

92, 6, 25

Data mining

جلد ۲

دوران فصل TAN آفره فصل های دیگر را معرفی کرده و خلاصه ای گفته

جلد ۲ در مورد این چرا ما به هافارین داده ما نیاز داریم صحبت کنیم. اینجاست که آن مثال قبلی مان را کامل تر کرده و از چند دیگه که این مثال را بررسی کرده، مثالهایی که بهش اکتفا داده مسائلی است که ما هم بصورت روزمره با آنها سروکار داریم.

Data Mining: Introduction

Lecture Notes for Chapter 1

Introduction to Data Mining

by

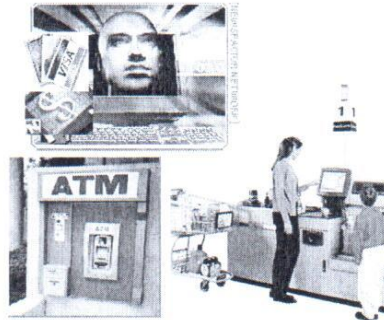
Tan, Steinbach, Kumar

دلیل که Data mining آسان را میگویند

Why Mine Data? Commercial Viewpoint

- Lots of data is being collected and warehoused

- Web data, e-commerce
- purchases at department/ grocery stores
- Bank/Credit Card transactions



- Computers have become cheaper and more powerful

- Competitive Pressure is Strong

- Provide better, customized services for an edge (e.g. in Customer Relationship Management)

خرید از اینترنت خریدگاه خیلی بزرگ

مسائل پیچیده

از آن شدن کامپیوترها

بسیار قدرت کامپیوترها زیادتر

شد و نرم افزارهای قویتری

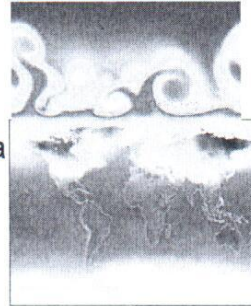
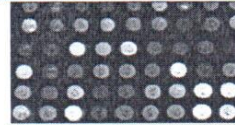
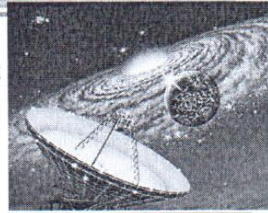
برای کامپیوترها نصب شد

بسیار در نرم افزار

بسیار در صنعت قرار

Why Mine Data? Scientific Viewpoint

- Data collected and stored at enormous speeds (GB/hour)
 - remote sensors on a satellite
 - telescopes scanning the skies
 - microarrays generating gene expression data
 - scientific simulations generating terabytes of data
- Traditional techniques infeasible for raw data
- Data mining may help scientists
 - in classifying and segmenting data
 - in Hypothesis Formation

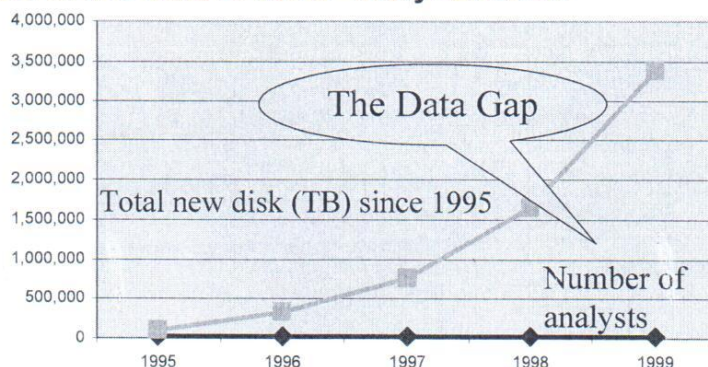


حجم داده ها بسیار زیاد شده
و به صورت اتوماتیک دارین تولید
میشوند و در لحظه میخوانند
اطلاعات پنهانی که در دل
این داده ها نهفته است را
کشف کنیم. البته از نقطه نظر
اقتصادی یا تجاری به این
مسئله نگاه شده که ما از نقطه
نظرات استرینت هوانیم به این
مسئله نگاه کنیم.
اطلاعاتی که در درون موجود
دارد و ما می توانیم اطلاعات

جدید جهانی از آنها کشف کنیم، آن ها هم به داده های مرتبط می شود. یعنی دانش کاربردی را می توانیم محدود کنیم
از بعد علمی مثال hyperspectral را مثال در جمله پیش زدیم، ما وقتی یک دوربین دیجیتال داریم و یک عکس از زمین می
گیریم، این دوربین در ۳ مایکرو موج از آن عکس میگیریم. این سه رنگ را با هم انضمام میکنند و یک عکس رنگی به ما
نشان می دهد. ولی دوربین
هایی که در ماهواره ها
نصب هستند که برای مقیوس
سنجی اند و در یک طیف رزده
این دوربین ها در 244
طول موج از زمین عکس
میگیرند و آن ها طول موجی
که از یک شیء سطح می شود را
داشته ایم، می توانیم بفهمیم
چیزی آن چیست. حالا
وقت کشد کوچه ترن عکس
فرام hyperspectral

