \leq 0,05. \quad (12)$$

Volba kritické oblasti nám určuje podmínku zamítnutí nulové hypotézy (je to $m \leq 4$ na hladině významnosti $\alpha = 0,05$). Zamítneme-li nulovou hypotézu, přijímáme hypotézu alternativní. Jinak ponecháváme platnost nulové hypotézy. Pozn. Platnost nulové hypotézy tedy nepotvrzujeme experimentem, nulovou hypotézu volíme na základě výzkumného programu, není to záležitost jen statistická, je to záležitost zde i pedagogických zkušeností, bez apriorních informací tedy k testování hypotéz přistupovat nemůžeme.

Hodnota β je závislá jak na hodnotě α , tak i na alternativní hypotéze. Volme proto alternativní hypotézu $H: p = 0,6$. Pak $\beta = \beta(\alpha = 0,05; H: p = 0,6) = P(m \notin W_{0,05} | H) = 0,83$. (Zmenšujeme-li α , roste nám β . Pro $\alpha = 0,10$ dostaneme v našem případě hodnotu stejnou, $\beta = 0,83$.) Takový test má špatnou vypovídající hodnotu, v této formulaci ho nemůžeme používat.

Poznamenejme ještě, že jednostranný interval spolehlivosti pro pravděpodobnost p , vyplývající z našeho předpokládaného výsledku $m = 7$ je pro $\alpha = 0,05$ roven $0,39 \leq p < 1$ a pro $\alpha = 0,10$ je $0,44 \leq p \leq 1$. (V tomto případě je však $\beta = \beta(\alpha = 0,00001; H: p = 0,6) > 0,9$.)

Pro $\alpha = 0,00001$ dostaneme „podmínku“, která p omezuje jen na interval $0,10 \leq p < 1$; z výsledku $m = 7$ v tom případě získáme velmi málo informace o pravděpodobnosti p .

Z uvedených úvah vyplývá, že některá použití statistických metod kladou na test požadavky, které jsou v rozporu s požadavky dobré informovanosti o žákově výkonu (zde to byl rozsah testu, rovnocennost položek a jejich statistická nezávislost).

Podobně jako o podmínkách testování statistických hypotéz, je třeba uvažovat i o interpretacích různých statistických ukazatelů. Příkladem jsou korelační koeficienty (jako uvažované míry závislosti, případně síly vztahu či shody) a různé ukazatele „spolehlivosti“ (Cronbachova alfa, KR_{24} , Spearman-Brownův vzorec, ...). Ukazatele ve svých hodnotách kombinují více vlastností měřeného než je „jen“ určitá závislost nebo spolehlivost (Řehák, 1998, 51–60; Zvára, 2002, 13–20; Půlpán, 2004, 40–43; Hendl, 2006, 297–335).

„... při důkladnějším rozboru libovolné metody teorie pravděpodobnosti zjistíme, že její použití vyžaduje častokrát nevšední úsilí, chceme-li získat věrohodné výsledky.“

V. N. Tutubalin ([12], str. 27)

4. Závěr

Autor (statistik) chtěl ukázat, že i elementární aplikace statistických metod vyžadují konzultace se statistikem. Upozorňuje na to, že učitel ve třídě (ne výzkumný pracovník) má daleko lepší možnost získat informace o znalostech informace o znalostech žáků než je „testování“. Informačně nejúčinnější je metoda dialogu se žáky a pozorování (např. jak děti reagují na otázky učitele).

Chce-li výzkumný pracovník použít statistických metod, pak je třeba již při přípravě výzkumného plánu s tím počítat. Je třeba si uvědomit, že měřícími prostředky jsou na jedné straně test a na druhé straně statistická technika, která „normu“ spoluvytváří. Test vytváří učitel, statistickou techniku si jen volí. Proto právě test by měl být vytvářen s největší pečlivostí. Větší zájem se většinou soustřeďuje jen na statistickou techniku (a to ještě jen opravdu na „techniku“, nikoliv smysl). A to je špatné. Užití statistických metod je třeba hlouběji analyzovat; výsledkem statistické analýzy musí být zasvěcený komentář interpretace, ne jen „proklikaná“ čísla z počítače. Přesto statistice nepřísluší poslední slovo.

Literatura

- [1] Beneš, P.; Janoušek, R.; Novotný, M.: Hodnocení obtížnosti textu středoškolských učebnic, *Pedagogika*, roč. LIX, 2009, s. 291–297.
- [2] Eysenck, M. W.; Keane, M. T.: *Kognitivní psychologie*, Academia, Praha 2008, 748 s., ISBN 978-80-200-1559-4.
- [3] Hendl, J.: *Přehled statistických metod zpracování dat*, Portál, Praha 2006, ISBN 80-7367-123-9, 583 s.
- [4] Hladík, J.: Vztah kognitivní a afektivní složky multikulturních kompetencí, *Pedagogika*, roč. LXI, 2011, s. 53–65.
- [5] Kuřina, F.: Tři pokusy řešit neřešitelné, *Pedagogika*, roč. LXI, s. 5–12.
- [6] Lindquist, E. F.: *Statistická analýza v pedagogickém výzkumu*, SPN, Praha 1967.
- [7] Průcha, J. (ed): *Pedagogická encyklopedie*, Portál, Praha, 2009, 935 s.
- [8] Půlpán, Z.: *Odhad informace z dat vágní povahy*. Academia, Praha, 2012.
- [9] Půlpán, Z.: *K problematice hledání podstatného v humanitních vědách*, Academia, Praha 2001, 135 stran, s. 124–126. ISBN 80-200-0855-1.
- [10] Půlpán, Z.: *K problematice zpracování empirických šetření v humanitních vědách*, Academia, Praha 2004, 181 stran. ISBN 80-200-1221-4.

- [11] Řehák, J.: Kvalita dat I. Klasický model měření reality a jeho praktický aplikační význam, *Sociologický časopis*, 34 (1), 1998, 61–60.
- [12] Tutubalin, V. N.: *Teorie pravděpodobnosti*, SNTL, Praha 1978, 216 s., DT 519.2.
- [13] Zvára, K.: Měření reability aneb bacha na Cronbacha, *Informační Bulletin ČStS*, č. 2, roč. 13, 2002, 13–20, ISSN 1210-8022.
- [14] Záhorec, J.: Konceptné a metodické východiská hodnotenia stavu vyučovania informatiky na vyššom sekundárnom stupni vzdelávania v Slovenskej Republike, *Obzory matematiky, fyziky a informatiky*, 1/2011 (40), s. 25–30, EV 915/08. ISSN 1335-4981.

HUDEBNÍ

Jméno: Prouha

Kontrolní práce

Co si pamatuješ o Bedřichu Smetanovi? Doplň:

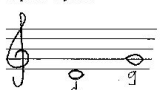
- Bedřich Smetana byl umělcem skládatel skladby
- Narodil se v Lydomyšli
- Čím byl jeho otec? Kde pracoval? na pivovaru Sládkem
- Na které hudební nástroje hrál? na housle a na klavír
- V které cizí zemi několik let působil? +
- Pro slavnostní otevření Národního divadla složil operu +
- Jaká nemoc postihla Smetanu 10 let před smrtí? ochluchl
- Přestal skládat svá díla, když onemocněl? ne, fort skládal
- Jak se jmenuje cizí fonických básní, v kterém opěvuje naši vlast? +
- Vyjmenuji Prouha

1. Poslechni si skladbu a napiš jméno Bedřich Smetana

2. Kolik základních tónů tvoří hudbu? Vyjmenuj je všechny: C, D, E, F, G, A, B

3. Která stupnice začíná i končí t +

4. Zapiš názvy not:



5. Co k sobě patří? Spojuj:



4. dubna

1. Čím měříme délku? metr

2. Čím měříme hmotnost? kg

3. Čím měříme teplotu? °C

4. Čím měříme rychlost? m/s

5. Čím měříme energii? J

6. Čím měříme výkon? W

7. Čím měříme tlak? Pa

8. Čím měříme sílu? N

9. Čím měříme práci? J

10. Čím měříme čas? s

11. Čím měříme délku? m

12. Čím měříme hmotnost? kg

13. Čím měříme teplotu? °C

14. Čím měříme rychlost? m/s

15. Čím měříme energii? J

16. Čím měříme výkon? W

17. Čím měříme tlak? Pa

18. Čím měříme sílu? N

19. Čím měříme práci? J

20. Čím měříme čas? s

31. března

1. Co tvoří neuvěřitelnou přírodu? +

2. Co jsou minerály a horniny? +

3. Co je magma? +

4. Jaké síly, horniny a minerály pade? +

5. Jaké síly, horniny a minerály pade? +

6. Jaké síly, horniny a minerály pade? +

7. Jaké síly, horniny a minerály pade? +

8. Jaké síly, horniny a minerály pade? +

9. Jaké síly, horniny a minerály pade? +

10. Jaké síly, horniny a minerály pade? +

11. Jaké síly, horniny a minerály pade? +

12. Jaké síly, horniny a minerály pade? +

13. Jaké síly, horniny a minerály pade? +

14. Jaké síly, horniny a minerály pade? +

15. Jaké síly, horniny a minerály pade? +

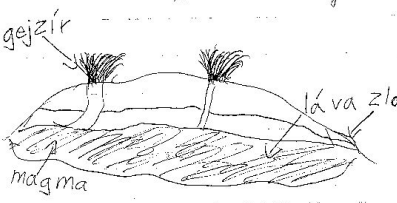
16. Jaké síly, horniny a minerály pade? +

17. Jaké síly, horniny a minerály pade? +

18. Jaké síly, horniny a minerály pade? +

19. Jaké síly, horniny a minerály pade? +

20. Jaké síly, horniny a minerály pade? +



3

MOŽNOSTI APLIKÁCIE ZHLUKOVEJ ANALÝZY V MANAŽÉRSKÝCH PODNIKOVÝCH ANALÝZACH

POSSIBLE APPLICATIONS OF THE CLUSTER ANALYSIS IN THE MANAGERIAL BUSINESS ANALYSIS

Ladislav Mura

Adresa: Ing. et Bc. Ladislav Mura, PhD., Ekonomická
fakulta, Univerzita J. Selyeho v Komárne,
Bratislavská 3322, 945 01 Komárno

E-mail: ladislav.mura@gmail.com

Abstrakt: Zhluková analýza sa používa v rôznych oblastiach ekonomických vied. Štatistická analýza formou zhľukovania je súborom metód, ktoré umožňujú hľadať v empirických údajoch zoskupenia podobných objektov. Použitie metód zhľukovej analýzy je obzvlášť vhodné najmä v štádiu explorácie problému. Pri analýze a komparácii rôznych ukazovateľov úspešnosti podnikateľských subjektov sa okrem iných štatistických metód často využíva práve zhľuková analýza. Cieľom článku je poukázať na možnosti aplikácie zhľukovej analýzy a jej využitia v manažérskych podnikových analýzach.

Kľúčové slová: Podniková analýza, zhľuková analýza, štatistické metódy.

Abstract: The cluster analysis is used in various fields of economics. The clustering form of statistical analysis is a set of methods that allow finding in the empirical data groups of similar objects. Using this cluster analysis method is particularly useful in the stage of exploration of the problem. By the analysis and comparison of various business success indicators is between other statistical methods often cluster analysis used. The aim of this paper is to point out possible applications of the cluster analysis and its use in the managerial business analysis.

Keywords: Business Analysis, Cluster Analysis, Statistical Methods.

1. Úvod

Zhľuková analýza sa zaoberá tým, ako by mali byť objekty teda štatistické jednotky zaradené do skupín tak, aby bola čo najväčšia podobnosť v rámci skupín a čo najväčšia rozdielnosť medzi skupinami. Zhľuková analýza sa používa napr. pri segmentácii trhu, pričom klasifikácia spotrebiteľov je založená

na kombinácii viacerých premenných. Premennými, teda segmentačnými kritériami môžu byť: pohlavie, vek, vzdelanie, životný štýl, náboženstvo, skúsenosti s produktom, veľkosť spotreby, frekvencia spotreby a pod. [5]

Štatistická analýza formou zhľukovania je súborom metód, ktoré umožňujú hľadať v empirických údajoch zoskupenia podobných objektov. Použitie metód zhľukovej analýzy je obzvlášť vhodné najmä v štádiu explorácie problému. Na rozdiel od faktorovej analýzy sa pri uplatnení zhľukovej analýzy zväčša nevenuje pozornosť základným dimenziám popisu sledovaných javov (premenným), ale základným typom sledovaných javov ako takých (objektom), ich podobnosti a nepodobnosti (hoci je niekedy problematické vymedziť, čo vlastne podobnosť je). Zhľuková analýza sa teda využíva na hľadanie typov objektov, ktoré sa vyznačujú špecifickými charakteristikami, tzn. slúži ako primeraný nástroj generovania predpokladov o klasifikácii objektov. [4]

Pri analýze a komparácii rôznych ukazovateľov úspešnosti podnikateľských subjektov sa okrem iných štatistických metód často využíva práve zhľuková analýza. Pomocou tejto metódy je možné podniky začleniť do zhľukov tak, aby v rámci skupiny podnikateľských subjektov bola čo najväčšia podobnosť. Možnosť aplikácie zhľukovej analýzy v podnikových analýzach na úrovni rôznych stupňov manažmentu je teda mnoho.

Pri zohľadnení iba jednej premennej je nájdenie zhľukov jednoduché: hodnoty premennej sa nanesú na číselnú os a zhľuky sa identifikujú vizuálne (napr. podľa veku nájdeme v súbore 2 skupiny respondentov: jednu okolo napr. 15 rokov a druhú okolo 40 rokov). Podobne použitím X-Y grafu možno jednoducho identifikovať zhľuky pri zohľadnení 2 premenných. V priestore (trojrozmerné usporiadanie) sa pomocou interaktívneho X-Y-Z grafu tiež dajú nájsť zhľuky vizuálne. Vizuálne identifikovať zhľuky pri zohľadnení viac ako 3 premenných súčasne sa však už nedá. Práve vtedy je možné použiť zhľukovú analýzu. Zhľuková analýza zahŕňa množstvo metód. Rozlišujeme dve základné skupiny:

- Hierarchické zhľukovacie metódy.
- Nehierarchické zhľukovacie metódy.

Hierarchické zhľukovacie metódy vychádzajú z jednotlivých objektov, ktoré reprezentujú zhľuky. Ich spájaním sa v každom kroku počet zhľukov postupne zmenšuje až sa nakoniec všetky zhľuky spoja do jedného celku. Hierarchické metódy vedú k hierarchickej (stromovej) štruktúre, ktorá sa graficky zobrazuje ako stromový diagram (dendrogram). Stromové zhľukovacie metódy začínajú výpočtom vzdialenosti medzi objektmi. [2] Euklidovská vzdialenosť medzi objektmi i a j s n charakteristikami (premennými) sa vypočíta:

$$d_{ij} = \sqrt{\sum_{k=1}^n (x_{ik} - x_{jk})^2}$$

Alternatívnu vzdialenosť predstavuje vzdialenosť Manhattan (City-block):

$$d_{ij} = \sum_{k=1}^n |x_{ik} - x_{jk}|$$

Euklidovská vzdialenosť vyjadruje vzdušnú vzdialenosť medzi dvoma objektmi a Manhattanovská vzdialenosť najkratšiu vzdialenosť, ktorú musí objekt prejsť, aby sa dostal z jedného miesta na druhé. Výhoda vzdialenosti Manhattan spočíva v znížení dopadu extrémnych prípadov na výsledky. Existujú ešte viaceré iné typy vzdialeností, ktoré sa používajú napr. pri kategorických premenných.

Nehierarchické zhukovacie metódy nevytvárajú stromovú štruktúru. Najznámejšia nehierarchická zhukovacia metóda je metóda k -priemerov. Táto metóda sa vyznačuje tým, že vyprodukuje presne k -zhlukov tak, aby bol vnútroskupinový súčet štvorcov minimálny. [2] Najvhodnejšia je na formovanie malého počtu zhlukov z veľkého počtu pozorovaní. Vyžaduje však intervalové premenné bez extrémnych hodnôt. Nominálne premenné sa dajú použiť, ale môžu spôsobovať problémy. Užitočnou metódou je neurčité zhukovanie (fuzzy zhukovanie), ktoré na rozdiel od ostatných zhukovacích metód, umožňuje čiastočné zaradenie objektu do viacerých zhlukov a to pomocou pravdepodobnosti. Cieľom je zabrániť skresleniu zhukovania kvôli prítomnosti nezaraditeľných objektov. Takéto individuum sa nepriradí ku žiadnemu zhuku (od každého sa príliš odlišuje), ale priradia sa mu pravdepodobnosti s ktorými sa bude nachádzať v jednotlivých zhukoch. Metóda sa často používa pri odhaľovaní podvodov v rôznych oblastiach. Napr. v bankovníctve sa bez vopred formulovanej definície podozrivej operácie z miliónov operácií klientov identifikuje pár desiatok takých, ktoré sa od zvyšných (zoskupených do niekoľkých zhlukov) pri použití viacerých premenných (napr. obrat, typ operácie, konštantný symbol, čas od zadania po jej splatnosť atď.) výrazne odlišujú.

Analýza pomocou zhukovania je technika pre zgrupovanie dát a objavovanie štruktúr v dátach. Najpoužívanejšou aplikáciou zhukovacích metód je deľba dátových množín na zhluky, alebo triedy, kde sú podobné dáta priradené do spoločného zhuku, a kde by rozdielne dáta mali patriť do rozdielnych zhlukov. Teda úlohou metód zhukovania je vhodne číselne vyjadriť vlastnosti objektov a zoskupiť podobné objekty do zhlukov. V reálnych aplikáciách nie

je veľmi často ostrá hranica medzi zhlukmi, teda fuzzy zhlukovanie sa často lepšie hodí pre dáta. Funkcie príslušnosti medzi nulou a jednotkou sú používané vo fuzzy zhlukovaní namiesto „crisp“ (pevná hodnota) pridelení dát pre zhluky. [7] Klasické typy zhlukovania pracujú s objektami, ktorých popis v sebe neobsahuje neurčitosť, vágnosť. Existujú objekty, ktoré sa nedajú popísať inak, než s určitou mierou vágnosti. Na tieto objekty nie je možné použiť klasické zhlukovacie metódy.

Fuzzy zhlukovanie môže byť použité ako stratégia učenia bez učiteľa v nahliadnutí na zgrupovanie dát. Ale fuzzy zhlukovanie je tiež veľmi užitočné pre zostavovanie fuzzy if – then pravidiel pre dáta. Štruktúra pravidiel závisí na požadovanej aplikácii. Pre chybnú diagnózu a iné klasifikačné úlohy, sa pravidlá zameriavajú na rozhodovanie do ktorej triedy v konečnej množine tried (ako v poriadku/prijateľné/chybné) by zadaný údaj mal byť zaradený. V systéme identifikácie, alebo aproximácie funkcií pravidlá popisujú väčšinou spojené prepojenie medzi rôznymi premennými (ako pri fuzzy riadení).

Ďalšia oblasť aplikácií analýzy fuzzy zhlukovania je analýza obrazu a rozpoznávania. Segmentácia a detekcia špeciálnych geometrických tvarov, ako sú kruhy a elipsy, môže byť dosiahnutá pomocou takzvaných zapúzdrovacích zhlukovacích algoritmov.

2. Materiál a metódy

Cieľom článku je poukázať na možnosti aplikácie zhlukovej analýzy a jej využitia v manažérskych podnikových analýzach. Parciálnymi cieľmi sú: objasniť vlastnosti hierarchických a nehierarchických zhlukovacích metód, objasniť možnosti fuzzy zhlukovej analýzy a ukázať praktickú aplikáciu zhlukovej analýzy s využitím počítačového programu. Zo softvérových produktov sme pre názornú ukážku využili program Matlab. Pri spracovaní príspevku sme využili sekundárne literárne pramene, predovšetkým príspevky vo vedeckých časopisoch a zborníkoch z medzinárodných konferencií. Z vedeckých metód boli využité logicko-poznávacie metódy.

3. Výsledky a diskusia

Je zrejmé, že ako ktorákoľvek metóda manažérskej podnikovej analýzy i zhluková analýza má okrem vyššie uvedených pozitív aj svoje nedostatky. V nasledujúcej časti článku vymedzíme niektoré z nich.

3.1. Negatíva klasického zhlukovania

Objekty, ktoré zhlukujeme pomocou klasického zhlukovania, sú popísané pomocou príznakov. Príznačky objektov môžu nadobúdať 3 základné typy:

- **Kvantitatívne:** hodnota znaku vyjadruje množstvo. Najčastejšie je znak tohto typu popísaný číslom patriacim do spočítateľnej, alebo nespočítateľnej množiny.
- **Kvalitatívne:** hodnota znaku je vybraná z konečnej množiny možných stavov. Hodnoty týchto znakov môžu byť aj disjunktné intervaly.
- **Binárne:** hodnota znaku je vybraná z dvojprvkovej množiny, kde jeden prvok znamená, že objekt nemá požadovanú vlastnosť a druhý znamená, že danú vlastnosť má. Najčastejšie sa táto dvojprvková množina definuje v tvare $\{0, 1\}$.

Popis objektu môže obsahovať aj kombináciu menovaných znakov. Vyššie uvedené typy príznakov objektu predpokladajú, že pre daný objekt je potrebné vybrať jednu konkrétnu hodnotu. Pokiaľ vybranému príznaku zvoleného objektu priradíme hodnotu, potom daný objekt vo vybranom príznaku už nemôže nadobúdať ďalšiu hodnotu. Príznačky objektu a tým i zvolený objekt sú presne definované a nad týmito objektmi je definované zhlukovanie.

Takéto objekty budeme nazývať klasické objekty a zhlukovanie nad nimi zhlukovanie klasických objektov. V praxi sa často vyskytujú objekty, ktoré nie je možné popísať vyššie uvedenými typmi znakov. Takýto objekt obsahuje znak, ktorého hodnoty nie je možné presne definovať, t.j. existuje znak objektu, ktorý môže súčasne obsahovať viac hodnôt alebo pre daný znak existuje „neurčitosť“, „vágnosť“ vo vyjadrení hodnôt tohoto znaku. Potom klasické zhlukovanie nie je možné aplikovať priamo na takéto typy objektov. V klasickom zhlukovaní sa tento prípad rieši tým, že danej „vágnej“ hodnote znaku priradíme hodnotu, ktorá najlepšie vystihuje daný znak objektu. Vyberie sa tzv. „hlavná hodnota“. Tým, že z celej „vágnej“ hodnote vyberieme iba túto „hlavnú hodnotu“, alebo zo všetkých možných hodnôt vyberieme iba jednu hodnotu, strácame informáciu, ktorá je obsiahnutá vo „vágnosti“ a ktorá môže mať na výsledok zhlukovania vplyv.

3.2. Zovšeobecnenie typov znakov a objektov

Bolo by vhodné zaviesť také typy príznakov objektov a zhlukovaní nad týmito objektmi, ktoré by brali do úvahy aj túto „vágnosť“. Je teda užitočné

„vágnosť“ použiť pri popise hodnôt príznaku a tým ju zapojiť do zhlukovacieho procesu. Existuje viac možností, ako „vágnosť“ popísať. Pretože človek dokáže triediť i objekty, ktorých popis znakov je „vágny“, je vhodné vybrať takýto popis „vágnosti“ hodnôt znakov, ktorý sa najviac približuje ľudskému uvažovaniu. Ako sa už ukázalo v podobných oblastiach, kde sa snažíme nahradiť ľudský vplyv na riešení problémov je vhodné túto „vágnosť“ popísať pomocou fuzzy množín. Takto definované hodnoty príznakov objektov, okrem popisu vágnosti, zovšeobecnenie hodnoty príznakov klasických objektov.

Objekty popísané fuzzy množinami budeme nazývať fuzzy objekty. Tieto fuzzy objekty môžu byť popísané tiež pomocou jazykových hodnôt predom definovaných jazykových premenných. Pri zhlukovaní sa využíva zovšeobecnenie štandardných zhlukovacích metód, kde sa miesto podobnosti (nepodobnosti) objektov zavádza pojem fuzzy podobnosti (fuzzy nepodobnosti), fuzzy objektov a s využitím tohoto pojmu je definované zhlukovanie nad fuzzy objektami.

3.3. Teória zhlukovania

Intuitívne chápeme, že dva objekty sú si podobné, keď majú niektoré vlastnosti rovnaké. Čím viac rovnakých vlastností majú, tým sú si podobnejšie. Z hľadiska množinovej teórie podobnosti môžeme objekt P_1 charakterizovať istou množinovou vlastnosťou $\{V_1\}$ a objekt P_2 množinou $\{V_2\}$.

Vzájomnú podobnosť môžeme určiť pomocou koeficientu podobnosti, ktorý definujeme nasledovne:

$$p(P_1, P_2) = \frac{|P_1 \cap P_2|}{|P_1 \cup P_2|} = \frac{|\sum V_1 \cap \sum V_2|}{|\sum V_1 \cup \sum V_2|}$$

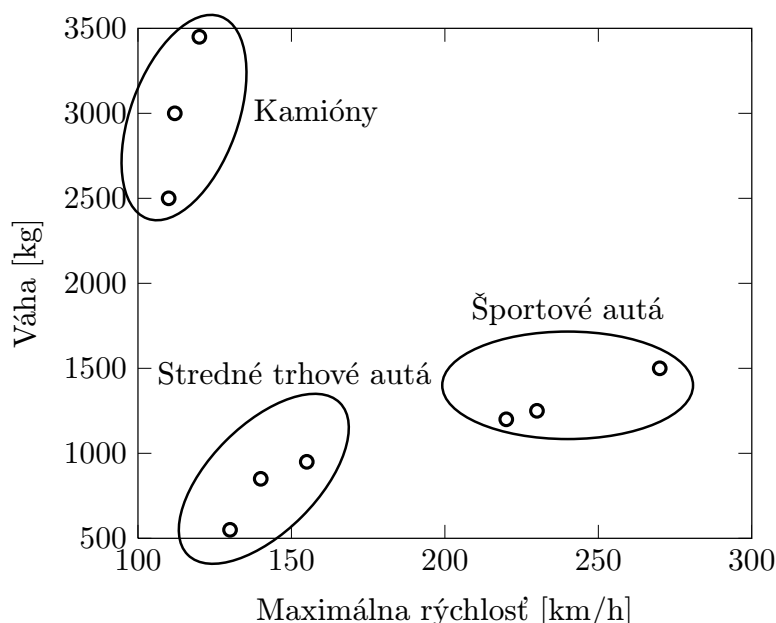
t.j. podielom mohutnosti prienikov a mohutnosti zjednotenia množín. Podobne môžeme využiť aj koeficient rozdielu definovaný ako:

$$d(P_1, P_2) = 1 - p(P_1, P_2)$$

Zhluk môžeme definovať ako množinu objektov, pri ktorých je hodnota koeficientu podobnosti vyššia, ako je hodnota istého prahu.

Vo všeobecnosti sa analýza dát týka objektov, ktoré sú popísané pomocou príznakov. Uvážme príklad, ktorý obsahuje dáta niektorých dopravných prostriedkov. Každý riadok takejto tabuľky popisuje objekt, a každý stĺpec popisuje príznak. Príznak môže byť považovaný za súbor hodnôt, z ktorého vystupujú skutočné hodnoty v danom zobrazenom stĺpci. Príznačky z osí abstraktného príznakového priestoru, v ktorom je každý objekt reprezentovaný

pomocou bodu. Ak aplikujeme naše teoretické vedomosti na podmienky priemyselného podniku, môžeme ako príklad uviesť napríklad zhlučky motorových vozidiel (obrázok 1).



