

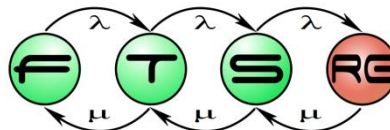
Az adatelemzés alapfeladatai

2017 ősz 6./7. alkalom

Kocsis Imre

ikocsis@mit.bme.hu

**Budapesti Műszaki és Gazdaságtudományi Egyetem
Hibatűrő Rendszerek Kutatócsoport**



MACHINE LEARNING / DATA MINING: ALAPFELADATOK

Adatelemzés: legfontosabb problémák

Csoportosítás
(*clustering*)

Osztályozás
(*classification*)

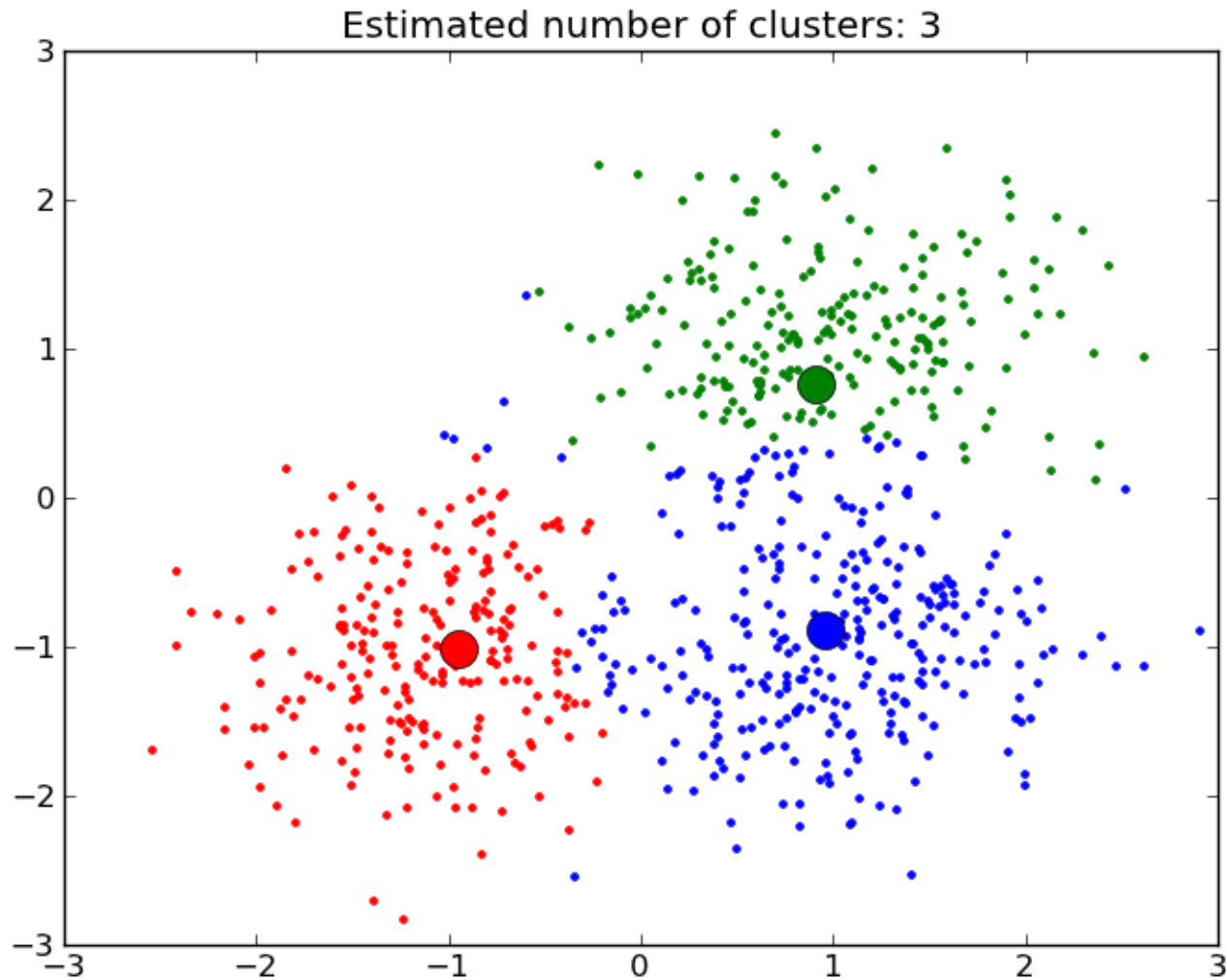
Asszoc. szabályok
(*assoc. rules*)

Regresszió
(*regression*)

Megközelítés

- Probléma-osztály;
 - Szemléltetés;
 - (Egy algoritmus kivonata)
-
- Fő cél: orientáció
 - ~~Deep learning, TensorFlow, ...~~
-
- N.B. egyre kevésbé „kézműves” tevékenység
 - „10 Algorithms every Data Scientist has to know”
 - SaaS/PaaS
 - Lásd később

Csoportosítás



K-means

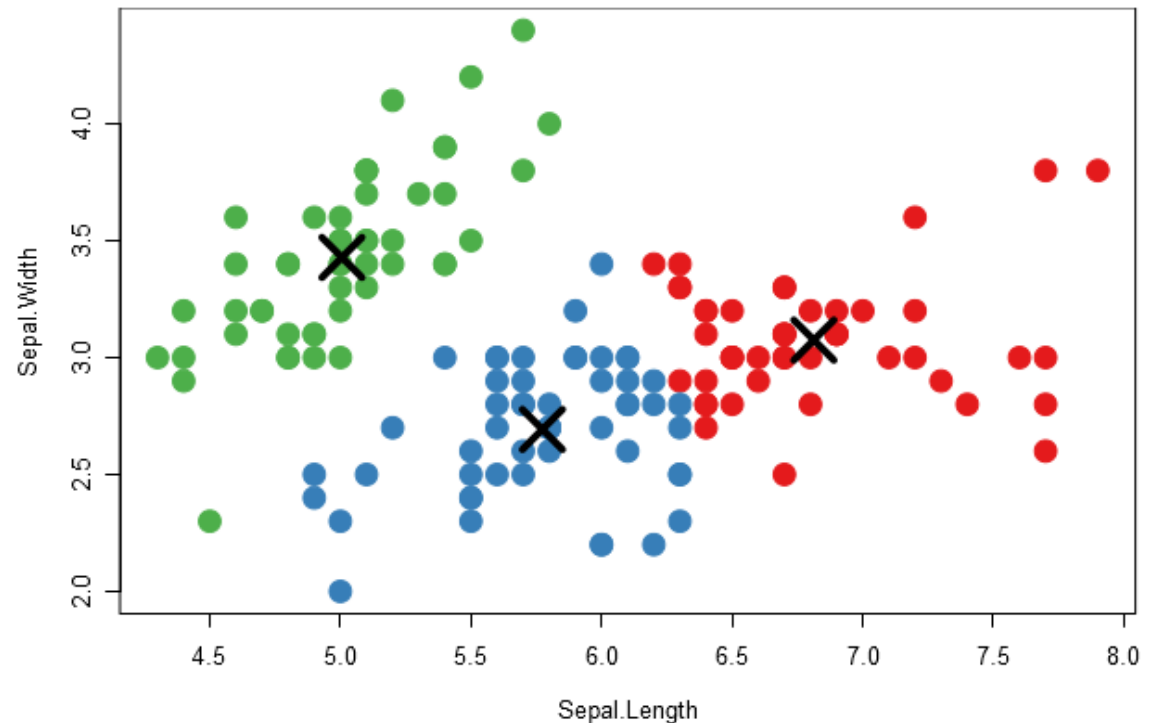
- Adatpontok: vektortér
- Klaszter reprezentációja: súlyponttal / "középponttal" (vektor-átlag)
- $r(C_i)$: i -edik klaszter reprezentánsa
- Minimalizálandó a négyzetes távolságösszeg, mint hiba:

$$E(C) = \sum_{i=1}^k \sum_{u \in C_i} d(u, r(C_i))^2$$

Demo

- <https://github.com/rstudio/shiny-examples/tree/master/050-kmeans-example>

Iris k-means clustering



Egy megoldás

```
{ $r(C_1), r(C_2), \dots, r(C_k)$ }  $\leftarrow$  repr. kezdeti halmaza  
while  $r(C_i)$  változik  
  do for  $\forall u \in D$  adott sorrendben  
    do  
       $h \leftarrow u$  klaszter-indexe  
       $j \leftarrow \operatorname{argmin}_i d(u, r(C_i))$   
      if  $h \neq j$  then {  
         $C_j \leftarrow C_j \cup \{u\}$   
         $C_i \leftarrow C_i \setminus \{u\}$   
         $r(C_j) \leftarrow \frac{1}{|C_j|} \sum_{v \in C_j} v$   
         $r(C_h) \leftarrow \frac{1}{|C_h|} \sum_{v \in C_h} v$   
      }  
return  $\mathcal{C}$ 
```

Régi klaszter

Új klaszter

Itt rögtön újra is számoljuk

Alapkérdések

- Klasztereken belül maximális homogenitás
 - „nagy hasonlóság”
 - „kis távolság”
- Klaszterek között „nagy távolság”
 - „kis hasonlóság” különböző klaszterek elemei között
- Hasonlósági mérce: „similarity metric”
 - Kategorikus változókra nehéz
 - Skálatranszformáció kellhet
 - Lehet választani...
 - Hierarchikus klaszterezés: dissimilarity metric
- Hasonlósági küszöb választása?
- Optimális klaszterszám?

Néhány távolság-mérték

Names	Formula
Euclidean distance	$\ a - b\ _2 = \sqrt{\sum_i (a_i - b_i)^2}$
Squared Euclidean distance	$\ a - b\ _2^2 = \sum_i (a_i - b_i)^2$
Manhattan distance	$\ a - b\ _1 = \sum_i a_i - b_i $
maximum distance	$\ a - b\ _\infty = \max_i a_i - b_i $

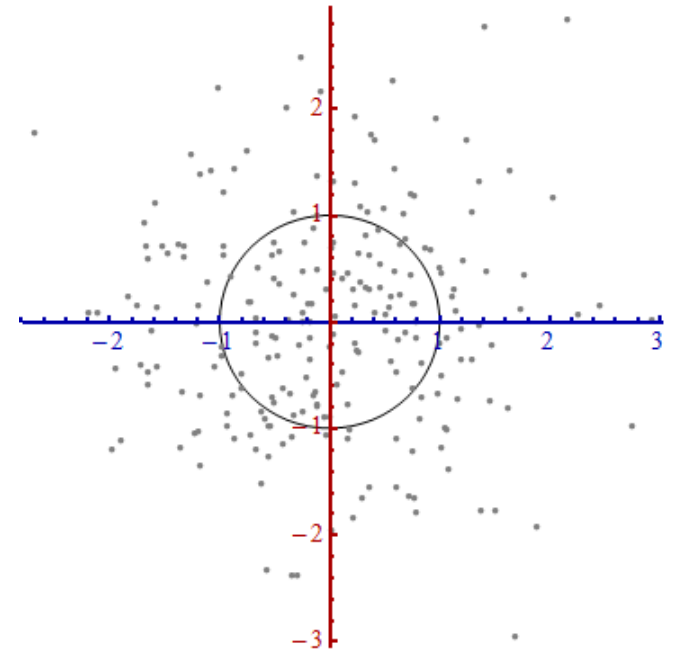
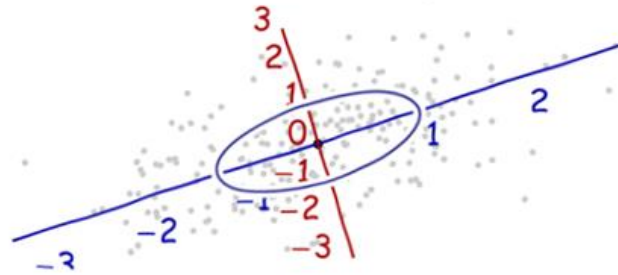
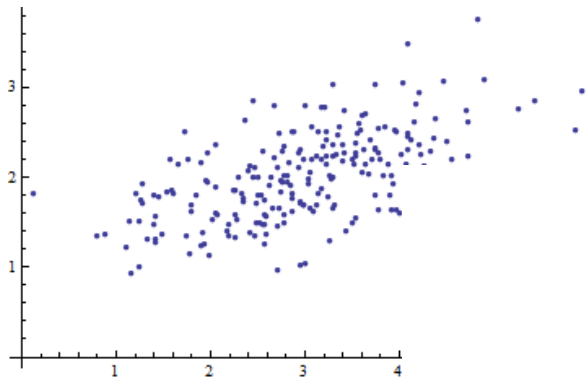
- (0,0) és (1,1) távolsága?

Mahalanobis távolság...?