Mining Large Data Sets - Motivation

- There is often information "hidden" in the data that is not readily evident
- Human analysts may take weeks to discover useful information
- Much of the data is never analyzed at all



From: R. Grossman, C. Kamath, V. Kumar, "Data Mining for Scientific and Engineering Applications"

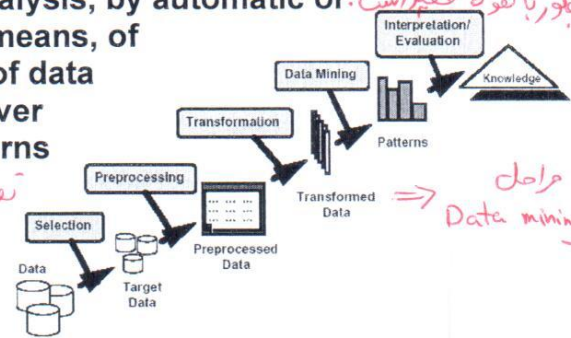
۱۰۲۴x۱۰۲۴x۱۰۲۴ یک عکس از عکس داریم، آیا می تونیم این را ذخیره کرد؟! آیا اطلاعاتی که درون اینها است را می تونیم دقیق در آورده؟!
چون بالاخره نور وجود دارد، در هوا ابر است، بخار آب است و... که باعث میشه طول موج دقیق مشخص نشود و ما باید سعی کنیم با روش
های داده کاوی بی بینیم که این چه شیئی بود؟! چه جیسی هواست؟! هم شیئی را مشخص کنیم و هم جیسیش را. در واقع ما می توانیم ۲ تا اعتبار
از داده کاوی در remote sensing بیرون بیرون کنیم. یکی به معنی اینکه عکس ما را Digitize کنیم یا حذف نور
در داده کاوی که در remote sensing است. C. Kamath, V. Kumar, R. Grossman, "Data Mining for Scientific and Engineering Applications"

مثال‌های زیادی می‌توان در این زمینه زد، اما هدف ما استخراج برسی داده‌های مشخص کنج، این است که ما هدف اصلی مان، کار کردن با حجم بزرگی از داده‌هاست و روش‌های Data Analyze نتوانند معمولاً این کار را پاسخ دهند، حالت پیشرفت Data Analyze و Data mining چندانی اش را نیفتیم. همین جمع بزرگ داده‌هاست که ما نتوانیم بر روی Data Analyze انجام دهیم علاوه بر این در Data Analyze هدفمان از قبل مشخص است، یعنی ما به دنبال کشف یک چیز جدید در داده‌ها هستیم، یک فرضیه یا ادعای تازه را در مورد آن فرضیه ادعا می‌کنیم و این تصمیم را با آمار برای آن داده‌ها می‌گیریم.

What is Data Mining?

Many Definitions

- Non-trivial extraction of implicit, previously unknown and potentially useful information from data
- Exploration & analysis, by automatic or semi-automatic means, of large quantities of data in order to discover meaningful patterns



چیز بی‌بهره داده‌های دیگر
کس از یک دید داده‌های را بی‌بهره
کرده است. بعضی داده‌ها را
را همان تحلیل داده‌ها تعیین کرده‌اند
اما بعداً مشخص شده‌اند.

تعریف ۲: توصیف، تحلیل، برسی و کشف الگوهای
Data mining برای یک جمع بزرگی از داده‌ها و کشف الگوهای
معنی

ولی در Data mining شاید در نهایت منجر به یک تصمیم شود، ولی آن مرحله دوم است، مرحله اول این است که اصلاً سوال ما چیست یعنی حتی ما در داده‌ها ممکن است سوالی هم پیدا کنیم و بعد بهش پاسخ دهیم و با سوال پاسخ را بگیریم.

What is (not) Data Mining?

What is not Data Mining?

- Look up phone number in phone directory
- Query a Web search engine for information about "Amazon"

What is Data Mining?