Obr. 1: Zhlučky motorových vozidiel v príznakovom priestore.

Zdroj: Výstup zo softvéru Matlab.

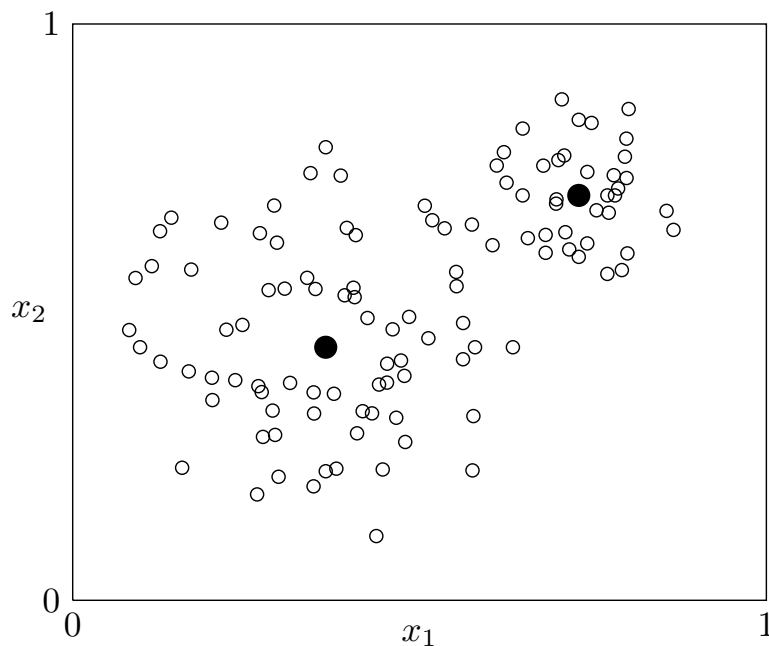
Napríklad, v štvorrozmernom súradnicovom systéme ohraničeného pomocou osí *maximálna rýchlosť*, *odpor vzduchu*, *farba* a *váha*, dopravný prostriedok V_1 je reprezentovaný pomocou bodu $(x_1, x_2, x_3, x_4) = (220; \text{červená}; 0,30; 1300)$. Farebná os je odlišná od ostatných osí pretože jej hodnoty sú vykresľované z diskkrétnej domény skôr ako z domény reálnych čísiel. Za účelom udržania myšlienky príznakového priestoru, musí byť diskrétna doména z tohto dôvodu usporiadaná.

Príznačky môžu byť vybrané priamo, alebo môžu byť generované pomocou vhodnej kombinácie príznakov. Druhá možnosť je nutná v prípade širších číselných príznakov, ktoré je nutné zredukovať na menší počet príznakov. Objekty sú fuzzy, ak jeden, alebo viac príznakov sú popísané fuzzy výrazmi. Ako príklad zoberme dopravný prostriedok s „veľmi rýchlym“ motorom auta, radšej ako maximálnu rýchlosť rovnú nejakému „crisp“ číslu. Zhlučky sú fuzzy, keď každý objekt je asociovaný so stupňom príslušnosti skôr ako s „crisp“ príslušnosťou. Príkladom je zhluček „športových áut“, konkrétne auto bude príslušnosť určitého stupňa závisiaceho na jeho maximálnej rýchlosti, odporu vzduchu a váhe.

Hard c -means algoritmus sa snaží lokalizovať zhluky v mnohorozmernom príznakovom priestore. Cieľom je priradiť každý bod v príznakovom priestore do konkrétneho zhľuku. Základný prístup je nasledovný [1]:

- Manuálne označiť „ c “ centrá zhľukov pre algoritmus, jedno centrum pre každý zhľuk, ktorý hľadáme. Toto požaduje predošlé informácie z vonkajšieho sveta o počte rozdielnych zhľukov do ktorých budú body rozdelené, takže algoritmus patrí do triedy algoritmov s učiteľom – supervízorom.
- Každý bod je priradený do zhľuku podľa toho, ku ktorému centru zhľuku je najbližšie.
- Nové centrum zhľuku je vypočítané pre každú triedu vzatím priemer-
ných hodnôt koordinátov bodov ktoré sú mu priradené.
- Ak neskončí v zhode s nejakou zastavovacou podmienkou, je potrebné ísť na krok 2.

Môžu byť pridané aj niektoré doplnkové pravidlá, pre odstránenie potreby poznať presne koľko sa tam nachádza zhľukov. Pravidlá dovoľujú susedným zhľukom spájanie a zhľuky ktoré majú široké štandardné odchýlky v koordinátoch dovoľujú delenie. Na obrázku 2 môžeme sledovať príklad dvoch zhľukov. Centrá zhľukov sú označené pomocou plných kruhov.



Obr. 2: Príklad s dvomi zhľukmi.
Inšpirácie: Jang & Gulley (1995).

V ďalšej časti článku uvažujme nad dátovými bodmi z obrázku. Zhlu-
vací algoritmus nájde jedno centrum v ľavom dolnom rohu a druhé v pravom
hornom rohu. Počiatočné centrá sú viac menej v strede obrázku a počas
iterácií sa pohybujú k ich konečným pozíciám. Každý bod patrí do jednej,
alebo druhej triedy, takže zhlu-
ky sú „crisp“, teda ostré. Vo všeobecnosti, dáta
príznakov musia byť normalizované správne pre vzdialenostné porovnávanie,
aby to pracovalo správne.

Hard c -means algoritmus má päť krokov. Ide o nasledovné:

1. Inicializácia centier zhlu-
kov c_i ($i = 1, 2, \dots, c$). To je typicky dosiahnuté
náhodne selektovaním c bodov z bodov dát.
2. Určenie matice susednosti M prostredníctvom:

$$m_{ik} = \begin{cases} 1 & \text{keď } \|u_k - c_i\|^2 \leq \|u_k - c_j\|^2, \text{ pre každé } j \neq i \\ 0 & \text{inak} \end{cases}$$

3. Výpočet funkcie vhodnosti:

$$J = \sum_{i=1}^c J_i = \sum_{i=1}^c \left(\sum_{k_i u_k \in C_k} \|u_k - c_i\|^2 \right)$$

4. Aktualizácia centier zhlu-
kov podľa:

$$c_i = \frac{1}{C_i} \sum_{k_i u_k \in C_k} u_k$$

5. Späť na krok číslo 2.

Algoritmus je iteratívny, a nie je tu žiadna garancia, že skonverguje k opti-
málnemu riešeniu. Správanie závisí na počiatočných pozíciách centier zhlu-
kov, a je odporúčané použiť nejakú metódu na nájdenie dobrých počiato-
čných centier zhlu-
kov. Je tiež možné najprv inicializovať náhodnú maticu M
a potom pokračovať v iteračnej procedúre.

3.4. Fuzzy zhlu- kovanie

Fuzzy zhlu-
kovanie zovšeobecňuje všetky zhlu-
kovacie metódy tým, že umo-
žňuje zhlu-
kovanie jedného objektu do viacej než jedného zhlu-
ku, pričom
v bežnom zhlu-
kovaní je každý objekt členom iba jedného zhlu-
ku. Predpo-
kladajme, že máme K zhlu-
kov a budeme definovať súbor premenných, ktoré
predstavujú pravdepodobnosť, že objekt je klasifikovaný do k -teho zhlu-
ku. [1]

V bežnom zhukovacom algoritme je jedna z týchto premenných rovná jednej a zbytok rovný nule. To predstavuje skutočnosť, že taký algoritmus klasifikuje každý objekt do jedného a práve jedného zhuku. Vo fuzzy zhukovaní „účasť objektu“, čiže prítomnosť objektu, je rozdelená do všetkých zhukov. Premenná sa môže rovnať 1 alebo 0 a suma týchto hodnôt musí byť rovná 1. Tento proces nazveme fuzzifikáciou zhukovej konfigurácie. Proces má veľkú výhodu, že nenúti objekt aby bol zaradený iba do jediného špecifického zhuku. Nevýhodou však je, že sa tu objavuje oveľa viac informácií, ktoré musia byť vysvetlené. Fuzzy algoritmus minimalizuje účelovú funkciu, ktorá je funkciou neznámych účastí v zhuku a ďalej funkciou vzdialenosti. Účasti v zhuku sú predmetom obmedzení a musia byť nezápornými číslami a ďalej účasti pre jeden objekt musia byť v sume rovné 1. To znamená, že účasti majú rovnaké obmedzenia, ako by to boli pravdepodobnosti, že individuum patrí do istej skupiny.

Miera vierohodnosti: jednou z najobtiažnejších úloh v zhukovej analýze je nájdenie vhodného počtu zhukov. Veľkosť „fuzzifikácie“ v riešení sa dá zmerať Dunnovým rozdeľovacím koeficientom, ktorý predstavuje mieru, ako tesne padne fuzzy riešenie na odpovedajúci pevný zhuk. Za pevné zhuky budeme považovať klasifikáciu každého objektu do zhuku, ktorý má najväčšiu účasť. [3]

Je pochopiteľné domnievať sa, že body v strede regiónu medzi dvoma centrami zhukov, majú postupnú príslušnosť *oboch* zhukov. Prirodzene je to zahrnuté pomocou fuzzifikácie definícií ako „nízky“ a „vysoký“. Fuzzifikovaný *c-means* algoritmus dovoľuje každému bodu aby patril zhuku podľa stupňa špecifikovaného stupňom príslušnosti a teda každý bod môže patriť niekoľkým zhukom.

Fuzzy c-means algoritmus delí kolekciu K bodov dát špecifikovaných pomocou m -rozmerných vektorov u_k ($k = 1, 2, \dots, K$) do c fuzzy zhukov, a nachádza centrum zhuku v každom, minimalizovaním funkcie vhodnosti. Fuzzy *c-means* je odlišná od hard *c-means*, hlavne pretože používa *fuzzy delenie*, kde bod môže prislúchať niekoľkým zhukom so stupňami príslušnosti. Na pojatie fuzzy delenia, matica príslušnosti M má dovolené obsahovať prvky v rozsahu $(0, 1)$. [6] Celková príslušnosť bodu zo všetkých zhukov, jednako, musí byť stále zhodná so súladom udržiavať vlastnosti že:

- Súčet každého stĺpca je jedna.
- Súčet všetkých elementov je K matice M . Funkcia vhodnosti je generalizácia J .

$$J(M, c_1, c_2, \dots, c_c) = \sum_{i=1}^c J_i = \sum_{i=1}^c \sum_{k=1}^K m_{ik}^q d_{ik}^2$$

V sériovom móde operácií, fuzzy c -means algoritmus určí centrá zhlukov c_i a maticu príslušnosti M pomocou nasledujúcich krokov:

1. Inicializuje sa matica príslušnosti M s náhodnými hodnotami medzi 0 a 1 v rámci ohraničení matice.
2. Vypočítajú sa c centrá zhlukov c_i ($i = 1, 2, \dots, c$):

$$c_i = \frac{\sum_{k=1}^K m_{ik}^q u_k}{\sum_{k=1}^K m_{ik}^q}$$

3. Vypočíta sa funkcia vhodnosti. Zastaví sa ak je buď pod určitým stupňom prahu, alebo jej nárast od predchádzajúcej interácie je pod istou toleranciou.
4. Vypočíta sa nová matica M :

$$m_{ik} = \frac{1}{\sum_{j=1}^c \left(\frac{d_{ik}}{d_{jk}} \right)^{2/(q-1)}}$$

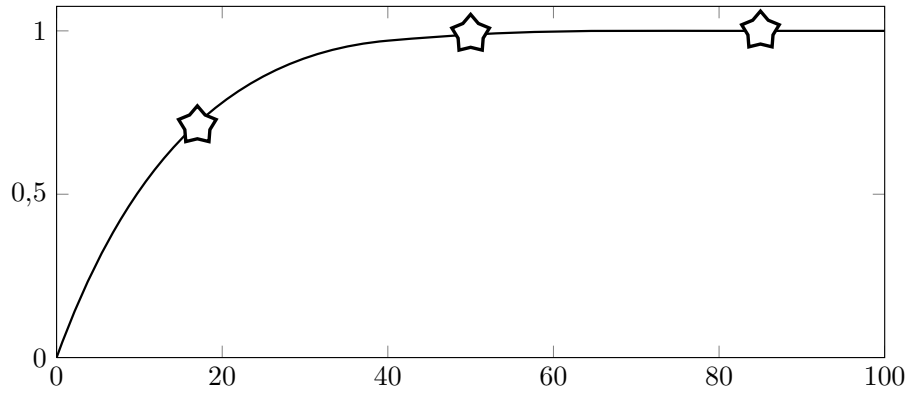
5. Späť na krok 2.

Centrá zhlukov môžu byť alternatívne najprv inicializované, pred uverejnením iteratívnej procedúry. Algoritmus nemusí skonvergovať k optimálnemu riešeniu a správanie závisí na inicializácii centier zhlukov, presne tak ako v prípade hard c -means algoritmu.

Dostávame sa späť k problému modelovania testovacích dát z predchádzajúceho príkladu, dáta boli preložené do FCM funkcie pomocou Matlab Fuzzy Logic Toolbox-u. Okrem požiadavky troch zhlukov, boli všetky ostatné nastavenia predvolené, výsledok: našli sa tri centrá zhlukov.

Tieto môžu byť použité na indikáciu kam umiestniť vrcholy troch fuzzy funkcií príslušnosti na vstupnú os.

Fuzzy c -means je kontrolovaný algoritmus – algoritmus s učiteľom – supervízorom, pretože je potrebné povedať koľko zhlukov c sa má hľadať. [3]



Obr. 3: Centrá zhlukov postredníctvom FCM algoritmu, indikované pomocou hviezd. Zdroj: Výstup zo softvéru Matlab.

Ak c nieje vopred známe, je potrebné aplikovať nekontrolovaný algoritmus. Subtraktívne zhľukovanie je založené na porovnávaní hustoty dátových bodov v príznakovom priestore. Myšlienka je nájsť regióny v príznakovom priestore s veľkou hustotou dátových bodov. Bod s najväčším počtom susedov je vybraný ako centrum pre zhľuk. Dátové body vo vnútri predšpecifikovaného fuzzy polomeru sú potom odstránené teda substraktované a algoritmus sa pozerá po novom bode s najvyšším počtom susedov. To pokračuje pokým všetky dátové body nie sú prehľadané. [7]

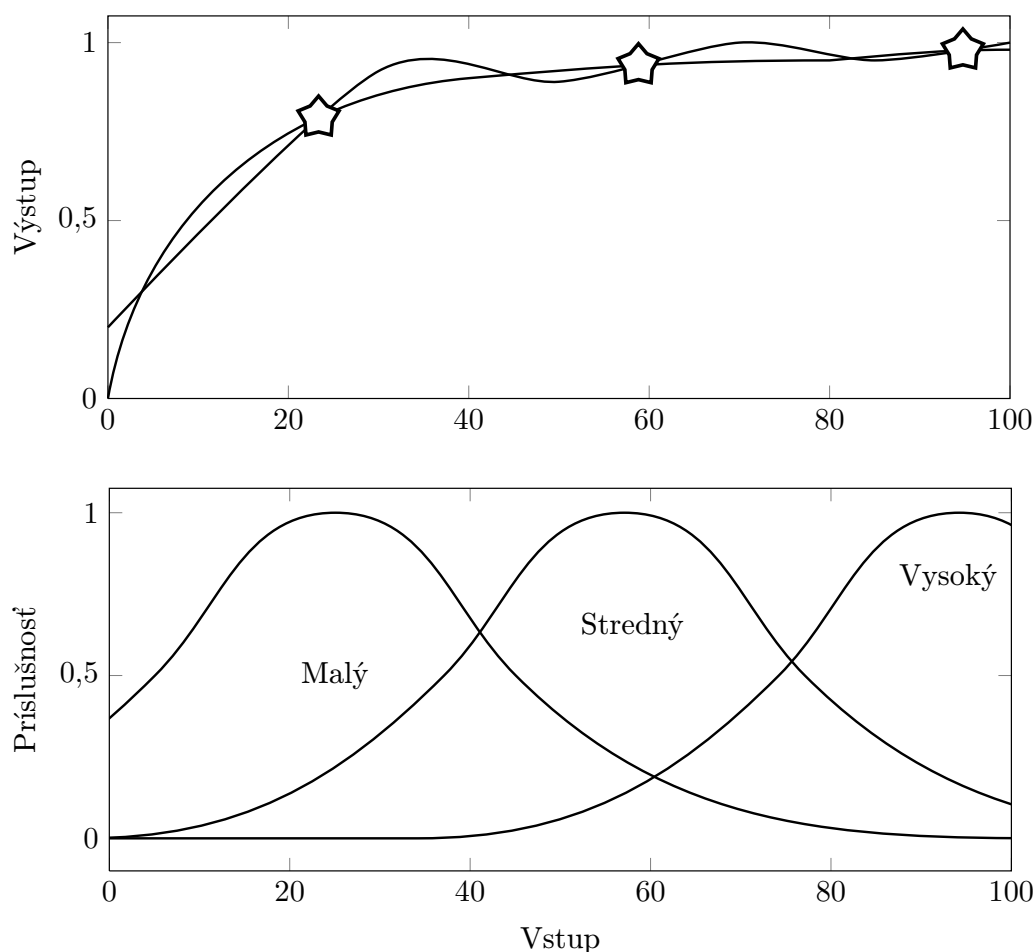
Ako príklad si vezmeme súbor K dátových bodov špecifikovaných pomocou m -rozmerných vektorov u_k , $k = 1, 2, \dots, K$. Bez straty všeobecnosti, dátové body nadobudnú normalizáciu. Od každého kandidáta na centrum zhľuku, meranie hustoty – denzita v dátovom bode u_k je definovaná nasledovne:

$$D_k = \sum_{j=1}^K \exp\left(-\frac{\|u_k - u_j\|}{(r_a/2)^2}\right),$$

kde r_a je pozitívna konštanta. Z tohto dôvodu, dátový bod bude mať vysokú hodnotu hustoty, ak má veľa susedných dátových bodov. Iba fuzzy susedstvo vo vnútri polomeru r_a prispieva k meraní hustoty.

Po vypočítaní hustoty pre každý dátový bod, bod s najväčšou hustotou je vybraný ako prvé centrum zhľuku (angl. cluster). Nech u_{cl} je vybraným bodom a D_{cl} je jeho hodnotou hustoty. V ďalšom, výpočet hustoty pre každý dátový bod u_k je korigovaný nasledovným vzorcom:

$$D'_k = D_k - D_{cl} \exp\left(-\frac{\|u_k - u_{cl}\|}{(r_b/2)^2}\right),$$



Obr. 4: Výsledok aplikácie subzhlukovania na tréovanie dáta. Zdroj: Výstup zo softvéru Matlab.

kde r_b je pozitívna konštanta. Pretože dátové body blízko prvého centra zhuku $\mathbf{u}_{cl}$ budú mať podstatne redukovanú hustotu, tým umožňujú nezvyčajne vyberať body ako ďalšie centrá zhukov. Je to zvyčajne viac ako r_a pre predchádzanie bližšie vzdialených centier zhukov, typicky $r_b = 1,5 \cdot r_a$.

Po porovnaní hustoty pre každý bod a jeho korekciu, ďalšie centrum zhuku $\mathbf{u}_{cl}$ je vyselektované a všetky porovnania hustoty sú korigované znova. Proces je opakovaný pokým nie je vygenerovaný dostatočný počet zhukov.

Horný obrázok ukazuje cieľovú funkciu (prerušovanú) a odhad (plná) tak isto, ako aj centrá zhukov ($\star$). Spodný obrázok zobrazuje tri funkcie príslušnosti korešpondujúce so zhukmi.

Keď aplikujeme substraktívne zhukovanie na množinu vstupno – výstupných dát, každé centrum zhuku reprezentuje pravidlo. Na generáciu pravidiel sú použité centrá zhukov ako centrá pre východiskové množiny na báze pravidiel typu singleton (alebo báze funkcií polomeru v báze polomeru funkcií neurónových sietí).

4. Záver

V článku sme orientovali našu pozornosť na zhodnotenie možností štatistickej analýzy prostredníctvom zhlukovania a jej využitia v manažérskych podnikových analýzach ako v teoretickej, tak i metodickej rovine. Poukázali sme na niektoré praktické sféry jej využitia aj s podporou vybraného počítačového programu Matlab. Záverom konštatujeme, že zhluková analýza má široké uplatnenie nielen vo vnútropodnikových analýzach, ale aj pri medzipodnikovej alebo medzisektorovej analýze odhaľovaní spoločných znakov.

Použitá literatúra

- [1] Jang, J. S. R.; Gulley, N. The Fuzzy Logic Toolbox for use with Matlab. The MathWorks, Inc., Massachusetts, 1995. [cit. 2011-08-16]
Dostupné na: <http://www.cs.nthu.edu.tw/~jang/publication.htm>
- [2] Kovářík, M.; Klímek, P. Shluková analýza v Matlabu. In: *Informační Bulletin ČStS*, roč. 22., č. 1/2011, s. 20–27, ISSN 1210-8022.
- [3] Meloun, M.; Militký, J. *Statistická analýza experimentálních dat*. Praha: Academia, nakladatelství AV ČR, 2004, 953 s. ISBN 80-200-1254-0.
- [4] Mura, L. Zhlukovanie ako štatistická metóda v ekonomických vedách. In: *Forum Statistium Slovacum*, č. 5/2010, s. 165–170, ISSN 1336-7420. Dostupné na:
<http://www.ssds.sk/casopis/archiv/2010/fss0510.pdf>
- [5] Rimarčík, M. *Zhluková analýza*. [online] [cit. 2011-08-11]
Dostupné na internete: <http://rimarcik.com/navigator/ca.html>
- [6] Řezanková, H.; Löster, T.; Húsek, D. 2011. Evaluation of Categorical Data Clustering. In: *Advances in Intelligent Web Mastering – 3*. Berlin: Springer Verlag, 2011, s. 173–182. ISBN 978-3-642-18028-6.
- [7] Stehlíková, B. Fuzzy zhluková analýza ako prostriedok hodnotenia rozdielov. In: *Štatistika a integrácia – 12. slovenská štatistická konferencia*. Bratislava: Slovenská štatistická a demografická spoločnosť. 2004, s. 165–169. ISBN 80-88946-37-9.

A USE OF OPTIMIZATION TOOLS AND GENETIC ALGORITHMS IN THE PRODUCTION PROCESS CONTROL

VYUŽITÍ NÁSTROJŮ OPTIMALIZACE A GENETICKÝCH ALGORITMŮ V PROCESU ŘÍZENÍ VÝROBY

Martin Kovářík, Petr Klímek

Adresa: Tomas Bata University in Zlín, Faculty of Management and Economics, nám. T. G. Masaryka 5555, 760 01 Zlín

E-mail: m1kovarik@fame.utb.cz, klimek@fame.utb.cz

Abstract: The following article deals with up-to-date field of genetic algorithms use in the production process control. The first part of this article is dedicated to introduction of optimization and demonstration of exercises in Matlab. Theory of genetic algorithms is mentioned in the next part of this paper. The last part describes two specific applications of genetic algorithms in the company production optimization. Matlab and Evolver software tools were used for the data computation. Data preparation for Evolver was done in Microsoft Excel.

Keywords: Evolver, Matlab, Genetic Algorithms, Selection, Optimization, Crossing, Mutation, Microsoft Excel.

Abstrakt: Následující článek se zabývá aktuální problematikou využití genetických algoritmů v procesu řízení výroby. První část tohoto článku se věnuje úvodu do optimalizace s ukázkami řešených příkladů v programovém prostředí Matlab. V další části je stručně uvedena teorie ke genetickým algoritmům. Poslední část článku následně popisuje konkrétní aplikace genetických algoritmů v procesu řízení výroby. K výpočtům byly použity programy Matlab a Evolver firmy Palisade. Pro program Evolver byla data zpracována v Microsoft Excelu.

Klíčová slova: Evolver, Matlab, genetické algoritmy, selekce, optimalizace, křížení, mutace, Microsoft Excel.

1. Úvod do optimalizace

Optimalizace je matematická disciplína, ve které hledáme minimum (resp. maximum) dané funkce $f(x)$ na dané množině M . Tato funkce se nazývá účelová či cílová. Množina (nazývá se množina přípustných řešení) bývá typicky popsána nějakými omezeními, nejčastěji soustavou rovnic nebo nerovnic apod. Minimum je matematická funkce, jejíž funkční hodnota představuje

nejnižší hodnotu ze všech vstupních parametrů. Funkce provádí porovnání jednotlivých parametrů a výsledkem je hodnota toho parametru, který se při porovnání se všemi ostatními jeví jako nejnižší. Maximum je matematická funkce, jejíž funkční hodnota představuje nejvyšší hodnotu ze všech vstupních parametrů. Funkce provádí porovnání jednotlivých parametrů a výsledkem je hodnota toho parametru, který se při porovnání se všemi ostatními jeví jako největší.

Poznámka: Funkce v matematice je definována jako zobrazení z nějaké množiny M do množiny čísel (většinou reálných nebo komplexních), nebo do vektorů (pak se mluví o vektorové funkci). Je to tedy předpis, který každému prvku z M jednoznačně přiřadí nějaké číslo nebo vektor (hodnotu funkce). Někdy se však slovo funkce používá pro libovolné zobrazení.

1.1. Globální optimalizace

Při řešení problému globální optimalizace hledáme pro danou účelovou funkci

$$f : D \rightarrow \mathbb{R}, \quad D \subset \mathbb{R}^d, \quad (2)$$

bod x^* tak, aby

$$x^* = \arg \min_{x \in D} f(x). \quad (3)$$

Bod x^* se nazývá globální minimum. Pro jednoduchost budeme uvažovat, že D je souvislá množina tvaru

$$D = \prod_{i=1}^d \langle a_i, b_i \rangle, \quad a_i < b_i, \quad i = 1, \dots, d. \quad (4)$$

1.2. Metody klasické optimalizace

Klasickými metodami optimalizace se rozumí takové, které se opírají o podmínky optimality. Pro diferencovatelnou funkci f stanovíme stacionární body, tj. všechna řešení rovnice $\text{grad } f(x) = 0$. Ty z nich, v nichž je Hessova matice pozitivně definitní, jsou body minima. Hessova matice má tvar

$$H = \left(\frac{\partial^2 f(x)}{\partial x_i \partial x_j} \right)_{i,j=1,\dots,d}. \quad (5)$$

Pozitivně definitní je taková matice $\mathbf{A}$, pro kterou platí:

- $x'Ax > 0$ pro každý nenulový vektor $x \in \mathbb{R}^d$,
- všechna vlastní čísla matice $\mathbf{A}$ jsou kladná,

- všechny hlavní minory $M_1, M_2, \dots, M_d$ matice $\mathbf{A}$ jsou kladné.

Poznámka: Minor příslušný k prvku $a_{i,j}$ matice $\mathbf{A}$ je determinant matice $A_{i,j}$, kterou získáme z matice $\mathbf{A}$ odstraněním i -tého řádku a j -tého slupce. Pokud nemáme podrobnější informace o účelové funkci f , každá z následujících metod určuje v případě optimalizace bez omezení lokální řešení. Metody pro řešení úloh nepodmíněné optimalizace rozdělujeme do tří kategorií podle požadavků na hladkost funkce f .

Metody přímého výběru – nevyžadují výpočty derivací.

Speciální spádové metody – vyžadují výpočet gradientů (např. metoda největšího spádu, kvazinevtonské metody).

Newtonova metoda – vyžaduje výpočet druhých derivací.