- Pont és eloszlás távolsága (S: kov-mátrix)

$$D_M(\vec{x}) = \sqrt{(\vec{x} - \vec{\mu})^T S^{-1} (\vec{x} - \vec{\mu})}$$

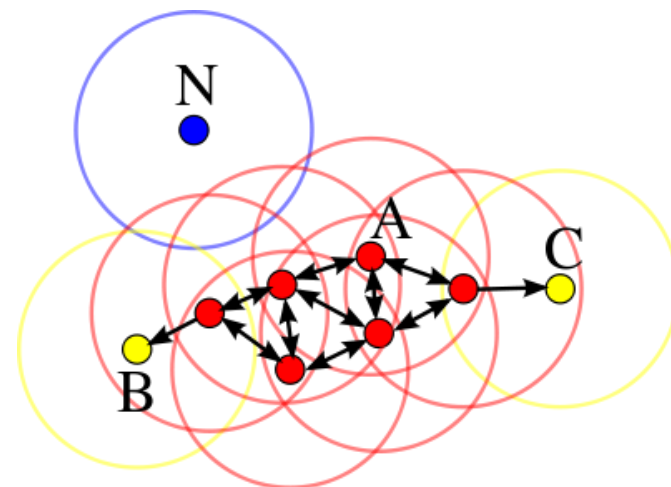
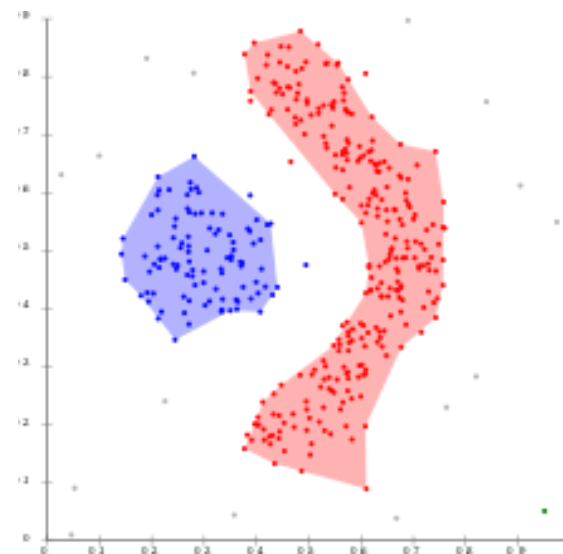
- “Szempléletesen”:



<https://stats.stackexchange.com/questions/62092/bottom-to-top-explanation-of-the-mahalanobis-distance>

További változatok

- Lineárisan nem szeparálható klaszterek: sűrűség alapú klaszterezés
- “packing together closely grouped points”
- PI. DBSCAN (“Density-based spatial clustering of applications with noise”)
 - Magpontok: legalább minPts pont e távolságon belül
 - Sűrűség-elérhető pontok
 - Outlierek



(Agglomeratív) hierarchikus klaszterezés

- CSAK KITEKINTÉS
- <https://github.com/joyofdata/hclust-shiny>

Hierarchical Clustering in Action

[joyofdata.de](#) / [@joyofdata](#) / [Hierarchical Clustering with R](#) / [github.com/joyofdata/hclust-shiny](#)

method

single

distance

euclidian

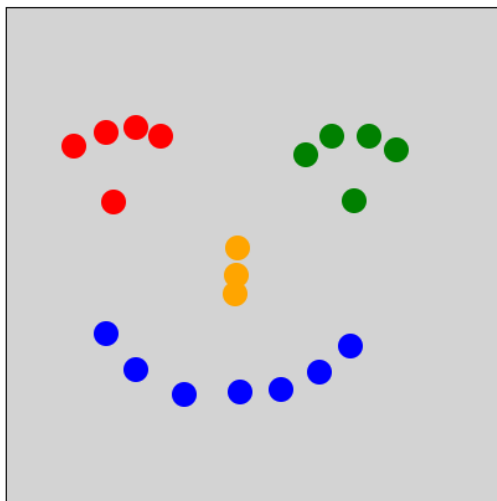
apply heuristic

min. max. branching gap

1

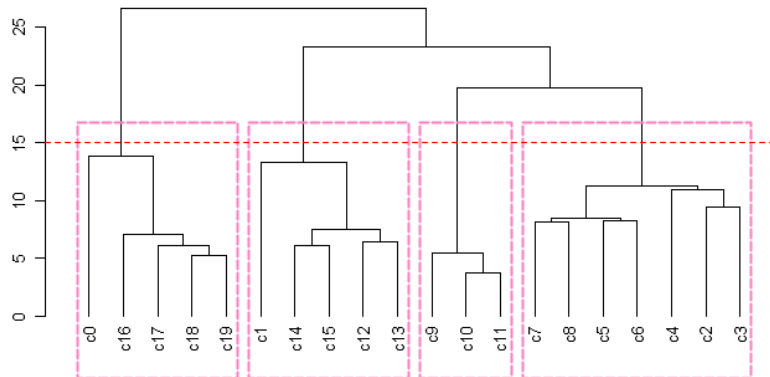
split tree at

15

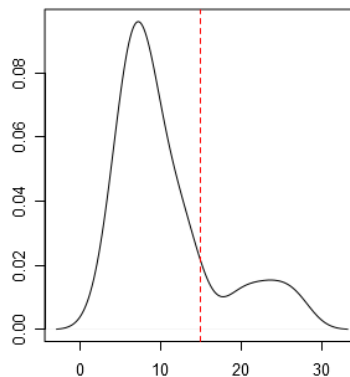


Left-Click to add or remove a point. Points can be dragged.

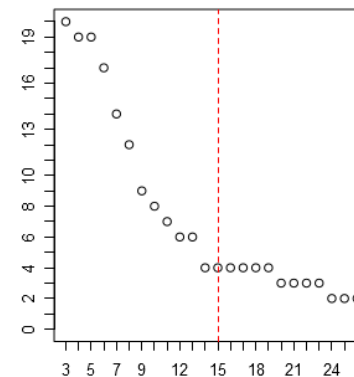
tree split at 15.00 - maximum branching height gap is 5.97



density of branching heights



num of clusters (y) when cutting at height (x)

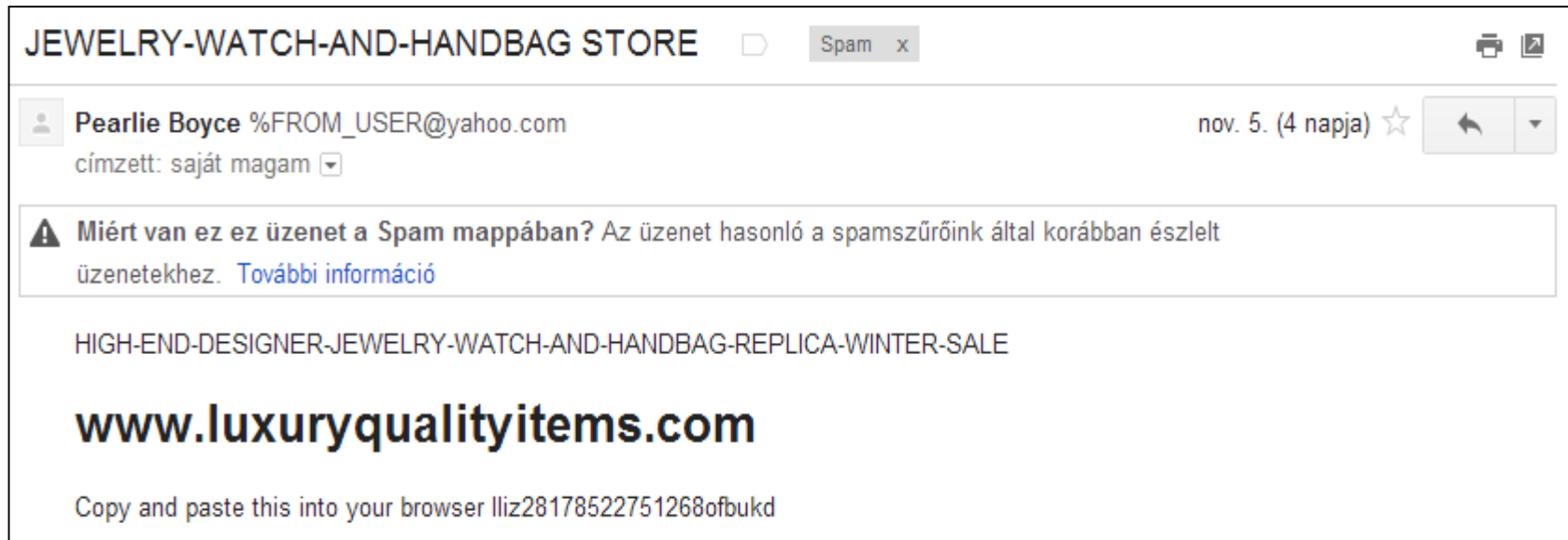


Osztályozás



Képosztályozás: a képen látható objektum madár vagy repülő?

Osztályozás



Levelek osztályozása: SPAM vagy nem SPAM?

Osztályozás

Netcool/OMNIBus Event List : Filter="Sun Events", View="Default"

File Edit View Alerts Tools Help

Sun Events Default

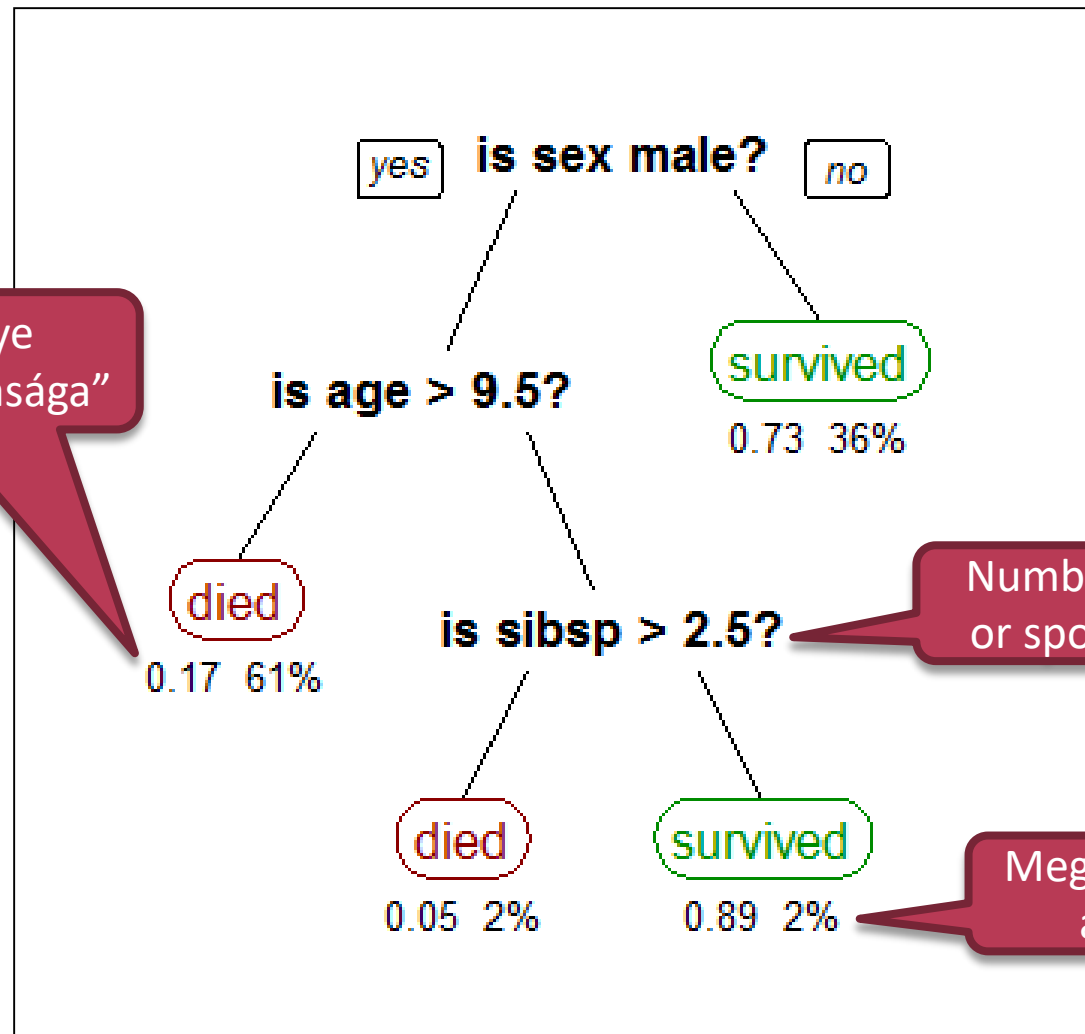
Node	Alert Group	Summary
wgs97-210	petTrapEntityPresenceDevice	Sensor io.id0.prstnt : A device is absent or has been removed.
wgs97-210	SunHwTrapPowerSupplyError	PowerSupply component /SYS/PS0/PWROK has detected an error (
wgs97-210	SunHwTrapPowerSupplyError	PowerSupply component /SYS/PS1/PWROK has detected an error (
wgs40-01	sunPlatStateChange	Device CH/FANBD1/FM2/SERVICE Operational Status has changed
wgs40-01	sunPlatStateChange	Device CH/MB/CMP0/BR3/CH0/D0/SERVICE Operational Status ha
wgs40-01	sunPlatStateChange	Device CH/SYS/PS_FAULT Operational Status has changed to disab
wgs40-01	sunPlatStateChange	Device CH/SYS/SERVICE Operational Status has changed to disable
wgs40-01	sunPlatStateChange	Device CH/SYS/TEMP_FAULT Operational Status has changed to di
wgs97-216	petTrapFanLowerCritical	Sensor Axial Fan 0 : Fan speed has decreased below lower critical thre
wgs97-210	SunHwTrapPreOSError	PreOS Error. Additional Info : System boot initiated
wgs97-210	SunHwTrapPreOSError	PreOS Error. Additional Info : Motherboard initialization
wgs97-210	petTrapPowerSupplyState	Sensor ps0.pwrok : Power Supply is connected to AC Power.
wgs97-210	petTrapPowerSupplyState	Sensor ps1.pwrok : Power Supply is connected to AC Power.