- Certain names are more prevalent in certain US locations (O'Brien, O'Rourke, O'Reilly... in Boston area)
- Group together similar documents returned by search engine according to their context (e.g. Amazon rainforest, Amazon.com,)

چیزی نیست؟
چیزی که تلفظ شماره و مانند اینها
این واقعاً Data mining نیست که
از آن بجای هم اطلاعاتی بگیریم شاید
هم بتوان اطلاعاتی بدست آورد اما اینجا هم این طور Data mining نیست
Deterministic است، داده‌های
stochastic نیست،
با یک query در دیتابیس کار می‌کنیم
بدست می‌آید که اسکنش کارش را بر سر
نمی‌آورد اما ما لا هم چینی می‌فروشیم و می‌خریم

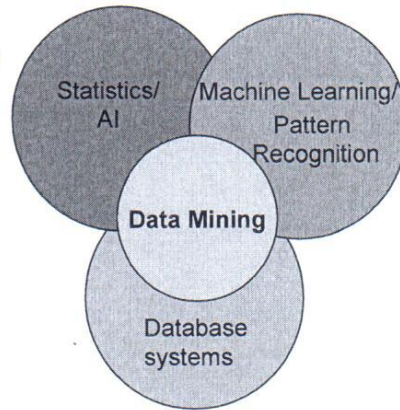
چیزی نیست؟
چیزی که تلفظ شماره و مانند اینها
این واقعاً Data mining نیست که
از آن بجای هم اطلاعاتی بگیریم شاید
هم بتوان اطلاعاتی بدست آورد اما اینجا هم این طور Data mining نیست
Deterministic است، داده‌های
stochastic نیست،
با یک query در دیتابیس کار می‌کنیم
بدست می‌آید که اسکنش کارش را بر سر
نمی‌آورد اما ما لا هم چینی می‌فروشیم و می‌خریم

این ۳ آمار ما با هم ادغام کنیم، به Data mining می‌رسیم.
چرا آمارها حاصلی در داده‌های فوق نیستند؟! چون Database system بلد نیستند.
Machine learning همان خیلی قوی نیست. چرا که ماشین‌های
قوی نیستند؟! چون از روش‌های آمار گریزان هستند، اینها را نتوان
از هم جدا کرد، چون یک آمار که به ما می‌دهد، هیچ حاشیه‌ای
نمی‌دهد.

یک کامپیوتری هم آنرا کار می‌نماید، کار خوبی نمی‌تواند ارائه دهد، پس لازم است این است که ما آمار برانیم Database
 برانیم، Pattern Recognition، Machine learning، هم برانیم. تمام درس‌های Pattern Recognition را می‌توان در دو درس
 داد: با Base آمار و بدون Base آمار که با Base آمار این خیلی بهتر است چون برای حل یک مسأله معمولاً سه درس
 داریم: تحلیل، بررسی و استفاده از بیس آمار است. ما معمولاً هر دو درس اول و دوم را برای حل کردن به مانده، تمام شده، صحتی بوده
 حل کردن، یا آن‌ها هم می‌تواند خیلی حل کردنش سخت است. این روش‌های آمار معمولاً عرض‌های Flexible برای دارند و راحت می‌توانند با مسائل
 برخورد کنند و حل‌شان کنند.

Origins of Data Mining

- Draws ideas from machine learning/AI, pattern recognition, statistics, and database systems
- Traditional Techniques may be unsuitable due to
 - Enormity of data
 - High dimensionality of data
 - Heterogeneous, distributed nature of data



© Tan, Steinbach, Kumar

Introduction to Data Mining

4/18/2004

7

Data mining مانند Data Analyze دو کار انجام می‌دهد: در Data Analyze ما آمار را به دو بخش تقسیم می‌کنیم: آمار توصیفی و آمار استنباطی
 در داده‌های آماری Inference از تکنیک Prediction استفاده می‌کنیم، به معنای پیش‌بینی است البته پیش‌بینی آینده را
 forecasting می‌گویند با prediction تفاوت است. Prediction به معنی Interpolation است یا آینده نزدیک یا گذشته نزدیک.
 داده‌های آماری برخلاف آمار به معنی Inference است اینجا بررسی prediction است. حالا ممکن است جای هم Inference انجام دهد، یعنی استنباطی انجام دهد.

Data Mining Tasks

- Prediction Methods
 - Use some variables to predict unknown or future values of other variables.
- Description Methods
 - Find human-interpretable patterns that describe the data.

کار راهبردی تر بررسی داده‌ها و کشف الگوها است، یعنی واقعیات Base را می‌بینیم اما باید شکل کلی آن را می‌توانیم خیلی
 اطلاعات بگیریم و یک سبکی هست به نام gapminder، که یک میلیون داده‌ها را نشان می‌دهد. مثلاً خط فقر، می‌تواند نقشه GAS
 این از جهان به ما نشان می‌دهد، یا یک کشور و قسمت‌های دیگر به نشان ران جان نشان می‌دهد. ما با چند آمار می‌توانیم فهمیم که
 From [Fayyad, et al.] Advances in Knowledge Discovery and Data Mining, 1996

© Tan, Steinbach, Kumar

Introduction to Data Mining

4/18/2004

8

فقرترین قاره که است، قوی‌ترین قاره که است و...، یعنی اطلاعاتی که ما باید روزها دنبالش بگردیم مثلاً برای GPA، نمرات دانش‌آموزان، آنجا
 یک نمودار آنی به ما نشان می‌دهد. نقشه GAS این، یعنی یک نقشه را می‌تواند در آن را ببرد، آماری می‌کند، آن رنگ‌هایی که به یک سبک است
 حرارت‌های آن می‌تواند.