Protože případy globální optimalizace řešíme často především v praxi, je velmi důležité zabývat se také teoretickou stránkou, všeobecnou úlohu globální optimalizace. Mnohdy velmi jednoduchá formulace celého optimalizačního problému může být matoucí a vyvolávat dojem, že určení deterministického algoritmu řešícího všeobecnou úlohu globální optimalizace je jednoduché. Ve skutečnosti analýza ukazuje, že takovýto deterministický algoritmus neexistuje.

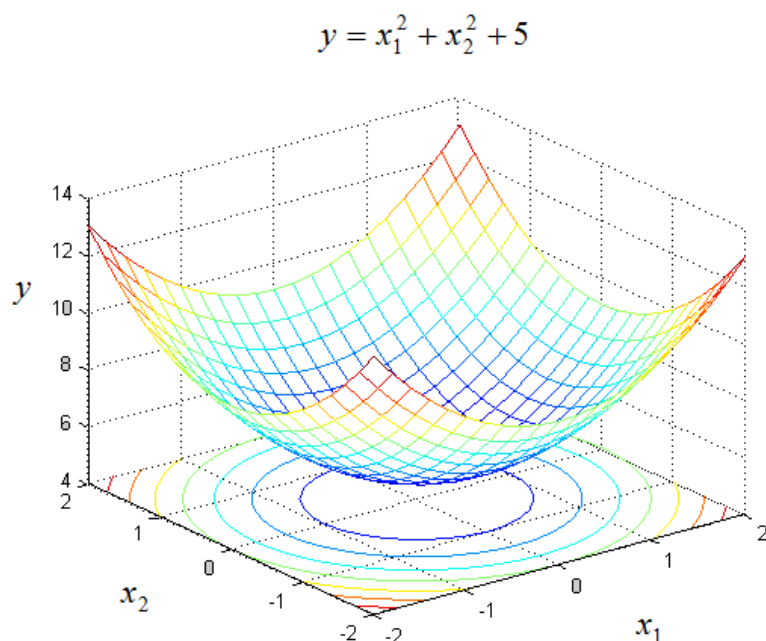
Příklad 1

Úkolem je najít minimum funkce zapsané ve tvaru $y = x_1^2 + x_2^2 + 5$, viz obrázek 1). Tento problém bude řešen v M-souboru s názvem *Opt1.m* vypsáním pod Programem 1.

```
fun = @(x) x(1)^2 + x(2)^2 + 5;
x0 = [1; 1];
options = optimset('LargeScale','off');
[x,fval,exitflag,output] = fminunc(fun,x0,options);
x
fval
output.funcCount
```

Program 1: Optimalizace neohraničené funkce

V prvním řádku je vytvořena proměnná *fun*, která obsahuje vzorec testované funkce. Do proměnné x_0 zapíšeme hodnotu možného řešení $[1; 1]$. Příkaz *options* obsahuje parametry optimalizace, které nastavujeme příkazem



Obrázek 1: Vykreslená funkce v programu Matlab

optimset v případě, že nechceme použít přednastavených hodnot. V našem případě vypneme volbu *LargeScale* na *off*. Výsledné nastavení si lze prohlédnout po vypsání proměnné *options*. Klíčovým příkazem je funkce *fminunc*, která hledá minimum neohrazené funkce. Parametrem je testovaná funkce *fun*, nabídnuté možné řešení v proměnné x_0 a navolené parametry optimalizace *options*.

Výstupem je proměnná *x* obsahující výsledek, *fval* proměnná obsahující hodnotu účelové funkce, v proměnné *exitflag* je informace o podmínce ukončení výpočtu a v proměnné *output* informaci o optimalizaci. Poslední tři příkazy provádějí výpis výsledných hodnot *x*, hodnotu účelové funkce *fval*, proměnná *output.funcCount* vypíše počet iterací optimalizace, viz Výsledek 2.

```
x = 0
    0
fval = 5
ans (output.funcCount)= 6
```

Výsledek 2: Výsledky optimalizace neohrazené funkce

U této rovnice, která představuje paraboloid posunutý ve směru souřadnice *y* o hodnotu 5, je patrné, že výsledek je správný.

Příklad 2: Sestavení výrobního programu

Uvedeme příklad, kdy strojírenský závod vyrábí čtyři výrobky V_1, V_2, V_3, V_4 . Při sestavování výrobního programu je třeba počítat s omezenou kapacitou výrobního zařízení, které je 1200 hodin a s omezeným množstvím suroviny S , které je 1400 tun. Výrobky V_1 a V_2 jsou polotovary potřebné pro výrobu výrobků V_3 a V_4 a mohou být též prodávány. Odbytové ceny těchto výrobků jsou 300, 600, 1000 a 3000 na jednu tunu. Úkolem je stanovit výrobní program, kterým firma dosáhne maximální hodnoty produkce.

Potřebné údaje jsou uvedeny v následující tabulce.

Tabulka 1: Tabulka zadaných hodnot

	V_1	V_2	V_3	V_4	Kapacita
Z	1,5	0,0	2,0	2,5	1200
S	2,0	1,5	2,0	0,0	1400
V_1		0,5	0,0	1,0	
V_2			0,5	2,0	
Cena	300	600	1000	3000	

Pro výše uvedený příklad můžeme psát soustavu nerovnic ve tvaru

$$\begin{aligned} 1,5x_1 &+ 2,0x_3 + 2,5x_4 \leq 1200 \\ 2,0x_1 + 1,5x_2 + 2,0x_3 &\leq 1400 \\ -1,0x_1 + 0,5x_2 &+ 1,0x_4 \leq 0 \\ &-1,0x_2 + 0,5x_3 + 2,0x_4 \leq 0 \\ 300x_1 + 450x_2 + 700x_3 + 1500x_4 &= z \end{aligned}$$

Způsob sestavení soustavy lineárních nerovnic je předmětem operační analýzy a nebude zde probírán. Podstatné pro náš příklad je skutečnost, že chceme maximalizovat zisk, který je dán účelovou funkcí z . Pro optimalizaci můžeme využít např. simplexové metody, kterou lze volat příkazem *linprog*.

Protože hledáme maximum a funkce *linprog* hledá pouze minimum, musíme uvedenou úlohu převést na řešení minima, účelová funkce bude mít záporná znaménka $f(x) = -300x_1 - 450x_2 - 700x_3 - 1500x_4$. Rovnice omezení zůstanou beze změny. Pro řešení úlohy vytvoříme M-soubor, například jako soubor *optimalizace.m*. Viz následující výpis Programu 2.

```

f = [-300; -450; -700; -1500];
A = [1.5 0 2 2.5; 2 1.5 2 0; -1 0.5 0 1; 0 -1 0.5 2];
b = [1200; 1400; 0; 0];
lb = zeros(4,1);
options = optimset('LargeScale','off','Simplex','on');
[x,fval,exitflag,output,lambda] = linprog(f,A,
    b,[],[],lb,[],[],options);
x
fval
lambda.ineqlin
lambda.lower

```

Program 2: Optimalizace simplexovou metodou

Popis příkazů je následující. První řádek definuje vektor $\mathbf{f}$ konstant účelové funkce. Druhý řádek definuje matici $\mathbf{A}$ konstant nerovnic omezujících podmínek, třetí řádek definuje vektor pravých stran nerovnic omezujících podmínek. Další řádek definuje vektor lb , levostranné omezení, tj. že řešení nesmí být menší jak nula. Příkaz *optimset* nastavuje proměnnou *options*, kdy pro simplexovou úlohu musí být parametr *LargeScale* vypnut *LargeScale* na *off* a simplexová metoda zapnuta *Simplex* na *on*. Spuštěním příkazu *linprog* s výstupními a vstupními parametry se provede optimalizace. Dalšími řádky vyvoláme výpis vektoru $\mathbf{x}$ obsahující výsledné hodnoty, hodnotu účelové funkce *fval*.

Nenulové hodnoty proměnných *lambda.ineqlin* a *lambda.lower* nás informují, zda omezení mělo vliv na optimální proces či nikoliv. Po spuštění souboru *optimalizace.m* v prostředí Matlab obdržíme výsledek, viz následující výpis Výsledek 3.

```

Optimization terminated.
x =
    400
    400
     0
    200
fval =
   -600000
ans =
     0
   428.5714
   557.1429

```

```
471.4286
ans =
    0
    0
392.8571
    0
```

Výsledek 3: Výsledky optimalizace

V našem konkrétním případě tedy bude maximální zisk 600 000 Kč v případě, že vyrobíme 400 ks výrobku V_1 , 400 ks výrobku V_2 a 200 ks výrobku V_4 , výrobek V_3 nebudeme vyrábět. Dále jsou vypsány omezující hodnoty, které měly vliv na výpočet optimalizace.

2. Genetické algoritmy

Genetické algoritmy (GA) jsou používány tam, kde by přesné řešení praktických úloh cestou systematického prozkoumávání trvalo téměř nekonečně dlouho. Ve své podstatě vychází genetické algoritmy z genetických procesů probíhajících v přírodě, kde se při evolučním vývoji nebo při šlechtění rostlin a živočichů prosazují jedinci disponující jistými žádoucími charakteristikami, které jsou na genetické úrovni determinovány kombinováním rodičovských chromozomů. Do manažerské oblasti, řízení organizací, se genetické algoritmy začaly prosazovat teprve v poměrně nedávné době. U zrodu genetických algoritmů stála myšlenka, že při hledání lepšího řešení složitých rozhodovacích problémů by bylo možné obdobným způsobem kombinovat části existujících řešení. Stejně jako v oblasti genetiky, i v terminologii genetických algoritmů je používán pojem chromozom, přičemž tento chromozom se skládá ze sekvence uspořádaných genů. Každý gen řídí dědičnost jednoho nebo několika znaků. Většina implementací genetických algoritmů pracuje s binární reprezentací chromozomu, tj. původní reprezentací chromozomu pomocí nul a jedniček, takže chromozomy představují binární řetězce, které většinou reprezentují zakódovaná dekadická čísla, v aplikacích parametry optimalizované funkce. Pro manipulaci s chromozomy bylo navrženo několik genetických operátorů. Z nich nejpoužívanější jsou tři, a to:

- selekce,
- křížení a
- mutace.

Při selekci jde o výběr chromozomů, které se stanou rodiči. Při výběru alespoň jednoho z rodičů je zde důležitým hlediskem jeho tzv. zdatnost. Zjed-

nodušeně řečeno, princip selekce spočívá v tom, že silnější rodičovský chromozom přechází do další generace. Křížení, při němž vzniká jeden nebo více potomků, představuje výměnu částí dvou či více rodičovských chromozomů, která způsobuje modifikaci chromozomů. Mutace potom představuje takovou modifikaci chromozomu, při níž dojde k náhodné změně. Tato eventualita se ale v přírodě vyskytuje jen zřídka. Genetické algoritmy následně pracují tím způsobem, že se nejprve vytvoří počáteční populace m chromozomů a následně se tato populace pomocí genetických operátorů mění tak dlouho, dokud není proces ukončen. Proces reprodukce, který sestává ze tří výše uvedených kroků, tedy selekce, křížení a mutace, je označován jako epocha evoluce populace. Počáteční generace je obvykle získána náhodným generováním.

Pokud jde o velikost populace, pak je logické, že příliš malá populace může zavinit špatné pokrytí prostoru řešení, zatímco naopak příliš velká populace zvyšuje výpočetní náročnost. Experimentální práce naznačují, že v mnoha případech je dostačující velikost populace kolem stovky, resp. mezi n a $2n$, kde n je délka binárního řetězce. Jednou z možností změny m -členné populace je vygenerovat pomocí křížení a mutace novou generaci m potomků a nahradit tak najednou celou rodičovskou generaci.

Jiné způsoby naopak umožňují jistou míru překrývání generace rodičů a potomků. Například vygenerovaný potomek nenahrazuje přímo svého rodiče, ale náhodně vybraného příslušníka aktuální populace. Při řešení optimalizačních problémů se ale objevuje ten požadavek, aby se nejlepší člen aktuální populace objevil v populaci nové. To je možné zajistit například tak, že nový chromozom nahrazuje jedince vybraného pouze z těch, kteří mají podprůměrnou kvalitu. Při aplikaci genetických algoritmů na problémy řízení organizací nebo výroby (viz dále), resp. při nevratném strategickém rozhodování, každý chromozom kóduje nějaké řešení problému. Chromozom zde tedy představuje genotyp s vlastností, jejímž nositelem je tento chromozom, je jemu odpovídající řešení.

V genetických algoritmech jsou přitom preferovány chromozomy s vyšší hodnotou zdatnosti, neboli fitness. A funkce zdatnosti (též také *fitness funkce*) musí být konstruována tak, že její hodnota je tím vyšší, čím lepší je hodnota účelové funkce. Praktické aplikace genetických algoritmů obvykle spočívají v řešení složitých optimalizačních úloh. Tyto jsou obsahem následující druhé části článku. V praxi se pomocí genetických algoritmů řeší úlohy optimalizace, využívají se k vyhledávání nejlepší topologie, v technologii, výrobě a v průmyslové automatizaci a jako alternativní metody učení umělých neuronových sítí.

Na rozdíl od gradientních metod, které reprezentují hledání lokálního minima nebo maxima pomocí jednoho zpřesňujícího se řešení, představují gene-

tické algoritmy jiný přístup, který používá populaci prozatímních řešení, jež paralelně procházejí parametrický prostor a navzájem se ovlivňují a modifikují pomocí genetických operátorů. Tím se dosahuje toho, že populace jedinců najde správné řešení rychleji, než kdyby se prohledával prostor izolovaně.

Nalezení minimální hodnoty multimodální funkce

Uvedený příklad ukazuje implementaci GA pro nalezení minimální hodnoty multimodální funkce označované jako Rastriginova funkce dané následujícím vztahem:

$$Ras(x) = 20 + x_1^2 + x_2^2 - 10(\cos 2\pi x_1 + \cos 2\pi x_2). \quad (1)$$

Náročnost úlohy nejlépe demonstruje zobrazení funkce pro dvě proměnné dle obrázku 2 a 3 zobrazených na další straně. Představa komplikovanosti optimalizační úlohy, například pro osm proměnných dle vztahu (1) při použití klasických gradientních metod je zřejmá.

Procedura evolučního algoritmu

vygenerujeme P (populaci o N bodech náhodně vybranou z D)

repeat

najdeme nejhorší bod v P , $\mathbf{x}_{WORST}$, s nejhorší hodnotou $\mathbf{f}$

zkopírujeme M nejlepších bodů z P do nové populace Q , $1 \leq M < N$

repeat

repeat

vygenerujeme nový zkušební bod $\mathbf{y}$ podle heuristiky aplikované na P

until $f(\mathbf{y}) \leq x_{WORST}$

zapisujeme nové zkušební body do Q

until Q není tvořeno N body

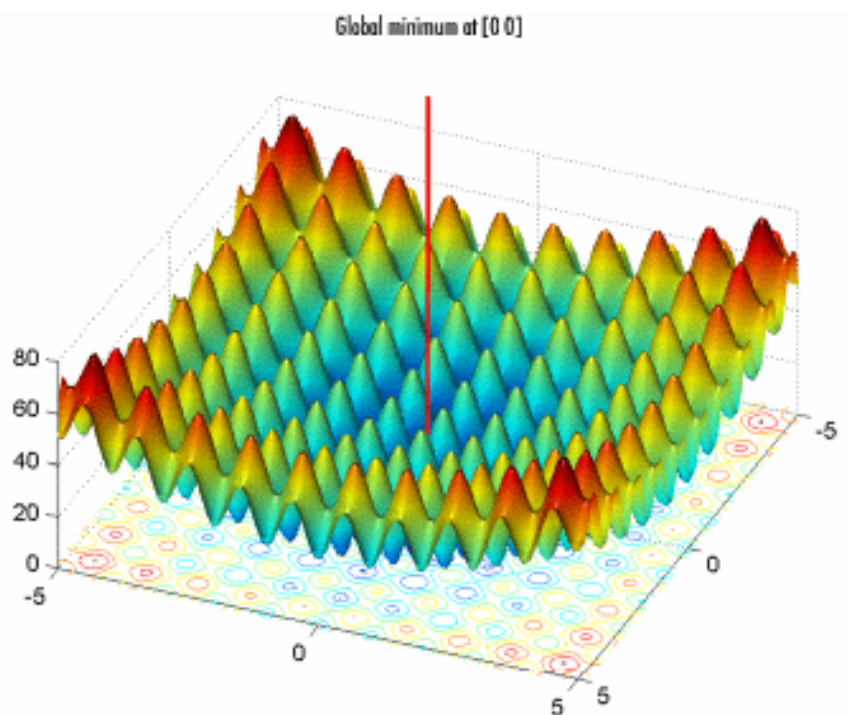
nahradíme P za Q **dokud** není splněna ukončovací podmínka

end

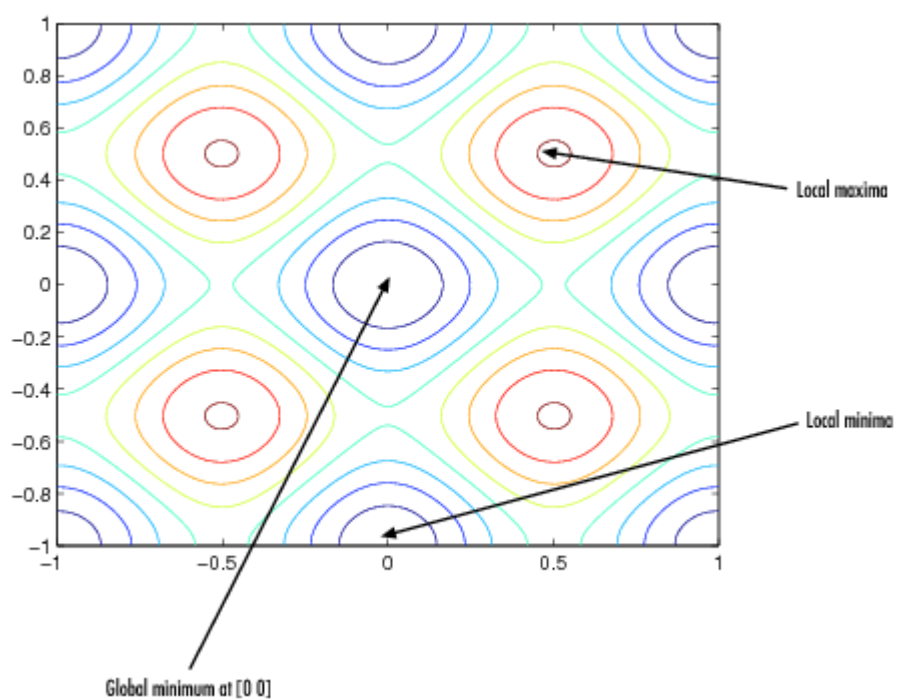
3. Praktické využití genetických algoritmů

Příklad 3

(Matlab – Genetic Algorithm Toolbox) Genetických algoritmů lze použít pro stanovení optima výroby. Na ukázkou uvedeme následující příklad. Máme vyrobit 12 výrobků, z toho 5 výrobků typu A, 3 výrobky typu B, 2 výrobky typu C, 1 výrobek typu D a E. K výrobě je potřeba 6 různých pracovišť. Počet strojů na pracovištích je v levé části tabulky 2 a doba ke zpracování různých



Obrázek 2: Multimodální Rastriginova funkce



Obrázek 3: Vrstevnicový graf nalezení alternativ maxima a minima

typů výrobků na různých pracovištích je různá a je dána křížovou tabulkou, viz tab. 2. Pokud bychom výrobky zadávali libovolně, nebyla by výroba optimální (minimální doba výroby). Proto je vhodné provést optimalizaci pořadí zadávání výrobků do výroby. Volit lze jak počet strojů na jednotlivých pracovištích, tak dobu zpracování různých výrobků na různých pracovištích.

Průběh samotného výpočtu lze sledovat např. za pomoci grafů *Best fitness* a *Score diversity*, viz [1], [2], [6].

Výsledky řazení součástek zobrazíme zvolením menu *File – Export to workspace* a zaškrtnutím menu *Export results to a Matlab structure named*. Název *optimresults* můžeme ponechat nebo jej změnit. Volbu potvrdíme tlačítkem *OK*. Výsledek řazení součástek obdržíme vypsáním příkazu *garesults* na pracovní ploše a příkazem *fval* obdržíme hodnotu doby výroby všech součástek po optimalizaci. Viz následující výpis Výsledek 1.

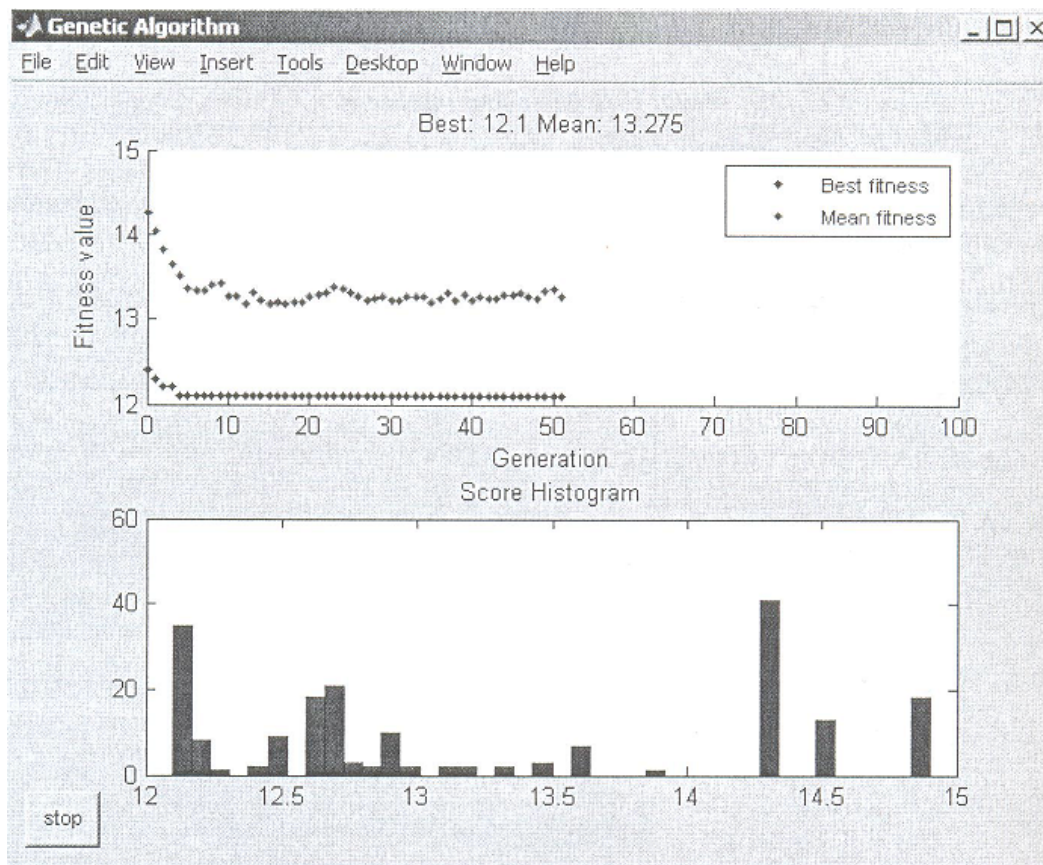
Příklad 4: Aplikace genetických algoritmů v procesu plánování výroby

Definice problému

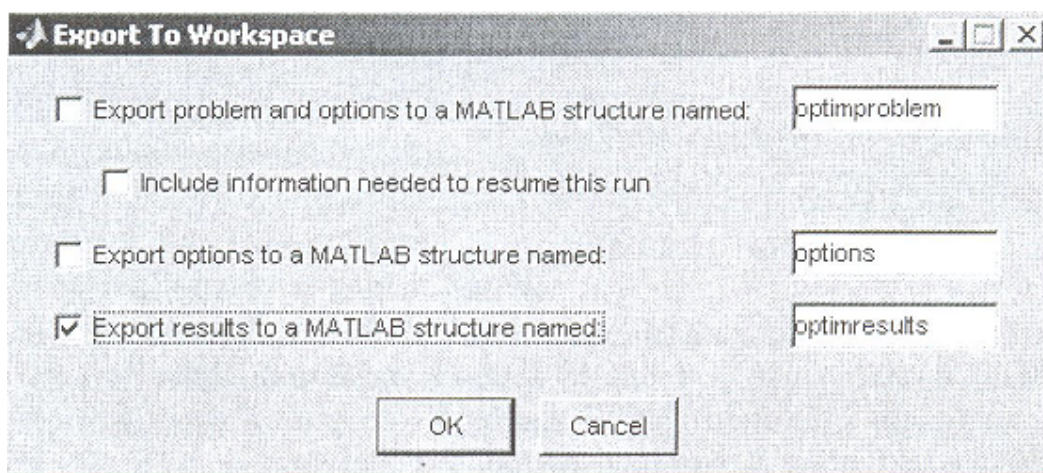
Podúkol byl zpracovaný ve firmě Mitas Zlín (dále představený formou posteru na konferenci Robust 2006 [5]), jako součást rozsáhlejšího projektu zabývajícího se aplikací statistických metod v procesu plánování výroby motoplášťů v rozsahu 117 rozměrů [4]. Konkrétně se jednalo o maximalizaci produkce při daných podmínkách. Nejdříve bylo nutné detailně popsat výrobní halu, viz obrázek 7. Analýza potvrdila fakt, že množství produkce je závislé výhradně na výkonové neekvivalenci jednotlivých částí výrobní linky. Úzkým místem se ukázaly být lisy, tedy produkce byla závislá výhradně na využití 21 instalovaných lisů. Výkony a kapacity ostatních částí výrobní linky nejsou rozhodující. Mají svůj systém optimalizace spočívající pouze v jednoduchém ručním rozhodování úsekových nejvyšších pracovníků na základě výroby a odbytu. Problém se tedy týkal výhradně maximálního využití lisů, a to na jednom lisu (A), na kterém se vyrábí 9 rozměrů, pěti lisech (B) s 30 rozměry a 15 lisech (C) se 78 rozměry. Lisy jsou v praxi využívány přibližně na 95 %. Cílem úkolu bylo maximální využití lisů v rozmezí 14 dnů, po které je plánování uskutečňováno.