0 2 0 2 8 1

No rows modified 12/11/2007 2:42:25 P root NCOMS [PRI]

Szabályok alapján Severity
osztályozása

Döntési fák (Titanic, túlélési esélyek)



Túlélés esélye
-> osztály "tisztasága"

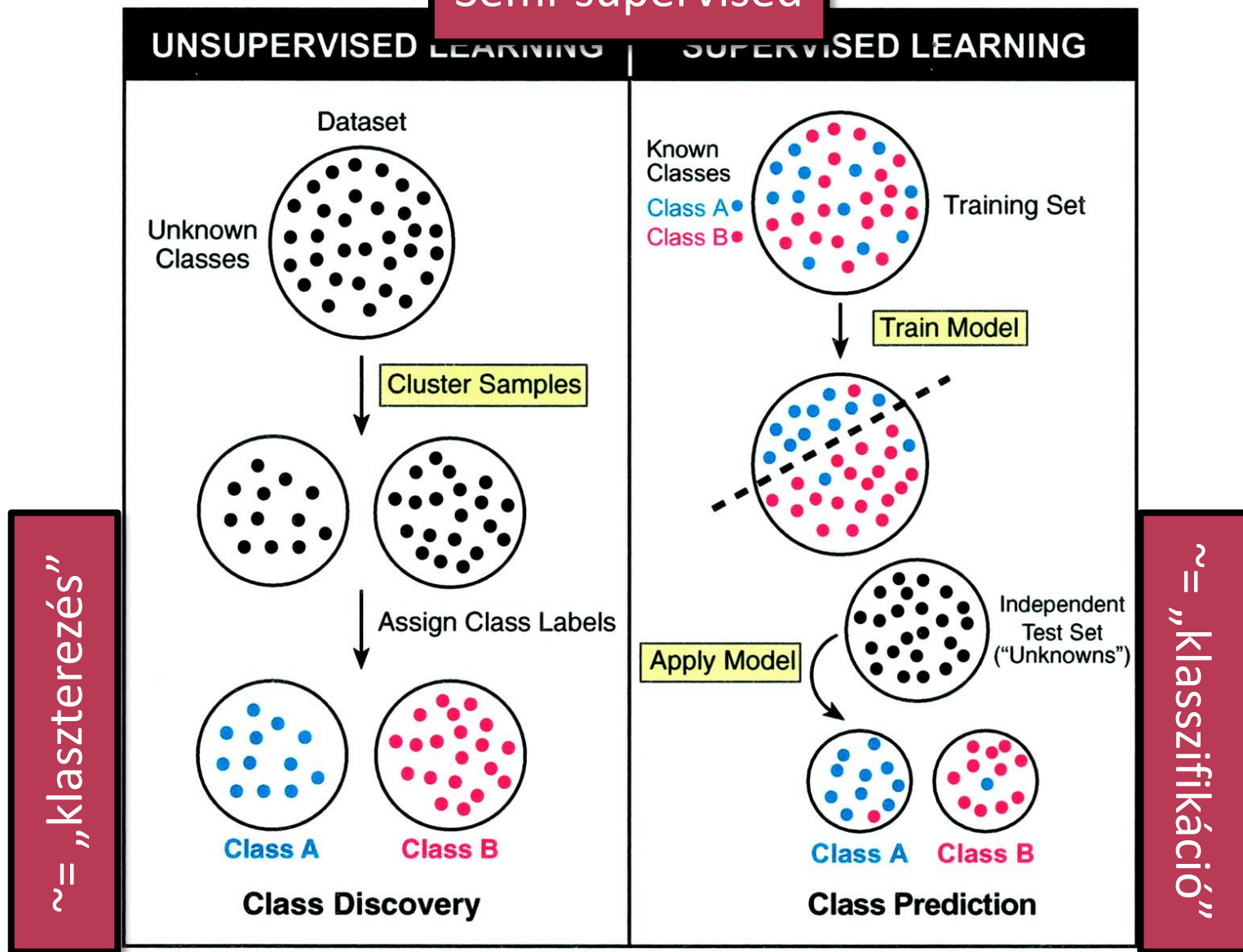
Number of siblings
or spouses aboard

Megfigyelések
aránya

https://en.wikipedia.org/wiki/Decision_tree_learning

Klaszterezés és klasszifikáció alapfeladat

Semi-supervised



Felügyelt és nem felügyelt tanulás

■ Felügyelt tanulás

- Adott néhány pontra az elvárt kimenet is
- $\approx$ a tanuló példákból való általánosítás
- Output: függvény
 - a meglévő mintapontokra jól képez le
 - megfelelően általánosítható

Tanulóhalmaz:
amin építjük a
modellt
Teszthalmaz:
amin ellenőrizzük

■ Nem felügyelt tanulás

- Nincs meg az elvárt kimenet
- Visszajelzés nélkül építi a modellt
- $\approx$ szabályok, összefüggések keresése (ismeretfeltárás)

Demo

- R „party” csomag
- Conditional Inference Tree, iris
 - Hogy ez mi, azt nem kezdjük itt nagyon részletezni
 - Party ctree dokumentáció: *“Roughly, the algorithm works as follows:*
 - *1) Test the global null hypothesis of independence between any of the input variables and the response (which may be multivariate as well).*
 - *Stop if this hypothesis cannot be rejected.*
 - *Otherwise select the input variable with strongest association to the response. This association is measured by a p-value corresponding to a test for the partial null hypothesis of a single input variable and the response.*
 - *2) Implement a binary split in the selected input variable.*
 - *3) Recursively repeat steps 1) and 2).”*

Demo

- R „party” csomag
- Conditional Inference Tree, iris
 - Hogy ez mi, azt nem kezdjük itt nagyon részletezni
 - Party ctree dokumentáció: *“Roughly, the algorithm works as follows:*

Kis érték (tip. 0.05) indikatív a (függetlenségi) hipotézis ellen *as well).*

- *Stop if this hypothesis cannot be rejected.*
- *Otherwise select the input variable with strongest association to the response. This association is measured by a p-value corresponding to a test for the partial null hypothesis of a single input variable and the response.*
- *2) Implement a binary split in the selected input variable.*
- *3) Recursively repeat steps 1) and 2).”*

Pl. entrópia-csökkenés maximalizálásával: a split csökkenti az össz-entrópiát (alágak entrópiájának súlyozott összege)

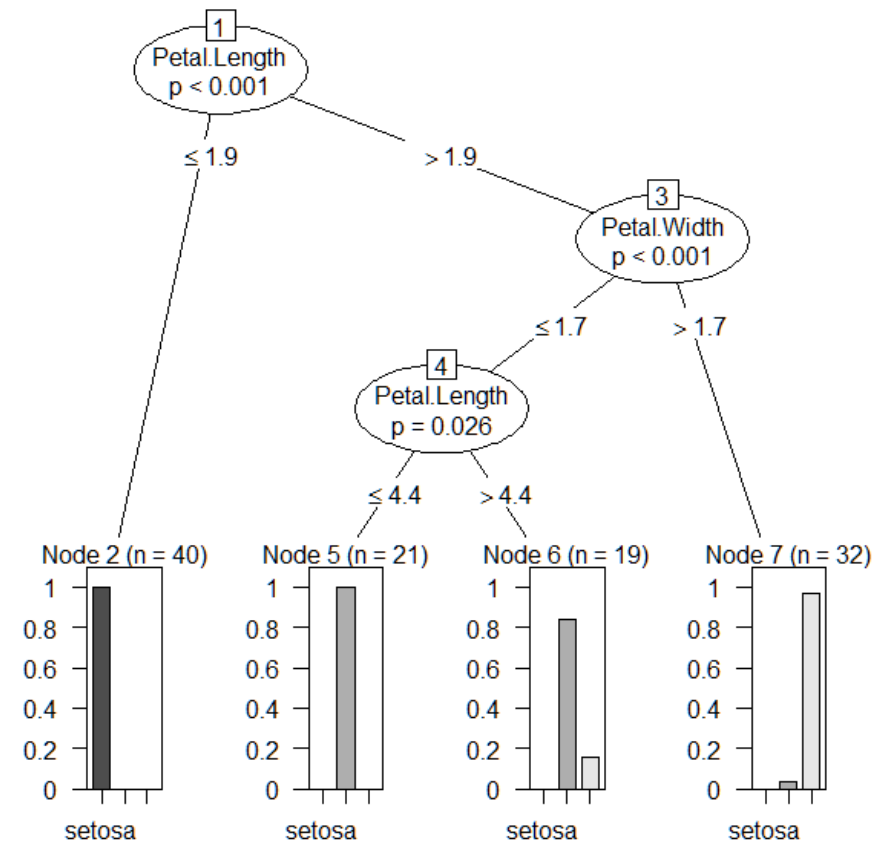
Demo

```
> myFormula <- Species ~ Sepal.Length + Sepal.Width + Petal.Length + Petal.Width
> iris_ctree <- ctree(myFormula, data = train.data)
> table(predict(iris_ctree), train.data$Species)
```

	setosa	versicolor	virginica
setosa	40	0	0
versicolor	0	37	3
virginica	0	1	31

```
> plot(iris_ctree)
> testPred <- predict(iris_ctree, newdata = test.data)
> table(testPred, test.data$Species)
```

	setosa	versicolor	virginica
testPred	10	0	0
setosa	10	0	0
versicolor	0	12	2
virginica	0	0	14



Bináris döntések jóságának mérése

Confusion
matrix

Jóslt állapot (predikció)

Valós állapot

	Jóslt állapot (predikció)	
	P	N
P	TP	FN (II. típusú hiba)
N	FP (I. típusú hiba)	TN

Érzékenység vagy specifikusság a fontos?

■ Érzékenység (sensitivity): $TP/(TP+FN)$

- érzékenység magas $\rightarrow$ FN alacsony
- Pl: Legyen kicsi az esélye, hogy egy fertőző betegség tesztje tévesen negatív $\rightarrow$ aki negatív eredményű az valójában is az.

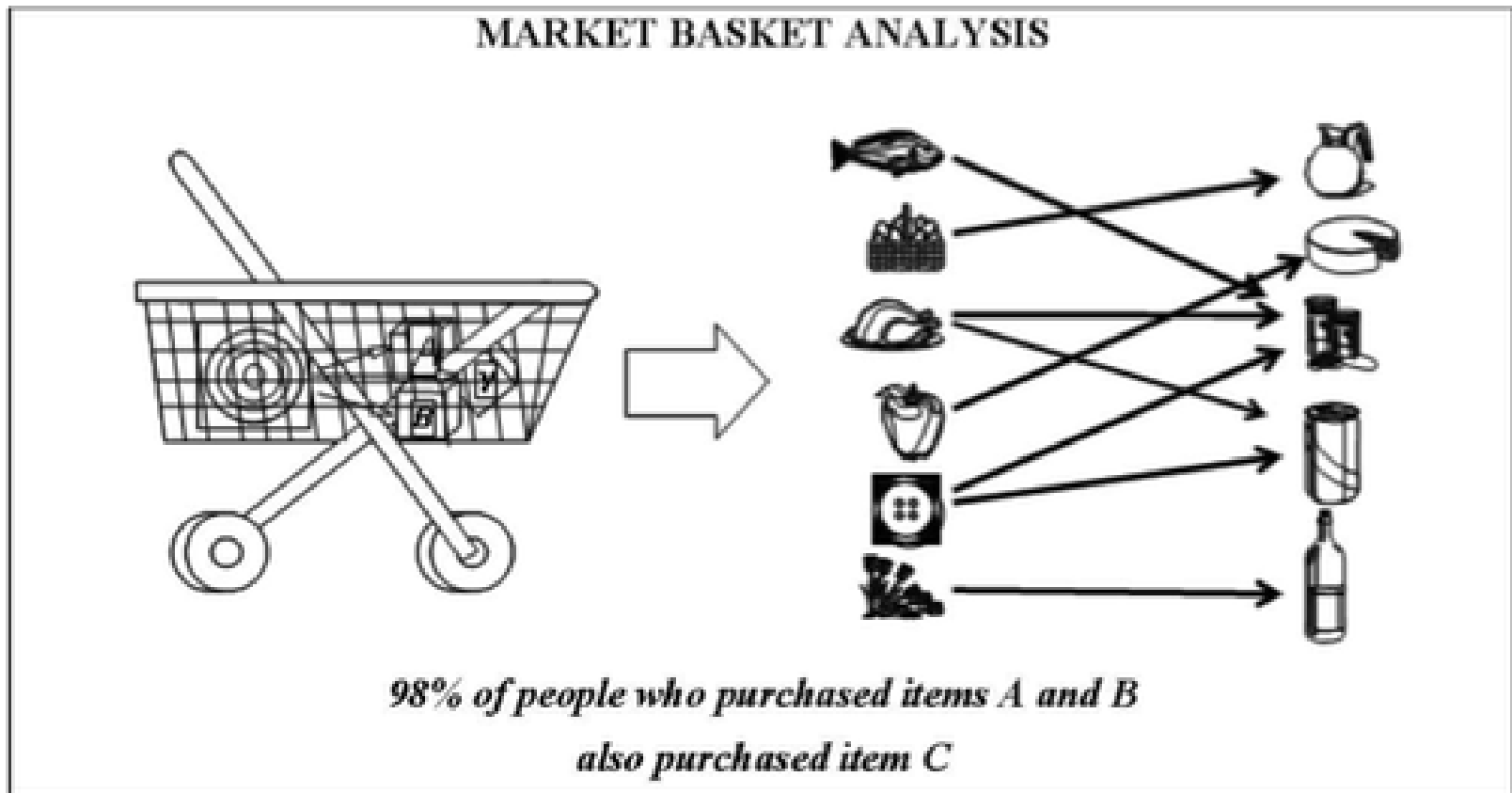
■ Specifikusság (specificity): $TN/(TN+FP)$

- specifikusság magas $\rightarrow$ FP alacsony
- Pl: Legyen kicsi az esélye, hogy egy genetikai teszt (egy kezelés megkezdéséhez) tévesen pozitív $\rightarrow$ aki pozitív eredményű az valójában is az.