همه مطالبی را می‌توانیم در Data mining آموزش دهیم! clustering / classification ... در دروس بعد (ارسلان)

Sequential Pattern Discovery و Association Rule Discovery در اینها به جستجو و تقریباً همیشه به یک طرفه می‌آید و اطلاعاتی از داده‌ها استخراج می‌کند که خیلی به ما کمک می‌کند. Regression برای Interpolation استفاده می‌شود. Extrapolation و Regression خیلی از هم جدا هستند. Predict کردن: ما مقدار وزن را داریم، می‌خواهیم ببینیم مقدار وزن چه رابطه‌ای با هم دارند، پس این می‌آید مقدار وزن برای افراد دیگر درست می‌کنیم. این مقدار وزن

Data Mining Tasks...

1- ● Classification [Predictive]

2- ● Clustering [Descriptive]

3- ● Association Rule Discovery [Descriptive]

4- ● Sequential Pattern Discovery [Descriptive]

5- ● Regression [Predictive]

6- ● Deviation Detection [Predictive]

Deviation Detection (Anomaly Detection) است، یعنی داده‌های نامرتب را می‌بینیم، آن‌ها سرتیغی‌تر از بقیه هستند.

تست‌های مختلفی برای تشخیص این موارد وجود دارد. Regression برای داده‌هایی که یک مقدار را می‌خواهیم پیدا کنیم استفاده می‌شود. Deviation یعنی آن‌ها از بقیه جدا هستند، این‌ها سرتیغی‌تر از بقیه هستند.

این‌ها مطالبی است که ما می‌خواهیم این‌ها را در این درس دهیم. تفاوت Classification و clustering: supervised classification است یعنی تعداد کلاس‌ها را از قبل می‌دانیم.

unsupervised clustering است یعنی تعداد کلاس‌ها را از قبل نمی‌دانیم.

در این‌ها هم داریم، در classification چون تعداد کلاس‌ها را می‌دانیم، هدفمان predictive است، یعنی پیش‌بینی.

© Tan, Steinbach, Kumar

Introduction to Data Mining

4/18/2004

9

جدیدتر می‌آید و می‌خواهیم بدانیم در کدام کلاس بگذاشتیم و این در clustering و Descriptive است یعنی توصیف می‌کنیم و می‌دانیم که این‌ها به چه کلاسی می‌آیند. اینها را می‌خواهیم تعیین کنیم. * در واقع آن‌ها این‌ها را می‌بینیم که به چه کلاسی می‌آیند و هدفمان predictive و آن‌ها را می‌خواهیم تعیین کنیم.

Classification: Definition

● Given a collection of records (training set)

– Each record contains a set of *attributes*, one of the attributes is the *class*.

● Find a *model* for class attribute as a function of the values of other attributes.

● Goal: previously unseen records should be assigned a class as accurately as possible.

– A *test set* is used to determine the accuracy of the model. Usually, the given data set is divided into training and test sets, with training set used to build the model and test set used to validate it.

© Tan, Steinbach, Kumar

Introduction to Data Mining

4/18/2004

10

Classification ما یک سری Data داریم که این Data موجود است، بعضی می‌گویند (داده‌های آموزشی) training set. هر داده یا رکورد شامل چند تا مقیاس است، به آن‌ها می‌گویند Attribute (هر رکورد شامل چند مقیاس است و یک Attribute نامیده می‌شود). Label می‌گویند، یعنی یک برچسبی دارد که این برچسب به ما می‌گوید این داده یا رکورد متعلق به چه کلاسی است. پس یک سری عدد یا برچسب داریم، برچسب‌ها را می‌خواهیم بگذاریم در چند تا کلاس. ما براساس داده‌های آموزشی مان کلاس‌ها را مشخص می‌کنیم، تعداد کلاس‌ها مشخص

است، حالاً روش ما بر مبنای این باشد که هر وقت داده جدیدی آمد بتواند این داده را در کلاس‌ها قرار دهد و طبقه‌بندی کند. در اینجا از *classification* نام می‌برند که هرگز در میان شایان انواع و اقسام تقسیم‌ها نیست. تقسیم‌های دیگر مثل *set* جدید من دهم، سوال اینجاست که اینها چیست؟ ما بر اساس اطلاعات قبلی که داریم، در مورد اینها نظر بدوید که اینها *label* آنها چیست. ما

Classification Example (yes, No)

این‌ها را به دو *cluster* تقسیم کنیم (yes, No)

حالا این کلاس‌ها را برای داده‌های جدید طبقه‌بندی کنیم، روش ما باید قادر باشد این کار را انجام دهد.