Řešení

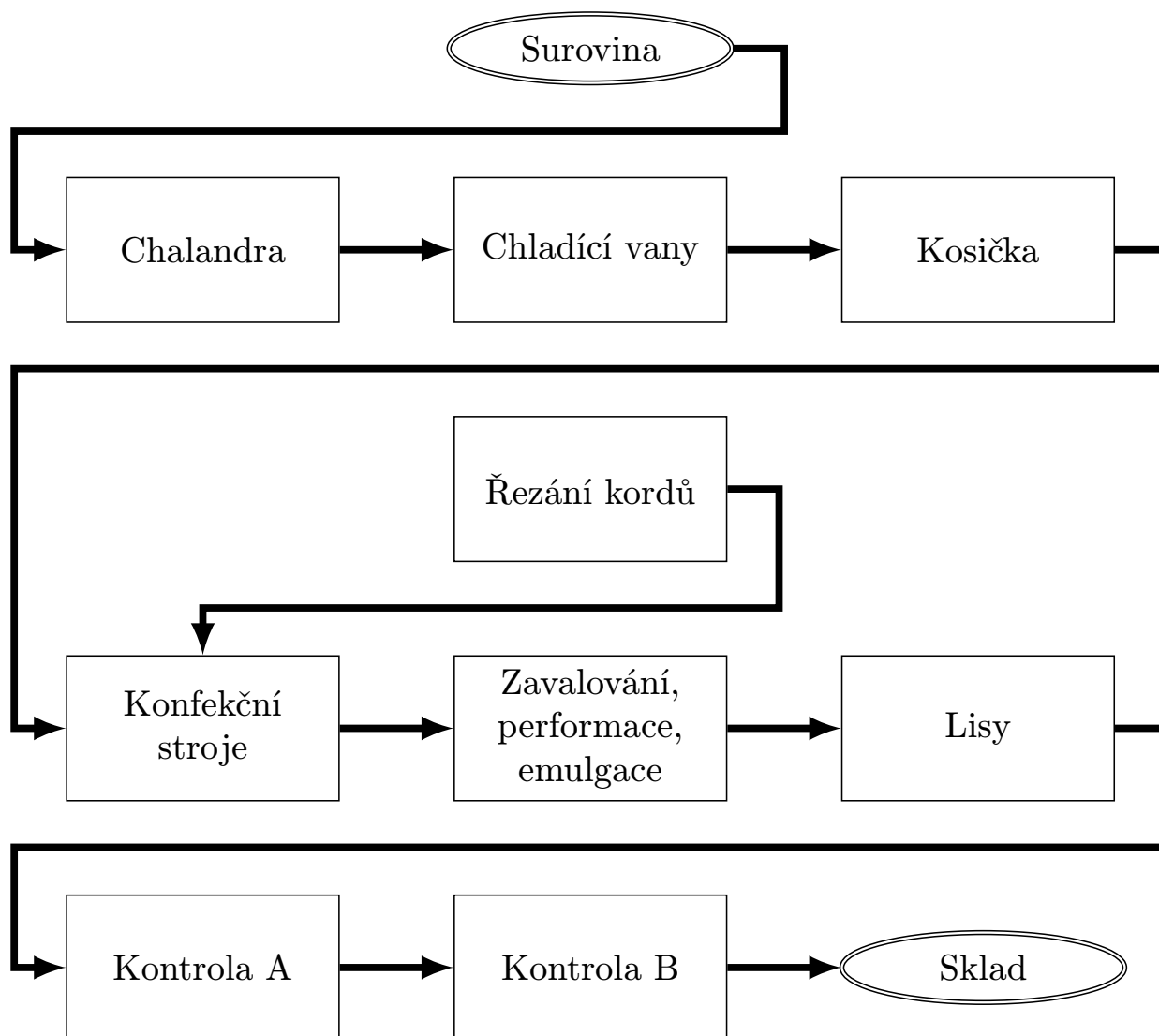
Pro řešení úkolu vzhledem k mnoha nezávisle proměnným byla zvolena metodika využívající genetické algoritmy. Úkol byl řešen za pomoci softwaru Evolver od společnosti Palisade, který pracuje jako doplněk (Add-In) v Microsoft Excelu. Řešení již nebude podrobněji rozebíráno, bude zdůrazněna pře-



Obrázek 5: Obrázovka grafů *Best fitness* a *Score diversity*
[Zpracováno v programu Matlab]



Obrázek 6: Obrázovka volby uložení výsledků
[Zpracováno v programu Matlab]



Obrázek 7: Schéma výroby [5]

devším praktická aplikace genetických algoritmů. Postup přípravy dat se odvíjel od zvolení účelové funkce vyjadřující využití lisovacích strojů. Bylo tedy nutné ke každému ze 117 výrobků přiřadit čas lisování včetně manipulace, viz sloupec C na obrázku 8. Další vstupní data byla počet jednotlivých lisů a jejich přiřazení k odpovídajícím rozměrům (sloupec D), skladové zásoby (P) a denní požadavky (U) na množství pro jeden každý rozměr.

Princip řešení spočíval v určení maximálního počtu vyrobených kusů při využití časového fondu 22,5 hodin, tedy tří směn po 7,5 hodinách. Jelikož jde právě o řešení na základě úzkého místa, lisů, které byly skutečně využity na téměř 100 %, uvažujeme další úseky provozu. Dále princip požaduje, aby bylo možné stanovit výrobu produktů, které je nutno bezpodmínečně vyrobit. To splňuje sloupec U („potřeba“). Programu Evolver byl nabídnut

sloupec Q pro stanovení počtu vyrobených kusů zvoleného produktu. Nyní přecházíme na konkrétní aplikaci genetických algoritmů, spustíme program Evolver a v nastavení (Settings), viz obr. 9, provedeme následující.

Do políčka „For the cell“ vložíme funkci, kterou chceme minimalizovat, maximalizovat nebo přiblížit konkrétní hodnotě (my budeme maximalizovat). K tomu je ale nutné také určit buňky, které bude algoritmus upravovat pro dosažení nejlepší kombinace pro výstup a stanovit podmínky a omezení. Buňky pro úpravu se vkládají do políčka „By adjusting the cells“, viz obrázek 10, podmínky a omezení se vkládají přes políčka „Subject to the Constrains“. Sloupec R (vyrobit tento výrobek – ano/ne) bude obsahovat pouze binární hodnoty a u sloupce Q (velikost výroby), který označíme jako sloupec pro úpravu, nastavíme také celá čísla, ale s hranicí (například $\{0; 500\}$) a metodu pro výpočet („Solving method“) v obou případech „Recipe“, viz obrázek 10. Program sice nabízí i jiné metody, ale pro naše účely je vhodná tato.

Mutační konstantu nastavíme na nízkou hodnotu 0,02, což je v praxi velmi častá hodnota, a práh křížení na hodnotu 0,9, která určuje pravděpodobnost, že se gen chromozomu nezmění. Software není omezen jen na pouhé základní (default) algoritmy pro proces selekce-křížení-mutace, poskytuje i modifikované algoritmy vhodnější pro specifické problémy, viz obrázek 11.

Je tedy možné použít následující postupy (algoritmy): arithmetic crossover, heuristic crossover, Cauchy mutation, Boundary mutation, non-uniform mutation, linear a local search. Pokud zvolíme všechny uvedené algoritmy (což program umožňuje), vybírá se při procesu řešení nejlepší možný postup. Problém na jeden den je ve velké většině vyřešen už po dvou minutách, a to i přesto, že byla zadána potřebná výroba využívající lisy z 90 % a současně nebyla algoritmu poskytnuta pomoc v podobě stanovení potřebné výroby již na počátku řešení. To je velmi dobrý výsledek.

4. Závěr

Pro aplikaci genetických algoritmů bylo rozhodnuto z důvodu složitosti výrobního procesu přípravy motoplášťů, který dokáže optimalizovat snad jen řízená metoda pokus-omyl. Genetické algoritmy tuto metodu více než vylepšily a jsou snad jediným možným nástrojem pro řešení velmi složitých problémů s velkým počtem proměnných. Po vytvoření predikce výběrem nejvhodnější matematické nebo statistické metody a po analýze výrobního procesu byl v Microsoft Excelu zkonstruován model, který v potřebném směru odpovídal skutečné povaze výroby. Efektivnost výroby lze interpretovat jako množství vyrobených kusů s maximálním využitím kapacit, proto zde byla jako úče-

	A	C	D	P	Q	R	S	T	U	V	W	X	Y	Z
1	1	počet lisů BOM 40"												
2	5	počet lisů BOM 30"							1	1				22,49
3	15	počet lisů BOM 40"							5	4				22,483
4									15	7	15	5	1	22,499
5														
6	č. výrobku	dobu lisování manipulace	+ počet lisů	sklad	GA výr.	GA ano/ne	nutná výroba	vyrobeno	potřeba	chybí	počet hodin na lis	počet hodin na lis	počet hodin na lis	celkem hodin
7	21212	15,3	15	50	1	0	0	0		0	0	0	0	0
8	22021	15,3	5	0	27	0	0	0		0	0	0	0	0
11	22041	15,3	5		335	0	1	335	320	0	0	42,71	0	42,713
12	22043	22,3	5		50	0	0	0		0	0	0	0	0
13	22049	22,3	5		56	0	0	0		0	0	0	0	0
14	22051	15,3	5		314	0	1	314	300	0	0	40,04	0	40,035
40	24432	32,3	15		16	0	0	0		0	0	0	0	0
41	24532	15,3	15		403	0	1	403	400	0	51,4	0	0	51,383
42	24544	15,3	15		66	0	0	0		0	0	0	0	0
44	24643	22,3	15		500	0	1	500	500	0	92,9	0	0	92,917
45	24649	32,3	15		17	0	0	0		0	0	0	0	0
46	24796	22,3	15		160	0	1	160	130	0	29,7	0	0	29,733
79	25141	17,3	1		68	0	0	0		0	0	0	0	0
80	25250	17,3	1		156	0	1	156	155	0	0	0	22,49	22,49
91	24087	22,3	15		27	0	0	0		0	0	0	0	0
92	24134	32,3	15		71	0	0	0		0	0	0	0	0
93	24167	22,3	15		301	0	1	301	300	0	55,9	0	0	55,936
94	24403	32,3	15		300	0	1	300	300	0	80,8	0	0	80,75
115	52008	15,3	5		30	0	0	0		0	0	0	0	0
116	52010	15,3	5		205	0	1	205	150	0	0	26,14	0	26,138
117	61011	15,3	5		55	0	0	0		0	0	0	0	0
121	52020	17,3	1		12	0	0	0		0	0	0	0	0

Obrázek 8: Úprava dat v Microsoft Excelu [4]

Evolver Settings

Find the

☐ Minimum
☒ Maximum
☐ Closest value to

For the cell

\$W\$:26

By Adjusting the Cells

Add ..

Delete

Edit...

\$Q\$7:\$Q\$123
\$R\$7:\$R\$23

Subject to the Constraints

Show:

☐ Hard
☐ Soft
☐ Range
☒ All

Add ..

Delete

Edit...

R: 7<=\$Q\$7:\$Q\$123<=FNN [INT]
R: 7<=\$R\$7:\$R\$123<=1 [INT]
T: 0 <= \$Z\$2:\$Z\$4 <= 22,5

Help..

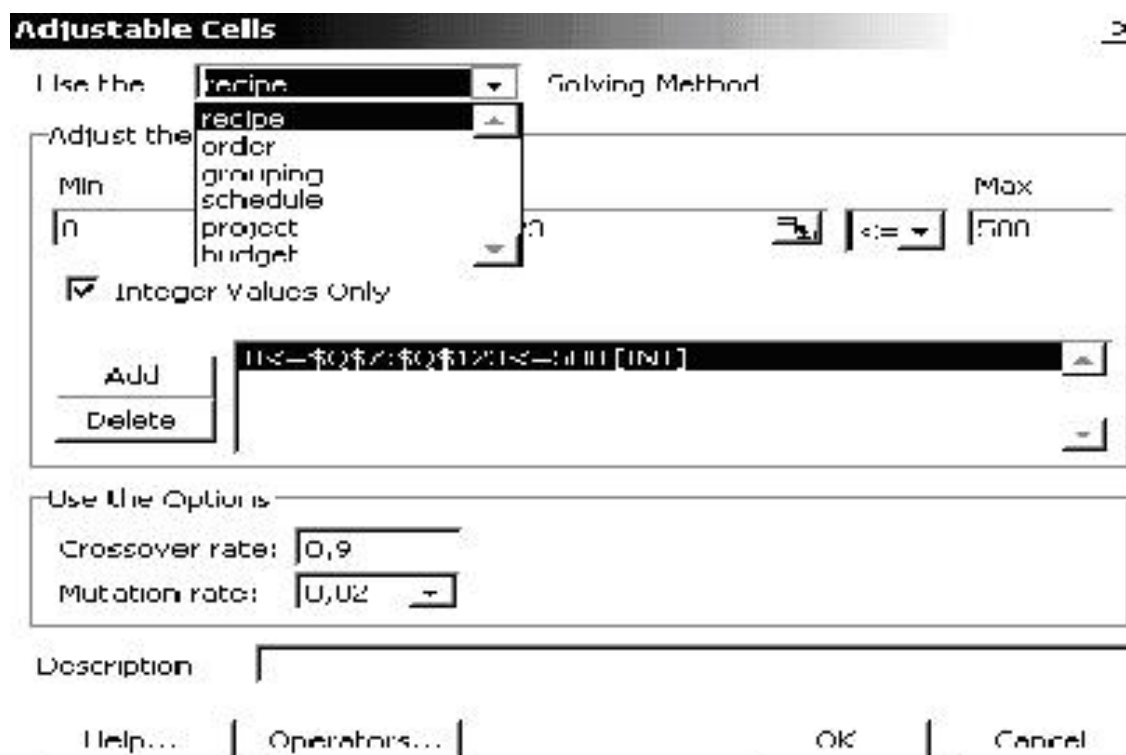
Options...

Macros..

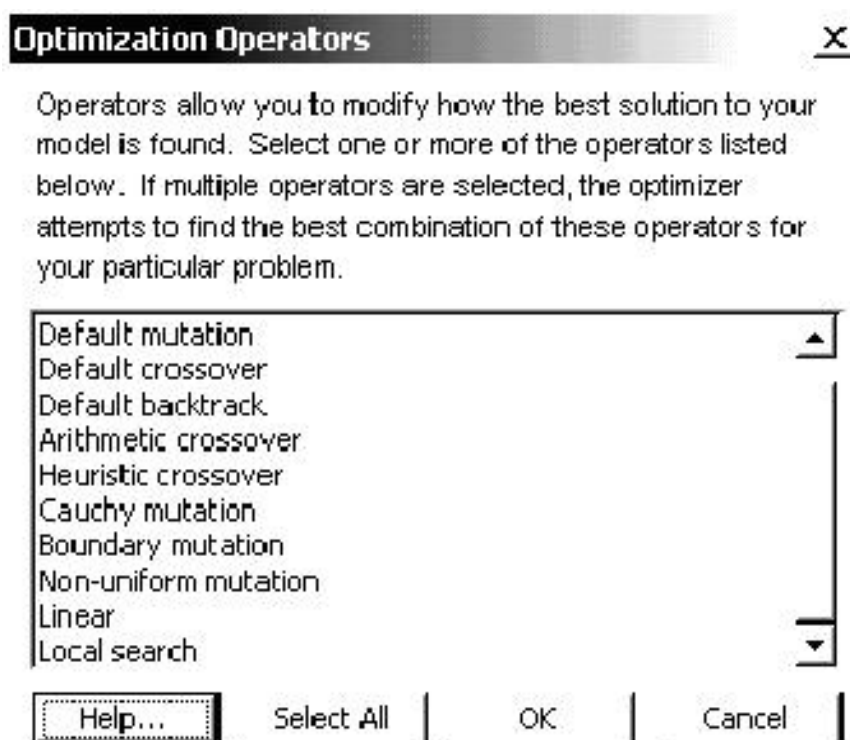
OK

Cancel

Obrázek 9: Nastavení v programu Evolver [5]



Obrázek 10: Další nastavení v programu Evolver [4]



Obrázek 11: Volby algoritmů [4]

lová funkce vybrána rovnice, jejíž výsledek určoval minimálně požadované množství vyrobených motoplášťů.

Data byla převedena do speciálního programu Evolver od společnosti Palisade, který pracuje jako doplněk (Add-In) v Microsoft Excelu. Mezi dílčí problémy patřilo určení vhodného pořadí receptur při vkládání do výroby. Úkol byl řešen rovněž pomocí evolučního (genetického) algoritmu, tentokrát ovšem v programu Matlab. Evoluční algoritmy jsou typické svojí robustností – schopností řešit obtížné optimalizační úlohy, nebo úlohy, ve kterých potřebujeme o něčem rozhodnout. Lze je charakterizovat vlastnostmi jako je multimodálnost a multikriteriálnost. Jejich nasazení je efektivní v úlohách, které lze definovat následovně:

- Prostor řešení je příliš rozsáhlý a chybí odborná znalost, která by umožnila zúžit prostor slibných řešení.
- Nelze provést matematickou analýzu problému.
- Tradiční metody selhávají.
- Úlohy s mnohačetnými extrémy, kritérii a omezujícími podmínkami.

Před užitím evolučního algoritmu bychom si měli uvědomit, že daný konkrétní typ algoritmu se hodí pro řešení jen určitého okruhu problémů. Evoluční algoritmy nejsou tedy vhodné pro řešení všech úloh. Neměli bychom proto také zapomínat na škálu klasických metod (metody přímého výběru, speciální spádové metody, Newtonova metoda), které jsou schopny vyřešit řadu optimalizačních úloh s relativně nízkou výpočetní náročností. Shrnutí nevýhod evolučních algoritmů:

- Pro mnohé úlohy je typická velká časová náročnost.
- Nelze otestovat, zda se jedná o globální optimum.
- Pro příliš rozsáhlé úlohy poskytují řešení příliš vzdálená od optima.
- Ukončení optimalizace je předem určené na základě časového limitu nebo stagnace kritériální funkce.

Použitá a doporučená literatura

- [1] Dostál, P. *Moderní metody ekonomických analýz*. Zlín: UTB, FaME, 2002. ISBN 80-7318-075-8.
- [2] Dostál, P.; Rais, K.; Sojka, Z. *Pokročilé metody manažerského rozhodování*. Grada, 2005. ISBN 80-247-1338-1.
- [3] Dostál, P. *Pokročilé metody analýz a modelování v podnikatelství a veřejné správě*. 1. vyd. Akademické nakladatelství CERM, Brno, 2008. ISBN 978-80-7204-605-8.

- [4] Kasal, R. *Projekt užití statistických metod v procesu plánování výroby firmy Mitas, a. s., Zlín*. Diplomová práce. Zlín: UTB, FaME, 2005. Bez přiřazeného ISBN.
- [5] Kasal, R.; Klímek P.; Stríž, P.; Říha, J. *Application of Genetic Algorithms to Planning Production in the Firm Mitas, a. s. (in Czech). Aplikace genetických algoritmů v procesu plánování výroby ve společnosti Mitas, a. s.* Scientific poster presented on the conference Robust 2006, January 2006. A1 format.
- [6] Kovářík, M. *Počítačové zpracování dat v programu Matlab*. 1. vyd. Martin Stríž, Bučovice, 2008. 278 s. ISBN 978-80-87106-09-9.
- [7] The MathWorks. *Matlab – Genetic Algorithm Toolbox – User’s Guide*, The MathWorks, Inc., 2007.
- [8] The MathWorks. *Matlab – Optimization Toolbox – User’s Guide*, The MathWorks, Inc., 2008.
- [9] Hynek, J. *Genetické algoritmy a genetické programování*. 1. vydání. Nakladatelství Grada Publishing, a. s., Praha, 2008. 182 s. ISBN 978-80-247-2695-3.
- [10] Venkataraman, P. *Applied Optimization with Matlab Programming*. 2. vydání. Nakladatelství John Wiley & Sons, Inc, 2009. 526 s. ISBN 978-0-470-08488-5.
- [11] Tvrdík, J. Evoluční algoritmy a adaptace jejich řídicích parametrů. *Automa*. 2007. č. 7–8, str. 453–457, [online]. [cit. 2010-7-17]. Dostupný na WWW: <http://www.automatizace.cz/article.php?a=1818>
- [12] Tvrdík, J. *Porovnání heuristik v algoritmech globální optimalizace*. Sborník konference Robust, JČMF, 2002, str. 321–332. ISBN 80-715-900-6.
- [13] Singiresu, S. Rao. *Engineering Optimization: Theory and Practice*. John Wiley & Sons, 1996. ISBN 978-0471550341.
- [14] Jablonský J. *Operační výzkum*, Professional Publishing, Praha, 2002. ISBN 80-86419-23-1.
- [15] Žižka M. *Vybrané statě z operačního výzkumu*, HF TUL, Liberec, 2003. ISBN 80-7083-691-1.
- [16] Kořenář, V.; Lagová, M. a kol. *Optimalizační metody*. Oeconomia, VŠE, Praha 2003. ISBN 80-245-0609-2.
- [17] Padberg, M. *Linear Optimization and Extensions*. Berlin: Springer, 1999. ISBN 3-540-65833-5.
- [18] Dantzig, G. B.; Thapa, M. N. *Linear Programming. 2. Theory and Extensions*. Springer, 2003. ISBN 0-387-98613-5.

VERIFICATION OF THE EFFICIENCY OF COMBINED TEACHING OF ENGLISH

Zuzana Kurucová, Beáta Stehlíková, Anna Tirpáková

Address: PaedDr. Zuzana Kurucová, Faculty of Law, Pan European University, Tomášikova 20, 851 05 Bratislava, Slovak Republic

E-mail: zuzana.kurucova@uninova.sk

Adresa: prof. RNDr. Beáta Stehlíková, CSc., Faculty of Economy and Business, Pan European University, Tematínska 10, 851 05 Bratislava

E-mail: bstehlikova@ukf.sk

Adresa: prof. RNDr. Anna Tirpáková, CSc., Faculty of Natural Sciences, Constantine the Philosopher University in Nitra, Trieda A. Hlinku 1, 949 74 Nitra, Slovak Republic

E-mail: atirpakova@ukf.sk

Abstract: This paper presents an experiment by which effectiveness of teaching foreign languages is verified by using three different forms: in-class model, the e-learning model and the combined model (combining classic teaching/learning and e-learning). The test used to confirm or reject the basic presumptions of the experiment was the Kruskal-Wallis' test. It has been proved that the most effective method was the combined method of foreign language learning.

Keywords: In-class Learning, Combined Methods of Teaching, E-learning, Teaching Methods.

1. Introduction

Within the university environment, e-learning is going through, or more precisely, has gone through various stages in the world wide measure. In this connection, it is possible to identify several common features that, paradoxically, rest upon diversity of progress. Meredith and Newton [10] state that e-learning is well-developed at some universities, whereas according to Marshall and Mitchell [9] other universities are facing problems achieving at least a minimum level of e-learning. According to Bleimann [3], knowledge society requires creating new approaches within educational process, what in final consequence – although from another perspective – confirm also Webster and Sudweeks [13]. In their opinion, the overall development in the society demands from learners to be independent (autonomous) and it subsequently

leads teachers to perceive particular aspects of education even more – including new forms of e-learning. Within the context of the research, it is necessary to point out that at university in Chile, an e-learning project was performed, which was aimed at the English language teaching as a foreign language. The author Emerita Bañados in her work [2] describes the execution of this project and its particular outcomes.

Currently, when the possibilities to prove successful on the labor market across Europe have grown, the foreign languages have become a necessity for everyone who seeks the application in its field. Preparation for future career does not only include the classical teaching in the classroom any more, but more and more space is devoted to alternative models of teaching. Such alternative models that complemented and in some cases replaced in class teaching approaches include e-learning. E-learning is in itself a form of virtual education, which may be under certain circumstances either fully or partially supplemented by traditional forms of teaching by means of a teacher and direct (not reproduced) communication with him. If the so-called classic e-learning model of education is examined, where the computer is understood in the strict sense (i.e. the hardware), then it can be identified in particular following basic sub-modules of e-learning in this system:

- a) Professional software (interactive CD, DVD and so on, which are not linked to access to the Internet).
- b) On-line learning via the Internet and dedicated programs, or web pages.
- c) On-line learning via the Internet, but without web pages being directly designed therefore (e.g. e-mail tutorials, sending of documents, tasks, etc. via e-mail).

2. Methodology of Research

The aim of the study was to verify effectiveness of three teaching methods in language learning, especially in English language teaching.

An experiment was conducted to verify the effectiveness, with the aim to confirm the benefits of e-learning in practice. The experiment was divided into three groups of students: students who were educated only in the classical mode of instruction (Standard), followed by students who were educated only in the on-line mode (On-line), and the third group consisted of students who attended the “traditional” hours and in a supplementary position applied and studied by means of e-learning (Combined). Teaching was focused on technical English – English in the Media, i.e. the technical language used in the mass media and for the needs of people working in this field. In the results achieved by respondents included in each of this group.

At the beginning experiment a so called pre-testing was conducted, i.e. the verification of language proficiency of members of all three groups before the experiment, and post-testing after the experiment, i.e. the verification of language proficiency of members of all three groups. Pre-tests and post-tests focused on verification of the same skills in all three groups so that these were quantifiable and comparable. Tests verifying the written knowledge (Writing) consisted of the following type of exercises: Grammar (G), Terminology (T) and Structure (S). The verified ability of students was to form a coherent text that is grammatically correct and uses the appropriate terminology in it. A student could get a maximum of 100 points (100 %) for each of the listed items, while the total for a section was 300 points.

Verification of student knowledge in the part “Listening” focused on the following sub-system of language and perception: Reproduction (R) Questioning (Q) and Terminology (T). A test should ascertain that the student understood the text as a whole (Reproduction), and was able to perceive the specificity (to answer specific questions of teacher) and that the student used the correct terminology (T) in the answers.

Test in the part “Speaking” should verify the fluency of speech – language skills (Fluency = F), aptness of the used vocabulary (Terminology = T) and the adequacy of grammatical constructions (G = Grammar). Also in this case a student could get up to 100 points (100 %) for each of the items, with a total of 300 points for this part.

Vocabulary as the final component of the test (“Vocabulary”) examined the knowledge of general vocabulary (Universal = U), specialized vocabulary (Pro = P), and the appropriateness of their application (Application = A). As in the previous sections, a student could get the maximum of 100 points (100 %) for each of these items, with a total of 300 points for this part. It follows that the student could get a total of 1,200 points (4 sections per 300 points) within the pre-test or post-test.

The assumption that the measured values of the observed characters (Writing, Listening, Speaking, Vocabulary) are normally distributed, is not justified in our case, because these values in the individual groups or samples (of pre-test and post-test) will be compared by means of the non-parametric one-sample Wilcoxon test. The Wilcoxon signed-rank test is a non-parametric alternative to the one sample t test. Its use is mainly motivated by uncertainty concerning the assumption of normality made in the t test [4].

When the statistical significance of differences between the three analysed groups in view of the achievements in pre-test and post-test, we have used the Kruskal-Wallis test. The Kruskal-Wallis test is a non-parametric test that

compares three or more unpaired groups. The calculations were performed in the program STATISTICA.

The statistical significance of the difference between the results achieved by respondents in pre-test and post-test in each group, the tested hypothesis is the following null hypothesis:

H_0 : Both selections come from the same basic set, in other words there is **no** statistically significant difference between the results achieved by respondents in the pre-test and post-test given the measured values (Writing, Listening, Speaking, Vocabulary). Compared to the alternative hypothesis:

H_1 : Both selections do not come from the same basic set, in other words there **is** statistically significant difference between the results achieved by respondents in the pre-test and post-test given the measured values (Writing, Listening, Speaking, Vocabulary).

3. Results and discussion

The results achieved by students in the pre-test and the post-test are summarized in the following table.

From Table 1, we can see that the respondents assigned to groups have reached different results in pre-test and post-test. In further analysis, we interested whether the statistically significant difference between both results achieved by respondents in pre-test and results in post-test will be confirmed,

Table 1: The average success rate in pre-test and post-test of each group (in %)

Object of Research	Writing				Listening				Speaking				Vocabulary				Test
	G	T	S	Σ	R	Q	T	Σ	F	T	G	Σ	U	P	A	Σ	Total
On-line (pre-test)	61	66	63	63	63	64	65	64	69	71	65	68	69	63	67	66	65
On-line (post-test)	64	73	67	68	67	67	72	68	70	77	66	71	71	73	76	73	70
Standard (pre-test)	53	51	59	54	58	58	58	58	64	68	60	64	73	71	79	74	63
Standard (post-test)	55	67	65	62	66	66	70	67	68	80	63	71	77	84	86	82	71
Combined (pre-test)	56	50	57	54	62	62	64	63	64	68	60	64	67	64	69	67	62
Combined (post-test)	81	84	82	82	79	80	83	81	78	87	87	84	90	94	94	93	85

Remark: G–grammar, T–terminology, S–structure, R–reproduction, Q–questioning, F–fluency, U–general vocabulary, P–specialized vocabulary, A–application.

and whether there are statistically significant differences between the groups with regard to the results achieved in pre-test and post-test.