További jellemzők:

https://en.wikipedia.org/wiki/Confusion_matrix
<http://people.inf.elte.hu/kiss/11dwhdm/roc.pdf>

Asszociációs szabályok



Alapfogalmak

- Asszoc. szabályok: elemhalmazok közötti asszociáció vagy korreláció
 - if LEFT then RIGHT
- Pl. Tx DB-k: „sör + pelenka + fejfájáscsillapító”
- Adatkeretre is működik

„Sokszor igaz”

$$\text{support}(LEFT \rightarrow RIGHT) = \frac{N_{BOTH}}{N_{TOTAL}}$$

„Ritkán téved”

$$\text{confidence}(LEFT \rightarrow RIGHT) = \frac{N_{BOTH}}{N_{LEFT}}$$

„Nem véletlen”

$$\text{lift}(LEFT \rightarrow RIGHT) = \frac{\text{confidence}(LEFT \rightarrow RIGHT)}{N_{RIGHT}}$$

Alapfogalmak

$$\blacksquare \text{ lift}(B \rightarrow 1) = \frac{P(\{B;1\})}{P(\{B;* \})P(\{*;1\})} \approx$$

1.56

- >1 : “többször fordulnak elő együtt, mint várható”
 - “ha pelenka, akkor többször sör”
- <1 : “kevesebbszer fordulnak elő együtt, mint várható”

Antecedent	Consequent
A	0
A	0
A	1
A	0
B	1
B	0
B	1

Lásd még: [https://en.wikipedia.org/wiki/Lift_\(data_mining\)](https://en.wikipedia.org/wiki/Lift_(data_mining))

Néhány megfontolás

- Potenciálisan rengeteg szabály
- Amit szeretnénk: „elég magas” confidence **és** support
- Redundanciák is lehetnek
- A „legérdekesebbeket” bányásszuk ki
 - Valamilyen „érdekességi” metrika alapján

Demo (Titanic, ismét)

- <http://brooksandrew.github.io/simpleblog/articles/association-rules-explore-app/>

Association Rules

Choose # of rule clusters

1 15 150

Sample

☒ All Rules ☐ Sample

Choose variables:

☒ Class
☒ Sex
☒ Age
☒ Survived

Support:

0 0.3965 1

Confidence:

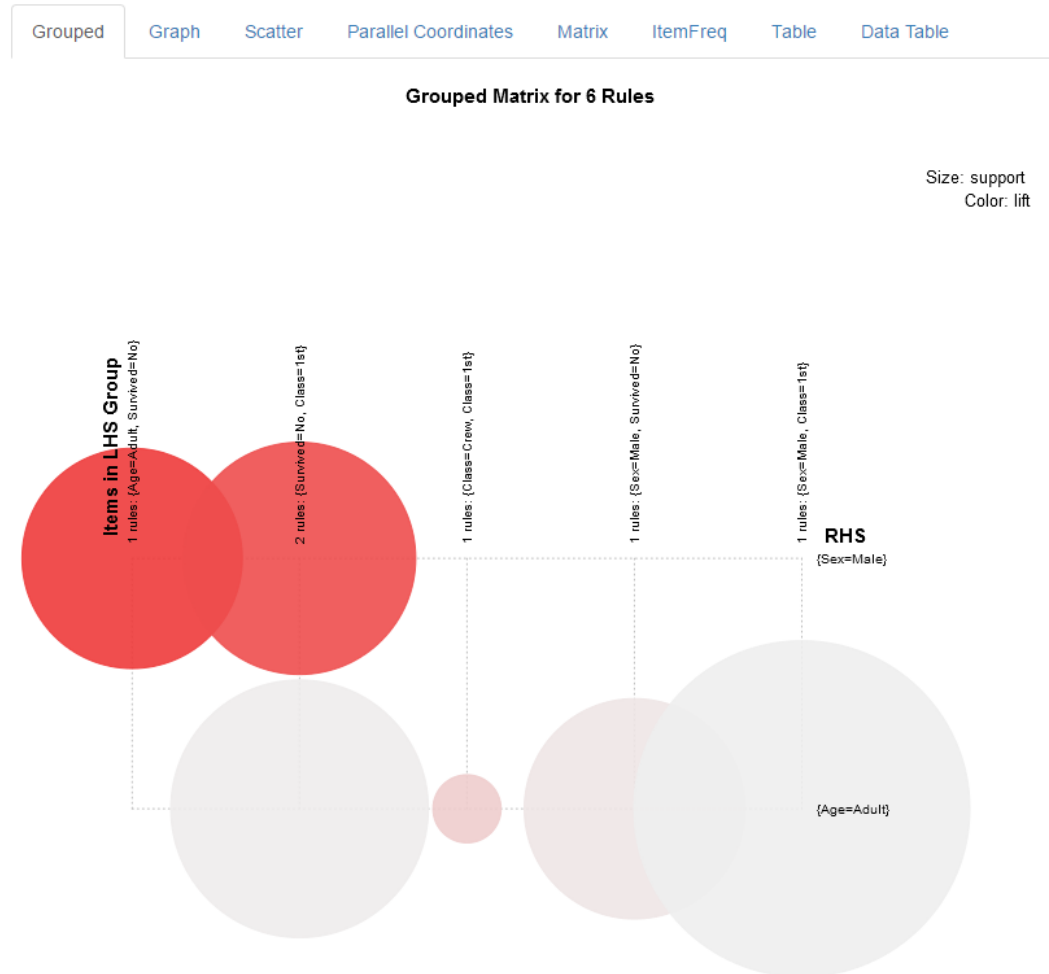
0 0.8983 1

Sorting Criteria:

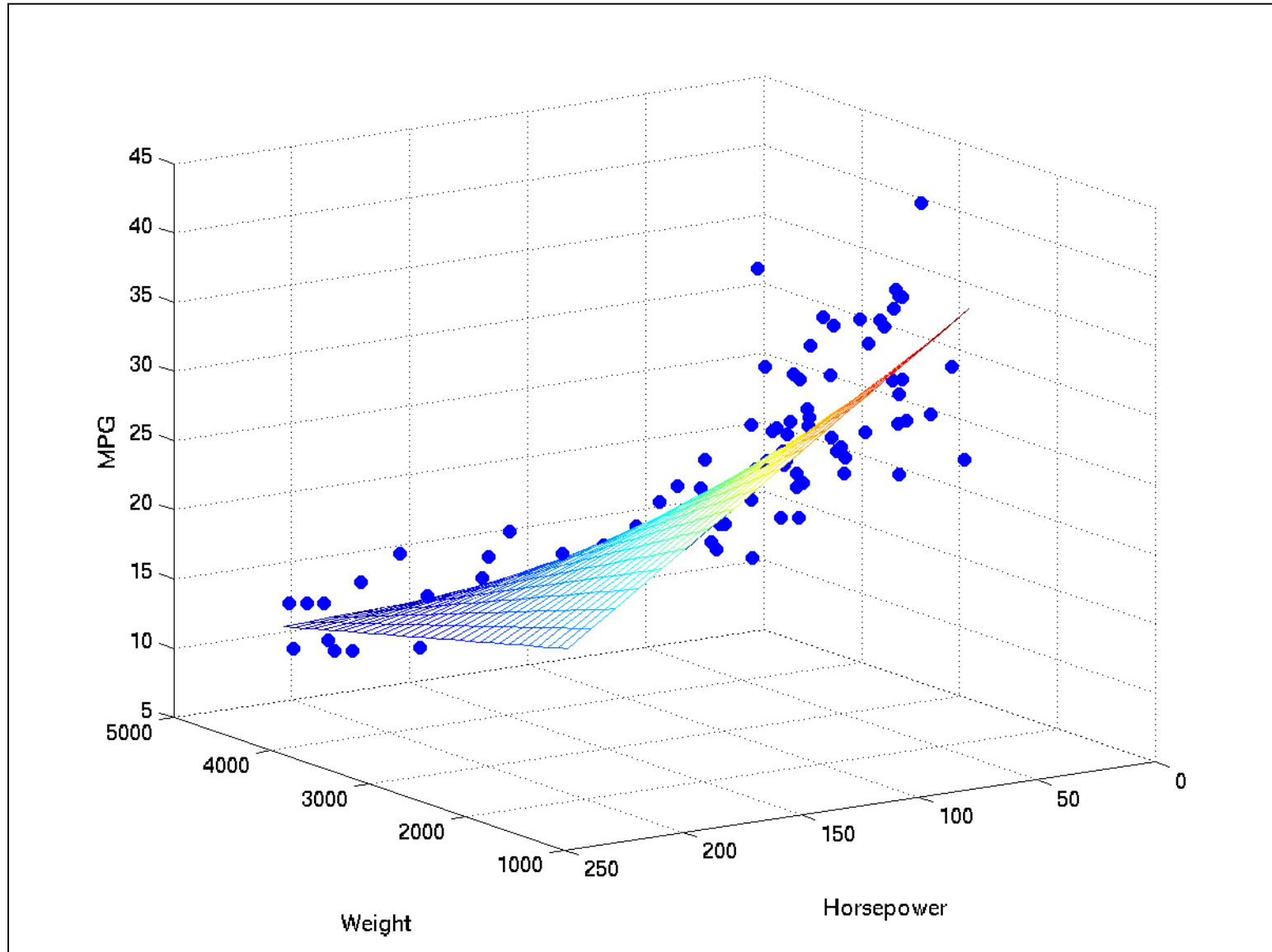
lift

Min. items per set:

2



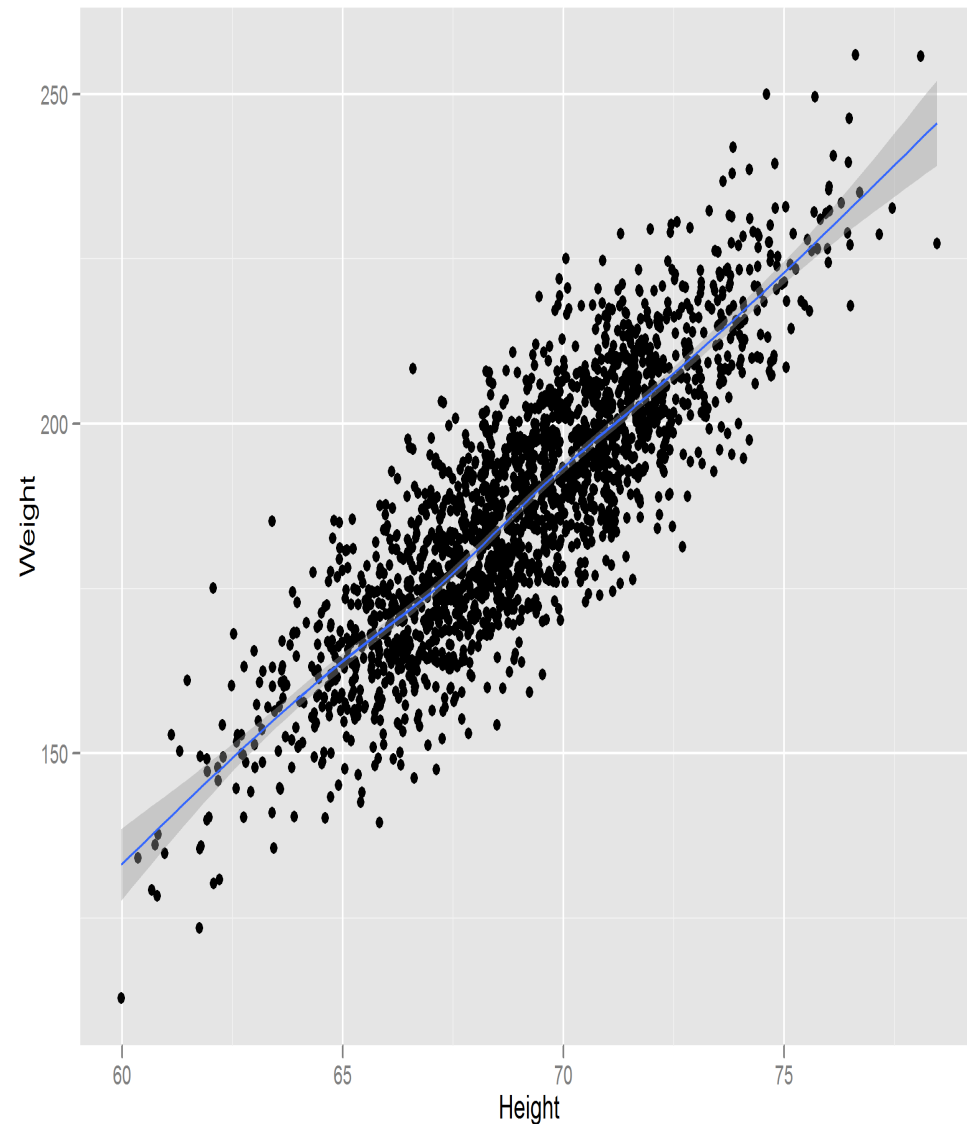
Regresszió



Regresszió

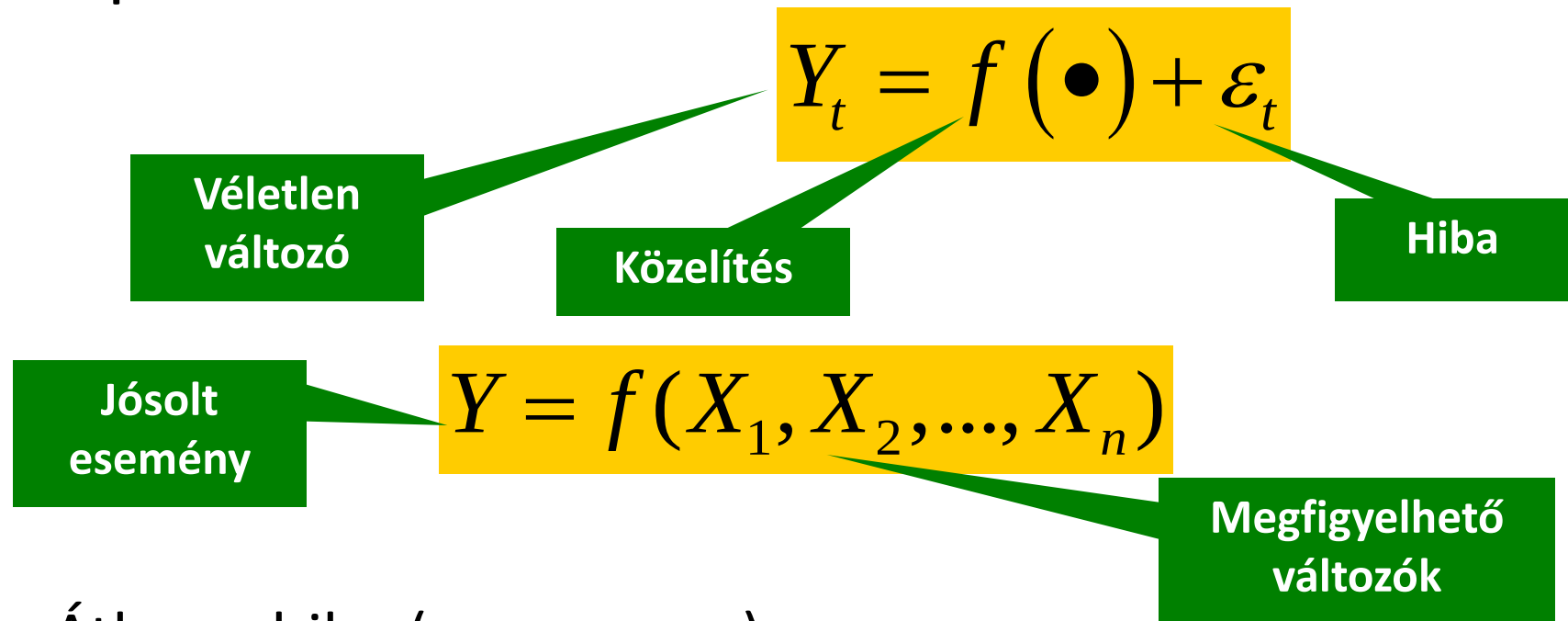
f függvény,

- bemenet:
az attribútumok értéke,
- kimenet: megfigyelések
legjobb közelítése
- „ökölszabály”
- Példa:
testtömeg/magasság
együttes eloszlás
valójában egyenesre
illeszthető

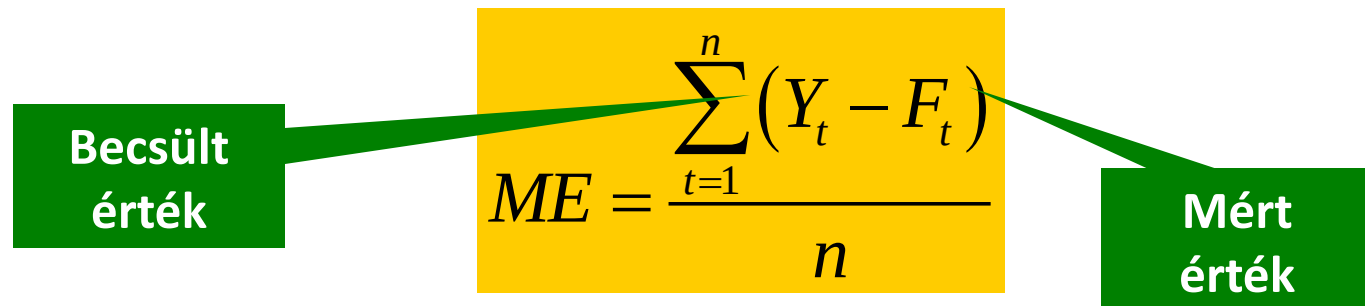


Regressziós módszerek

- Alapelv:



- Átlagos hiba (mean error)



Lineáris regresszió

- Egyszerű lin. függvény illesztése az adatokra
 - “nem vár alapvető változást a rendszer viselkedésében”

$$Y = a + bX$$

- Legkisebb négyzetek módszere
 - keressük azokat az a, b paramétereket, amelyekre

$$SSE = \sum_{t=1}^n \varepsilon_t^2 = \sum_{t=1}^n (Y_t - F_t)^2 \quad \text{minimális (Sum of Squared Errors)}$$

- cél:

$$\sum_{t=1}^n (Y_t - F_t)^2 = \sum_{t=1}^n [Y_t - (a + bX_t)]^2$$

Zárt alak: levezetés (parc. deriválás)

$$\frac{d \sum_{t=1}^n [Y_t - (a + bX_t)]^2}{da} = \sum_{t=1}^n (-2) [Y_t - (a + bX_t)] = 0$$

$$na = \sum_{t=1}^n (Y_t - bX_t)$$

$$a = \bar{Y} - b\bar{X}$$

**X_i, Y_i a mért értékpárok
(pl. idő, terhelés)**

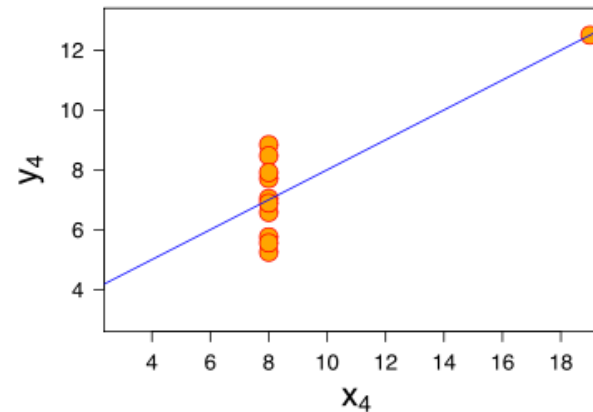
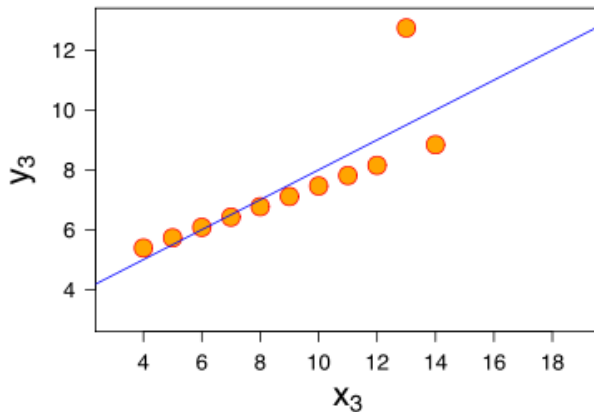
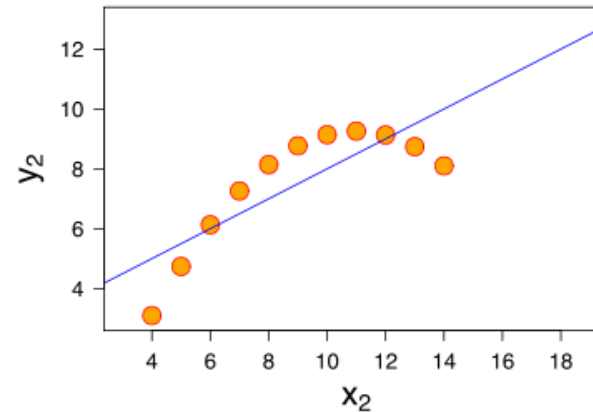
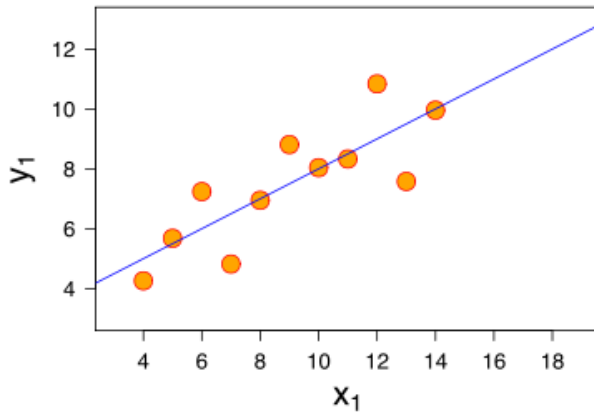
$$\frac{d \sum_{t=1}^n [Y_t - (a + bX_t)]^2}{db} = \sum_{t=1}^n X_t [Y_t - (a + bX_t)] = 0$$

$$\sum_{t=1}^n X_t \left[Y_t - \frac{1}{n} \sum_{t=1}^n (Y_t - bX_t) - bX_t \right] = \sum_{t=1}^n X_t Y_t - \frac{1}{n} \left(\sum_{t=1}^n X_t \right) \left(\sum_{t=1}^n Y_t \right) + \frac{1}{n} b \left(\sum_{t=1}^n X_t \right) \left(\sum_{t=1}^n X_t \right) - b \sum_{t=1}^n X_t^2 = 0$$

$$b = \frac{n \sum_{t=1}^n X_t Y_t - \left(\sum_{t=1}^n X_t \right) \left(\sum_{t=1}^n Y_t \right)}{n \sum_{t=1}^n X_t^2 - \left(\sum_{t=1}^n X_t \right)^2}$$

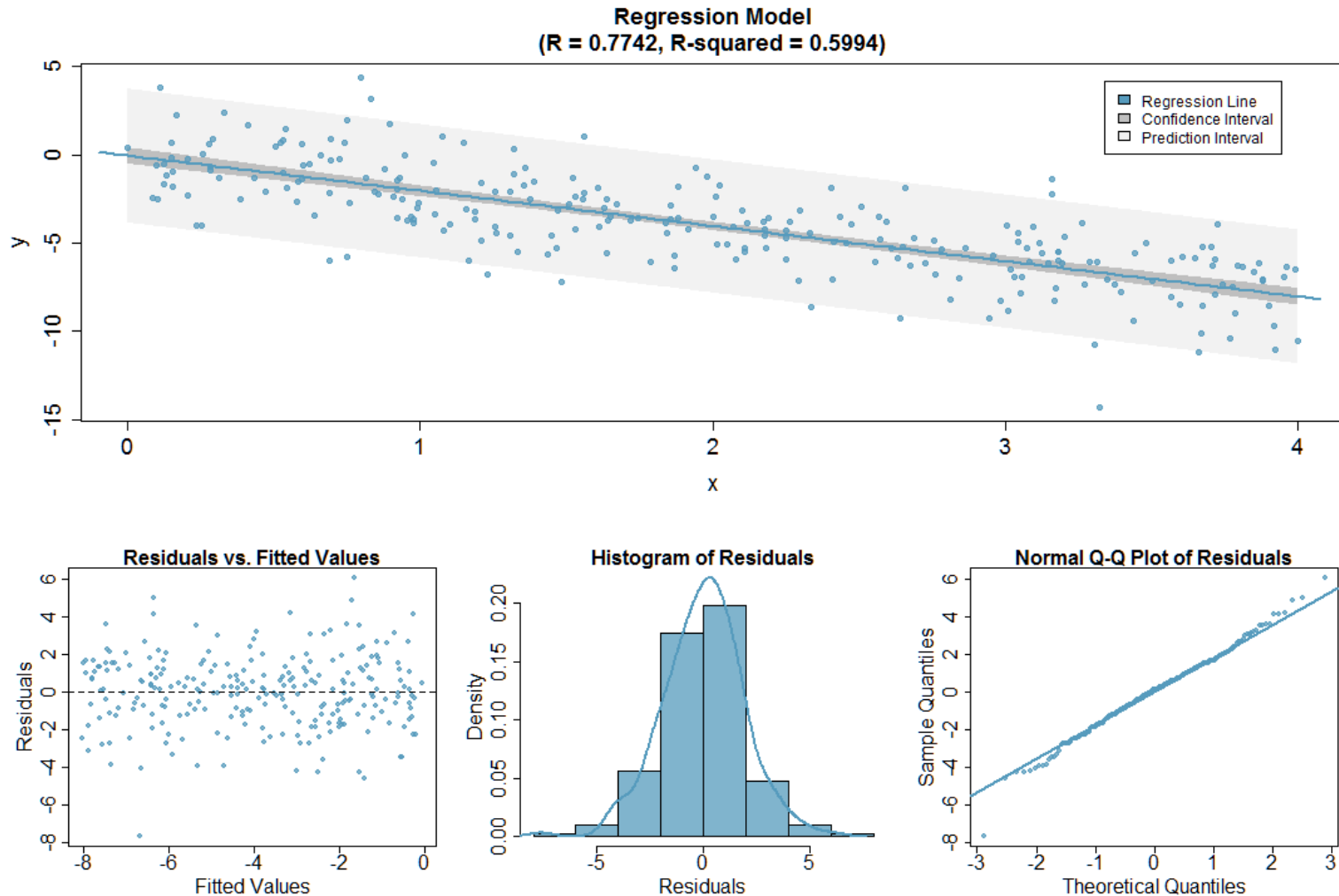
Ismétlés: Anscombe négyese

- Legjobban illeszkedő egyenes mindenre van
- Minőségileg különböző adatpontokra is