Tid	Refund	Marital Status	Taxable Income	Cheat
1	Yes	Single	125K	No
2	No	Married	100K	No
3	No	Single	70K	No
4	Yes	Married	120K	No
5	No	Divorced	95K	Yes
6	No	Married	60K	No
7	Yes	Divorced	220K	No
8	No	Single	85K	Yes
9	No	Married	75K	No
10	No	Single	90K	Yes

Refund	Marital Status	Taxable Income	Cheat
No	Single	75K	?
Yes	Married	50K	?
No	Married	150K	?
Yes	Divorced	90K	?
No	Single	40K	?
No	Married	80K	?

Training Set → Learn Classifier → Model → Test Set

© Tan, Steinbach, Kumar Introduction to Data Mining 4/18/2004 11

کاربردهایی که *classification* دارد خیلی خیلی زیاد است، مثلاً یک مثال که از کورس خودمان داریم، مثال ما می‌خواهیم استان‌ها را بر اساس چندین حرفه‌ای بودن یا نبودن، درآمد، آمار فروش، فرهنگ، نتیجه این خیلی جالب می‌شود. یک نقشه *GAS* این می‌تواند خودمان که رتبه‌بندی کرده‌ایم، آن را به‌کار می‌بریم و می‌توانیم ببینیم که این استان‌ها را چه کسانی هستند که به‌کار می‌بریم.

Classification: Application 1

● Direct Marketing

- Goal: Reduce cost of mailing by *targeting* a set of consumers likely to buy a new cell-phone product.
- Approach:
 - ◆ Use the data for a similar product introduced before.
 - ◆ We know which customers decided to buy and which decided otherwise. This {buy, don't buy} decision forms the *class attribute*.
 - ◆ Collect various demographic, lifestyle, and company-interaction related information about all such customers.
 - Type of business, where they stay, how much they earn, etc.
 - ◆ Use this information as input attributes to learn a classifier model.

From [Berry & Linoff] Data Mining Techniques, 1997

اگر ما می‌خواهیم دسته‌بندی کنیم خیلی جالب است مثلاً ممکن است به نتایجی برسیم که باورهای سنتی باشد. مثلاً ما می‌توانیم از قبل که می‌خواهیم بدانیم که در میان این ۳۰ میلیون نفر، چه کسانی هستند که در این کلاس قرار می‌گیرند و چه کسانی قرار نمی‌گیرند. به نایب خیلی تفاوت زیادی دارد نسبت به بقیه.

Classification: Application 2

خودمان بخوان

- Fraud Detection
 - Goal: Predict fraudulent cases in credit card transactions.
 - Approach:
 - ◆ Use credit card transactions and the information on its account-holder as attributes.
 - When does a customer buy, what does he buy, how often he pays on time, etc
 - ◆ Label past transactions as fraud or fair transactions. This forms the class attribute.
 - ◆ Learn a model for the class of the transactions.
 - ◆ Use this model to detect fraud by observing credit card transactions on an account.

Classification: Application 3

- Customer Attrition/Churn:
 - Goal: To predict whether a customer is likely to be lost to a competitor.
 - Approach:
 - ◆ Use detailed record of transactions with each of the past and present customers, to find attributes.
 - How often the customer calls, where he calls, what time-of-the day he calls most, his financial status, marital status, etc.
 - ◆ Label the customers as loyal or disloyal.
 - ◆ Find a model for loyalty.

Classification: Application 4

- Sky Survey Cataloging

- Goal: To predict class (star or galaxy) of sky objects, especially visually faint ones, based on the telescopic survey images (from Palomar Observatory).

- 3000 images with 23,040 x 23,040 pixels per image.

- Approach:

- ◆ Segment the image.
 - ◆ Measure image attributes (features) - 40 of them per object.
 - ◆ Model the class based on these features.
 - ◆ Success Story: Could find 16 new high red-shift quasars, some of the farthest objects that are difficult to find!

From [Fayyad, et.al.] Advances in Knowledge Discovery and Data Mining, 1996

Classifying Galaxies

Courtesy: <http://aps.umn.edu>

Early



Class:

- Stages of Formation

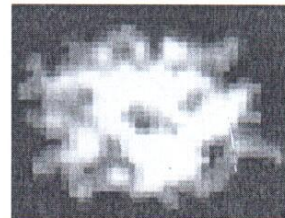
Attributes:

- Image features,
- Characteristics of light waves received, etc.