First, we verified the statistical significance of the difference in the results achieved by respondents in pre-test and post-test in all three groups in each study area.

We proceeded analogously when verifying the statistical significance of differences in the results of pre-test and post-test in other parts (Writing, Listening, Speaking, Vocabulary) not only together, but also separately for each group (On-line, Standard, Combined). The results are clearly recorded in the following table.

Table 2: Results of Wilcoxon test (pre-test vs. post-test) in each group (values of Z and p -values)

Object of Research	Writing		Listening		Speaking		Vocabulary	
	v. of Z	p -value	v. of Z	p -value	v. of Z	p -value	v. of Z	p -value
On-line	3.92	0.0001	3.81	0.0030	3.82	0.0015	3.92	0.0031
Standard	3.92	0.0001	3.92	0.0011	3.92	0.0010	3.86	0.0015
Combined	3.06	0.0010	3.06	0.0012	3.06	0.0001	3.06	0.0001
Total	6.27	0.0001	6.25	0.0001	6.21	0.0001	6.26	0.0001

Since p -value is smaller than significance level alpha 0.01 in all cases, we reject the tested hypothesis α at the significance level 0.01 in favour of the alternative hypothesis H_1 . We have shown a statistically significant increase in knowledge level of respondents in all areas – Writing, Listening, Speaking, as well as Vocabulary.

We analyzed the differences among the three groups of respondents in the next step, with regard to achievements in pre-test and post-test. The results of pre-test and post-test achieved by the respondents in each group are illustrated graphically in Fig. 1 and Fig. 2 on the next page.

There are differences in the achieved results of pre-test and post-test as may be seen in Figure 1 and 2. The use of statistical methods should ascertain which of these differences are statistically significant, i.e. whether the level of factor that represents different groups of respondents has an impact on the average observed character. If the number of files was being compared (the number of factor levels) is k , then the problem can be formulated in the form of a statistical hypothesis as follows. The tested hypothesis is

H_0 : Distributions of k essential files are identical. When the tested hypothesis is compared with an alternative hypothesis H_1 : Not all k distributions

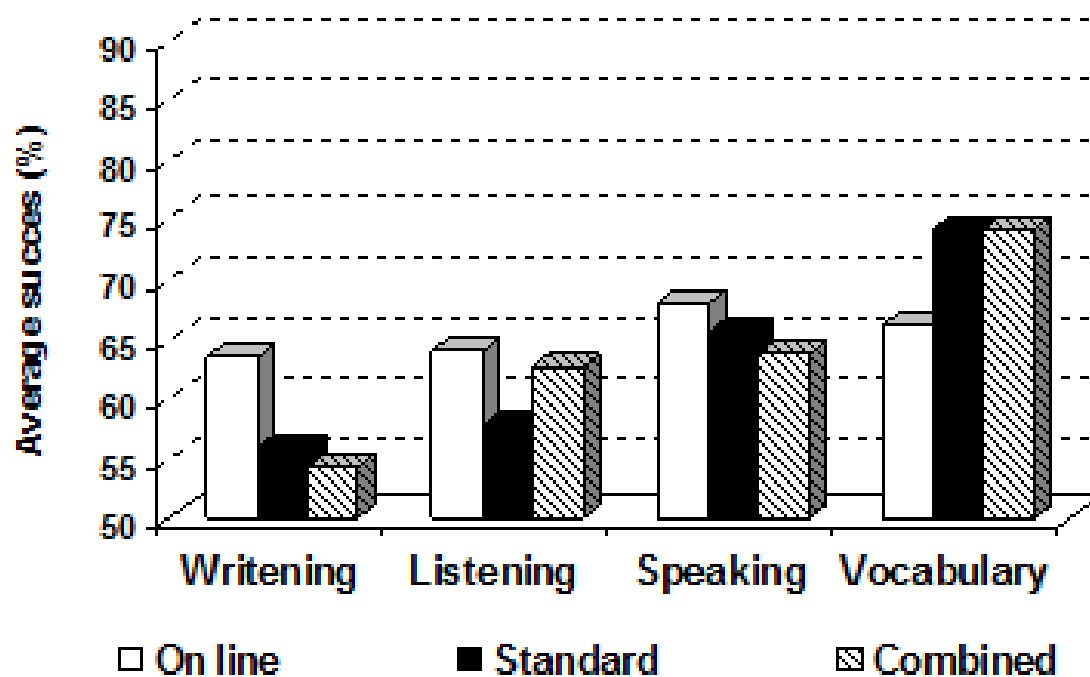


Figure 1: Results of pre-test in the groups of respondents from individual areas

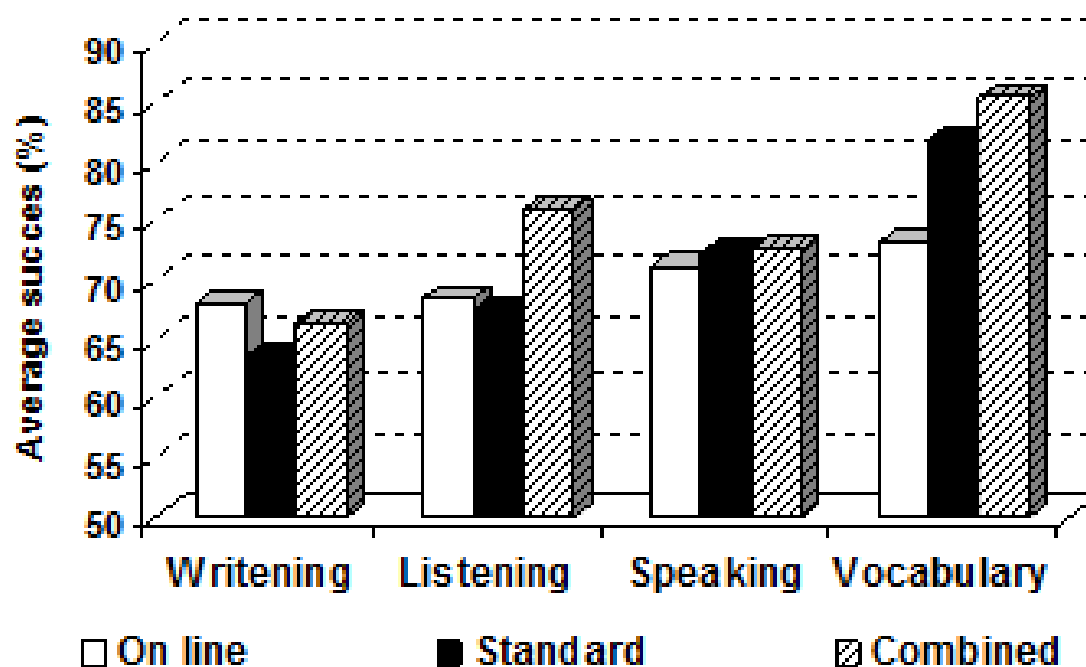


Figure 2: Results of post-test in the groups of respondents from individual areas

are identical. The same problem can be formulated as follows H_0 : Average values of the observed character do not depend on the levels of factor H_1 : Average values depend on the levels of factor.

Since the assumption of normal distribution of the observed characters for testing the null hypothesis is not justified, the non-parametric Kruskal-Wallis test was used.

The Kruskal-Wallis test should verify whether the three subgroups formed according to the levels of factor – the type of group (On-line, Standard, Combined) significantly differ in the observed character – pre-test success rate (= number of points). Therefore $k = 3$, whereby $n_1 = 20$, $n_2 = 20$, $n_3 = 12$, $n = n_1 + n_2 + n_3 = 32$ respondents.

The value of the test criterion $K = 0.433566$ and the probability value $p = 0.8051$ was calculated after the entry of the pre-test results using the Kruskal-Wallis test. Since the calculated value of probability p is greater than 0.05, the null hypothesis H_0 cannot be rejected at the significance level $\alpha = 0.05$. This means that the groups were not statistically significantly different in the pre-test results. All three groups reached approximately the same percentage in the pre-test results.

The results of post-test were evaluated in the same way. The value of the test criteria $K = 11.27376$ and the level of significance p -value of 0.0036 was calculated by means of the Kruskal-Wallis test. As the calculated probability value p is smaller than 0.05, the null hypothesis H_0 was rejected at the significance level of 0.01 in favor of the alternative hypothesis. Tests confirmed that each group of respondents (On-line, Standard, Combined) is statistically significantly different with regard to success in the post-test.

Given that the tests confirmed a statistically remarkable difference among the observed groups, it is obvious which groups are statistically significantly different from each other. This question was answered by the Kruskal-Wallis test of multiple comparisons (Table 3).

Table 3: p -values of Kruskal-Wallis test of multiple comparisons in post-test

	On-line	Standard	Combined
On-line		0.997000	0.015937*
Standard	0.997000		0.004079*
Combined	0.015937*	0.004079*	

It is visible from the table that they are statistically significantly different between the On-line group and the Combined group, given the results in post-test, ($p = \mathbf{0.015937^*}$), as well as between the Standard and the Combined group ($p = \mathbf{0.004079^*}$), but between the On-line and the Standard group ($p = 0.997000$) there is no statistically significant difference given the results in post-test.

It is obvious that the isolated application of classical learning or the isolated application of e-learning teaching model have approximately similar or comparable effects and conclusions. Conversely, the combination of modes has resulted in significant improvement of language training and language skills of students.

4. Conclusion

Before summing up some of the results of the research, it must be stressed that the statistical method used in this research was not new. The aim of the research was not to introduce new statistical methods but to apply an existing one (Kruskal-Wallis test) in order to test and confirm three models of education (the in-class model, the e-learning model and the combined model).

The authors believe that the efficiency of the combined model of education is based on the following facts:

- a) the use and exercise of all senses within the process of learning,
- b) the fact that although an e-learning model was tested, students still had the possibility of one-to-one communication and tuition with a person,
- c) the fact that e-learning enabled students to study continuously also in time period in-between the classes (in-class education).

The use of the statistical method confirmed that the respondents' level of knowledge has increased statistically significantly in all groups and in all areas – Writing, Listening, Speaking, and Vocabulary.

The Kruskal-Wallis test showed that the combination of in-class teaching and e-learning has resulted in significant improvement of students' language training and language skills.

References

- [1] Anděl, J. (2003). *Statistické metody = Statistical methods*. MatfyzPress, Praha.
- [2] Bañados, E. (2006). *A Blended-learning Pedagogical Model for Teaching and Learning EFL Successfully Through an Online Interactive Multimedia Environment*. *Calico Journal*, Vol. 23, No. 3.
- [3] Bleimann, U. (2004). *Atlantis University – A New Pedagogical Approach Beyond-Learning*. Campus-Wide Information Systems, Emerald, Bradford, S., Vol. 21, No. 5, pp. 191–195.
- [4] Hollander, M. and Wolfe, D. A. (1999). *Nonparametric Statistical Methods*. Second Edition, New York: John Wiley & Sons.
- [5] Kurucová, Z. (2010). *E-learning ako prostriedok a zdroj hybridných dopadov na výchovu a vzdelávanie = E-learning as a means of power and hybrid effects on education*. In *Notitiae ex Academia Bratislavensi Iurisprudentiae*, Bratislava, I/2010, pp. 56–60.
- [6] Kurucová, Z. (2010). *Vybrané modely E-learningu vo vyspelých spoločnostiach = Selected models of E-learning in advanced societies*. In *Notitiae ex Academia Bratislavensi Iurisprudentiae*, Bratislava, III/2010, pp. 69–80.
- [7] Kurucová, Z. (2011). *Úvod do tvorby efektívneho e-learningového kurzu = Introduction to creating an effective e-learning course*. In *Notitiae ex Academia Bratislavensi Iurisprudentiae*, Bratislava, IV/2010, pp. 88–95.
- [8] Kurucová, Z. (2011). *Aplikácia nástrojov e-learningu s využitím IKT vo vzdelávaní*. Dizertačná práca, Paneurópska vysoká škola, Fakulta masmédií, Bratislava 2011.
- [9] Marshall, S. and Mitchell G. (2002). *Ane-Learning maturity model*. Proceedings of Ascilite 2002, Auckland, New Zealand.
- [10] Meredith, S. and Newton, B. (2003). *Models of e-learning: Technology promise vs learner needs literature review*. In: *The International Journal of Management Education*, pp. 43–56.
- [11] Stehlíková, B., Tirpáková, A., Poměnková, J. and Markechová, D. (2009). *Metodologie výzkumu a statistická inference = Research Methodology and Statistical Inference*. Brno: Mendel University in Brno, 325 p.
- [12] Tirpáková, A. and Markechová, D. (2008). *Štatistika v praxi = Statistics in practice*. FPV UKF in Nitra, 390 p.
- [13] Webster, R. and Sudweeks, F. (2006). *Teaching for e-Learning in the Knowledge Society: Promoting Conceptual Change in Academics' Approaches to Teaching*. *Current Developments in Technology-Assisted Education*. [online]
- [14] <http://www.formatex.org/micte2006/pdf/631-635.pdf>
[accessed 15th May 2007]

EKONOMETRICS IN OPEN SOURCE SOFTWARE

EKONOMETRIE V OPEN SOURCE SOFTWARE

Petr Klímek, Martin Kovářík

Adresa: doc. Ing. Petr Klímek, Ph.D., Ing. Martin Kovářík, Ph.D.

Univerzita Tomáše Bati ve Zlíně, Fakulta managementu
a ekonomiky, nám. T. G. Masaryka 5555, 760 01 Zlín

E-mail: klimek@fame.utb.cz, mlikovarik@fame.utb.cz

Abstract: This paper introduces software Gretl which is an open-source software. It is focused on time series analysis. Gretl is a suitable software tool for econometrics education at universities. Programme installation is described at the beginning of the paper followed by pointing at built-in manuals, functions and different data formats description. Total software evaluation can be found at the end of this article.

Keywords: Gretl, Econometrics, Time Series, ARIMA.

Abstrakt: Tento článek představuje program Gretl, který je volně dostupný na Internetu. Je zaměřen především na analýzu časových řad. To jej činí vhodným nástrojem pro použití v kurzech ekonometrie na vysokých školách. Po stručném úvodu je zde popsána instalace programu, jeho vestavené manuály a funkce, práce s různými formáty dat. Na závěr následuje celkové zhodnocení programu.

Klíčová slova: Gretl, ekonometrie, časové řady, ARIMA.

1. Úvod

Program Gretl (z anglického GNU **R**egression, **E**conometrics and **T**ime-series **L**ibrary) je open-source statistický softwarový paket, který se hlavně zaměřuje na statistické metody, které se používají v ekonometrii. Původně existoval jako paket **E**conometrics **S**oftware **L**ibrary (ESL), který vyvinul profesor Ramu Ramanathan na University of California v San Diegu. Ve vývoji dále pokračoval profesor Allin Cottrell na Wake Forest University.

Po dalším vývoji se objevila první stabilní verze programu Gretl 1.0, která byla uveřejněna 15. listopadu 2002 (<http://gretl.sourceforge.net/ChangeLog.html>). Tato verze je v neustálém vývoji a tento článek se zabývá poslední verzí 1.8.0, která byla vyvinuta v roce 2009. Kromě profesora Allina Cottrella je dalším hlavním vývojářem programu Riccardo Jack Lucchetti (docent ekonometrie na Università Politecnica delle Marche, Ancona, Itálie).

Motto Gretlu je anglicky „by econometricians, for econometricians“, tedy „od ekonometrů pro ekonometry“. Původně byl program určen pro systém Linux, ale dnes jsou dostupné i verze pro Microsoft Windows a Mac OS. Je napsán v programovacím jazyce C a je volně dostupný podle obecné veřejné licence GPL (anglicky General Public Licence; <http://www.gnu.org/licenses/gpl.html>). Také společnost Free Software Foundation (<http://www.fsf.org/>) přijala program Gretl jako program GNU.

Kromě anglické verze si můžeme vybrat další jazykové mutace v baschičtině, francouzštině, němčině, italštině, polštině, portugalštině a španělštině.

Jako mnoho dalších open-source programů můžeme Gretl najít také na oficiálních stránkách SourceForge:

<http://sourceforge.net/>

Zdrojový kód programu lze nalézt na stránce:

<http://gretl.cvs.sourceforge.net/gretl/gretl/>

Pro diskusi o vývoji Gretlu byla zřízena následující adresa:

<http://lists.wfu.edu/pipermail/gretl-devel/>

2. Instalace

Program Gretl se dá stáhnout ze serveru <http://gretl.sourceforge.net/> jako samospustitelný instalační soubor (jedná se o jeden soubor o cca 10 MB) pro OS Microsoft Windows, dále pak jako disc image pro Mac OS nebo Debian GNU/Linux paket pro OS Linux. Instalace pro Windows XP je automatická a bezproblémová. Také instalace na PC s OS Windows Vista je stejně snadná.

Pro zobrazení grafů Gretl používá jiný open-source program pro kreslení Gnuplot. Tento program je zahrnut v instalačním souboru pro Microsoft Windows a je automaticky nainstalován s Gretlem.

Po instalaci Gretlu si uživatel může dále doinstalovat dva volitelné speciální softwarové pakety: X-12-ARIMA od US Census Bureau z roku 2007 a nebo TRAMO/SEATS (anglicky **T**ime series **R**egression with **A**rima **N**oise, **M**issing values and **O**utliers/**S**ignal **E**xtraction in **A**rima **T**ime **S**eries). Tyto pakety jsou vhodné pro časové řady s měsíční nebo nižší frekvencí sledování. Tyto dva pakety se dají stáhnout z webových stránek programu Gretl jako samospustitelné soubory pro OS Microsoft Windows.

3. Manuály a podpora

Program Gretl obsahuje v Help menu dva dokumentační soubory v pdf formátu: Gretl User's Guide [1] a Gretl Command Reference [2]. Oba dokumenty

jsou psány v angličtině. Druhý jmenovaný je také dostupný přímo v helpu programu jako hypertext (viz obrázek 1). První dokument – The User’s Guide – se skládá ze čtyř částí. V první části jsou popsány základy práce s programem včetně práce s grafickým rozhraním a s příkazovým řádkem, práce s daty spolu se skriptovacím jazykem Gretlu. Druhá část je věnována aplikacím ekonometrických metod v Gretlu spolu s detailním popisem různých metod odhadu parametrů, které Gretl umožňuje. Třetí část se zabývá technickými detaily, především propojením s L^AT_EXem. Poslední část obsahuje několik příloh s dalšími informacemi o různých formátech dat v Gretlu, o numerické přesnosti programu apod. Dokument Gretl User’s Guide v pdf dává uživateli dobrý a podrobný přehled o práci s programem a o jeho možnostech různých typů analýz. Jsou zde uvedeny mnohé příklady, které jsou podrobně doplněny screenshoty a komentáři k nim. Dokument Gretl User’s Guide je užitečnou pomůckou jak pro začátečníky, tak i pokročilé uživatele.

4. Jak program pracuje

V porovnání s jinými open-source programy (jako například R) je program Gretl mnohem uživatelsky příjemnější a jeho ovládání je intuitivní. Uživatel má pohodlný přístup skrz grafické prostředí programu k jednotlivým analýzám a funkcím programu. Na obrázku 2 vidíme hlavní okno programu Gretl pro naši analýzu. Abychom odhadli například lineární regresní model pomocí metody nejmenších čtverců, zvolíme v programu Model — Ordinary Least Squares. . . na hlavní liště (viz obrázek 3) a dále v dialogovém okně vybereme proměnné, které chceme zahrnout do modelu. Výsledky metody nejmenších čtverců získáme ve výstupu na obrázku 4 jako samostatné okno. Odtud si uživatel může zvolit z několika nových nabídek. V nich se dají provádět další testy a analýzy dat, vykreslit grafy výsledků, uložit výsledky analyzovaných dat (například grafy reziduí nebo vyrovnaných hodnot) pro pozdější použití.

Například z menu Analysis — Bootstrap. . . se dá vstoupit do dialogového okna pro bootstrap analýzu modelu (viz obrázek 5). Všechny analýzy, které jsou prováděny přes klasické menu, se ukládají rovněž jako script v příkazovém řádku (viz obrázek 6). Ten může být znovu spuštěn nebo uložen pro pozdější použití. Kromě klasické menu lišty program Gretl obsahuje tlačítka pro speciální příkazy a analýzy, která jsou umístěna v dolní části hlavního okna programu (viz obrázek 2). Ta dávají uživateli rychlejší přístup k často používaným příkazům jako například vytvoření nového scriptu, otevření okna konsole Gretlu (viz obrázek 7) a další. Jestliže chceme data zobrazit nebo editovat, lze otevřít vybrané proměnné v samostatném okně (viz obrázek 8).

5. Formáty

Gretl používá svůj vlastní formát založený na XML, ale lze samozřejmě **importovat** i jiné formáty dat jako například ASCII, Open Document Format (ODF) pro tabulkové procesory, Microsoft Excel, Gnumeric, EViews work, GNU Octave, Stata, SPSS a JMulTi. Gretl umožňuje dále **export** dat do programů CSV, R, GNU Octave, JMulTi and PcGive. Program Gretl podporuje širokou škálu statistických metod hlavně se zaměřením na ekonometrii, a v ní především na modely regresní. Kromě známé metody nejmenších čtverců (OLS – Ordinary Least Squares) zde nalezneme i váženou metodu nejmenších čtverců, dvoustupňovou metodu nejmenších čtverců a další (viz obrázek 7). Rovněž jsou k dispozici i nelineární modely jako jsou logit, probit, tobit, heckit, Poisson, logistic a nelineární nejmenší čtverce (viz obrázek 9). Dále pod hlavním menu v modelech (Models) lze najít také metodu maximální věrohodnosti, simultánní rovnice (SUR), dvou a třístupňové nejmenší čtverce a další. Pro panelová data lze odhadnout například modely s pevnými nebo náhodnými efekty a dynamické modely. Pro analýzu časových řad nalezneme v programu Gretl širokou nabídku metod včetně filtrování.

6. Co program neobsahuje

Ačkoliv poskytuje Gretl mnoho možností pro ekonometrickou analýzu a analýzu časových řad, pár drobností mu chybí. Například analýza přežívání (Survival Analysis). Také editování dat by mohlo být snadnější. Například v současné verzi lze kopírovat a vložit pouze jedno pozorování v rámci jedné proměnné. Dále nelze vkládat data v textovém formátu (pouze v číselném).

7. Závěr

Program Gretl je statistický software, který se velice snadno používá. Poskytuje mnoho různých druhů analýz, které jsou důležité v ekonometrii a v analýze časových řad. Jeho přívětivost k uživateli a také skutečnost, že program je k dispozici zdarma, ho činí ideálním softwarovým nástrojem pro výuku ekonometrie a analýzu časových řad na vysokých školách. Protože si studenti mohou program zdarma stáhnout do svých domácích počítačů, nejsou závislí v přípravě na počítačových učebnách, které jsou neustále obsazeny.

Program je zcela určitě také vhodný pro profesionální ekonometry a ekonomy. Skutečnosti, že program je open-source a že je naprogramován v jazyce C, který je hodně znám, jej činí vhodným pro další vývoj, kterému napomohou statistikové a ekonometři, kteří jsou v tomto programovacím jazyce zbláhli. Jistě v budoucích verzích programu lze očekávat vložení dalších

modelů a analýz. Budoucnost Gretlu je proto dle našeho názoru velmi slibná. Lze jej tedy zcela bez obav doporučit jak studentům, tak i profesionálním uživatelům.

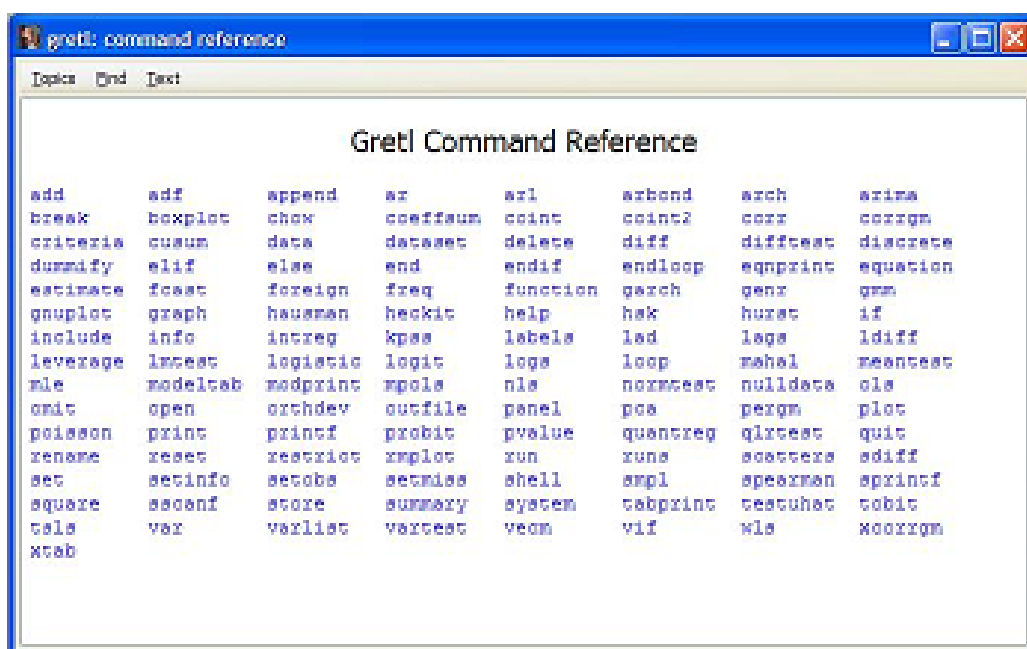
Poděkování

Tento příspěvek vznikl za podpory Interní grantové agentury UTB, projekt č. IGA/73/FaME/10/D, pod názvem *Rozvoj využívání matematicko-statistických metod v řízení kvality*.

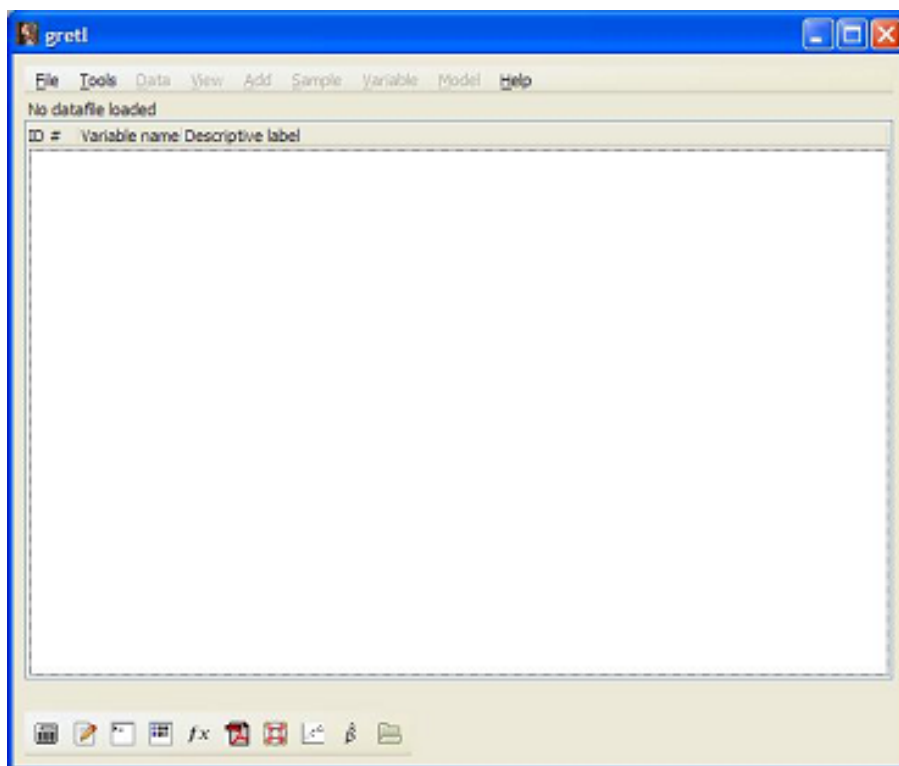
Seznam použité literatury

- [1] Cottrell, A., Lucchetti, R. *Gretl User's Guide*. 2008.
- [2] Cottrell, A., Lucchetti, R. *Gretl Command Reference*. 2008.

