



OPERAČNÍ PROGRAM
LIDSKÉ ZDROJE
A ZAMĚSTNANOST

Statistika v evaluacích



Mgr. Lubomíra Červová
21.6.2012



OPERAČNÍ PROGRAM
LIDSKÉ ZDROJE
A ZAMĚSTNANOST

POPORUČUJEME
VŠE BUDOUCNOSTI
www.esf.cz

Program

- základní metody analýzy číselných proměnných (*průměr, kvantily, směrodatná odchylka, intervaly spolehlivosti, histogram, boxplot*)
- porovnání průměrů ve skupinách
- korelační analýza
- lineární regrese

Používaný software: IBM SPSS Statistics



OPERAČNÍ PROGRAM
LIDSKÉ ZDROJE
A ZAMĚSTNANOST

POPORUČUJEME
VŠE BUDOUCNOSTI
www.esf.cz

2

Data

- informace z databáze Albertina o firmách v letech 2008 a 2009
 - základní informace z účetních výkazů firem
- data z monitorovacího systému ESF Monit 7+
 - informace o tom, zda firma získala podporu z grantu

+ další datové soubory pro doplnění ...



OPERAČNÍ PROGRAM
LIDSKÉ ZDROJE
A ZAMĚSTNANOST

POPORUČUJEME
VŠE BUDOUCNOSTI
www.esf.cz

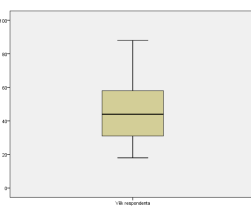
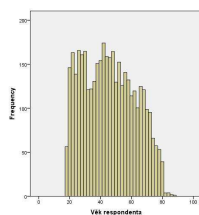
3

Základní metody analýzy číselných proměnných

Popis souboru

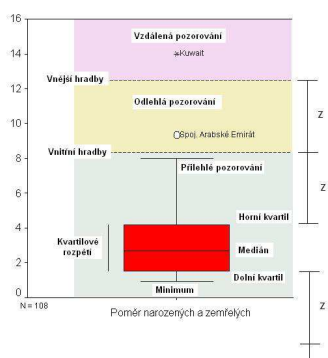
číselná proměnná

Descriptive Statistics					
	N	Minimum	Maximum	Mean	Std. Deviation
Věk respondenta	4011	18	88	45,34	16,714
Valid N (listwise)	4011				



5

Boxplot



6

Komparace skupin

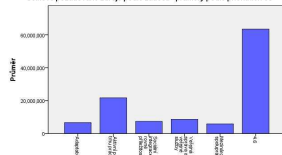
Porovnání jedné nebo více číselných proměnných ve skupinách

Přibližné hodnoty požadovaných finančních prostředků dle žádosti

Statistika: Průměr

	Celkové zdroje - žádosti	Souhrnné financování - žádosti	Veřejná fn. podpora celkem - žádosti
Příjemci osa	5603492	84014	6488073
Adaptabilita	21650898	9593	21660491
Aktivní politický trh práce	7371197	5820	7377017
Sociální integrace a rovné příležitosti	6888513	28084	6916597
Veřejná správa a veřejné služby	5799127	77	5799204
Mediální spolupráce	63517095	0	63517095
CELKEM	8261469	40913	8302382

Celkové požadované zdroje podle žádosti - průměry podle prioritních os



7

Testy hypotéz



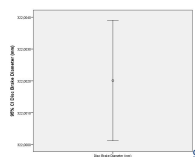
Interval spolehlivosti

bodový odhad x intervalový odhad

interval spolehlivosti

- vyjadřuje oblast, kde by se s danou spolehlivostí (obvykle 95%) měla nacházet hodnota statistiky pro celou populaci
- předpokládá pravděpodobnostní výběr (obvykle prostý náhodný výběr)

Descriptives		Statistic	Std. Error
Disc Brake Diameter (mm)	Mean	322,002009	,0009568
	85% Confidence Interval for Mean	322,000116	
	Lower Bound	322,000116	
	Upper Bound	322,001895	
	5% Trimmed Mean	322,001553	
	Median	322,001811	
	Variance	,000	
	Std. Deviation	,0108224	
	Minimum	321,9759	
Range	Maximum	322,0306	
	Range	,0547	
	Interquartile Range	,0133	
	Skewness	,200	,214
Kurtosis	Kurtosis	-,126	,425



Úvod do testování statistických hypotéz

Výběrová šetření

- Lze určitou vlastnost výběrového souboru zobecnit s danou spolehlivostí na celou populaci? *(například průměrný plat zaměstnanců v určitém oboru)*
- Mohou být zjištěné rozdíly na výběrovém souboru způsobené náhodou? *(například rozdíly mezi regiony)*



Princip testování statistických hypotéz (1)

H_0 : nulová hypotéza (na začátku předpokládáme, že platí)

H_A : alternativní hypotéza

T ... testová statistika (*kritérium pro rozhodování*)

Hladina spolehlivosti (volíme předem, obvykle 95%)

Pro testování použijeme testovou statistiku T a nalezneme vhodnou podmnožinu W , tzv. **kritický obor**

Rozhodování:

Pokud $\{T \in W\} \Rightarrow$ zamítáme H_0 , v opačném případě H_0 nezamítáme

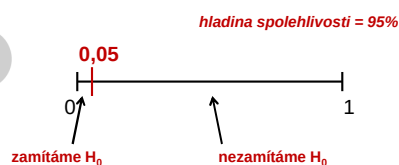
	rozhodnutí		
platí		H_0	H_A
	H_0	OK	chyba I. druhu
	H_A	chyba II. druhu	OK

Princip testování statistických hypotéz (2)

Significance (dosažená hladina významnosti)

- pravděpodobnost, že zamítneme nulovou hypotézu ačkoliv platí

Rozhodování na základě significance:



Princip testování statistických hypotéz (3)

síla testu: pravděpodobnost, že statistický test zamítne nulovou hypotézu v případě, že neplatí

síla testu = $1 - \beta$, kde

β je pravděpodobnost chyby II. druhu

Sílu testu ovlivňuje:

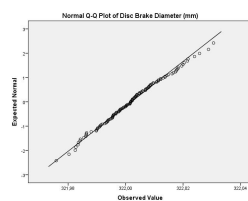
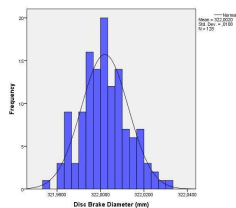
- **použitý statistický test** (některé testy jsou silnější než jiné)
- **velikost výběru** (obecně čím větší je velikost výběru, tím větší je síla testu)
- zvolená **hladina spolehlivosti**
- **effect size** (zjištěný rozdíl od nulové hypotézy)
- jednostranná x oboustranná **alternativní hypotéza**
- **chyba měření** (jakákoliv chyba, která narušuje přesnost měření, snižuje sílu testu)

13

Testování normality

Tests of Normality						
	Kolmogorov-Smirnov ^a			Shapiro-Wilk		
	Statistic	df	Sig.	Statistic	df	Sig.
Disc Brake Diameter (mm)	,051	128	,200 [*]	,993	128	,785

a. Lilliefors Significance Correction
*. This is a lower bound of the true significance.