Intermediate



Late



Data Size:

- 72 million stars, 20 million galaxies
- Object Catalog: 9 GB
- Image Database: 150 GB



صورتِ مُدراستفاده می‌شود را اول شتاب می‌گیریم، بر این اساس سعی می‌کنیم اول آنها را کلاستر کنیم، پس در ادامه
منطقی‌ترین روش را برای یافتن خوشه‌ها می‌گیریم که در این ارتباط دانه باشد با آن داکومننت‌ها را جدا می‌کنیم.

Clustering: Application 1

● Market Segmentation:

- Goal: subdivide a market into distinct subsets of customers where any subset may conceivably be selected as a market target to be reached with a distinct marketing mix.
- Approach:
 - ◆ Collect different attributes of customers based on their geographical and lifestyle related information.
 - ◆ Find clusters of similar customers.
 - ◆ Measure the clustering quality by observing buying patterns of customers in same cluster vs. those from different clusters.

Clustering: Application 2

● Document Clustering:

- Goal: To find groups of documents that are similar to each other based on the important terms appearing in them.
- Approach: To identify frequently occurring terms in each document. Form a similarity measure based on the frequencies of different terms. Use it to cluster.
- Gain: Information Retrieval can utilize the clusters to relate a new document or search term to clustered documents.

یادآوری از این Document clustering را در این اسلاید 21 آورده که گفته 3204 مقالات از Los Angeles
 را در نظر بگیریم، Similarity measure مان را تعریف می‌کنیم (این تعریف مان Document صاف است که مان)

Illustrating Document Clustering

- Clustering Points: 3204 Articles of Los Angeles Times.

- Similarity Measure: How many words are common in these documents (after some word filtering). برخی از کلمات را در نظر نمی‌گیریم.

	Category	Total Articles	Correctly Placed
→	Financial	555	364
→	Foreign	341	260
→	National	273	36
→	Metro	943	746
→	Sports	738	573
→	Entertainment	354	278

Clustering of S&P 500 Stock Data

- Observe Stock Movements every day.
- Clustering points: Stock-{UP/DOWN}
- Similarity Measure: Two points are more similar if the events described by them frequently happen together on the same day.
- We used association rules to quantify a similarity measure.

	Discovered Clusters	Industry Group
1	Applied-Matl-DOWN, Bay-Network-DOWN, 3-COM-DOWN, Cabletron-Sys-DOWN, CISCO-DOWN, HP-DOWN, DSC-Comm-DOWN, INTEL-DOWN, LSI-Logic-DOWN, Micron-Tech-DOWN, Texas-Inst-DOWN, TeIllabs-Inc-DOWN, Natl-Semiconduct-DOWN, Oracle-DOWN, SGI-DOWN, Sun-DOWN	Technology1-DOWN
2	Apple-Comp-DOWN, Autodesk-DOWN, DEC-DOWN, ADV-Micro-Device-DOWN, Andrew-Corp-DOWN, Computer-Assoc-DOWN, Circuit-City-DOWN, Compaq-DOWN, EMC-Corp-DOWN, Gen-Inst-DOWN, Motorola-DOWN, Microsoft-DOWN, Scientific-Atl-DOWN	Technology2-DOWN
3	Fannie-Mac-DOWN, Fed-Home-Loan-DOWN, MBNA-Corp-DOWN, Morgan-Stanley-DOWN	Financial-DOWN
4	Baker-Hughes-UP, Dresser-Inds-UP, Halliburton-HLD-UP, Louisiana-Land-UP, Phillips-Petro-UP, Unocal-UP, Schlumberger-UP	Oil-UP

اسلاید 23: ارزش های Descriptive است نه prective قیل که تسبیب. مثلاً که جلب پیش را هم
 این که یونیک به معنی در یک چیز دیگر میخورد در یک مغازه بزرگ. کفایت نشان دادند که حتی ما سیر هم میخورد.

Association Rule Discovery: Definition

- Given a set of records each of which contain some number of items from a given collection;
 - Produce dependency rules which will predict occurrence of an item based on occurrences of other items.

TID	Items
1	Bread, Coke, Milk
2	Beer, Bread
3	Beer, Coke, Diaper, Milk
4	Beer, Bread, Diaper, Milk
5	Coke, Diaper, Milk

Rules Discovered:

{Milk} --> {Coke}

{Diaper, Milk} --> {Beer}

Association Rule Discovery: Application 1

- Marketing and Sales Promotion:
 - Let the rule discovered be
 $\{Bagels, \dots\} \rightarrow \{Potato\ Chips\}$
 - Potato Chips as consequent => Can be used to determine what should be done to boost its sales.
 - Bagels in the antecedent => Can be used to see which products would be affected if the store discontinues selling bagels.
 - Bagels in antecedent and Potato chips in consequent => Can be used to see what products should be sold with Bagels to promote sale of Potato chips!