Demo

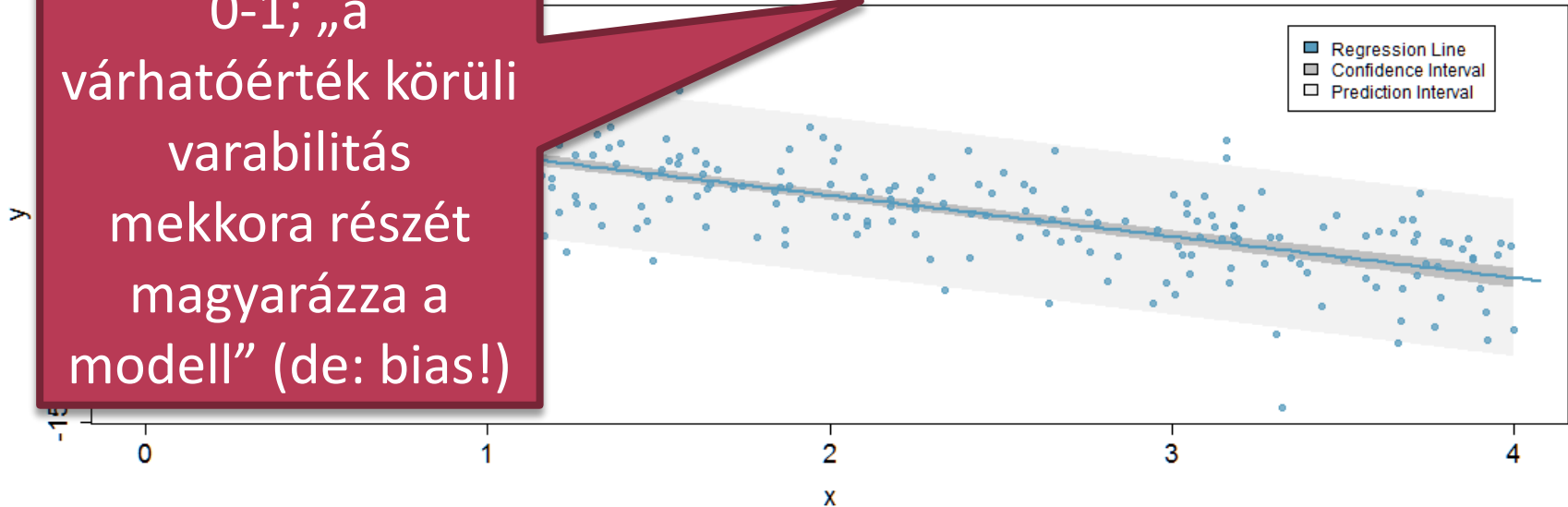
- https://github.com/ShinyEd/ShinyEd/tree/master/slr_diag



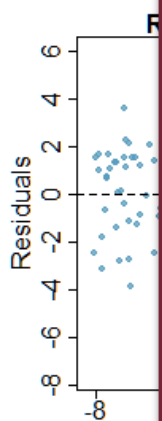
Néhány tulajdonság

0-1; „a várhatóérték körüli variabilitás mekkora részét magyarázza a modell” (de: bias!)

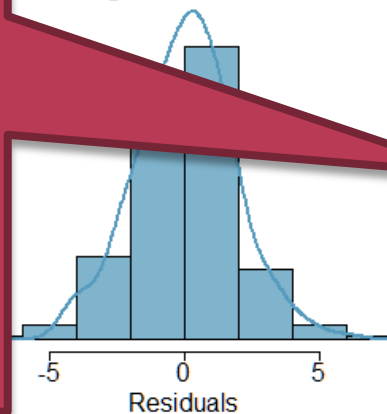
Regression Model
($R = 0.7742$, $R\text{-squared} = 0.5994$)



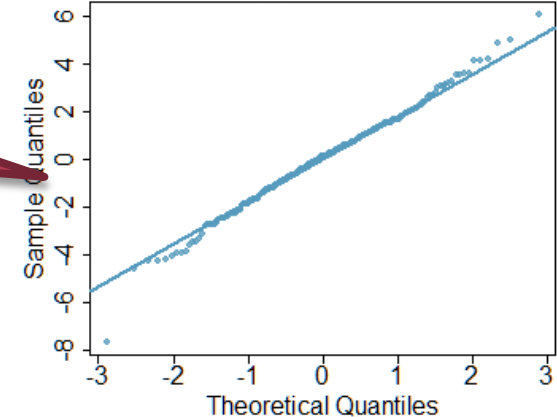
Kvantilisok: azonos valószínűségű intervallumokat adó vágáspontok (2-kvantilis: medián)



Histogram of Residuals



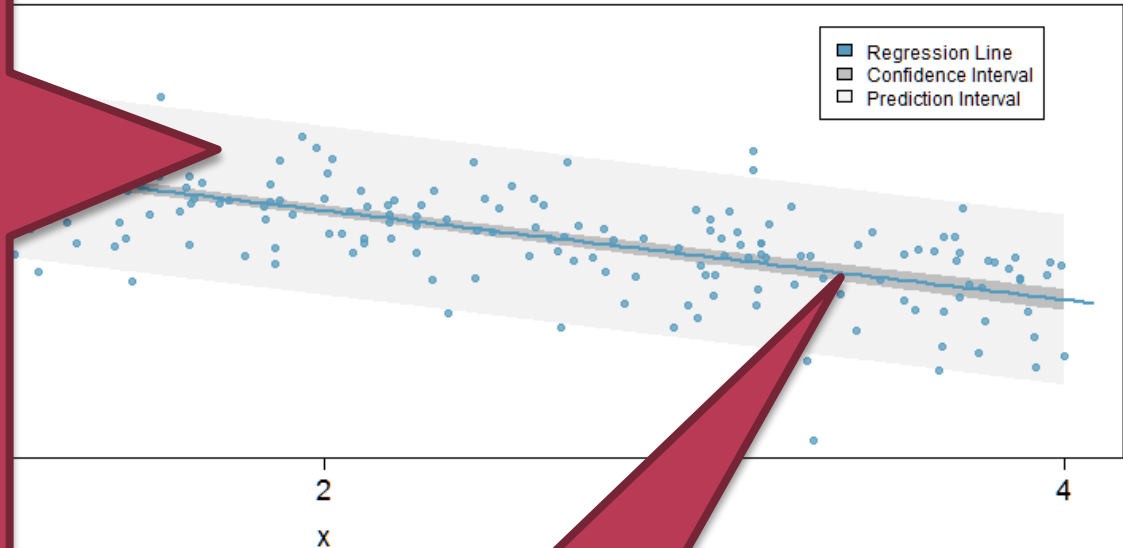
Normal Q-Q Plot of Residuals



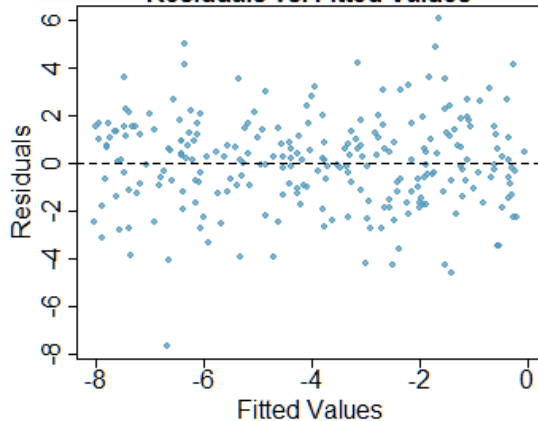
Néhány tulajdonság

Tfh. A hiba ftlen,
normál eloszlású, 0
várhatóértékű és
konstans szórású.
„Szinte biztos,
hogy ide esik az
előrejelzés” (ált.
95%)

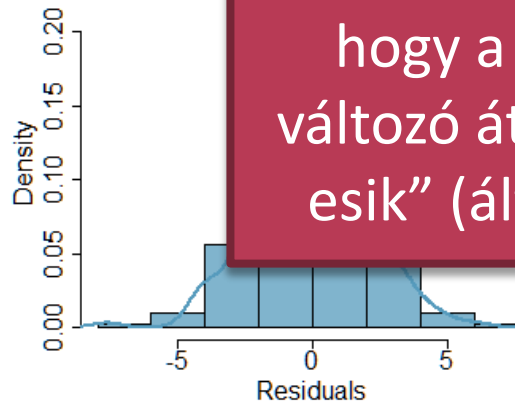
Regression Model
($R = 0.7742$, $R\text{-squared} = 0.5994$)



Residuals vs. Fitted Values

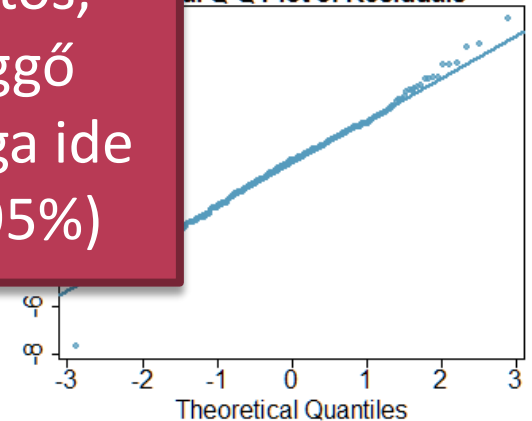


Histo



„Szinte biztos,
hogy a függő
változó átlaga ide
esik” (ált. 95%)

Normal Q-Q Plot of Residuals



Adatelemzés: legfontosabb problémák

Idősorelemzés- és
előrejelzés

Anomália-detektálás
(*anomaly detection*)

„*Structured
prediction*”

Osztályozás
(*classification*)

Asszoc. szabályok
(*assoc. rules*)

Csoportosítás
(*clustering*)

Mintakeresés (*freq.
pattern mining*)

Regresszió
(*regression*)

Graph analysis

Feature
selection

Dimensionality
reduction

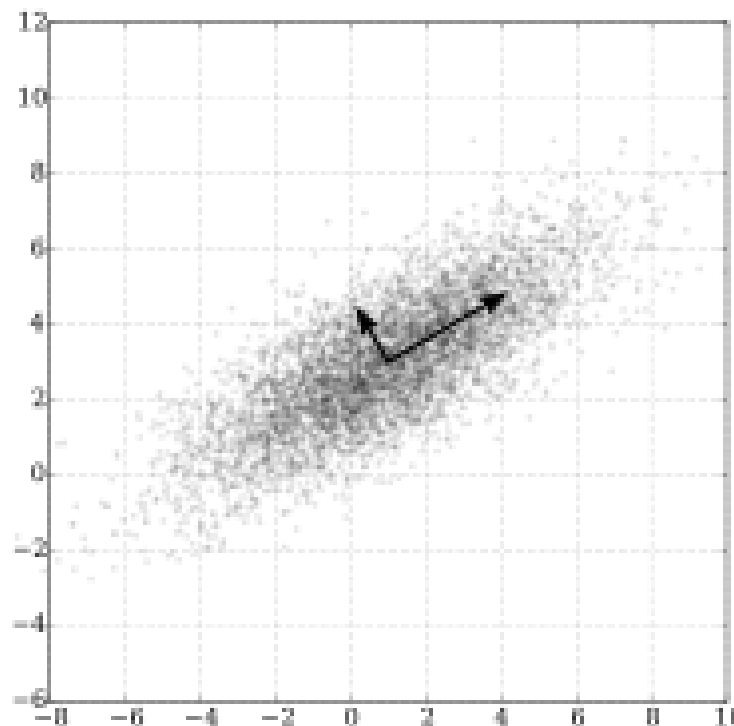
Feature
extraction

Principal Component Analysis

- Ortogonális transzformáció olyan ortogonális bázisra (lineárisan független változók; “főkomponensek”), ahol
 - Az első komponens varianciája a lehető legnagyobb
 - Az i -edik varianciája a lehető legnagyobb úgy, hogy merőleges legyen az eddigiekre.
- Értelme: lehet, hogy összvarianciát leíró komponens jóval kevesebb lesz, mint változó
- Nem faktoranalízis; az látens változók által okozott közös varianciát keres
 - És ezzel felteszi egy mögöttes modell létezését
 - Az “Exploratory Factor Analysis” tartalmazhat PCA-szerű lépést...