H_0 : proměnná má normální rozložení

14

T-testy, ANOVA

Testy hypotéz o průměrech

- **Jednovýběrový t-test** – porovnání průměru s konstantou
- **Párový t-test** – porovnání průměrů dvojice proměnných
- **Dvouvýběrový t-test** – porovnání průměrů ve dvou skupinách
- **Jednoduchá ANOVA** (analýza rozptylu) – porovnání průměrů v několika skupinách

Předpoklady:

- nezávislost pozorování
- nezávislost skutečných hodnot a chyb
- **normální rozložení** (testovaná proměnná/rozdíl proměnných/testovaná proměnná ve skupinách)
- dvouvýběrový t-test, ANOVA: **skupiny se nepřekrývají**
- dvouvýběrový t-test, ANOVA: **shoda rozptylů ve skupinách**

15

Jednovýběrový t-test

H_0 : průměr (střední hodnota) proměnné odpovídá testované konstantě, tj. $\mu = 322$

One-Sample Statistics				
	N	Mean	Std. Deviation	Std. Error Mean
Disc Brake Diameter (mm)	128	322,00209	,0108224	,0009566

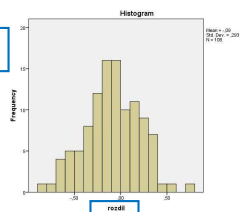
One-Sample Test						
	Test Value = 322				95% Confidence Interval of the Difference	
	t	df	Sig. (2-tailed)	Mean Difference	Lower	Upper
Disc Brake Diameter (mm)	2,101	127	.038	-.0020694	-.000116	-.003962

Párový t-test

H_0 : průměry (střední hodnoty) proměnných se rovnají, tj. $\mu_X = \mu_Y$

		Mean	N	Std. Deviation	Std. Error Mean
Pair 1	první měření	17,0119	108	1,05382	,10140
	druhé měření	17,0982	108	1,06586	,10254

		N	Correlation	Sig.
Pair 1	první měření & druhé měření	108	,962	,000



Paired Samples Test									
		Paired Differences					t	df	Sig. (2-tailed)
		Mean	Std. Deviation	Std. Error Mean	95% Confidence Interval of the Difference				
					Lower	Upper			
Pair 1	prvi mjeseci - drugi mjeseci	-.08624	,29259	,02815	-,14205	-,03043	-3,063	107	,003

Dvouvýběrový t-test

H_0 : průměry (střední hodnoty) ve dvou skupinách se rovnají, tj. $\mu_1 = \mu_2$

Group Statistics					
	Měřicí přístroj	N	Mean	Std. Deviation	Std. Error Mean
Síla	A	37	19,0192	,32840	,05399
	B	35	19,0411	,23192	,03920
Průměr vývodu	A	37	1,5055	,09345	,01536
	B	35	1,4889	,11545	,01952

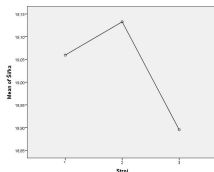
Independent Samples Test									
	Levene's Test for Equality of Variances			t-Test for Equality of Means					
	F	Sig.	t	df	t Sig. (2-tailed)	Mean Difference	Std. Error Difference	95% Confidence Interval of the Difference	
								Lower	Upper
Silika	3,000	.037	-.340	71	.735	-.02201	.06735	-.15724	.11324
			-.343	64,874	.733	-.02201	.06735	-.15616	.11404
Pada-mirinda	2,877	.108	.872	70	.504	.01650	.02469	-.03268	.06580
			.866	65,457	.507	.01650	.02464	-.03301	.06610

Jednoduchá analýza rozptylu (ANOVA)

H_0 : průměry (střední hodnoty) ve dvou nebo více skupinách se rovnají, tj.
 $\mu_1 = \mu_2 = \dots = \mu_K$

Test of Homogeneity of Variances				
Síla	Levene Statistic	df1	df2	Sig.
Průměr výkonu	1,544	2	69	,059

ANOVA						
		Sum of Squares	df	Mean Square	F	Sig.
Síla	Between Groups	,709	2	,354	4,880	,010
	Within Groups	5,012	69	,073		
	Total	5,721	71			
Průměr výkonu	Between Groups	,082	2	,041	3,020	,055
	Within Groups	,710	69	,010		
	Total	,792	71			



Typy testů

Parametrické testy

- testy o parametrech známých pravděpodobnostních modelů, předpokládají určitý typ rozdělení, ze kterého výběr pochází (obvykle normální)
- drobné odchylky od předpokladů většinou nevedí
- při výrazném porušení předpokladů můžeme získat nekorektní výsledky

Neparametrické testy

- určeny pro situaci, kdy nejsou splněny předpoklady parametrických testů
- nemají předpoklady o typu rozdělení dat
- vyváženo menší citlivostí (vyšší pravděpodobnost nezamítnutí testované hypotézy)
- testují jinou hypotézu o rozdělení základního souboru než je hypotéza o jeho parametru

Klasifikace neparametrických testů

- jednovýběrové testy
- dvouvýběrové testy
 - nezávislé výběry
 - závislé výběry
- vícevýběrové testy
 - nezávislé výběry
 - závislé výběry

Neparametrické testy pro 2 závislé výběry

Wilcoxonův test (*Wilcoxon*)

- test hypotézy o shodě mediánů ve dvou závislých výběrech
- silnější než znaménkový test
- předpokládá stejný tvar rozdělení
- princip: celkové pořadí na základě absolutních hodnot rozdílů testovaných proměnných, průměrné pořadí $+i - i$ - by měla být podobná

Znaménkový test (*Sign test*)

- test hypotézy o shodě mediánů ve dvou závislých výběrech
- princip: podobný jako u W. testu, ale bere v úvahu pouze znaménka rozdílů, aplikuje binomický test

McNemarův test (*McNemar*)

- pouze pro dichotomické proměnné
- zjišťuje, zda po určité události došlo ke změně výskytu určitého jevu

Neparametrické obdoby párového t-testu

22

Neparametrické testy pro dva nezávislé výběry

Mann-Whitneyův test (*Mann-Whitney U*)