Association Rule Discovery: Application 2

- Supermarket shelf management.
 - Goal: To identify items that are bought together by sufficiently many customers.
 - Approach: Process the point-of-sale data collected with barcode scanners to find dependencies among items.
 - A classic rule --
 - ◆ If a customer buys diaper and milk, then he is very likely to buy beer.
 - ◆ So, don't be surprised if you find six-packs stacked next to diapers!

Association Rule Discovery: Application 3

- Inventory Management:
 - Goal: A consumer appliance repair company wants to anticipate the nature of repairs on its consumer products and keep the service vehicles equipped with right parts to reduce on number of visits to consumer households.
 - Approach: Process the data on tools and parts required in previous repairs at different consumer locations and discover the co-occurrence patterns.

اسلاید 27: خیلی وقتها داریم با Sequential تولید می‌شوند یا به صورت دنباله‌ای نیستند ولی

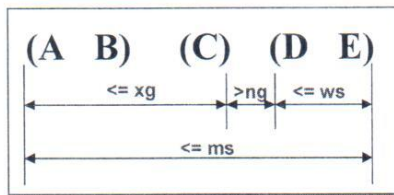
اگر ما آن داده‌ها را پشت سر هم قرار دهیم می‌توانیم به نتیجه‌ای برسیم. Sequential Pattern در واقع بررسی می‌کند که اینها استفاده می‌شود. ما می‌خواهیم ببینیم چه ارتباطی بین تغییرها است. معمولاً تغییرهای مستقل با یکدیگر می‌توانند وجود داشته باشند، به هم می‌خورند و این وابستگی نشان می‌دهد. ممکن است زمان هم

Sequential Pattern Discovery: Definition

- Given is a set of objects, with each object associated with its own timeline of events, find rules that predict strong sequential dependencies among different events.

(A B) (C) → (D E)

- Rules are formed by first discovering patterns. Event occurrences in the patterns are governed by timing constraints.



در بحث Sequential که در کتاب Time series می‌بینیم به داده‌های ما یک بدنه می‌دهیم اما نه شده و آن بدنه زمان است. اگر داده‌ای به زمان وابسته باشد می‌توانیم آن را به صورت Time series قرار دهیم. وقت کنید که یک داده‌ای می‌خواهیم با یک Time series قرار دهیم چون غیر قابل شمار هستند و چون در زمان دارند تو می‌گیریم. در سری‌های زمانی ما معمولاً علامت‌ها می‌گذاریم یا آن‌ها را می‌گذاریم.

Sequential Pattern Discovery: Examples

- In telecommunications alarm logs,
 - (Inverter_Problem Excessive_Line_Current) (Rectifier_Alarm) → (Fire_Alarm)
- In point-of-sale transaction sequences,
 - Computer Bookstore: (Intro_To_Visual_C) (C++_Primer) → (Perl_for_dummies, Tcl_Tk)
 - Athletic Apparel Store: (Shoes) (Racket, Racketball) → (Sports_Jacket)

داده‌های ما در آنجا خواهند داشت. chang point می‌توانیم به تغییرات زمانی داریم. اینها را می‌توانیم predict کنیم، کار می‌کند یا نه. Descriptive است. مثال‌هایی از Pattern را می‌بینیم که اگر مثلاً در فروشگاه کتاب، کسی این دو کتاب را بخرد، کتاب بعدی که می‌خواهد بخرد این است، پس آن‌ها در آمازون این دو کتاب را می‌خریم، خود آمازون چند کتاب دیگر را پیشنهاد می‌دهد و بعد آنرا می‌بخریم. تخفیف هم به ما می‌دهد. چون دیدیم مثلاً آن کتاب C و کتاب C++ را بخرد کتاب برود فوراً می‌توانیم به ما خودمان بگویم پیشنهاد دهیم.

Regression خطی قوی است، همانند دارنوزدان در قبل گفتیم، می‌توانیم تصادف یا آنفوزار از جملای مختلف که یک سن دارند را انتخاب کنیم، در یک سن معنوی یا درس‌های مختلف و بعد می‌توانیم به این داده‌ها خط برازش دهیم، خط امکان ندارد، یک معنی برای این دهیم، روش‌های ریاضی خطی خوب نیستند چون این داده‌ها داریم خطی مدعی نیستیم از آن‌ها رد کردیم مدعی با درم خودشان که ما به نادر و خطی یا رانده‌های زاناس داردیم خاصه مدعی دل‌آگاه

Regression

- Predict a value of a given continuous valued variable based on the values of other variables, assuming a linear or nonlinear model of dependency.
- Greatly studied in statistics, neural network fields.
- Examples:
 - Predicting sales amounts of new product based on advertising expenditure.
 - Predicting wind velocities as a function of temperature, humidity, air pressure, etc.
 - Time series prediction of stock market indices.

Time series با Regression می‌تواند است، در Time series Extrapolation است. در Time series prediction است، در Time series forecasting است.