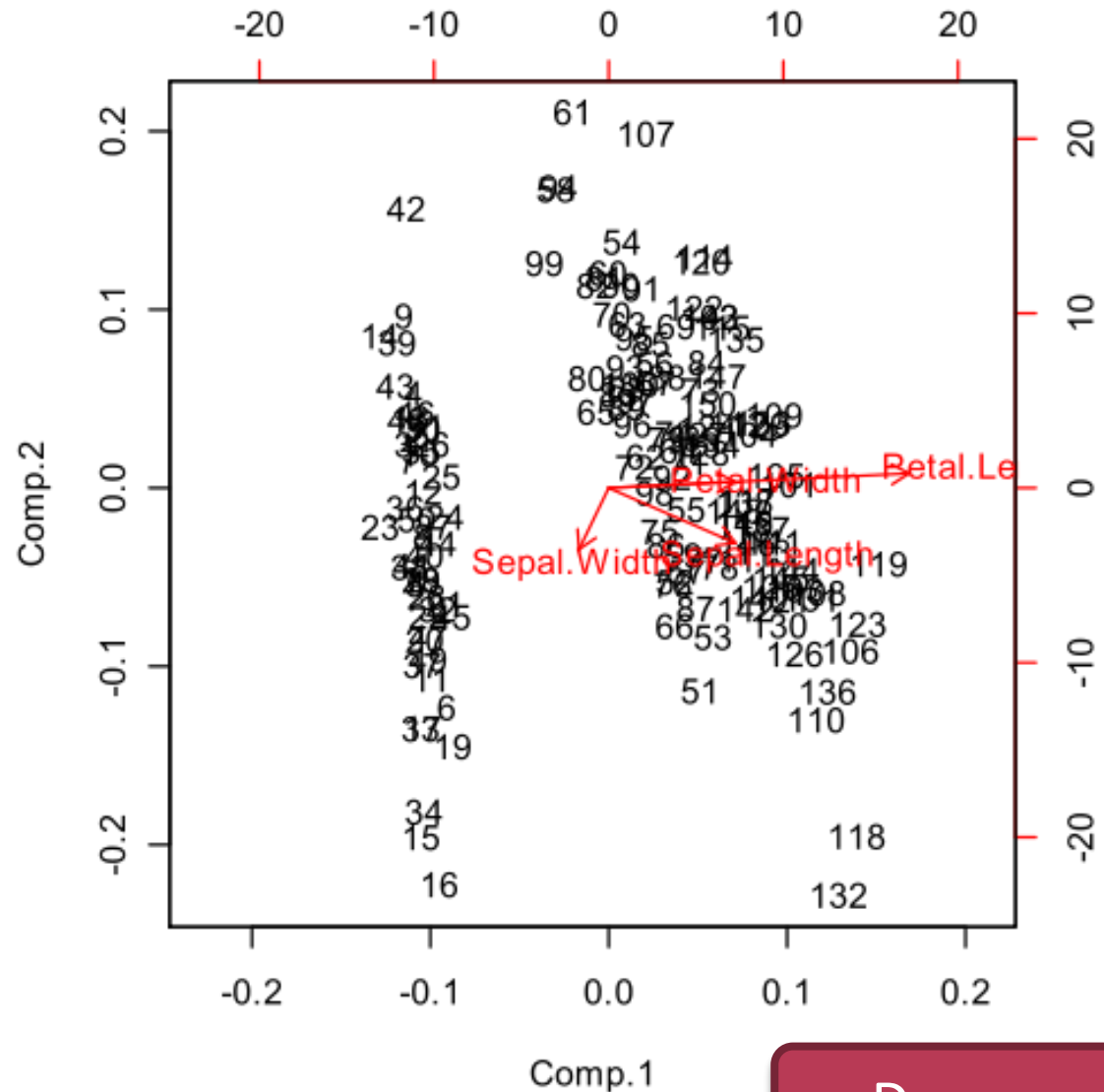
Principal Component Analysis

- Eltolás a várhatóértékbe + utána forgatás
- Skálaérzékeny, de lehet normalizálva is csinálni (pl. Z-score-ra)
- Ennek persze megint lehetnek nem kívánt hatásai
- Matematikáját nem tárgyaljuk



Biplot

- (Implicit)
feltételezés: 2
komponens
“elégésgesen leírja”
- Komponensenkénti
“koordináták”
(score-ok)
- + Az eredeti
változók súlya
(“loading”) az első
két faktorban



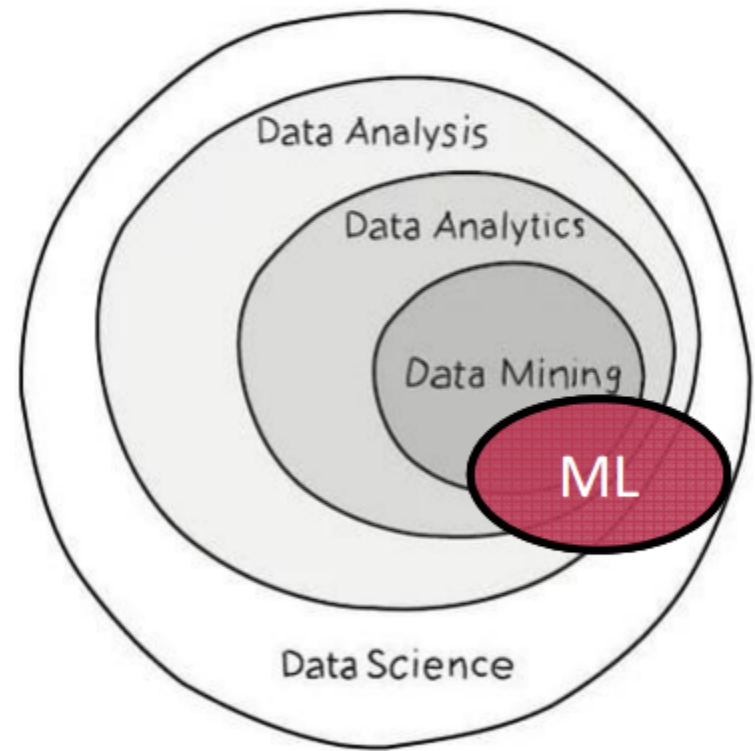
Demo

Gépi tanulás

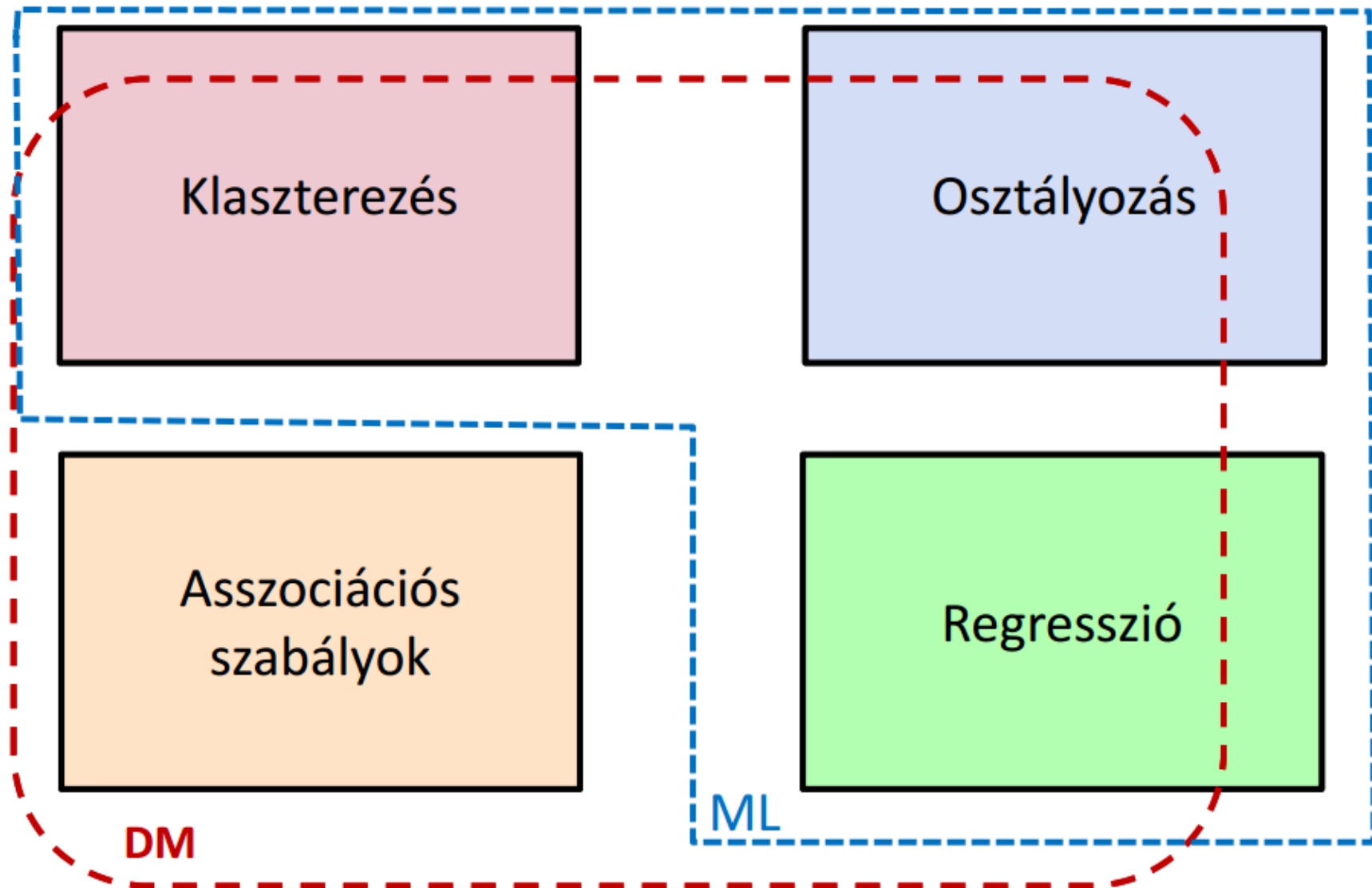
- **Gépi tanulás** (Machine Learning - ML)
 - ~ Olyan módszerek és algoritmusok, melyek a rendelkezésre álló adathalmazból egy **modellt tanulnak**, és ezt felhasználva **következtetést (predikciót) tesznek lehetővé** (a meglévő vagy más/újabb adatokról).
- Adatorientált módszereket használ, tehát lényeges a
 - Változósám
 - Mintaszám
- Ezek aránya jelentősen befolyásolja a tanult modell minőségét

Adatbányászat

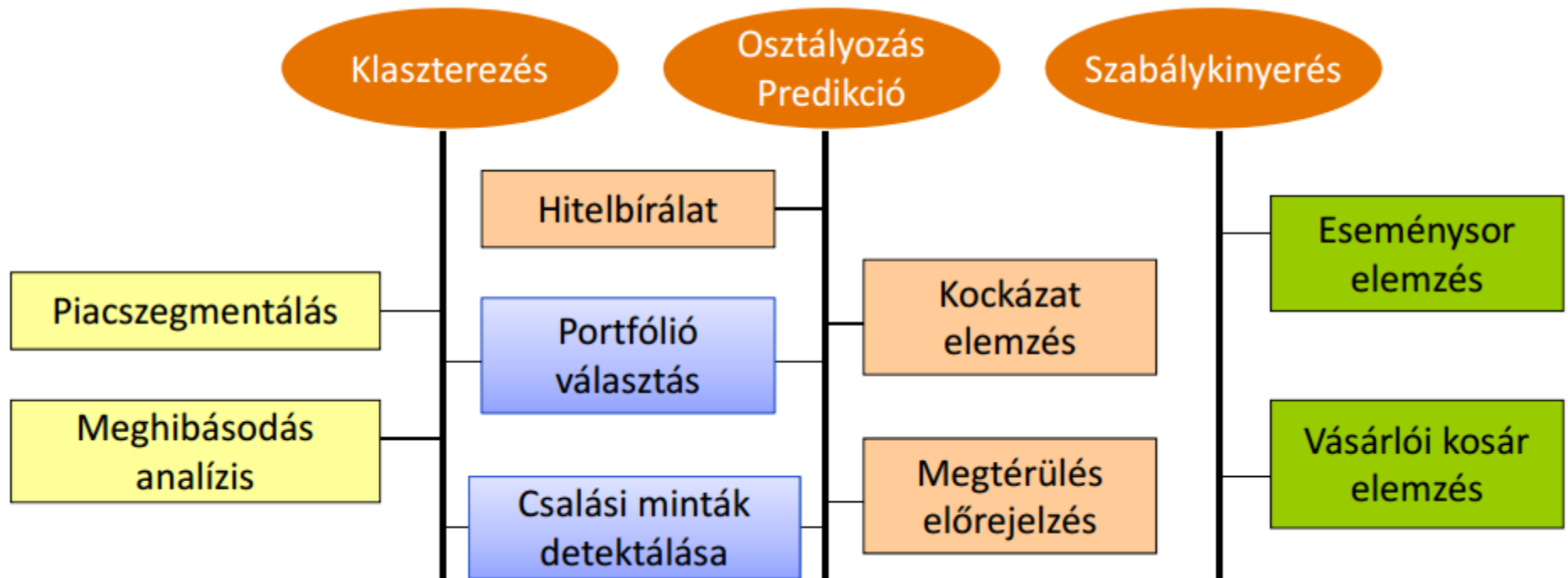
- Gépi tanulás
 - Előrejelző modellt épít
- Adatbányászat
 - ML módszereket felhasználva nyer ki új információt
 - Inkább feltáró jellegű
 - Inkább nem felügyelt tanulás



ML vs DM

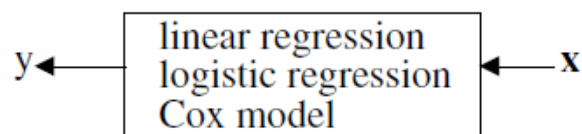


ML alapú üzleti intelligencia alkalmazások



Machine Learning, Data Mining, statisztika

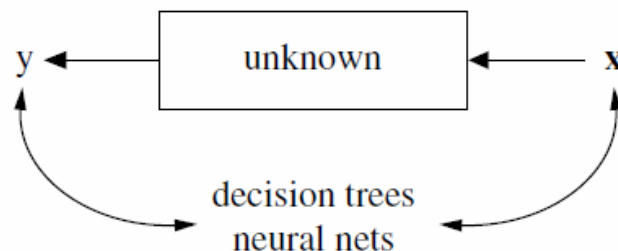
- Az eszközök és a „DNS” ugyanaz; különböző kultúrák
- Statisztika: (stochasztikus) adatmodellezés
- ML: informatikusok tanuló programot akartak
- DM: tudás kinyerése az adatokból (EDA...stat. Modellezés)



Model validation. Yes–no using goodness-of-fit tests and residual examination.

Estimated culture population. 98% of all statisticians.

Lásd [2][3]



Model validation. Measured by predictive accuracy.
Estimated culture population. 2% of statisticians, many in other fields.