- testuje hypotézu, že obě skupiny pocházejí ze stejného základního souboru
- předpokládá stejný tvar rozložení dat ve skupinách
- princip: společné pořadí, průměrné pořadí v obou skupinách by měla být podobná

Dvouvýběrový Kolmogorov-Smirnovův test (*Kolmogorov-Smirnov Z*)

- obecnější a slabší než Mann-Whitneyův test
- založen na porovnání empirických distribučních funkcí
- testová statistika vychází z jejich maximální absolutní odchylky
- citlivý na polohu a škálu měření

Wald-Wolfowitzův test (*Wald-Wolfowitz runs*)

- nemá žádné dodatečné předpoklady
- princip: společné pořadí, v setříděné posloupnosti by se měla náhodně střídát pozorování z obou skupin → Runs test

Neparametrické obdoby dvouvýběrového t-testu

23

Neparametrické testy pro K nezávislých výběrů

Kruskal-Wallisův test (*Kruskal-Wallis H*)

- testuje hypotézu, že všechny skupiny pocházejí ze stejného základního souboru
- předpokládá shodu tvaru rozdělení ve skupinách
- test neříká nic o tom, které skupiny se od sebe odlišují
- princip: společné pořadí všech skupin, průměrná pořadí ve skupinách by měla být podobná

Mediánový test (*Median*)

- testuje hypotézu o shodě mediánů ve dvou a více nezávislých výběrech
- vhodný v situacích, kdy předpokládáme ve výběrech různá rozdělení
- princip: společný medián, kontingenční tabulka kategorií $\times$ kategorie pod/nad mediánem, test Chí-kvadrát

Neparametrické obdoby ANOVA

24

Neparametrické testy – ukázka (1)

Mann-Whitneyův test

Ranks			
	Měřicí přístroj	N	Mean Rank
Hloubká	A	37	35,73
	B	35	36,26
	Total	72	
Výška	A	37	36,62
	B	35	36,37
	Total	72	
Síla	A	37	34,86
	B	35	39,23
	Total	72	
Průměr vývodu	A	37	38,46
	B	35	34,43
	Total	72	

H₀: obě skupiny pocházejí ze stejného základního souboru

Test Statistics ^a				
	Hloubká	Výška	Síla	Průměr vývodu
Mann-Whitney U	639,000	643,000	587,000	575,000
Wilcoxon W	1269,000	1273,000	1290,000	1205,000
Z	-.096	-.851	-.682	-.817
Asymp. Sig. (2-tailed)	.924	.860	.485	.414

a. Grouping Variable: Měřicí přístroj

Neparametrické testy – ukázka (2)

Kruskal-Wallisův test

Ranks			
	Straj	N	Mean Rank
Hloubká	1	24	34,58
	2	24	36,13
	3	24	38,79
	Total	72	
Výška	1	24	25,42
	2	24	28,79
	3	24	55,39
	Total	72	
Síla	1	24	38,25
	2	24	44,13
	3	24	26,13
	Total	72	
Průměr vývodu	1	24	29,68
	2	24	44,33
	3	24	36,68
	Total	72	

H₀: všechny skupiny pocházejí ze stejného základního souboru

Test Statistics ^{a,b}				
	Hloubká	Výška	Síla	Průměr vývodu
Chi-Square	,497	29,336	8,498	6,386
df	2	2	2	2
Asymp. Sig.	,780	,000	,009	,041

a. Kruskal-Wallis Test
b. Grouping Variable: Straj

Korelační analýza

Statistika

- **statistika** zkoumá variabilitu dat:
 - popisuje ji
 - vysvětluje ji
 - predikuje ji
- **korelační analýza** zkoumá společnou variabilitu (kovariabilitu):
 - popisuje ji
 - používá ji pro vysvětlení
 - používá ji pro predikci

Úlohy a otázky:

A) Souvisí spolu výskyt proměnné X a proměnné Y tak, že s vyššími hodnotami X se pojí vyšší hodnoty Y a (a nižšími nižší), či naopak s vyššími hodnotami X se pojí nižší hodnoty Y (a s nižšími X vyšší Y)?

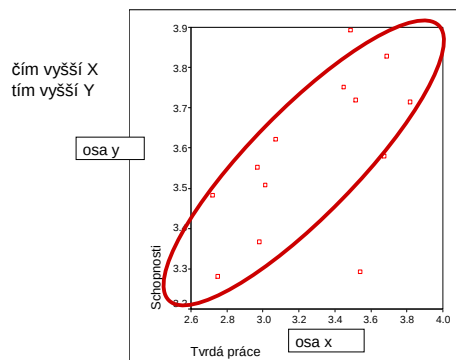
B) Můžeme v datech zjistit souběžnost resp. protiběžnost hodnot dvou číselných proměnných?

C) Je hodnota Y důsledkem hodnoty X? Reprezentuje proměnná X příčinu pro důsledek Y?

D) Jsou X a Y nositeli (částečně) stejné informace?

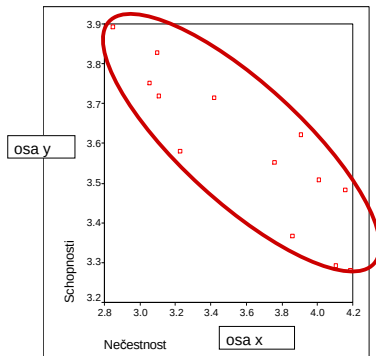
E) Vylučují se (resp. doplňují se) X a Y nebo naopak jedno předpokládá druhé?

Souběh variabilit



Protiběh variabilit

čím vyšší X
tím nižší Y
čím nižší X
tím vyšší Y



31

Korelační analýza

- korelace = vztah dvou **číselných** proměnných
- vztah:
 - **empirický**: korelace je zachycení vztahu mezi vzniklými čísly, statistickými řadami, bez ohledu na to, co znamenají
 - **kauzální**: korelace je odraz geneze příčinného procesu