© Tan, Steinbach, Kumar

Introduction to Data Mining

4/18/2004

29

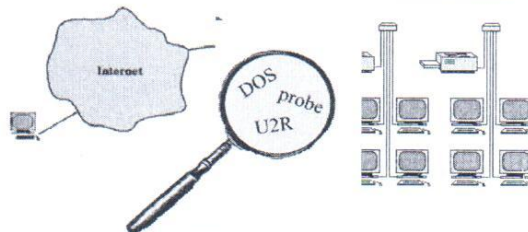
Regression یک ابزار خطی خوبی است که به کمک آن می‌توانیم learning را انجام دهیم و در واقع statistical learning است. فرض می‌آید که در داده‌ها یک رابطه وجود دارد، ما این مدل را با neural network حل می‌کنیم. در واقع neural network یک ابزار برای انجام مدل‌های Regression است اما تفاوت Regression neural network چیست؟ در neural network می‌توانیم

Deviation/Anomaly Detection

- Detect significant deviations from normal behavior
- Applications:
 - Credit Card Fraud Detection



- Network Intrusion Detection



Typical network traffic at University level may reach over 100 million connections per day

© Tan, Steinbach, Kumar

Introduction to Data Mining

4/18/2004

30

انحراف = Anomaly یعنی نا به معیار - نا به معیار یعنی خرید نا به معیار، داده نام معیار، یعنی یک اتفاق در Database - همان اتفاق که به دنبال معمول نیست می‌خواهیم اینها را تشخیص کنیم، یا آنرا اتفاق معمول افتاده باشد، می‌خواهیم ببینیم آیا واقعاً این داده نا به معیار به صورت رندوم قرار پیش آمده یا این به صورت systematic پیش آمده است. اینها کارهای است که ما می‌خواهیم انجام دهیم و یکی از بجهت‌های روز است Anomaly Detection. اما در داده‌ها ما یا جالبه‌های جدیدی رو می‌شوریم، در داده‌ها

Challenges of Data Mining

- Scalability
- Dimensionality
- Complex and Heterogeneous Data
- Data Quality → یک اصل استخراج پذیر است
- Data Ownership and Distribution
- Privacy Preservation } دو مورد آخر به همراه خودمان مطالعه کنید
- Streaming Data

با داده های مواجه می شویم که اینها با مقیاس اسمی - فاصله ای و ... اندازه گیری شده اند. لزوماً با مقیاس نسبتی اندازه گیری نشده اند. به خاطر همین Scalable نیستند، تقسیم با هم جمع شان کردن، ضربشان کردن، قواعد خاصی بر روی آنها حاکم است. بعضی بگویند ما که می از موضوعات به مرور در دارند برایش راه حل جدیدی می دهند، Dimensionality یا بُعد داده ها ما است معنی تعداد متغیرات است که ما در داده ها داریم و وقتی از تعداد داده های ما بیشتر است. مثلاً ما میخواهیم به تعداد کم غرض داشته باشیم، به قدری که اگر بخواهیم بررسی انجام دهیم، میخواهیم آن ها را با هم مقایسه کنیم، چند تا از آن ها را جدا کنیم و در اینجا ما ۵۰ تا داده داریم، ۱۰ هزار متغیر، هیچ بررسی آجاری وجود ندارد که ما بتوانیم ۱۰ هزار را مقیاس ۵۰ تا داده بتوانیم از آن نتایج بگیریم. این ابعاد که بزرگ شود هیچ روشی نمیتواند این داده ها را تحلیل کند و باید روش جدیدی ابداع کنیم برای داده های High dimensional یا این بعد داده ها را کم کنیم مثلاً روش PCA، ICA و Random projection و ... که روش هایی هستند که بعد داده ها را کم میکنند، داده ها را می بیند و یک فضای دیگر در آن فضا مدل را حل میکند و جواب را بر می گرداند، هر کدام یک بهتری و یک محدودیت نسبت به دیگری دارند. ساختار داده های ما ممکن است خیلی پیچیده باشد، پیچیدگی به معنی نوع داده ها، توزیع داده ها، هم بستگی داده ها، اندازه داده ها و ... بعضی در داده ها Data quality است، وقتی یک بررسی شماره انجام میشود در یک کشور، صور سال یک به شماره انجام میشود، این داده های که از طریق شماره برداشته میشود، یک ال طول می دهد تا Quality شان را چک کنیم، افرادی را به طرقت کنند بخاطر همین است که نتایج سریعی را بعد از ارسال اعلام میشود، اگر کیفیت داده ها خوب نباشد ما به بی راهه می رویم در داده کاری. ما گفتم داده ها و توزیع داده ها خیلی سخت است، ما نمی توانیم در یک در زمان از یک Data بگیریم.