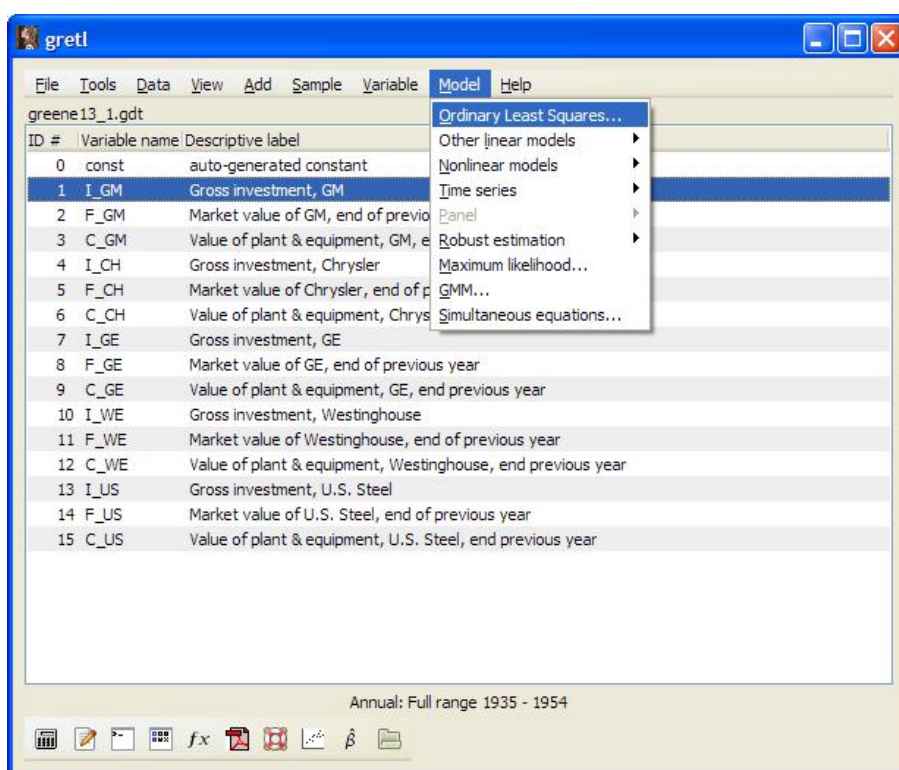
Obrazové přílohy



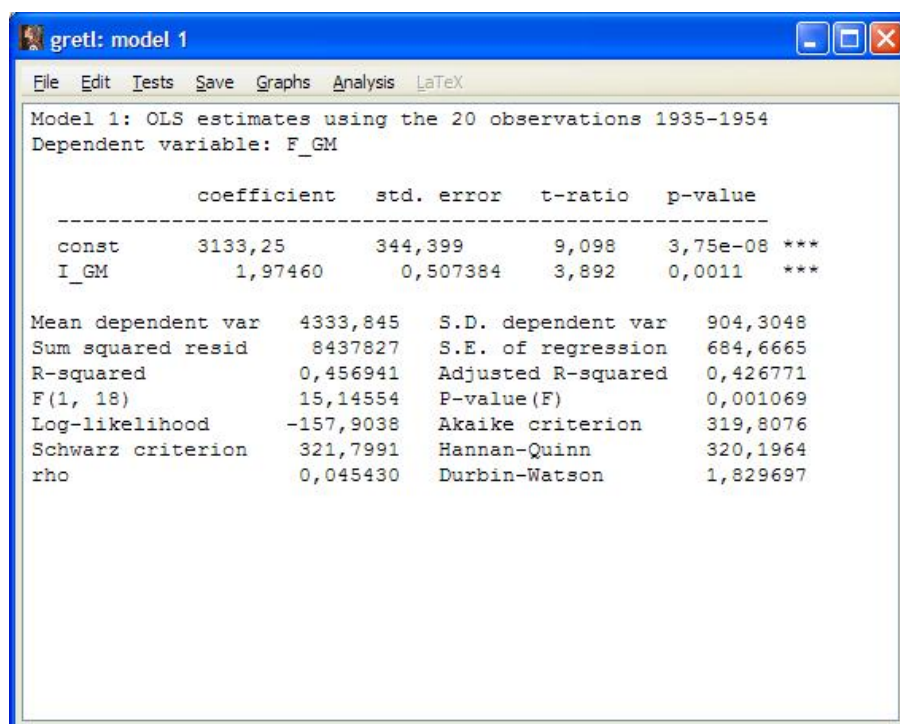
Obrázek 1: Odkazy na příkazy v programu Gretl.



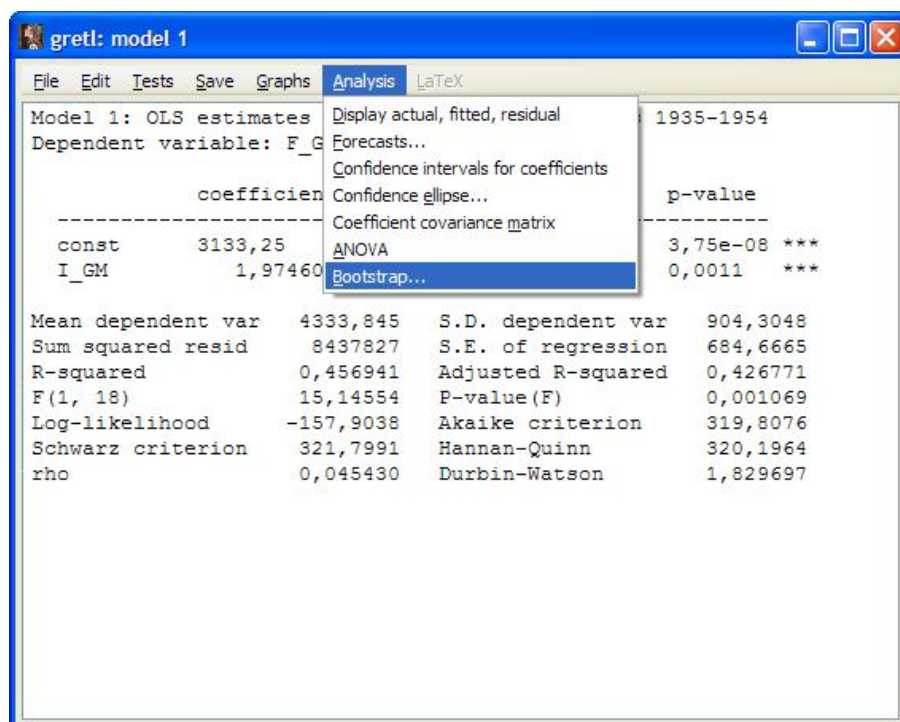
Obrázek 2: Úvodní okno při spuštění programu.



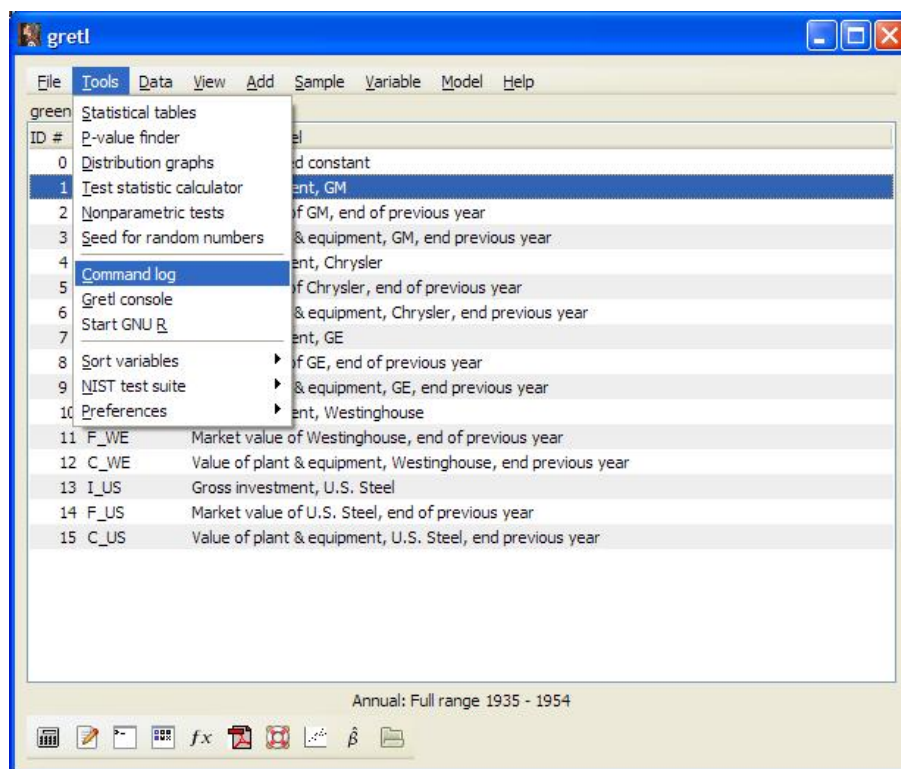
Obrázek 3: Odhad lineárního regresního modelu pomocí metody nejmenších čtverců v programu Gretl.



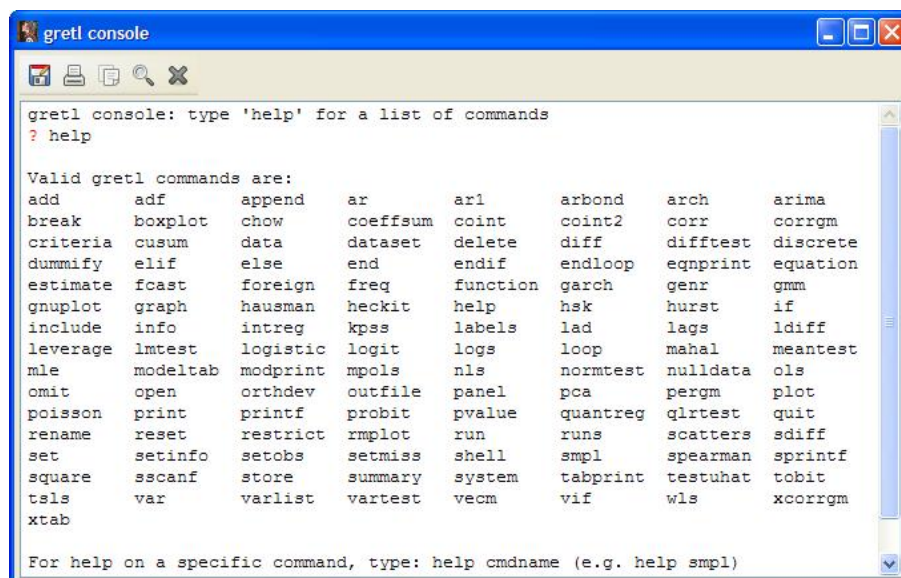
Obrázek 4: Výstupní diagnostika odhadu lineárního regresního modelu.



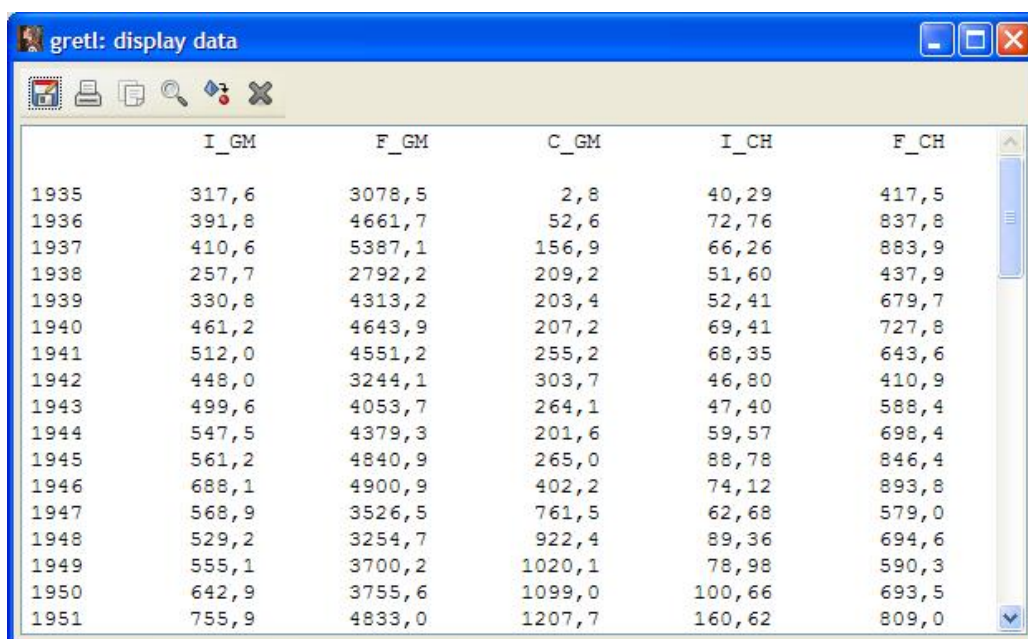
Obrázek 5: Dialogové okno pro bootstrap analýzu modelu.



Obrázek 6: Všechny analýzy, které jsou prováděny přes klasické menu, se ukládají rovněž jako script v příkazovém řádku.



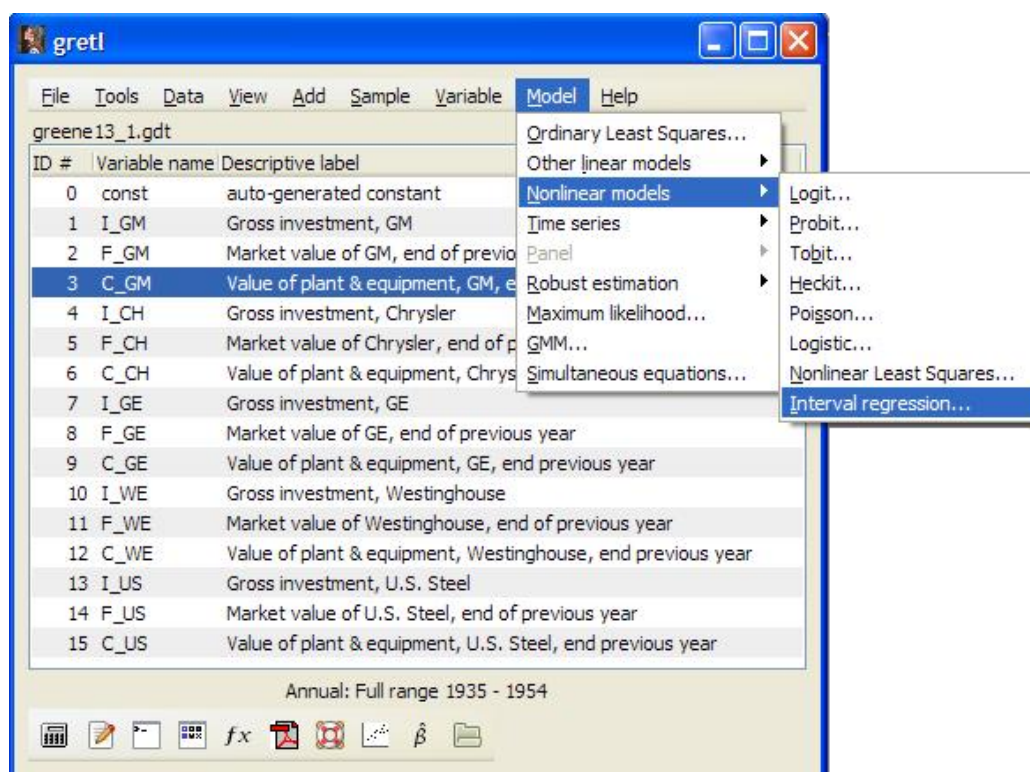
Obrázek 7: Vytvoření nového scriptu otevřením okna console Gretlu.



The screenshot shows the 'gretl: display data' window. It contains a table with 6 columns: I_GM, F_GM, C_GM, I_CH, and F_CH. The rows represent years from 1935 to 1951. The data values are as follows:

	I_GM	F_GM	C_GM	I_CH	F_CH
1935	317,6	3078,5	2,8	40,29	417,5
1936	391,8	4661,7	52,6	72,76	837,8
1937	410,6	5387,1	156,9	66,26	883,9
1938	257,7	2792,2	209,2	51,60	437,9
1939	330,8	4313,2	203,4	52,41	679,7
1940	461,2	4643,9	207,2	69,41	727,8
1941	512,0	4551,2	255,2	68,35	643,6
1942	448,0	3244,1	303,7	46,80	410,9
1943	499,6	4053,7	264,1	47,40	588,4
1944	547,5	4379,3	201,6	59,57	698,4
1945	561,2	4840,9	265,0	88,78	846,4
1946	688,1	4900,9	402,2	74,12	893,8
1947	568,9	3526,5	761,5	62,68	579,0
1948	529,2	3254,7	922,4	89,36	694,6
1949	555,1	3700,2	1020,1	78,98	590,3
1950	642,9	3755,6	1099,0	100,66	693,5
1951	755,9	4833,0	1207,7	160,62	809,0

Obrázek 8: Možnost editace vybrané proměnné v samostatném okně.



Obrázek 9: Možnost použití nelineárních modelů jako jsou logit, probit, tobit, heckit, Poisson, logistic a nelineární nejmenší čtverce.

TVORBA INTERNETOVÝCH UČEBNIC PRAVDĚPODOBNOСТИ A STATISTIKY JAKO SOUČÁST PROJEKTU MI21

Martina Litschmannová

Adresa: Ing. Martina Litschmannová, VŠB-TU Ostrava, Fakulta elektrotechniky a informatiky, Katedra aplikované matematiky

E-mail: martina.litschmannova@vsb.cz

Abstrakt: Jedním z výstupů projektu „Matematika pro inženýry 21. století – inovace výuky matematiky na technických školách v nových podmínkách rychle se vyvíjející informační a technické společnosti“ (zkráceně MI21) je vytvoření 15 základních matematických modulů, 14 specializovaných matematických modulů, 12 mezioborových modulů a 20 modulů pro pedagogy. Výukové moduly jsou vytvářeny jednak ve formě určené pro tisk, jednak v elektronické formě vhodné pro prohlížení na monitoru. Obrazovkové verze modulu jsou doplňovány řadou multimediálních prvků, jako jsou animace, výpočetní applety, interaktivní testy, videa či 3D grafika. Tento příspěvek popisuje dva moduly věnované vybraným kapitolám z pravděpodobnosti a úvodu do statistiky, jejichž tiskové verze, připravené animace a výpočetní applety byly aplikovány ve výuce během tzv. pilotních kurzů v letním semestru školního roku 2010/2011.

1. Úvod

V rámci projektu „Matematika pro inženýry 21. století – inovace výuky matematiky na technických školách v nových podmínkách rychle se vyvíjející informační a technické společnosti“, jehož hlavními řešiteli jsou Vysoká škola báňská – Technická univerzita Ostrava a Západočeská univerzita v Plzni, je realizováno 6 klíčových aktivit:

- Tvorba výukových modulů.
- Zavádění výukových modulů do výuky v pilotních kurzech.
- Oponentury výukových modulů a vyhodnocování pilotních kurzů.
- Konečné úpravy výukových modulů na základě výsledků oponentur a hodnocení studentů a pedagogů.
- Propagace vytvořených výukových modulů mezi studenty středních škol, vysokých škol a pedagogy.
- Vytvoření modulů pro pedagogy a jejich prezentace na Seminári o výuce matematiky.

V rámci první klíčové aktivity, tvorby výukových modulů, je vytvářeno 41 výukových modulů (15 základních matematických modulů, 14 specializovaných matematických modulů a 12 mezioborových modulů), které by se měly stát základem pro inovaci matematických a odborných kurzů prezenčního i kombinovaného studia. V současné době již je většina výukových modulů ve formě určené pro tisk vytvořena a jsou průběžně aplikovány ve výuce během tzv. pilotních kurzů. Tyto pilotní kurzy probíhají na VŠB-TU i na pracovišti partnera, na ZČU Plzeň, od letního semestru školního roku 2010/2011 do letního semestru školního roku 2011/2012. Na konci běhu pilotních kurzů vyplňují studenti evaluační dotazníky, v nichž hodnotí přínos inovovaných výukových modulů – srozumitelnost a dostatečnost textu, srozumitelnost a množství řešených příkladů, užitečnost úloh k samostatné práci, grafickou úpravu a míru používání textu. Statistické vyhodnocení dotazníků je následně předáváno garantům modulů, kteří tak získávají zpětnou vazbu od studentů.

V zimním semestru 2011/2012 je většinou autorů připravována elektronická forma výukových modulů, vhodná pro prohlížení na monitoru. Elektronická podoba výukových modulů (tzv. obrazovková verze) by měla textům přidat další dimenze – multimedialitu (zvuk, videa), hypertextové provázání textu, propojení s internetem, dynamiku a interaktivitu (animace, 3D grafika, interaktivní testy). Materiály v elektronické formě i formě pro tisk jsou v rámci realizace klíčové aktivity „Propagace vytvořených výukových modulů mezi studenty SŠ, VŠ a pedagogy“ umísťovány na webových stránkách projektu (<http://mi21.vsb.cz/>).

Zmíňme ještě klíčovou aktivitu „Vytvoření modulů pro pedagogy a jejich prezentace na Seminári o výuce matematiky“. Během realizace této klíčové aktivity je vytvářeno 20 modulů pro pedagogy. Tyto moduly se zaměřují na obtížnější aplikace, mezioborové a mezipředmětové vazby. Učitelé matematiky i odborných předmětů tím získávají možnost lépe pochopit možnosti využití moderní matematiky v odborných předmětech. Všechny moduly jsou průběžně prezentovány na Seminári o výuce matematiky a vyvěšovány na webové stránky projektu.

2. Výukový modul „Vybrané kapitoly z pravděpodobnosti“

Výukový modul „Vybrané kapitoly z pravděpodobnosti“ je určen pro studenty technických oborů vysoké školy. Měl by poskytnout výchozí představu o základních pojmech a úlohách spadajících do oblasti pravděpodobnosti (kombinatorika, úvod do teorie pravděpodobnosti, náhodná veličina a ná-



Obrázek 1: Náhled na webové stránky projektu „Matematika pro inženýry 21. století“

hodný vektor, vybraná rozdělení náhodné veličiny – cca 190 stran). Obtížnější části výkladu jsou prezentovány jen s nejnutnější mírou formálních prvků, mnohá odvození a důkazy jsou zařazeny pouze do kapitol určených pro zájemce o pozadí předkládaných vztahů. V současné době je dokončena tisková verze modulu ($\text{T}_{\text{E}}\text{X}$) a je připravena řada animací a výpočetních appletů pro verzi obrazovkovou. Animace jsou připraveny ve formátech SWF (Flash), JAR (Java) a XLSX (Microsoft Excel). Nejjednodušší varianta Flash animací obsahuje pouze sekvenci obrázků (vytvořených v Microsoft PowerPoint), mezi kterými se lze přepínat. Jde v podstatě o komentované řešené příklady. Tímto způsobem byly vytvořeny animace:

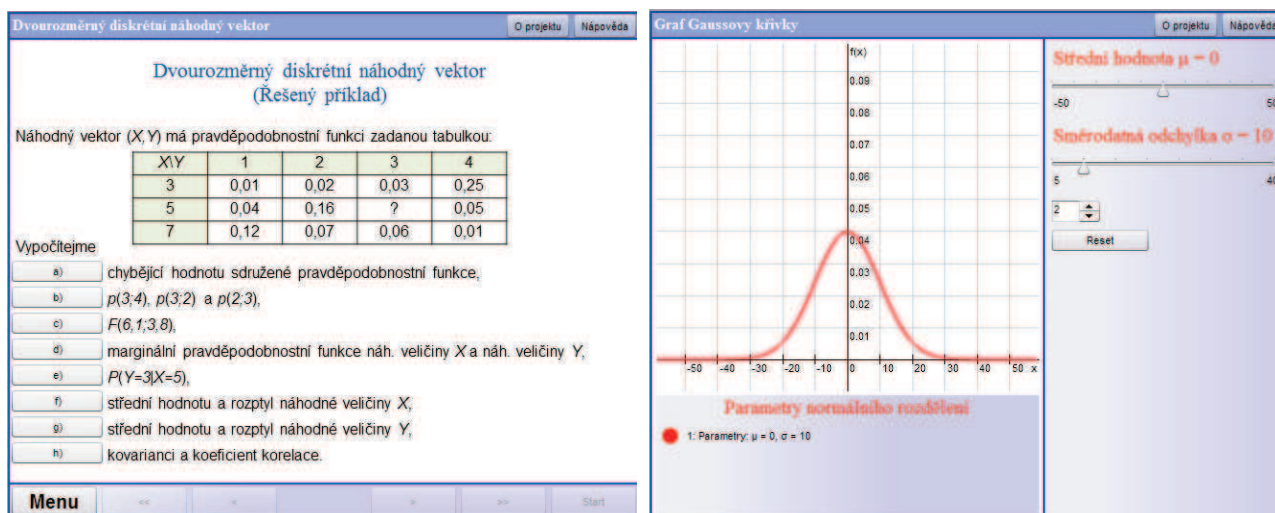
- Aplikace Bayesovy věty v biomedicíně.
- Diskrétní náhodná veličina.
- Dvourozměrný diskrétní náhodný vektor.

Druhým typem animací jsou interaktivní animace, tj. animace, které obsahují prvky pohybující se v čase:

- Klasická definice pravděpodobnosti (Flash, demonstrace výpočtu pravděpodobnosti pomocí klasické definice pravděpodobnosti).
- Korelační koeficient (Java, animace umožňuje graficky určit možné realizace náhodného vektoru (X, Y) a následně odhadnout korelaci jeho složek).
- Přejít mezi histogramem a hustotou pravděpodobnosti (Java).
- Graf Gaussovy křivky (Flash, vliv parametrů normálního rozdělení na křivku hustoty pravděpodobnosti).
- Spojitá rozdělení (Microsoft Excel, vliv parametrů vybraných osmi rozdělení na křivky hustoty pravděpodobnosti a distribuční funkce).