32

Popis kovariability

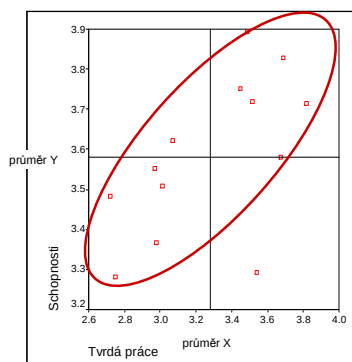
- **kovariance**: souběh variabilit dvou proměnných

$$c(X, Y) = \frac{1}{N-1} \sum (X_i - \bar{X}) * (Y_i - \bar{Y})$$

$$c(X, Y) = \frac{1}{N * (N-1)/2} \sum_{i \neq j} (X_i - X_j) * (Y_i - Y_j)$$

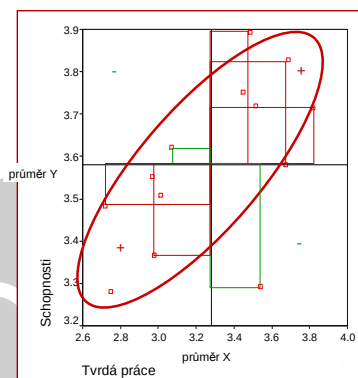
33

Kovariance = součet orientovaných plošek



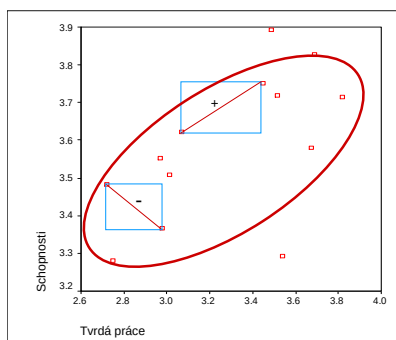
34

Kovariance = součet orientovaných plošek



35

Kovariance = součet orientovaných plošek



36

Korelace – kovariance v poměru k rozptylům

korelace = vztah dvou proměnných
= kovariance standardizovaná
k rozptylům obou proměnných
= měří **vztah** dvou variabilít,
nikoliv jejich velikost

Korelační koeficient

$$r = \frac{\text{cov}(X, Y)}{\sqrt{\text{var} X \cdot \text{var} Y}} \quad r = \frac{\text{cov}(X, Y)}{S_X \times S_Y}$$

- r je vypočten jako kovariance v poměru ke geometrickému průměru rozptylů
- jmenovatel je také součin směrodatných odchylek $s_x \cdot s_y$

Pearsonův lineární korelační koeficient

Vlastnosti:

- r definován, pro $n > 1$
- r definován pro nenulové variability;
nesmí platit $s_x = 0$ nebo $s_y = 0$
- $r = 1$, právě když body jsou seřazeny v nějaké přímce s nenulovým kladným spádem
- $r = -1$, právě když body jsou seřazeny v nějaké přímce s nenulovým záporným spádem
- čím více se r blíží k $+1$, tím více se body shlukují kolem stoupající přímky; čím více se r blíží k -1 , tím více se body shlukují kolem klesající přímky
- jestliže v mraku bodů nelze vystopovat žádný *lineární* trend, $r = 0$

Pearsonův lineární korelační koeficient

Další vlastnosti:

- **r** se nezmění, když se
 - **posune** škála jedné nebo obou proměnných o **libovolnou konstantu** (změna počátku)
 - **změní** škála jedné nebo obou proměnných **násobkem** libovolnými činiteli (změna měřítka)

Nulové r:

- mrak bodů tvoří pravidelný kruhový útvar
- přímka, kolem které se shlukují body, je vodorovná nebo svislá
- body leží symetricky kolem osy procházející průměrem X a to i když odpovídají úplné závislosti Y na X,
např. $Y = (X - 4)^2$
- silné shlukování kolem rostoucí/klesající přímky je zkresleno bodem vzdáleným od mraku
- kříží se kladný a záporný trend – překrytí dvou bodových mraků

Zkreslení koeficientu korelace

- **vzdálený bod** - mrak bodů ukazuje na silnou/slabou korelaci, ale vzdálený bod ji uměle sniž/zvýší
- **dvě skupiny** nulové korelace umístěné v rovině vykazují vyšší korelaci
- číselné proměnné mají **diskretní** povahu (škála celých čísel od 1 do K) a přesné seskupení hodnot kolem přímky není plně možné
- jsou-li **rozložení** X a/nebo Y výrazně **šikmá** s dlouhým koncem

Další míry korelace

- jiná data (pořadí)
- šikmá rozložení
- vzdálená pozorování
- zvyklosti oboru
- komparace s jinými výstupy

Pořadová korelace

- koeficienty pořadové korelace vycházejí při výpočtu koeficientu z pořadí a vzájemných pozic hodnot
- nejsou citlivé na vzdálená pozorování
- neměří linearitu vztahu, ale shodu pořadí obou statistických řad
- mohou se použít pro číselná data a pro data, která vyjadřují pouze pořadí případů z hlediska X a Y

Pořadová korelace - Spearmanovo ρ

- Spearmanův koeficient pořadové korelace ρ vznikne tak, že se **do vzorečku pro Pearsonův lineární korelační koeficient dosadí místo hodnot X a Y jejich pořadí v řadě**

$\rho = 1$, pokud jsou řady zcela shodné

$\rho = -1$, pokud jsou řady zcela protichůdné

$\rho = 0$, pokud mezi řadami není žádná tendence ke shodě či protichůdnosti, ale pořadí jsou k sobě zcela náhodně

Pořadová korelace - Kendallovo τ

- vypočte se z **počtu souhlasných a nesouhlasných dvojic** (shod a neshod v pořadí pro dvojice případů):
 - shoda dvojice: případy jsou u obou proměnných orientovány ve stejném směru, tj. když první je větší než druhý u X je větší i u Y (a naopak)
 - neshoda dvojice – případy opačně orientovány v pořadí, tj. když první je větší než druhý u X je menší u Y (a naopak)
- shody a neshody jsou znaménka při výpočtu kovariance vycházejícího ze všech dvojic**

$$\tau = (\text{počet shod} - \text{počet neshod}) / (\text{počet shod} + \text{počet neshod})$$

- vlastnosti stejné jako u Spearmanova koeficientu

Testy hypotéz

u korelačních koeficientů je základní dvojicí hypotéz, které testujeme:

H_0 : korelační koeficient je nulový
 H_A : korelační koeficient je nenulový

- prokazujeme, že naše spočtená míra je signifikantně nenulová, tedy, že v datech se projevuje nějaký vztah
- platí pro Pearsonův, Spearmanův i Kendallův koeficient

Použití korelací – konfirmační přístup

testování hypotéz a teorií: korelovanost prokazuje představu o kauzálním procesu

nebezpečí:

- korelovanost vznikne jiným procesem, než je předpokládáno
- data jsou nepřesná (nebo je jich málo) a korelovanost je proto rozumně a nepřesvědčivě nízká
- kauzální proces existuje, ale projevuje se jinými závislostmi než lineárními
- kauzální procesy se projevují jen na části souboru a korelace je tak překryta jinými vztahy

[illegible]

Parciální korelace

- X, Y, Z (tři známé proměnné v datech, jejichž vzájemné korelace jsou známy):

$$r(X, Y / Z) = \frac{r(X, Y) - r(X, Z) \times r(Y, Z)}{\sqrt{(1 - r(X, Z)^2) \times (1 - r(Y, Z)^2)}}$$