Machine learning	Statistics
network, graphs	model
weights	parameters
learning	fitting
generalization	test set performance
supervised learning	regression/classification
unsupervised learning	density estimation, clustering
large grant = \$1,000,000	large grant= \$50,000
nice place to have a meeting: Snowbird, Utah, French Alps	nice place to have a meeting: Las Vegas in August

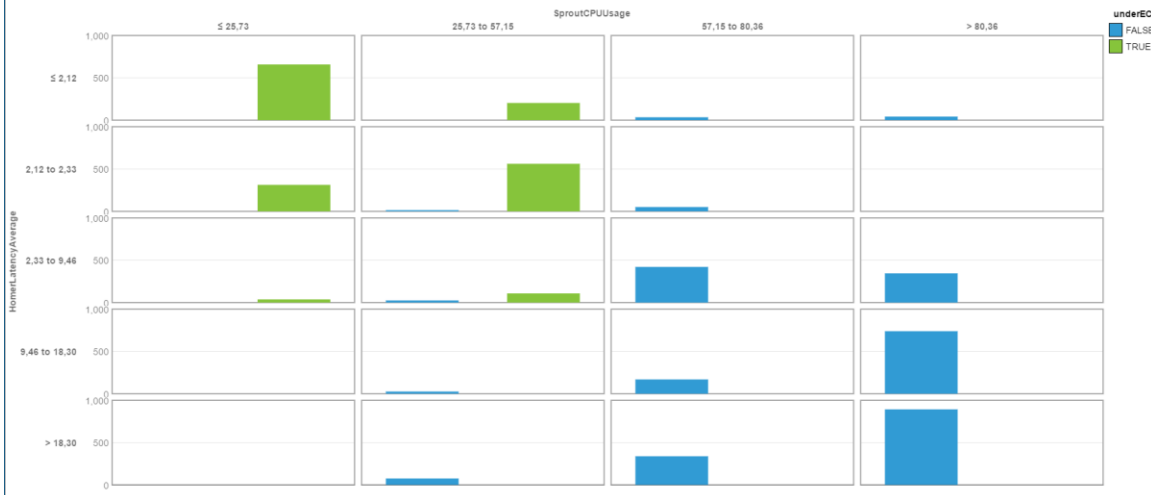
Lásd [4]

AUTOMATIZÁLT ADATELEMZÉS: IBM WATSON ANALYTICS

Automatizált adatelemzés – IBM Watson Analytics

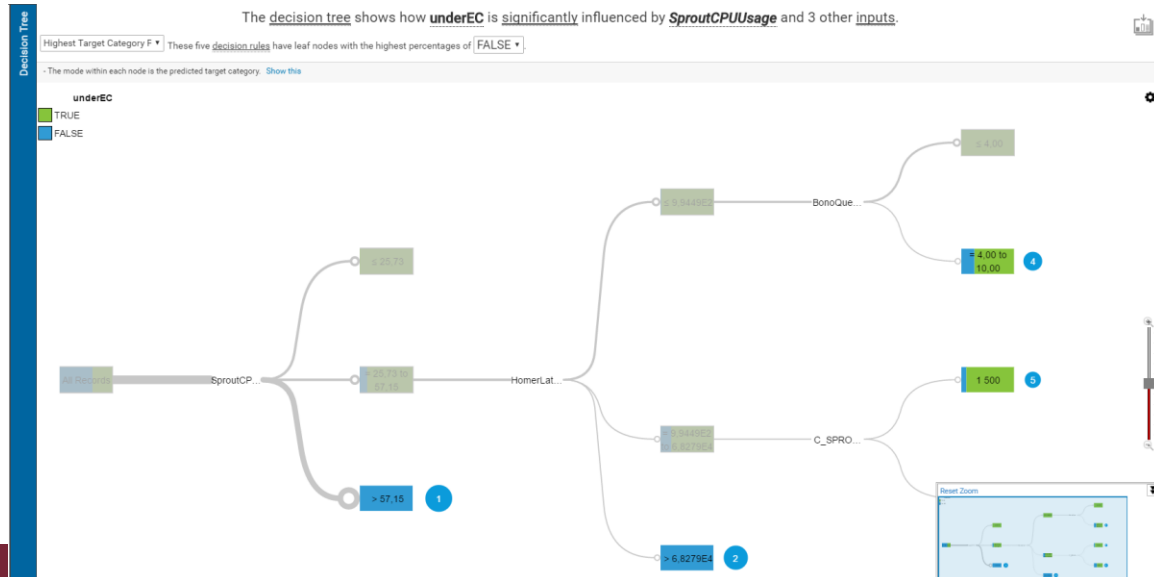
SproutCPUUsage and HomerLatencyAverage drive underEC. (Predictive Strength: 99%)

The distribution of underEC is shown at each combination of categories of SproutCPUUsage and HomerLatencyAverage.

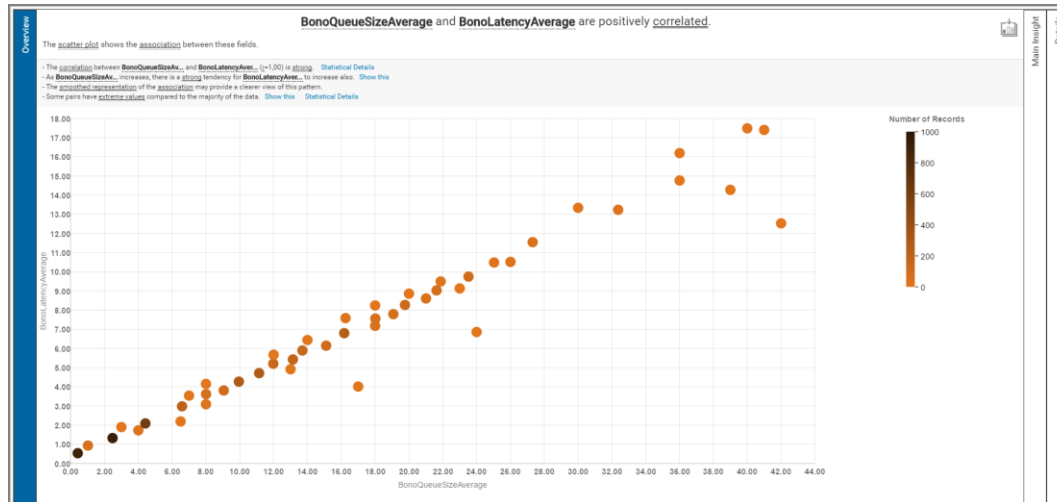


Automatikus
“predikció”

The decision tree shows how underEC is significantly influenced by SproutCPUUsage and 3 other inputs.



Automatizált adatelemzés – IBM Watson Analytics



“Érdekes
asszociációk”
automatikus feltárása

Average Quality Score

61

low quality high quality

sinusoida100_RTs.csv has 45 fields & 1367 records.

Results by Quality Measure

EXCLUDED FROM PREDICTIVE ANALYSIS

5 fields (11%) have constant values.

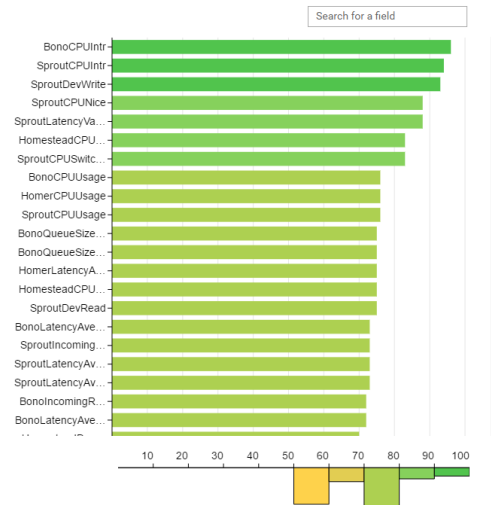
INTERESTING

25 fields (56%) have outliers.

27 fields (60%) have skewed distributions.

1 field (2%) is imbalanced.

Data Quality by Field



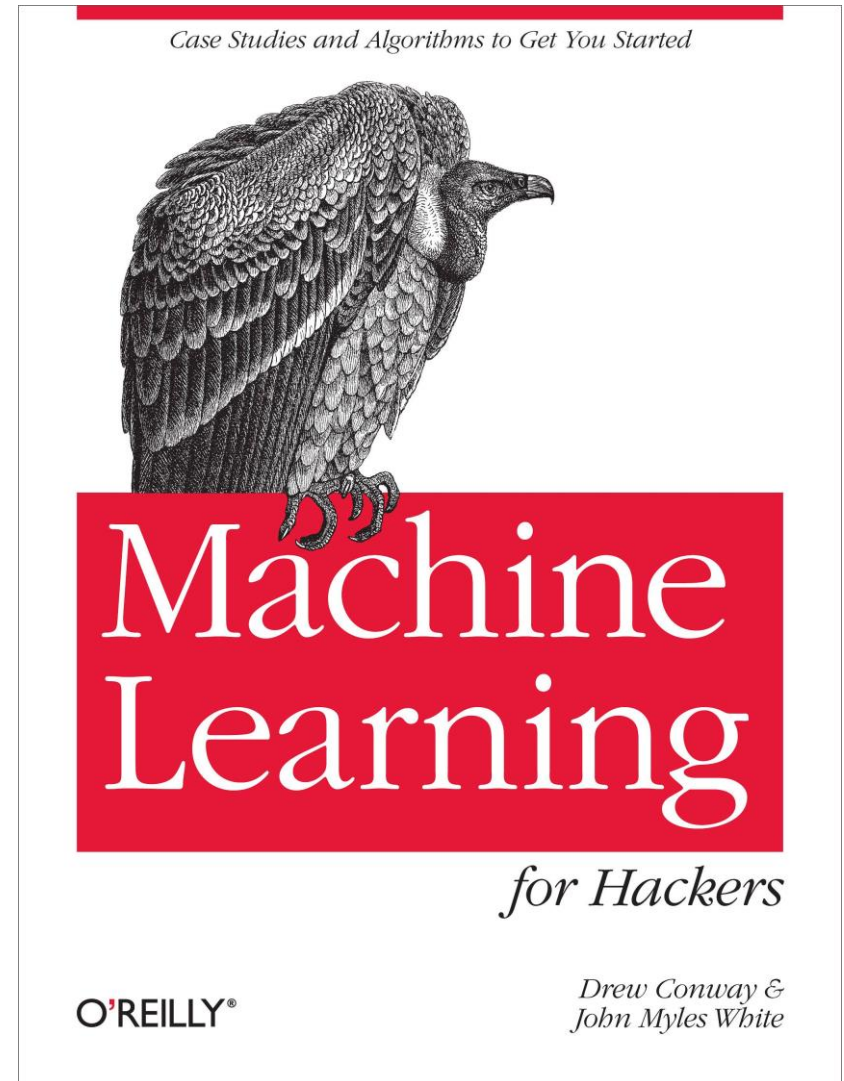
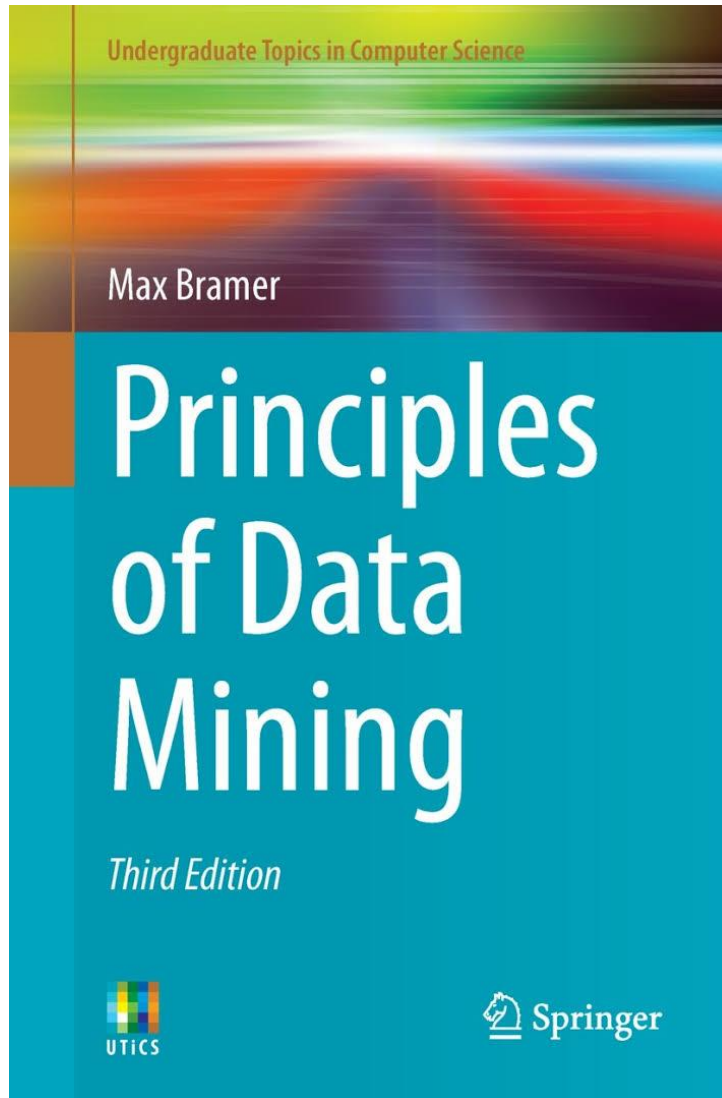
Automatizált
adatminőség-
értékelés

IRODALOM

Javasolt magyar nyelvű anyagok

- Helyenként nagyon mély, de kiváló tanulmány/jegyzet a Számításelméleti és Inf. Tud. Tanszékről:
 - <http://www.cs.bme.hu/nagyadat/bodon.pdf>
- Egyéb:
 - Dr. Abonyi János: Adatbányászat a hatékonyság eszköze
 - Iványi Antal: Informatikai Algoritmusok 2. kötet (28-29. fejezetek), ELTE Eötvös Kiadó

További javasolt kezdőirodalom



Hivatkozások

- [1] Theus, M., Urbanek, S.: Interactive graphics for data analysis: principles and examples. CRC Press (2011)
- [2] https://projecteuclid.org/download/pdf_1/euclid.ss/1009213726
- [3] <https://www.r-bloggers.com/whats-the-difference-between-machine-learning-statistics-and-data-mining/>
- [4] <http://statweb.stanford.edu/~tibs/stat315a/glossary.pdf>