Výše uvedené animace jsou doplněny výpočetním appletem:

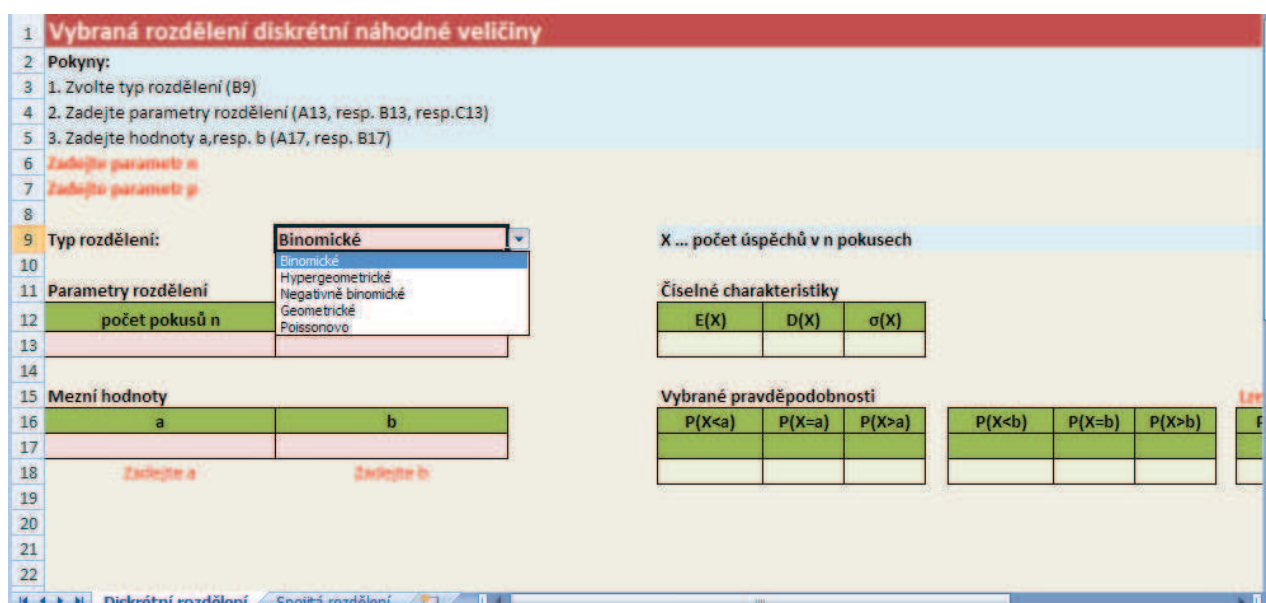
- Vybraná rozdělení pravděpodobnosti (Microsoft Excel, umožňuje výpočet základních číselných charakteristik, pravděpodobností, resp. hustoty pravděpodobností a kvantilů pro vybraná rozdělení).



Obrázek 2: Ukázka úvodních stránek Flash animací typu: sekvence obrázků (vlevo) a jednoduché interaktivní animace (vpravo)

3. Výukový modul „Úvod do statistiky“

Studiem výukového modulu „Úvod do statistiky“ navazují studenti na studium modulu „Vybrané kapitoly z pravděpodobnosti“. Modul by měl umožnit, aby si studenti učinili výchozí představu o základních pojmech a úlohách



Obrázek 3: Ukázka úvodní stránky výpočetního appletu (Microsoft Excel)

spadajících do oblasti statistiky (statistické zjišťování, explorační analýza, testování hypotéz, vybrané testy parametrických i neparametrických hypotéz, analýza závislosti, úvod do korelační a regresní analýzy – cca 330 stran). V současné době je dokončena tisková verze modulu (T_{EX}).

Zároveň je již vytvořena řada animací a výpočetních appletů pro obrazovkovou verzi:

- Výběrové charakteristiky (Java, demonstrace způsobu výpočtu výběrových charakteristik numerické proměnné a jejich souvislosti s histogramem a krabicovým grafem).
- Paretova analýza (Flash, řešený příklad).
- Rozdělení průměru (Java, demonstrace centrální limitní věty).
- Intervalové odhady (Java, demonstrace vlivu rozsahu výběru a spolehlivosti odhadu na interval spolehlivosti).
- Intervalové odhady jednovýběrové (Microsoft Excel, výpočetní applet).
- Intervalové odhady rozdílu, resp. podílu parametrů dvou populací (Microsoft Excel, výpočetní applet).
- Chyba I. a II. druhu (Flash, animace prostřednictvím testu střední hodnoty (při známém rozptylu) demonstruje vliv testované hodnoty, populačního rozptylu, hodnoty testované v jednoduché alternativní hy-

potéze, rozsahu výběru a kritické hodnoty průměru na velikost chyb I. a II. druhu).

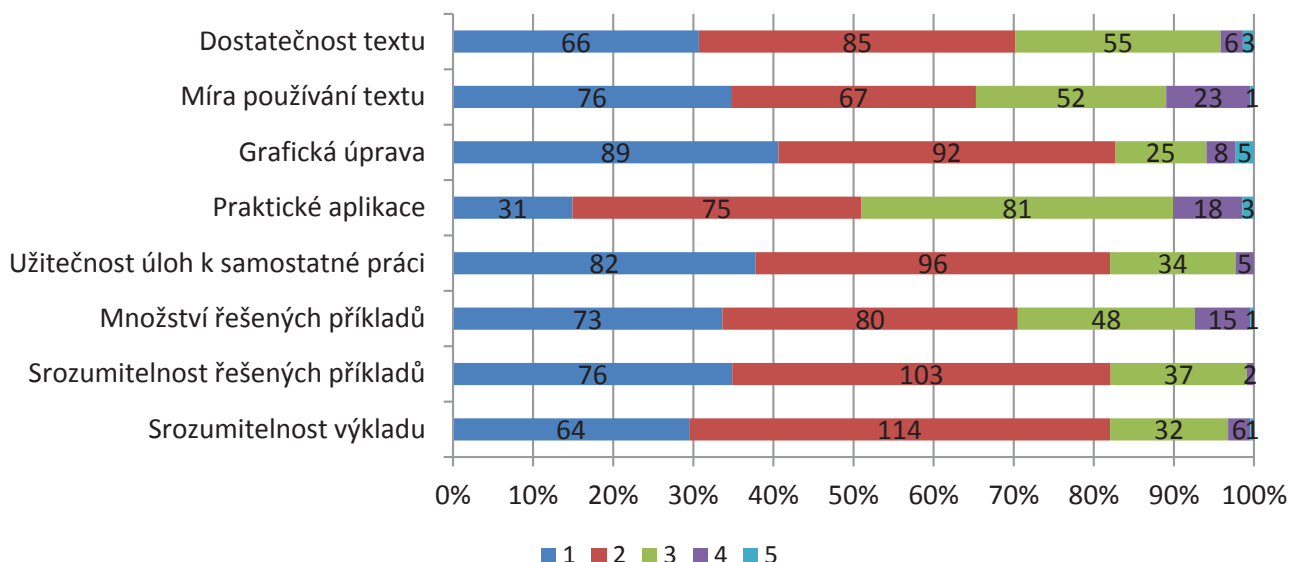
- Operační charakteristika (Flash, animace prostřednictvím testu střední hodnoty (při známém rozptylu) demonstruje vliv testované hodnoty, populačního rozptylu, rozsahu výběru a zvolené hladiny významnosti na velikost chyby II. druhu a tvar operační charakteristiky a silofunkci).
- p -hodnota (Flash, animace prostřednictvím testu střední hodnoty (při známém rozptylu) demonstruje výpočet p -hodnoty při různých typech alternativy, vliv testované hodnoty c , populačního rozptylu σ^2 a rozsahu výběru n na tvar křivky hustoty nulového rozdělení průměru (modrá křivka), tj. rozdělení průměru za předpokladu platnosti nulové hypotézy, vliv výběrového průměru na p -hodnotu a tím i na rozhodnutí o výsledku testu při dané hladině významnosti α a vliv zvolené hladiny významnosti α na rozhodnutí o výsledku testu).
- Testy dobré shody (Flash, řešený příklad).
- Kolmogorovův-Smirnovův test (Flash, řešený příklad).
- Analýza závislosti dvou kategoriálních veličin (Flash, řešený příklad).
- ANOVA (Java, interaktivní demonstrace principu ANOVY).
- Regrese (Java, umožňuje vygenerovat bodový graf náhodného výběru $z(X, Y)$, odhadnout od oka regresní přímku a srovnat vlastní odhad s odhadem metodou nejmenších čtverců).

Výše uvedené animace jsou doplněny výpočetními applety:

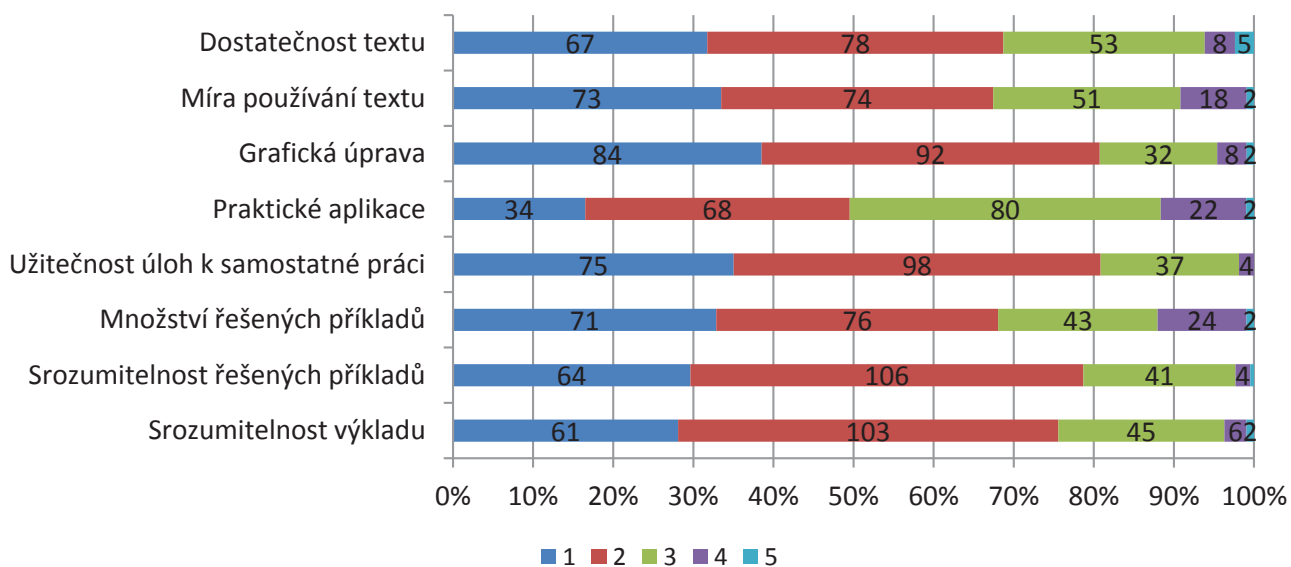
- Explorační analýza (Microsoft Excel), který pro numerickou proměnnou umožňuje výpočet základních číselných charakteristik, vykreslení krabicového grafu, kategorizaci proměnné do požadovaného počtu kategorií a následné vykreslení histogramu (je možno sledovat vliv počtu kategorií na tvar histogramu). Pro kategoriální proměnnou je konstruována tabulka četnosti, přičemž je zohledněno, zda jde o proměnnou nominální nebo ordinální.
- Kruskalův-Wallisův test (Excel, Kruskalův-Wallisův test včetně post hoc analýzy).
- Friedmanův test (Microsoft Excel, Friedmanův test včetně post hoc analýzy).

Obrazovkové verze jsou nyní k dispozici na <http://mi21.vsb.cz/>.

Hodnocení modulu Vybrané kapitoly z pravděpodobnosti (220 respondentů)



Hodnocení modulu Úvod do statistiky (219 respondentů)



Obrázek 4: Vyhodnocení evaluačních dotazníků z pilotních kurzů v letním semestru školního roku 2010/2011.

4. Vyhodnocení evaluačních dotazníků – pilotní kurzy v letním semestru školního roku 2010/2011

Výukové moduly „Vybrané kapitoly z pravděpodobnosti“ a „Úvod do statistiky“ patří mezi moduly, které byly aplikovány v pilotních kurzech předmětů Statistika (statistika pro nematematické obory), Statistika I. (úvodní kurz statistiky pro obor Výpočetní matematika) a Biostatistika (kurz statistiky pro obor Biomedicínský technik) v letním semestru školního roku 2010/2011.

Evaluační dotazník vyplnilo v závěru kurzu cca 220 studentů. Z výsledků (známkování jako ve škole, 1 ~ „Výborně“, viz Obr. 4) lze usuzovat, že studenti většinou nerozlišovali hodnocené moduly, které používali v rámci jednoho předmětu (stejná slovní hodnocení na dvojicích dotazníků).

V naprosté většině položek hodnotily více než 2/3 studentů předkládané materiály pozitivně (známka 1 nebo 2). Jako nejslabší se jeví výsledky v položkách Praktické aplikace a Množství řešených příkladů. Ze slovního hodnocení lze usuzovat, že část studentů by uvítala místo skript s výkladem teoretického základu rozsáhlou sbírku řešených příkladů. Tento požadavek lze považovat za „klasický“ požadavek studentů, pedagogové pilotních kurzů považují počet řešených příkladů za dostatečný.

Poděkování: Příspěvek byl realizován v rámci projektu „Matematika pro inženýry 21. století – inovace výuky matematiky na technických školách v nových podmínkách rychle se vyvíjející informační a technické společnosti“.

Registrační číslo projektu je CZ.1.07/2.2.00/07.0332. Řešitelem projektu jsou Vysoká škola báňská – Technická univerzita Ostrava a Západočeská univerzita v Plzni.

Reference

- [1] Litschmannová, M. (2011): *Úvod do statistiky*. Učební texty v elektronické podobě, dostupné z World Wide Web: <http://mi21.vsb.cz/modul/uvod-do-statistiky>
- [2] Litschmannová, M. (2011): *Vybrané kapitoly z pravděpodobnosti*. Učební texty v elektronické podobě, dostupné z World Wide Web: <http://mi21.vsb.cz/modul/vybrane-kapitoly-z-pravdepodobnosti>

UKÁZKA VYUŽITÍ PROGRAMU WINQSB PŘI ŘEŠENÍ ÚLOH OPERAČNÍ ANALÝZY

Alena Kolčavová

Adresa: Mgr. Alena Kolčavová, Univerzita Tomáše Bati ve Zlíně,
Fakulta managementu a ekonomiky, Ústav statistiky a
kvantitativních metod, Mostní 5139, 760 01 Zlín

E-mail: kolcavova@fame.utb.cz

Abstrakt: V úvodu příspěvku je nastíněna současná situace stavu připravenosti FaME UTB ve Zlíně na distanční formu vzdělávání. Nedílnou součástí distančního vzdělávání je samostudium a proto je příspěvek zaměřen na využitelnost programu WinQSB při výuce předmětu Kvantitativní metody v rozhodování. Použití programu je demonstrováno na příkladu případové studie.

Klíčová slova: Kvantitativní metody v rozhodování, lineární programování, sestavování matematických modelů, optimalizační metody, operační výzkum.

Abstract: The introductory part of the article outlines the state of readiness of the Faculty of Management and Economics, Tomas Bata University in Zlin for distance learning. Self-study is an integral part of distance learning, and for this reason the article concentrates on WinQSB programme usability in the subject Quantitative Methods in Decision Making. The programme application is demonstrated in a case study.

Keywords: Quantitative Methods in Decision Making, Linear Programming, Draw Up of Mathematic Models, Optimisation Methods, Operating Research.

1. Zkušenosti s výukou předmětu Kvantitativní metody v rozhodování

V příspěvku jsou na ukázce řešené případové studie zhodnoceny zkušenosti s využíváním programu WinQSB ve výuce předmětu Kvantitativní metody v rozhodování.

FaME UTB ve Zlíně nemá zatím akreditováno distanční vzdělávání, ale kurzy DiV, které absolvují někteří pedagogové, jsou předzvěstí trendu, kterým se fakulta hodlá ubírat. Byla zakoupena také licence pro výukový program EDEN (což je prostředí pro tvorbu distančních textů). V rámci rozvojového programu bylo zahájeno zpracování výukových a studijních materiálů

v prostředí Moodle. Jedná se o disciplíny, které jsou součástí kombinovaného studijního programu. Využívání systému Moodle je plánováno i na dalších fakultách UTB ve Zlíně.

Disciplína Kvantitativní metody v rozhodování se v současné době vyučuje v prezenční formě studia i v kombinované formě studia ve všech studijních programech akreditovaných na FaME UTB ve Zlíně v magisterském studiu (1. ročník magisterského programu). Kombinovaná forma studia vychází z učebního plánu studia prezenčního. Vykazuje již řadu distančních prvků, zejména uplatňovaných v kombinované formě studia (použití programu WinQSB, Microsoft Project, pomocí nichž si studenti ověřují správnost sestavení matematických modelů, ...). Nezbytná je také možnost komunikace s vyučujícím přes internet.

Každý pedagog, který vyučuje v kombinované formě studia, je nucen těmto studentům poskytnout speciální studijní materiály, které jim umožní samostudium. Ne každá problematika se dá jednoduše studovat distančně. V kurzu Kvantitativní metody v rozhodování jsou řešeny případové studie z praxe. Studenti musí na základě rozboru ekonomického problému sestavit matematický model. Ruční řešení tohoto modelu je pro složitost výpočtu a časovou náročnost téměř nemožné a v dnešní době informačních technologií i velmi neefektivní.

Naše fakulta využívá speciální výukový program WinQSB (verze 1.00 for Windows), který je i volně stažitelný na internetu. Obsahuje nejen moduly pro řešení úloh lineárního programování, ale i moduly z ostatních oblastí operačního výzkumu podle tohoto členění:

- Analýza rozhodování (Decision Analysis)
- Dynamické programování (Dynamic Programming)
- Plánování výrobních zdrojů (Facility Location and Layout)
- Předpovědi a lineární regrese (Forecasting and Linear Regression)
- Lineární programování (Linear and Integer Programming)
- Analýza Markovových procesů (Markov Process)
- Plánování materiálových zdrojů (Material Requirements Planning)
- Síťová analýza (Network Modeling)
- Nelineární programování (Nonlinear Programming)
- Analýza PERT_CPM
- Kontrola kvality (Quality Control Chart)
- Rozvrhování pracovních úkolů (Job Scheduling)
- Teorie front (Queuing Analysis)
- Modely hromadné obsluhy (Queuing System Simulation)

Práce s tímto programem je jednoduchá a na veškeré nejasnosti lze najít odpověď v nabídce Help.

Bez použití specializovaného programu by výuka tohoto předmětu byla pouhou ukázkou „zboží přes sklo výlohy, bez možnosti vyzkoušet si kvalitu a funkčnost“.

Využití tohoto programu ve výuce je doloženo ukázkou vyřešené případové studie. Je z tématického celku: Lineární programování – Sestavování matematických modelů – Kapacitní úlohy.

Snad tato ukázka bude přínosná i pro ostatní pedagogy, kteří se zabývají operačním výzkumem.

2. Případová studie: Plán rozšíření výroby pro firmu vyrábějící hračky

Podnikatelské prostředí: Výrobní firma.

Firma byla velmi úspěšná během prvních 6 měsíců působení a nyní hledá možnost přesunutí výroby do oblasti, kde výdaje na práci a materiál jsou podstatně levnější. Podařilo se jí nalézt oblast, kde je levná pracovní síla a uzavřít zde smlouvu s místním dodavatelem plastů. Dodavatel se zavázal zásobovat firmu 1500 kg plastů týdně za podstatně nižší ceny. Výroba by měla být efektivnější, došlo by k zdvojnásobení zisku z výroby V1 až na 480 Kč/kus a k ztrojnásobení zisku z výroby V2 až na 450 Kč/kus.

Nové prostory by byly vybaveny stroji a dělníky pracujícími 40 hodin v řádné pracovní době a navíc dělníci mohou odpracovat týdně 32 přesčasových hodin.

Z účetního hlediska vybavení pro dělníky, sociální pojištění, zdravotní pojištění, mzda a provozní výdaje za 1 přesčasovou hodinu budou stát společnost o 5400 Kč více než 1 řádná hodina pracovní doby.

Firma udělala marketingový průzkum trhu i pro dva nové produkty V3 a V4, které by chtěla uvést na trh. V tabulce jsou uvedeny požadavky na materiál, čas a zisk z jednotlivých druhů výrobků.

Firma podepsala smlouvu s odběratelskou firmou na dodávku 240 ks výrobků V2 týdně. Marketingové oddělení provedlo průzkum trhu a na základě tohoto průzkumu rozhodlo, že nejoblíbenější výrobek V1 bude zahrnovat 50 % celkové produkce a ostatní výrobky budou zahrnovat méně než 40 % produkce každý. Pod vlivem příznivých podmínek pro výrobu a odbytu oddělení vývoje navrhuje zvýšit celkovou výrobu hraček na 1000 ks týdně.

Management firmy chce určit týdenní plán produkce hraček za změněných podmínek (včetně všech přesčasových hodin, pokud jsou nezbytně nutné). Cílem je maximalizovat týdenní zisk.

Tabulka 1: Požadavky na materiál, čas a zisk z jednotlivých druhů výrobků

Výrobek	Spotřeba plastů (kg/kus)	Spotřeba času (min/kus)	Předpokládaný zisk (Kč/kus, EUR/kus)
V1	1,0	3	480 (15)
V2	0,5	4	450 (14)
V3	1,5	5	600 (19)
V4	2,0	6	660 (21)
Disponibilní množství	1500 kg	40 h řádná pr. doba 32 h přesčasová p. d.	

Rozbor problému:

1. Firma chce maximalizovat čistý týdenní zisk.
2. Musí být stanoven týdenní výrobní plán jednotlivých výrobků.
3. Existují následující omezení:
 - Dostupnost suroviny (plastů) – 1500 kg týdně.
 - Řádná pracovní doba dělníků ($40 \cdot 60 = 2400$ minut týdně).
 - Dostupnost přesčasové doby (32 h týdně).
 - Minimální množství vyrobených V2 týdně (240 ks týdně).
 - Vhodný výrobkový mix: V1 = 50 % z celkové produkce.
 - V2, V3, V4 menší než 40 % z celkové produkce.
 - Minimální celková produkce ($V1+V2+V3+V4 = 1000$ ks).

Proměnné:

Firma musí rozhodnout nejen o velikosti týdenní produkce jednotlivých výrobků, ale také určit výši přesčasových hodin týdně.

x_1 = počet ks výrobků V1 vyráběných za 1 týden.

x_2 = počet ks výrobků V2 vyráběných za 1 týden.

x_3 = počet ks výrobků V3 vyráběných za 1 týden.

x_4 = počet ks výrobků V4 vyráběných za 1 týden.

x_5 = počet přesčasových hodin za 1 týden.

Účelová funkce – chceme maximalizovat čistý týdenní zisk:

$$z_{\max} = 480x_1 + 450x_2 + 600x_3 + 660x_4 - 5400x_5$$

Omezení:

1. Surovina (kg): $x_1 + 0,5x_2 + 1,5x_3 + 2x_4 \leq 1500$
2. Čas (min.): $3x_1 + 4x_2 + 5x_3 + 6x_4 \leq 2400 + 60x_5$
3. Přesčas (hod.): $x_5 \leq 32$

4. Požadavky odběratele V2 (ks): $x_2 \geq 240$
Nyní zavedeme další proměnnou x_6 = celková týdenní produkce (v ks).
5. Potom platí: $x_6 = x_1 + x_2 + x_3 + x_4 + x_5$
Po úpravě: $x_1 + x_2 + x_3 + x_4 + x_5 - x_6 = 0$
6. Týdenní výroba V1 = 50 % z celkové produkce: $x_1 = 0,5x_6$
7. Týdenní výroba V2 menší než 40 % z celkové produkce: $x_2 \leq 0,4x_6$
8. Týdenní výroba V3 menší než 40 % z celkové produkce: $x_3 \leq 0,4x_6$
9. Týdenní výroba V4 menší než 40 % z celkové produkce: $x_4 \leq 0,4x_6$
10. Celková produkce větší než 1000 ks týdně: $x_6 \geq 1000$

Podmínka nezápornosti: $x_i \geq 0, i = 1, 2, \dots, 6$

3. Ekonomická interpretace problému

Po zadání dat v modulu Linear and Integer Programming (viz obr. 1) z menu ve WinQSB volíme *Solve and Analyze*, poté vybíráme *Solve the Problem*. Výstupní sestavu vidíme na obr. 2 na další straně. Tu si okomentujeme.

3.1. Řešení primárního problému

Optimální bude vyrábět týdně: 570 ks výrobků V1, 240 ks výrobků V2, 330 ks výrobků V3, výrobky V4 vůbec nevyrábět. Při této skladbě výroby bude počet přesčasových hodin týdně 32 a celková týdenní produkce bude 1140 ks výrobků týdně. Čistý týdenní zisk bude činit 406 800 Kč (12 712,5 EUR).

3.2. Řešení duálního problému

Neméně důležité je i řešení duálního problému.

1. K dispozici máme 1500 kg suroviny, spotřebujeme ji při výše uvedené výrobě pouze 1185 kg. Zbytek 315 kg nám zůstane na skladě. Proto je zde stínová cena nulová.

2. Máme k dispozici 2400 min. řádného prac. času, který je plně využit. Přidáním 1 min. řádného času do daného výrobního procesu by se hodnota čistého týdenního zisku zvýšila o 135 Kč = 4,2 EUR. Pro zvážení této eventuality je podstatné, kolik Kč nás stojí 1 min. řádné prac. doby, tzn. mzdové a provozní náklady. Pokud je to částka nižší než 135 Kč, stojí za to uvažovat o rozšíření výroby, za předpokladu, že ostatní ukazatele jsou také příznivé.

3. Máme k dispozici 32 hodin přesčasového prac. času, který je plně využit. Stínová cena je 2700 Kč = 84,4 EUR, tzn. přidáním 1 přesčasové hodiny do daného výrobního procesu by se hodnota čistého týdenního zisku zvýšila o 2700 Kč. Platí zde podobná úvaha jako v předešlém bodě.

Variable -->	X1	X2	X3	X4	X5	X6	Direction	R. H. S.
Maximize	480	450	600	660	-5400			
C1	1	0.5	1.5	2			<=	1500
C2	3	4	5	6	-60		<=	2400
C3					1		<=	32
C4		1					>=	240
C5	1	1	1	1		-1	=	0
C6	1					-0.5	=	0
C7		1				-0.4	<=	0
C8			1			-0.4	<=	0
C9				1		-0.4	<=	0
C10						1	>=	1000
LowerBound	0	0	0	0	0	0		
UpperBound	M	M	M	M	M	M		
VariableType	Continuous	Continuous	Continuous	Continuous	Continuous	Continuous		

Obrázek 1: Zadávání vstupních údajů v programu WinQSB

	12:34:25		Sunday	March	02	2003		
	Decision Variable	Solution Value	Unit Cost or Profit c(j)	Total Contribution	Reduced Cost	Basis Status	Allowable Min. c(j)	Allowable Max. c(j)
1	X1	570,0000	480,0000	273 600,0000	0	basic	120,0000	600,0000
2	X2	240,0000	450,0000	108 000,0000	0	basic	-M	465,0000
3	X3	330,0000	600,0000	198 000,0000	0	basic	582,8571	M
4	X4	0	660,0000	0	-75,0000	at bound	-M	735,0000
5	X5	32,0000	-5 400,0000	-172 800,0000	0	basic	-8 100,0000	M
6	X6	1 140,0000	0	0	0	basic	-180,0000	60,0000
	Objective	Function	(Max.) =	406 800,0000				
	Constraint	Left Hand Side	Direction	Right Hand Side	Slack or Surplus	Shadow Price	Allowable Min. RHS	Allowable Max. RHS
1	C1	1 185,0000	<=	1 500,0000	315,0000	0	1 185,0000	M
2	C2	2 400,0000	<=	2 400,0000	0	135,0000	1 840,0000	3 408,0000
3	C3	32,0000	<=	32,0000	0	2 700,0000	22,6667	48,8000
4	C4	240,0000	>=	240,0000	0	-15,0000	110,7692	480,0000
5	C5	0	=	0	0	-75,0000	-880,0000	112,0000
6	C6	0	=	0	0	150,0000	-132,6316	440,0000
7	C7	-216,0000	<=	0	216,0000	0	-216,0000	M
8	C8	-126,0000	<=	0	126,0000	0	-126,0000	M
9	C9	-456,0000	<=	0	456,0000	0	-456,0000	M
10	C10	1 140,0000	>=	1 000,0000	140,0000	0	-M	1 140,0000

Obrázek 2: Řešení problému pomocí programu WinQSB

4. Máme vyrobit 240 ks výrobků V2, které vyrobíme. Vyrobením 1 ks výrobku V2 navíc by se snížil týdenní zisk o 15 Kč = 0,7 EUR.