- je-li parciální korelační koeficient roven nebo blízko k nule (**nesignifikantní**), znamená to, že proměnná Z plně vysvětluje korelaci mezi X a Y
- pokud se parciální koeficient jen **podstatně redukuje**, znamená to, že Z ovlivňuje vztah X a Y, ale není samo



OPERAČNÍ PROGRAM
VÝSKUM, VÝVOJ
A ZAMĚSTNANOST

PODPORUJEME
VÁŠ BUDOUCNOST
www.mis.cz

50

Parciální korelace

Model 1: Společná příčina $X, Y, Z:$

$Z = \text{společná příčina}$

$r(X, Y) \neq 0, r(X, Z) \neq 0, r(Y, Z) \neq 0,$
 $r(X, Y/Z) = 0$

Model 2: Zprostředkující vlastnost

$X \longrightarrow Z \longrightarrow Y$ $r(X, Y) \neq 0, r(X, Z) \neq 0, r(Y, Z) \neq 0,$
 $r(X, Y/Z) = 0$

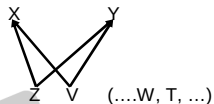
Modely 1 a 2 se nedají statisticky odlišit

ssf STATISTICKÝ ÚSTAV  EVROPSKÁ UNIE  OPERAČNÍ PROGRAM
LIDSKÉ ZDROJE
A ZAMĚSTNANOST PODPORUJEME
VAŠI BUDOUCNOST
www.mfcr.cz 51

Parciální korelace

model 3 – společné příčiny

X, Y, Z, V



Z = společná příčina

$r(X, Y) \neq 0$, $r(X, Z) \neq 0$, $r(Y, Z) \neq 0$, $r(X, V) \neq 0$, $r(Y, V) \neq 0$,
 $r(X, Y/Z) \neq 0$, $r(X, Y/V) \neq 0$
 $r(X, Y/Z, V) = 0$



52

Korelační analýza

- **korelační analýza je první stupeň analýzy vztahů, po ní následují:**
 - regresní analýza,
 - faktorová analýza,
 - analýza kovariančních struktur a kauzálních sítí
 - další speciální analýzy ...

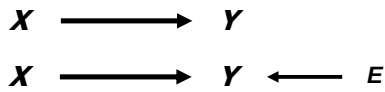


53

Lineární regrese



Regresní analýza: jednosměrný vztah



- nezávislá proměnná → závislá proměnná
- příčina → následek
- prediktor → predikant

u korelace: jednosměrný nebo symetrický vztah

u regrese: jednosměrný vztah

Popis vztahu rovnicí

- rovnice

model vztahu

$$Y = f(X) + \varepsilon$$

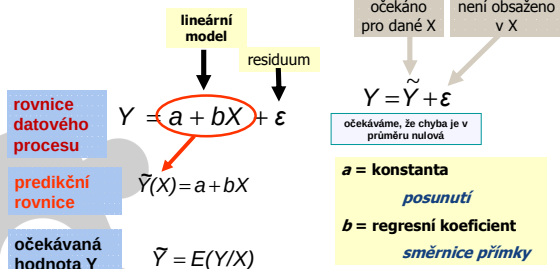
směr vztahu

forma převodu
chyba rovnice

rovnice rozloží hodnotu Y na dvě části:

- a) model = převod z X
- b) residuum/zbytek, které se neúčastní převodu

Dva významy rovnic: model vztahu/procesu a predikce



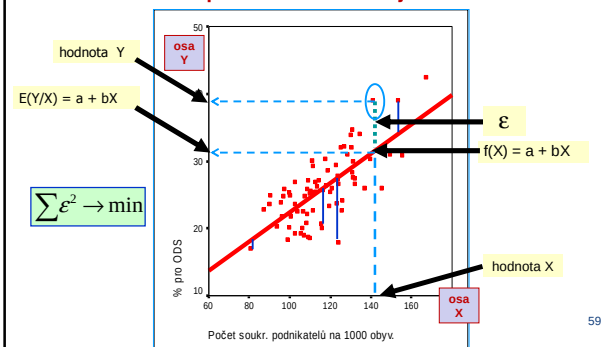
Významy parametrů: rovnoměrné posunutí a individuální změna

- **b** = regresní koeficient – koeficient úměry vlivu X na Y u každého jednoho případu
 - $b > 0$... přímka má růstový/stoupavý trend
kladný trend
s rostoucím X roste Y
 - $b < 0$... přímka má ztrátový/klesavý trend
záporný trend
s rostoucím X klesá Y
 - $b = 0$... přímka je rovnoběžná s osou X, absence trendu
s rostoucím X se Y nemění: nulový trend
hodnota Y na X nezávisí
- **a** = posunutí – hodnota Y pro nulové X nebo koeficient rovnoměrné změny pro každý případ bez ohledu na jeho X hodnotu

58

Lineární regrese: zachycení souběhu variabilit přímku

hledáme vhodnou přímku – metoda nejmenších čtverců



59

Dekompozice variability závislé (vysvětlované, cílové) proměnné Y

$$Y = a + bX + \varepsilon$$

$$\text{var}(Y) = \text{var}(E(Y/X)) + \text{var } \varepsilon$$

$$TSS = MSS + ESS$$

celkový (korigovaný k průměru) součet čtverců Y =
součet čtverců očekávaných/predikovaných
modelových hodnot +
součet čtverců odchylek (residuí, chyb)

$$TSS = \sum (Y_i - \text{ave}Y)^2$$

$$MSS = \sum (a + bX_i - \text{ave}Y)^2$$

$$ESS = \sum (Y_i - (a + bX_i))^2$$

60

Testování významnosti modelu: tabulka ANOVA

F - test – kritérium pro zjištění existence vztahu

$$F(1, n-2) = MSS / (ESS / n-2)$$

$$df = (1, n-2)$$

dosažená významnost F = alfa

$$F = R^2 / (1 - R^2) * (n-2)$$

n = počet případů

Dekompozice variability Y: koeficient determinace

variabilita Y je dekomponována do dvou částí korespondujících rozdělení hodnoty Y na dvě složky modelem:

$$Y = \text{očekávání/predikce} + \text{chyba/residuum}$$

$$\text{var } Y = \text{var}(E(Y/X)) + \text{var } \varepsilon$$

z toho logicky odvodíme míru determinace:

koeficient determinace:

$$R^2 = \text{var}(E(Y/X)) / \text{var}(Y)$$

$$= 1 - (\text{var } \varepsilon / \text{var}(Y))$$

$$= 1 - ESS / TSS = MSS / TSS$$

Koeficient determinace

koeficient determinace – míra explanační síly rovnice, intenzita působení X na Y lineárním vztahem

vyjádřeno %: $100 \cdot R^2 \%$

procento var Y vysvětlené modelovou rovnicí

$$R^2 = \text{čtverec korelačního koeficientu } r(X, Y)$$

$$R^2 = r(X, Y)^2$$

$$F = R^2 / (1 - R^2) * (n-2)$$

n = počet případů

Testy významnosti parametrů

$H_0: b = 0$ vs. $H_A: b \neq 0$

t - test – kritérium významnosti b

$$t(n-2) = b / \text{sterr}(b)$$

$$df = n-2$$

dosažená významnost $t = \alpha$

obdobně pro parametr a

pro rovnici s jedním prediktorem:

$$t^2 = F$$

Konfidenční intervaly pro b a a

konfidenční interval pro parametr – *kritérium přesnosti odhadu*

interval pro spolehlivost $\gamma = (1 - \alpha)$

$$b = \text{odhadnuté } b \pm t(n-2; \alpha/2) * \text{sterr}(b)$$

$$\gamma = 1 - \alpha$$

analogicky pro a

skutečná hodnota parametru není
intervalem zachycena s rizikem
 α

Konfidenční pás pro přímku (očekávané hodnoty k danému X)

konfidenční pás pro přímku – *kritérium pro přesnost určení přímky*

pás spolehlivosti pro $E(Y/X)$:

$$\text{odhad } (a + b * X) = \text{odhad } a + \text{odhad } b * X \\ \pm t(n-2; \alpha/2) * s * \sqrt{1/n + (X - \text{ave}X)^2 / \sum X_i^2}$$

$$\gamma = 1 - \alpha$$

skutečná přímka vztahu není
obsažena v pásu spolehlivosti s
rizikem α

Predikční pás pro jednotlivá pozorování (odhad Y pro dané známé X)

predikční pás – kritérium přesnosti predikce

predikční pás pro jednotlivé pozorování:

odhad $Y = a + bX$

$$\pm t(n-2; \alpha/2) * s * \sqrt{1 + 1/n + (X - \text{ave}X)^2 / \sum x_i^2}$$

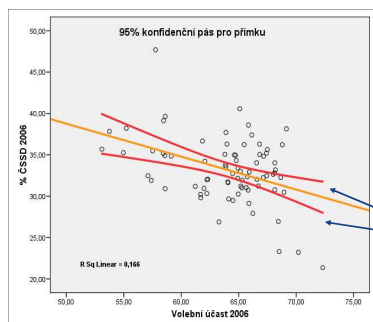
$\gamma = 1 - \alpha$

*predikovaná hodnota je mimo v
hranice pásu s rizikem α*



67

Regresní pás spolehlivosti

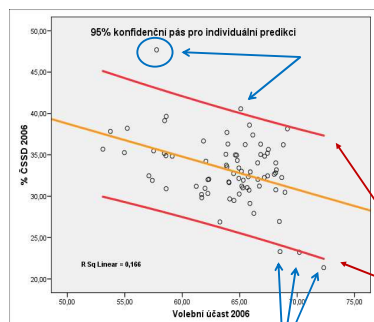


**intervaly spolehlivosti
pro parametry:**
a = (45.5; 71.89)
b = (-0.604; -0.193)

**pás spolehlivosti
pro regresní
přímku**

68

Predikční pás pro jednotlivá pozorování



**intervaly
spolehlivosti pro
parametry:**
a = (-5.09; 6.31)
b = (0.169; 0.266)

**pás predikce pro
jednotlivá X**

69

Chyba rovnice

- náhodné vlivy
 - při měření, zjišťování
- náhodné vlivy
 - při chování
- tvar funkce
- vlivy chybějících proměnných

regresní rovnice vyhlazuje tyto vlivy a očisťuje vztah o nepodstatné a nahodilé

Rozptyl reziduí - var ε

nevychýlený odhad chybové variance – kritérium
predikční přesnosti

$$s_r^2 = ESS/(n-2) = \sum(Y_i - \text{očekávané } Y_i)^2/(n-2)$$

směrodatná odchylka s_r
dosahuje minimální možné hodnoty pro lineární model – je to
kritérium pro výběr přímky

Předpoklady modelu

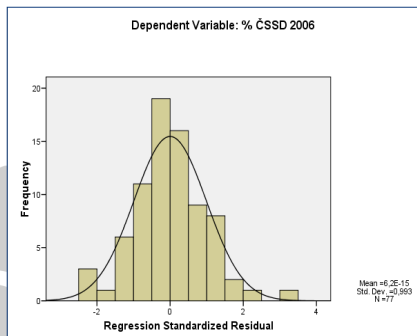
pro odhad metodou nejmenších čtverců:

- předpoklad tvaru – linearita a aditivita chyby
- nezávislost reziduí na prediktorech

pro testování:

- nezávislost reziduí mezi sebou
- homoscedasticita – rozptyl reziduí je konstantní vzhledem k X
- normalita reziduí

Normalita reziduí

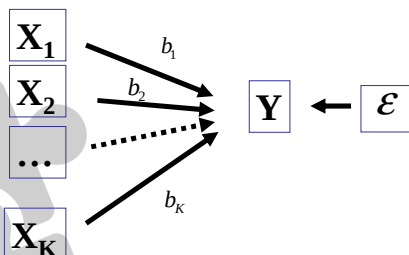


Vychýlená pozorování

- box plot standardizovaná rezidua
- standardizovaná rezidua (pokud jsou normálně rozložena) :
hodnoty přes ± 3.00 (or ± 2.00 , or $\pm c$)
- (seznamy extrémních příkladů)
- vizuální analýza – histogram reziduí
- Cookova distance (práh = $4/(n-p)$)
- Leverage (práh = $2p/n$; p = počet odhadovaných parametrů)
- dffbeta
- dffFit

Mnohorozměrná lineární regresní analýza

schéma lineárního regresního modelu:



Mnohorozměrná lineární regresní analýza

- je nalezení **predikční rovnice regresního modelu**:

$$Y = f(X_1, X_2, \dots, X_K) + \varepsilon$$

$$E(Y) = \tilde{Y} = f(X_1, X_2, \dots, X_K)$$

- hodnoty Y predikujeme pomocí proměnných $X_1, \dots, X_K$ až na chybu

Mnohorozměrná lineární regresní analýza

- lineární (nejčastější tvar) regrese**

$$Y = a + b_1 X_1 + b_2 X_2 + \dots + b_K X_K + e$$

převodní koeficient
 X_1 na Y

odhady z
dat

chyba rovnice
šum

Mnohorozměrná lineární regrese: rovnice s K nezávislými proměnnými

$$Y = a + b_1 X_1 + b_2 X_2 + \dots + b_K X_K + \varepsilon$$

$$Y = \tilde{Y} + \varepsilon$$

$$\tilde{Y} = a + b_1 X_1 + b_2 X_2 + \dots + b_K X_K$$

hodnota Y , kterou
očekáváme pro dané
hodnoty X_k

$$\tilde{Y} = E(Y / X_1, X_2, \dots, X_K)$$

$Y = \text{predikční model} + \text{residuum}$

predikční model = lineární kombinace prediktorů + konstanta

Vlastnosti rovnice

- Y je číselná proměnná, X jsou číselné proměnné
- koeficienty β jsou modelové hodnoty, koeficienty b jsou jejich odhady (pozor na dvojznačnost značení !!!)
- b_k je **převodní koeficient X_k na Y**
nazývá se **parciální regresní koeficient**
 b_k = **přírůstek Y při jednotkové změně X_k , jsou-li ostatní X beze změny**
= **čistý vliv X_k na Y , parciální, částečný**
(vliv očištěný od korelací X_k s jinými prediktory)
- obvyklé kritérium hledání modelu je **metoda nejmenších čtverců**:

$$\sum \varepsilon^2 \rightarrow \min$$

Mnohorozměrná lineární regrese: R^2 a R

Koeficient determinace = podíl/procento z var Y , které je vysvětleno množinou prediktorů pomocí lineárního modelu

$$R^2 = \frac{\text{var}(\tilde{Y})}{\text{var}(Y)} \\ = \frac{MSS}{TSS}$$

$$R^2 = (\text{var}(Y) - \text{var}(\varepsilon)) / \text{var}(Y)$$

Koeficient vícenásobné korelace R = korelační koeficient mezi Y a hodnotami predikcí/očekávání odvozených z modelové funkce (tj. z nalezené lineární kombinace nezávislých proměnných X).

Lineární kombinace modelu **maximalizuje** korelační koeficient s Y .

$$R = r(Y, a + b_1 X_1 + b_2 X_2 + \dots + b_K X_K)$$

80

Testování významnosti modelu

F - test – kritérium pro zjištění existence vztahu

kritérium toho, zda v model obsahuje významnou vztahovou informaci mezi množinou nezávislých proměnných a Y

$$F(k, n-k-1) = (ESS/k) / (MSS/n-k-1)$$

dosažená významnost F = α *

n = počet případů

k = počet regresních koeficientů

$$F = \frac{R^2 / (1 - R^2) * (n - k - 1)}{k}$$

n = počet případů

k = počet regresních koeficientů

Kolinearita

$$\hat{Y} = a + b_1 X_1 + b_2 X_2$$

- **kolinearita** – vysoká korelace mezi X_1 a X_2 (obecně vysoká korelovanost mezi nezávislými proměnnými) vede **k nestabilitě** odhadu koeficientů
- je **obtížné separovat** vlivy vysoce korelovaných proměnných
- přelévání dat mezi X_1 a X_2 pro nová pozorování a pro jiné datové soubory k témuž problému a z toho plynoucí změny v b_1 a v b_2
- **singularita** – kompletní korelovanost mezi skupinou prediktorů, tj. jeden lze plně vyjádřit jako lineární kombinaci ostatních
 - krajní případ kolinearity – *regresi nemůžeme počítat, jednu proměnnou je nutno vynechat*

Metody pro automatický výběr prediktorů

metody výběru prediktorů:

- prediktory jsou určeny - **ENTER** – všechny vstoupí do rovnice (rozhodnutí uživatele)
- 1. metoda **FORWARD** – postupné zařazování prediktorů
- 2. metoda **BACKWARD** – postupné vyřazování prediktorů
- 3. metoda **STEPWISE** – kombinace obou
- postupný vstup určených bloků proměnných (prediktorů)

Prokládání křivky

Prokládání křivky v IBM SPSS Statistics

Obecný model vztahu závislé proměnné y na nezávislé proměnné x :

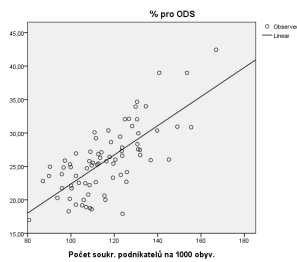
$$y = f(x; b) + \varepsilon$$

b ... vektor neznámých parametrů

ε ... náhodná chyba

- modely **vycházející z jednoduché lineární regrese**: funkce jsou lineární v parametrech nebo je lze na funkce lineární v parametrech převést vhodnou transformací (např. logaritmováním)
- vektor b odhadujeme **metodou nejmenších čtverců**, tj. na základě požadavku, aby součet čtverců chyb (reziduí) byl co nejmenší
- procedura *Curve Estimation* v *IBM SPSS Statistics Base* nabízí celkem 11 takových modelů

Model: Linear (přímka)



rovnice:

$$y = b_0 + b_1x + \varepsilon$$

posunutí ve směru osy y

sklon přímky

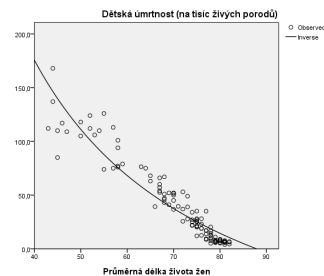
Model Summary and Parameter Estimates

Dependent Variable: % pro ODS

Equation	R Square	F	df1	df2	Sig.	Constant	b1
Linear	.515	79,696	1	75	.000	.609	.219

The independent variable is Počet soukr. podnikatelů na 1000 obyv.

Model: Inverse (hyperbola)



rovnice:

$$y = b_0 + \frac{b_1}{x} + \varepsilon$$

posunutí ve směru osy y

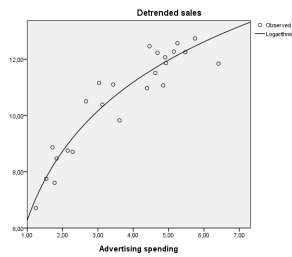
tvar křivky

Model Summary and Parameter Estimates

Equation	R Square	F	df1	df2	Sig.	Constant	b1
Inverse	.877	767,960	1	107	.000	-146,838	12908,889

The independent variable is Průměrná délka života žen.

Model: Logarithmic (logaritmus)



rovnice:

$$y = b_0 + b_1 \ln(x) + \epsilon$$

↑ ↑
posunutí ve tvar křivky
směru osy y

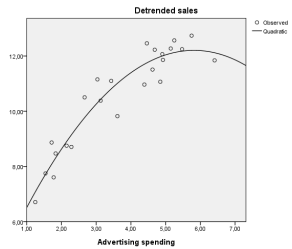
Model Summary and Parameter Estimates

Dependent Variable: Detrended sales

Equation	Model Summary				Parameter Estimates	
	R Square	F	df1	df2	Sig.	Constant
Logarithmic	.801	199.631	1	22	.000	6,274

The independent variable is Advertising spending.