5. Pokud bychom zvýšili počet vyrobených výrobků o 1 ks týdně, snížil by se týdenní zisk o 75 Kč = 2,3 EUR.

6. Máme vyrobit 570 ks výrobků V1, které vyrobíme. Vyrobením 1 ks výrobku V1 navíc by se zvýšil týdenní zisk o 150 Kč = 4,7 EUR. Uvažovat o zvýšení výroby V1 je možné pouze za předpokladu, že výrobní náklady na tento výrobek jsou nižší než 150 Kč/kus a všechny potřebné suroviny máme k dispozici v dostatečném množství.

7. Máme vyrobit alespoň 456 ks výrobků V2, vyrobíme jich o 216 ks více. Zde je stínová cena nulová, protože nemá smysl uvažovat o změně výše týdenního zisku, vyrobíme-li 1 ks V1 navíc. Totéž platí i o omezeních C8–C10.

Literatura

- [1] Jablonský, J. *Operační výzkum – kvantitativní modely pro ekonomické rozhodování*. 1. vydání. Praha: Professional Publishing, 2002. 323 s. ISBN 80-86419-23-1.
- [2] Lawrence, J.; Pasternack, B. *Applied Management Science*. New York: John Wiley, 1998. 665 s. ISBN 0-471-13776-6.
- [3] Gros, I. *Kvantitativní metody v manažerském rozhodování*. 1. vydání Praha, Grada publishing, a. s., 2003. 432 s. ISBN 80-247-0421-8.
- [4] WinQSB. Version 1.00. Copyright Yih-Long Chang. [cit 2004-04-11].

~ ~ ~

Pozvánky na akce

- 6.–7. 12. 2012, Bratislava 21. Medzinárodný seminár Výpočtová štatistika. Infostat. <http://www.ssds.sk/>
- 6. 12. 2012, Bratislava. Prehliadka prác mladých štatistikov a demograf. Bratislava, Infostat. <http://www.ssds.sk/>
- NTTS 2013 (New Techniques and Technologies for Statistics), 5 to 7 March 2013 in Brussels. <http://www.ntts2013.eu/>
- ICEEIM 2013, Peking, March 14–15, 2013. International Conference on e-Education, e-Business and Information Management. <http://www.iceepsd.org/icibet2013/>
- QRPC 2013, 4.–7. 6. 2013, New York. <http://www.trilobyte.cz/Oznamení/Pozvanka-na-konferenci-o-aplikovane-statistice-v-kvalite-prumyslu-a-vede.html>

THE FIRST CZECH TEXTBOOK OF STATISTICAL METHODS AND ITS AUTHOR – STANISLAV KOHN

PRVNÍ PŮVODNÍ ČESKÁ UČEBNICE STATISTICKÝCH METOD A JEJÍ AUTOR – STANISLAV KOHN

Prokop Závodský

Adresa: VŠE, FIS, KSTP, nám. W. Churchilla 4, 130 67 Praha 3

E-mail: zavodsky@vse.cz

Abstract: The first textbook of modern statistical methods written in Czech language *Základy teorie statistické metody* (Praha 1929) has been written by Stanislav Kohn (1888–1933), coming from Varsovian Jewish family. Finally, he got into Prague across St. Petersburg, Tiflis (Tbilisi) and Paris and was named here an associate professor of Russian faculty of law. The above mentioned textbook piqued interest even abroad, but severe disease and untimely death prevented Kohn in further work. His tomb can be found in New Jewish cemetery in Prague, in the vicinity of the tomb of Franz Kafka.

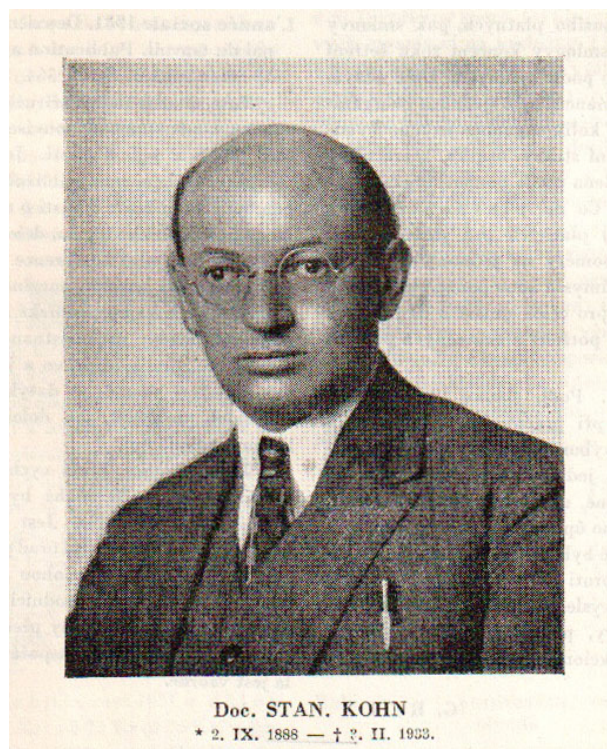
Keywords: History of Statistics, Stanislav Kohn.

Abstrakt: První česky psanou učebnici moderních statistických metod *Základy teorie statistické metody* (Praha 1929) napsal Stanislav Kohn (1888–1933), pocházející z varšavské židovské rodiny. Přes St. Petersburg, Tiflis (Tbilisi) a Paříž se nakonec dostal až do Prahy, kde byl jmenován docentem ruské právnické fakulty. Uvedená učebnice vzbudila velký zájem i v zahraničí, ale těžká choroba a předčasná smrt zabránily Kohnovi v další práci. Jeho hrob nalezneme na Novém židovském hřbitově v Praze, nedaleko hrobu Franze Kafky.

Klíčová slova: Dějiny statistiky, Stanislav Kohn.

V letech 1919–1920 se vedle dalších nových ústředních úřadů Československa konstituuje i Státní úřad statistický (SÚS), který má brzy 700–800 zaměstnanců. Na tradičních, ale zejména na nově vzniklých českých vysokých školách a fakultách (Vysoká škola speciálních nauk a Vysoká škola obchodní ČVUT, Přírodovědecká fakulta UK a další) patří statistika mezi významné předměty. Podívejme se, jaká je v této době situace s českou statistickou literaturou.

Učebnice moderních statistických metod, které se tehdy ve světě bouřlivě rozvíjely, u nás ještě na počátku 20. let zcela chyběla, neexistovala v tomto oboru ani česká odborná terminologie. Drobnými příspěvky zde byly *Vybrané kapitoly z matematické statistiky* od polyhistora v oblasti přírodních věd Václava Lásky (1862–1943)¹ a obsáhlý referát o francouzské knize Armanda Julina od předního odborníka SÚS Josefa Mráze (1882–1934).² Příručky „otce čs. statistiky“ Dobroslava Krejčího³ se zabývají především popularizací statistiky a organizací sběru statistických dat. Pro úplnost jmenujme ještě brožuru předního českého pojistného matematika té doby Josefa Beneše *O statistice a její teorii, o vědách a zájmech, s nimiž souvisí* (1920), stručně a poněkud nepřehledně uvádějící čtenáře do různých problémů statistiky, pravděpodobnosti a pojistné matematiky.



Velkým pokrokem bylo přeložení do češtiny obsáhlé publikace představiatele anglické biometrické školy George Udny Yuleho *Úvod do teorie statistiky* (Praha 1926). Překladaři byli: profesor české techniky v Brně Vladimír Novák a již zmíněný doc. Josef Mráz, překlad inicioval a vlastním nákladem vydal SÚS. Překladem této rozsáhlé knihy (přes 500 stran) se oba jmenovaní vědci zasloužili o vytvoření české statistické terminologie⁴, většinou používané dodnes – uveďme kupř. termíny: směrodatná odchylka, regrese, korelace, kontingence a další.

První původní česky psaná učebnice statistických metod vyšla rovněž péčí SÚS o tři roky později. Jsou to *Základy teorie statistické metody* (Praha 1929) od Stanislava Kohna. Mají přibližně stejný rozsah jako Yuleho práce, ale jsou

¹ Čs. statistický věstník, roč. II (1921), s. 225–258 a 313–342 i jako samostatná publikace. Obsahuje jen některé problémy statistiky a pro čtenáře bez matematického vzdělání je nesrozumitelná.

² Julinovy „Základy teoretické a praktické statistiky“. Čs. statistický věstník, roč. III (1922), s. 284–316.

³ *Základy statistiky, zvláště pro zemědělce a družstevníky* (1. vydání 1920, 2., podstatně rozšířené, 1923).

⁴ Jen částečně mohli navázat na výše uvedené publikace, zejména na Mrázův referát, citovaný v pozn. 2.

modernější, lépe zachycují nejnovější vývoj statistiky ve světě a obsahují i širší sortiment statistických metod.

Stanislav Kohn⁵ se narodil 2. září 1888 v židovské rodině ve Varšavě, kde pak navštěvoval střední obchodní školu. Po dvou semestrech na přírodovědecké fakultě v Krakově⁶ studoval na ekonomickém odboru (fakultě) polytechniky v Petrohradě, který měl tehdy v rámci Ruska špičkovou úroveň. Mezi jeho nejvýznamnější učitele, s nimiž se po absolvování polytechniky (1911) spřátelil, patřili profesori A. A. Čuprov⁷ a P. B. Struve⁸.

Již první Kohnova publikovaná práce o financích pojišťovacích společností vzbudila ohlas u odborné veřejnosti. Během první světové války pak Kohn pracoval v matematicko-statistickém oddělení ministerstva zemědělství, kde se zabýval problematikou výběrových šetření v zemědělské statistice.

Roku 1918 odjíždí Kohn z Petrohradu do Tiflisu (dnes Tbilisi), kde získal docenturu národního hospodářství a statistiky na tamější polytechnice a brzy i renomé výborného pedagoga. Litografické vydání Kohnových přednášek⁹ se dočkalo velmi příznivého hodnocení z pera A. A. Čuprova, který jinak chválou šetřil.

Likvidace polytechniky a bolševický režim přiměly Kohna opustit Rusko, od ledna 1921 žil v Paříži. Během následujících dvou let se zde vědecky i žurnalisticky zabýval otázkami národního hospodářství a demografickými i sociálními poměry v sovětském Rusku.

Po dvou letech pařížského pobytu se Kohnovi otevírají i jiné možnosti. Hlavní úřad statistický v rodné Varšavě (kde žije Kohnova matka) mu nabízí stálé místo s perspektivou dobré kariéry. Profesor Čuprov (tehdy působil v Drážďanech) ho chce doporučit na některou universitu v Německu. Kohn si ale vybírá Prahu, kam ho zve prof. Struve na ruskou právnickou fakultu.

Proč se Kohn roku 1923 rozhodl pro Prahu, kde nemohl počítat se solidním finančním zajištěním? Domnívám se, že hlavním důvodem k tomuto rozhodnutí byla skutečnost, že Kohn se již sžil s ruskými intelektuálními kruhy, zejména se statistiky a národohospodáři – toho času v exilu. Vyhovovala mu též perspektiva vysokoškolského působení a volnější vědecké práce.

⁵V ruské literatuře se píše Stanislav Salezijevič Kon.

⁶Snad se laskavý čtenář neurazí, připomenuli-li, že Varšava tehdy patřila do carské říše, zatímco Krakov do rakouské části habsburské monarchie.

⁷Aleksandr Aleksandrovič Čuprov (1874–1926) patřil v první čtvrtině XX. stol. k nejvýznamnějším světovým statistikům. Roku 1917 opustil Rusko a r. 1925 přesídlil do Prahy.

⁸Petr Berggardovič Struve (1870–1944) byl významný ruský ekonom, historik, politik a publicista. Roku 1920 emigroval z Ruska, v letech 1922–1924 žil v Praze, kde byl, stejně jako Čuprov, profesorem ruské právnické fakulty.

⁹*Kurs lekcij po teoriji statistiki* (Tiflis 1919).

Československo (a zejména Praha) se za podpory vlády i osobně presidenta Masaryka stalo ve 20. letech centrem vzdělaného ruského a ukrajinského exilu. Roku 1921 byla do Prahy z Vídně přeložena ukrajinská univerzita, následujícího roku zde byla založena ruská právnická fakulta. Praha se tak stala unikátním sídlem čtyř národních universit.

Kohn prožil v Praze deset neplodnějších let svého života. Jako soukromý docent přednášel statistiku na již zmíněné ruské právnické fakultě. Od roku 1924 byl též spolupracovníkem Národohospodářského ústavu prof. Prokopoviče.¹⁰ Přednášel zde o statistické teorii i o hospodářské a demografické situaci v SSSR a publikoval různé aktuality o SSSR i důkladné ekonomické a statistické analýzy.

Zároveň získal Kohn místo konsulenta (odborného poradce) v Brdlíkově Zemědělském ústavu účetnicko-správovědném.¹¹ Navázal zde na svou někdejší práci v petrohradském ministerstvu zemědělství a dále rozvíjel teorii i praktickou stránku výběrových metod v zemědělství. Na toto téma, v naší literatuře dosud téměř neznámé,¹² publikoval Kohn česky i rusky několik významných statí.¹³

Z Kohnovy publikační činnosti kolem poloviny 20. let¹⁴ připomeňme obsáhlou recenzi poslední významné práce A. A. Čuprova *Grundbegriffe und Grundprobleme der Korrelationstheorie* a zhodnocení jeho díla v nekrolozích po Čuprově úmrtí na jaře 1926. Na 3. Ruském akademickém sjezdu v Praze v říjnu 1924 vystoupil Kohn s pozoruhodným referátem *O pravděpodobnostně-statistickém přístupu k problémům ekonomické teorie*, tyto myšlenky dále

¹⁰Sergej Nikolajevič Prokopovič (1871–1955) – ruský ekonom, statistik, publicista i politik (ministr Kerenského prozatímní vlády). Roku 1922 se štěstím unikl výkonu trestu smrti a byl ze sovětského Ruska vypovězen. V Berlíně založil národohospodářský ústav nesoucí jeho jméno, roku 1924 se s ním přestěhoval do Prahy. Ústav dosahoval evropské úrovně, pořádal přednášky, semináře, vydával vědecké i popularizující publikace (především o hospodářském vývoji v SSSR) apod. Z důvodu hroící Hitlerovy agrese opustil Prokopovič roku 1938 Prahu a přenesl svou činnost do Ženevy.

¹¹Vladislav Brdlík (1879–1964) byl přední národohospodářský expert agrární strany, profesor zemědělství na ČVUT atd. Výše jmenovaný ústav, spojovaný vždy s jeho jménem, řídil od jeho založení (1919) až do roku 1946. Roku 1929 byl Brdlík zvolen zakládajícím členem Čs. statistické společnosti. Od roku 1948 aktivně působil v exilu.

¹²Např. v Krejčího publikaci (viz pozn. 3), věnované zejména organizaci statistických šetření v zemědělství, se možnost reprezentativního výběru vůbec neuvažuje.

¹³Kohnovu neúplnou bibliografii (39 položek bez drobných aktualit a jiných příspěvků) obsahuje nekrolog [3].

¹⁴Publikoval v ruském, českém, německém, polském, francouzském, anglickém i bulharském jazyce – v různých evropských zemích.

rozvinul následujícího roku ve dvou příspěvcích do vědeckých sborníků.¹⁵ Citované statě jsou patrně u nás (a snad i v rusky psané literatuře) prvními příspěvky k vědní disciplíně, která byla ve světě teprve v počátcích – k ekonometrii. O mezinárodním ohlasu Kohnových prací svědčí, že se stal r. 1930 zakládajícím členem The Econometric Society v USA.

Ve 20. letech byl Kohn zvolen za člena i několika dalších významných mezinárodních vědeckých společností, účastnil se řady vědeckých konferencí, např. 18. kongresu Mezinárodního statistického institutu ve Varšavě (1929).

Na jaře 1931 Kohn těžce onemocněl. Přestože byl dlouhodobě upoután na lůžko, neztrácel optimismus a napsal ještě několik pozoruhodných statí, anglická Royal Economic Society mu nabídla členství. 3. února 1933 však Stanislav Kohn v nedožitém 45. roce věku zemřel. Je pohřben na Novém židovském hřbitově v Praze, poblíž hrobu Franze Kafky (sektor 22, řada 6, hrob 8).¹⁶

Vraťme se ke Kohnově rozsáhlé učebnici *Základy teorie statistické metody* (17 + 483 stran). Její název vychází ze skutečnosti, že Čuprov a jeho žáci deklarovali statistiku jako zvláštní metodologickou vědu – teorii statistické metody.¹⁷ První část knihy je věnována popisné statistice, do níž autor zařadil i teorii indexů. Druhá, podstatně rozsáhlejší část, nadepsaná *Statistické bádání o příčinných spojeních*, pojednává především o teorii pravděpodobnosti a regresní analýze (včetně zkoumání závislostí mezi kvalitativními znaky) a též o analýze časových řad.

Zárodkem Kohnovy učebnice byly jeho přednášky v Gruzii (viz pozn. 9), ale ty zcela přepracoval a rozšířil. Zcela novou kapitolou je např. pojednání o časových řadách, kde Kohn vychází z nejnovějších prací harvardské školy, ale i jejích kritiků a probírá mimo jiné i problematiku korelace časových řad, opožděné korelace atd. Jedná se o první systematický výklad této problematiky v české literatuře – podobně jako v kapitole o indexech, kde je Kohnův výklad založen na nejnovějších publikacích amerických, německých a jiných statistiků.

Zejména v těchto originálních kapitolách přispěla Kohnova učebnice k rozvoji české statistické terminologie, zásluhu samozřejmě mají i doc. Mráz

¹⁵ O „statistifikacii“ političeskoj ekonomiji (*K kritike těoretičeskich postrojenij P. B. Struve*). In: Sbornik posvjaščennyj P. B. Struve, Praha 1925 a *Matěmaticeskoje i empiričeskoje napravlenije v těorii ceny*. In: Russkij ekonomičeskij sbornik II (1925), s. 5–51 a III (1925), s. 31–49.

¹⁶ Nápis na náhrobku dnes již není čitelný, na fotografii na následující straně je text rekonstruován podle mých starších záznamů.

¹⁷ Viz např. předmluvu ke Kohnově učebnici (s. V).



a prof. Schoenbaum¹⁸ (v menší míře i další osoby, převážně z řad pracovníků SÚS), kteří Kohnovi s přípravou českého textu knihy významně pomáhali.

Přestože Kohnova učebnice vyšla česky a nikoli v některém z jazyků tehdejší světové vědy, vzbudila v zahraničí velkou pozornost. Berlínský profesor a jinak obávaný kritik L. v. Bortkiewicz (sám se přes svou ohromnou erudici k sepsání souborné práce o statistice nikdy neodhodlal) vyjádřil svůj údiv nad množstvím probrané látky a doporučil přeložit práci do němčiny. Učebnici kladně hodnotili i profesori G. U. Yule, O. Anderson, v Polsku O. Lange ad. Jednání o německém a polském vydání bylo nakonec neúspěšné, především pro Kohnův zhoršený zdravotní stav a pozdější úmrtí.

Při odpovědi na otázku po příčině velkého mezinárodního ohlasu Kohnova díla je třeba připomenout následující skutečnost: jednotlivé statistické směry či školy se v Evropě tradičně vyvíjely dosti odděleně – lišily se obory, v nichž statistiku rozvíjely, a stejně tak přístupy k řešení problémů a metody, na něž se soustředily. Také vzájemná znalost výsledků práce anglických biometriků

¹⁸Emil Schoenbaum (1882–1967) byl tehdy profesorem pojistné matematiky na UK.

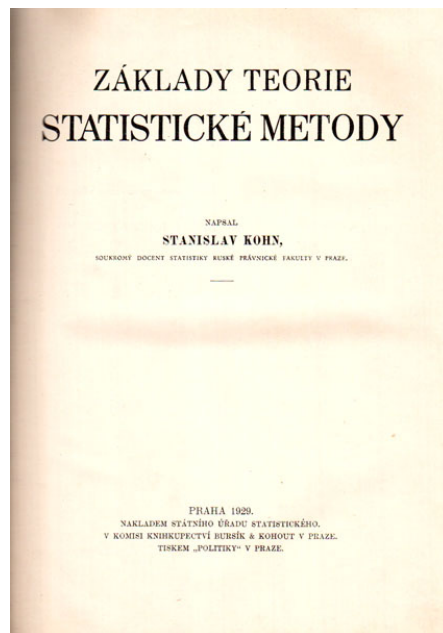
(Pearson, Yule ad.), německých sociálních statistiků (Lexis, Bortkiewicz), francouzských vědců (soustřeďujících se na filosofické a logické základy teorie pravděpodobnosti) a ruských matematiků a statistiků (A. A. Markov, A. A. Čuprov) nebyla obvykle na vysoké úrovni – výjimku tvořili zejména Bortkiewicz a Čuprov se svými žáky. Jak již uvedeno, Bortkiewicz se do souborné učebnice statistiky nepustil. Zato Ahasveru Kohnovi se podařilo (po smrti jeho guru a na jeho počest) nejen využít ve své práci nejnovější poznatky různých statistických škol, ale zvládl je skloubit do jednotného systému moderní statistiky.

Historik teorie pravděpodobnosti doc. Mačák v citované práci¹⁹ projevil určitý podiv nad skutečností, že u nás v meziválečném období vyšla jediná původní učebnice statistiky. Na otázku, kdo kromě Kohna by byl schopen takovou učebnici napsat, si odpovídá, že např. prof. Láska. K souhlasu s tímto názorem musím však podotknout, že by to ovšem od Lásky vyžadovalo důkladné studium rozsáhlé světové literatury oboru, který se ve 20. a 30. letech skutečně bouřlivě rozvíjel. Podobný předpoklad by samozřejmě musel splnit i prof. Schoenbaum, případně brněnský profesor Bohuslav Hostinský.

Ale neměli bychom tu opomenout ani doc. Mráze (od r. 1929 vicepresidenta SÚS), předního českého znalce statistické literatury,²⁰ který se k sepsání takové učebnice skutečně chystal. Bohužel odešel na věčnost již rok po Kohnovi.

Ve 30. letech již u nás vyrůstala nová generace fundovaných statistiků, z nichž především prof. Jaroslav Janko svou monografií *Základy statistické indukce* (1937) a dalšími publikacemi prokázal, že by byl schopen učebnici statistiky kvalitně napsat.

Faktem je, že Kohnovy *Základy* zůstaly po více než dvě desetiletí (!) jedinou českou publikací svého druhu – ve 30. letech žádná podobná nevyšla, za protektorátu okupanti dohlíželi, aby nebyla vydána žádná publikace připomínající vysokoškolskou učebnici. Únorem 1948 pak byl zahájen boj proti „matematickým formalismům ve statistice“, zatímco ekonometrie (s kybernetikou a genetikou) se ocitla v kategorii buržoazních pavěd.



¹⁹Mačák [2], s. 105.

²⁰Aby jako čtyřicetiletý statistik s právnickým vzděláním mohl porozumět moderní statistické literatuře, neváhal se zapsat k vysokoškolskému studiu vyšší matematiky.

Literatura

- [1] Dmitrijev A.: Iz plejady Čuprovcev: Stanislav Salezijevič Kon. Voprosy statistiki, roč. 1998, s. 81–82.
- [2] Mačák K.: Vývoj teorie pravděpodobnosti v českých zemích do roku 1938. Praha 2005.
- [3] Mráz J.: + Doc. Stanislav Kohn (Posmrtné vzpomínky). Čs. statistický obzor, roč. XIV (1933), s. 162–167.
- [4] Mráz, J.: Základy teorie statistické metody (recenze). Čs. statistický věstník, roč. XI (1930), s. 134–139.
- [5] Čuvakov V. N. (ed.): Něžabytnyje mogily. Rossijskoje zaruběžije: Někrologi 1917–1999. Tom 1–6. Moskva 1999–2007.
- [6] Ostrouchov P.: Pamjati S. S. Kona. Rossija i slavjanstvo, roč. 5 (1933), 18. II. 1933, s. 2.
- [7] Pamjati S. S. Kona. Bjulletěň Ekonomičeskogo kabiněta prof. S. N. Prokopoviča, roč. X (1933), No. 103, s. 1–2.
- [8] Práce ruské, ukrajinské a běloruské emigrace vydané v Československu 1918–1945. Díl I, sv. 1–3. Praha 1996.
- [9] Struve P.: Krupnyj učenyj i chorošij čelověk. Rossija i slavjanstvo, roč. 5 (1933), 18. II. 1933, s. 2.

~ ~ ~

Mikukláš 2012

Vážené kolegyně a kolegové,

rok se s rokem sešel a je tu zase – Mikukláš! To znamená, že se sejdeme opět v Karlíně v Praze v respiriu MFF UK, Sokolovská 83, tentokrát v den svátku sv. Mikuláše, 6. prosince 2012.

Na programu budou odborné přednášky, vystoupení skupiny FAB, s. r. o., a možná přijde i Mikukláš! Začátek předpokládáme v 9:00.

Všichni zájemci jsou vítáni!

Na shledanou se těší,

výbor společnosti

Obsah

Jiří Ivánek

Compound Quantification of Implications, Double-implications and Equivalency in Four-fold Tables	1
---	---

Zdeněk Půlpán

Jsou meze pro užití didaktických testů k odhadu vědomosti žáků a studentů?	14
---	----

Ladislav Mura

Možnosti aplikácie zhlukovej analýzy v manažérskych podnikových analýzách	27
--	----

Martin Kovářik, Petr Klímek

Využití nástrojů optimalizace a genetických algoritmů v procesu řízení výroby	41
--	----

Zuzana Kurucová, Beáta Stehlíková, Anna Tirpáková

Verification of the Efficiency of Combined Teaching of English	60
--	----

Petr Klímek, Martin Kovářik

Ekonometrie v Open Source Software	69
--	----

Martina Litschmannová

Tvorba internetových učebnic pravděpodobnosti a statistiky jako součást projektu MI21	78
--	----

Alena Kolčarová

Ukázka využití programu WinQSB při řešení úloh operační analýzy	86
--	----

Prokop Závodský

První původní česká učebnice statistických metod a její autor – Stanislav Kohn	93
---	----

Informační Bulletin České statistické společnosti vychází čtyřikrát
do roka v českém vydání. Příležitostně i mimořádné české a anglické číslo.

Časopis je zařazen do seznamu Rady pro výzkum, vývoj
a inovace, více viz server <http://www.vyzkum.cz/>.

Předseda společnosti: prof. RNDr. Gejza DOHNAL, CSc.

ÚTM FS ČVUT v Praze, Karlovo náměstí 13, 121 35 Praha 2

E-mail: gejza.dohnal@fs.cvut.cz

Redakční rada: prof. Ing. Václav ČERMÁK, DrSc. (předseda), prof. RNDr.
Jaromír ANTOCH, CSc., doc. Ing. Josef TVRDÍK, CSc., RNDr. Marek MALÝ,
CSc., doc. RNDr. Jiří MICHÁLEK, CSc., doc. RNDr. Zdeněk KARPÍŠEK,
CSc., prof. Ing. Jiří MILITKÝ, CSc., prof. RNDr. Gejza Dohnal, CSc.

Technický redaktor: Ing. Pavel STRÍŽ, Ph.D., pavel@striz.cz

Informace pro autory jsou na stránkách <http://www.statspol.cz/>

DOI: 10.5300/IB, <http://dx.doi.org/10.5300/IB>

ISSN 1210–8022 (Print), ISSN 1804–8617 (Online)