Model: Quadratic (parabola)



rovnice:

$$y = b_0 + b_1x + b_2x^2 + \epsilon$$

souřadnice vrcholu:

$$V = \left[-\frac{b_1}{2b_2}, b_0 - \frac{b_1^2}{4b_2} \right]$$

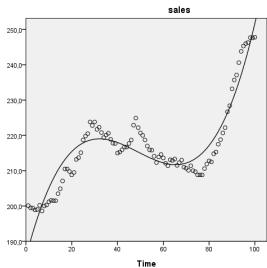
Model Summary and Parameter Estimates

Dependent Variable: Detrended sales

Equation	Model Summary				Parameter Estimates		
	R Square	F	df1	df2	Sig.	Constant	b1
Quadratic	.988	104,213	2	21	.000	3,903	2,854

The independent variable is Advertising spending.

Model: Cubic (kubická křivka)



rovnice:

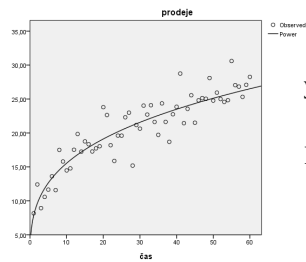
$$y = b_0 + b_1x + b_2x^2 + b_3x^3 + \epsilon$$

Model Summary and Parameter Estimates

Dependent Variable: sales

Equation	Model Summary				Parameter Estimates		
	R Square	F	df1	df2	Sig.	Constant	b1
Cubic	.877	227,808	3	96	.000	165,507	2,485

Model: Power (obecná mocnina)



rovnice:

$$y = b_0 x^{b_1} * (e^\varepsilon)$$

nebo

$$\ln(y) = \ln(b_0) + b_1 \ln(x) + \varepsilon$$

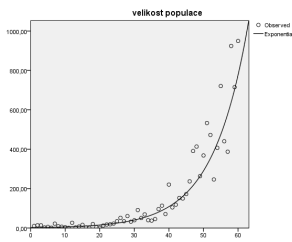
Model Summary and Parameter Estimates

Equation	Model Summary					Parameter Estimates	
	R Square	F	df1	df2	Sig.	Constant	b1
Power	.854	340.453	1	58	.000	7.832	.298

The independent variable is čas.

91

Model: Exponential (exponenciála)



rovnice:

$$y = b_0 \exp(b_1 x) * (e^\varepsilon)$$

nebo

$$\ln(y) = \ln(b_0) + b_1 x + \varepsilon$$

obecná exponenciála, růstová křivka a exponenciála se od sebe neliší typem průběhu křivky, ale pouze numerickými hodnotami a významem odhadovaných parametrů.

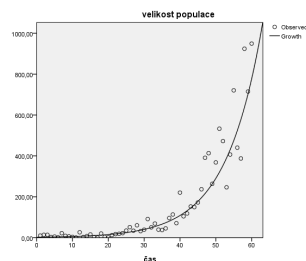
Model Summary and Parameter Estimates

Equation	Model Summary					Parameter Estimates	
	R Square	F	df1	df2	Sig.	Constant	b1
Exponential	.787	214.995	1	58	.000	2.105	.099

The independent variable is čas.

92

Model: Growth (růstová křivka)



rovnice:

$$y = \exp(b_0 + b_1 x) * (e^\varepsilon)$$

nebo

$$\ln(y) = b_0 + b_1 x + \varepsilon$$

obecná exponenciála, růstová křivka a exponenciála se od sebe neliší typem průběhu křivky, ale pouze numerickými hodnotami a významem odhadovaných parametrů.

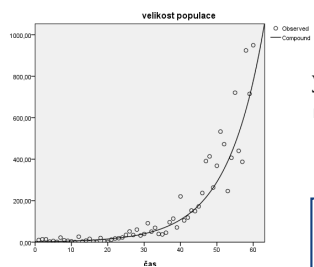
Model Summary and Parameter Estimates

Equation	Model Summary					Parameter Estimates	
	R Square	F	df1	df2	Sig.	Constant	b1
Growth	.787	214.095	1	58	.000	.744	.099

The independent variable is čas.

93

Model: Compound (obecná exponenciála)



rovnice:

$$y = b_0 b_1^x * (e^\varepsilon)$$

nebo

$$\ln(y) = \ln(b_0) + \ln(b_1)x + \varepsilon$$

obecná exponenciála, růstová křivka a exponenciála se od sebe neliší typem průběhu křivky, ale pouze numerickými hodnotami a významem odhadovaných parametrů.

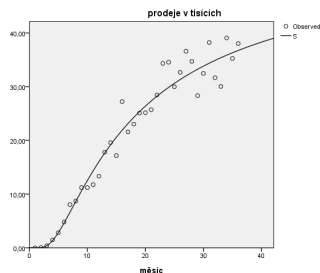
Model Summary and Parameter Estimates

Equation	Model Summary				Parameter Estimates	
	R Square	F	df1	df2	Sig.	Constant
Compound	.767	214.095	1	58	.000	2,105

The independent variable is čas.

94

Model: S (S-křivka)



rovnice:

$$y = \exp(b_0 + \frac{b_1}{x}) * (e^\varepsilon)$$

nebo

$$\ln(y) = b_0 + \frac{b_1}{x} + \varepsilon$$

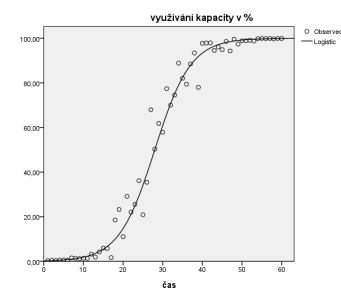
Model Summary and Parameter Estimates

Equation	Model Summary				Parameter Estimates	
	R Square	F	df1	df2	Sig.	Constant
S	.998	21324.839	1	34	.000	4,018

The independent variable is měsíc.

95

Model: Logistic (logistická křivka)



rovnice:

$$y = \frac{1}{\frac{1}{u} + b_0 b_1^x * (e^\varepsilon)}$$

nebo

$$\ln\left(\frac{1}{y} - \frac{1}{u}\right) = \ln(b_0) + \ln(b_1)x + \varepsilon$$

parametr u vyjadřuje hodnotu, kterou je funkce omezena shora (zde 100%)

Model Summary and Parameter Estimates

Equation	Model Summary				Parameter Estimates	
	R Square	F	df1	df2	Sig.	Constant
Logistic	.973	2093.785	1	58	.000	4,744

The independent variable is čas.

96
