



# Computational Communications and Data Analysis

## Lecture 9: Text Mining and Data Visualization

Ting Wang

- Text Mining
  - Data Visualization using Python
- Data Mining Essentials

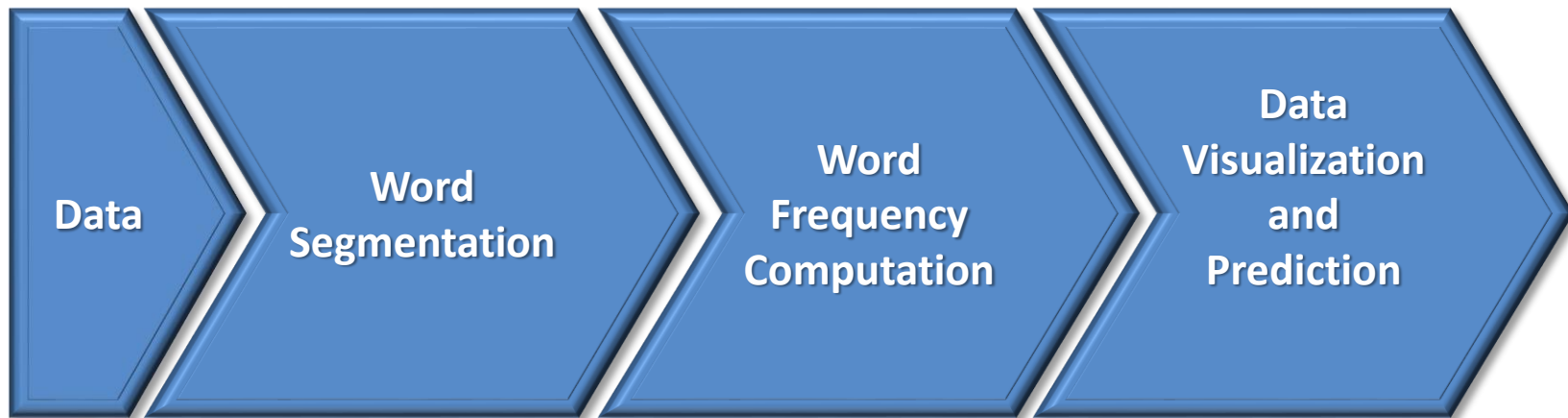




online text data mining based on natural language processing

# Text Mining

*Now, we have data, how to mining it?*



## *Case Description*

### *Motivations:*

- To measure a news objectively
- To obtain new information efficiently

### *Methodologies:*

- Describe a news report by quantitative method
- Technical integration by computer science, statistics and journalism



## *Steps:*

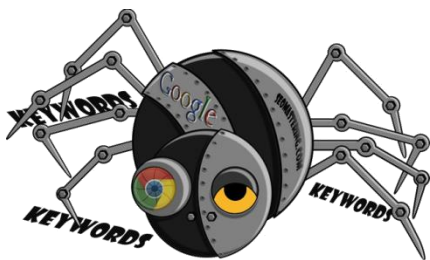
1. Download a news report
2. Word segmentation
3. Word tag extraction and statistical computing
4. Data visualization and news summarization



## Step 1: Download a News Report

- Example: <http://news.sina.com.cn/w/2018-05-21/doc-ihawmatz9906261.shtml>

*News also can be obtained by web crawler or databases*



原标题：德媒：五大国这次要被逼得联手了

据德国《星期日世界报》20日报道，来自德国、法国、英国、俄罗斯和中国的外交官员正在协商一项新协议，希望借此挽救2015年签署的伊核协议，并说服特朗普解除对伊朗的制裁。这些外交官还将于25日在维也纳就此举行会议。不过路透社20日援引3名欧盟消息人士的话否认与会各方将讨论新协议。分析认为，鉴于欧盟自知力量有限，因此有意与中俄共同商讨新协议，但短期内，这个目标并不现实。



《星期日世界报》从欧盟高层人士获得的消息称，德国、法国、英国、俄罗斯和中国间的会谈定在下周末，但美国不会出席，伊朗官员是否参加还不得而知。会谈的目的是商讨美国退出伊朗核协议后的下一步进程。

报道称，新协议和2015年的伊核协议相似，但新增限制伊朗弹道导弹和地区角色的条款，未来还可能增加对伊朗的财政援助内容。如果新协议能够达成，有助于说服特朗普解除对伊朗的制裁。

但3名曾参与阻止美国总统特朗普退出伊核协议谈判的欧盟消息人士20日晚些时候告诉路透社，上述消息并不正确，“本周五的维也纳会议将讨论伊核协议的实施问题和细节。”德国外交部目前尚未就有关消息予以回应。

## *Step 2: Word segmentation (1)*

### Database Preparation

- Word Dictionary (required)
- Stop Word Dictionary (required)
- Dictionaries of Terms (optional)
- Word Chains (required if using N-gram)
- Part of Speech (optional)
- Word Sentiment (optional for Sentiment Analysis)



## Step 2: Word segmentation (2)

### Chinese Word Segmentation

- FMM
- BMM
- N-gram



```
def word_seg_fmm(content): #正向匹配
    MaxLen=10 #最大词长
    Len=MaxLen #动态切割词长
    Seg_Content="" #返回的切割结果

    while len(content)>0:
        if content[0:Len] in WordMap: #词典中有匹配
            Seg_Content=Seg_Content+content[0:Len]+"|"
            content=content[Len:]
            Len=MaxLen
            #print("Seg_Content1:"+Seg_Content)
            continue
        else: #词典中无匹配
            Len=Len-1
            if Len==1: #仅剩一个词还没匹配到
                Seg_Content = Seg_Content + content[0:Len] + "|"
                content = content[Len:]
                Len = MaxLen
            #print("Seg_Content2:" + Seg_Content)
    return Seg_Content[:-1]
```

```
def word_seg_bmm(content): #逆向匹配
    MaxLen=10 #最大词长
    Len=MaxLen #动态切割词长
    Seg_Content="" #返回的切割结果

    while len(content)>0:
        if content[-Len:] in WordMap: #词典中有匹配
            Seg_Content=content[-Len:]+"|" + Seg_Content
            content=content[:-Len]
            Len=MaxLen
            #print("Seg_Content1:"+Seg_Content)
            continue
        else: #词典中无匹配
            Len=Len-1
            if Len==1: #仅剩一个词还没匹配到
                Seg_Content = content[-Len:] + "|" + Seg_Content
                content = content[:-Len]
                Len = MaxLen
            #print("Seg_Content2:" + Seg_Content)
    return Seg_Content[:-1]
```

## *Step 2: Word segmentation (3)*

- Tips for Chinese Word Segmentation
  - Initialization is very important
  - Segment in the memory (not hard disk or data bases) to accelerate the segmentation speed
  - Using “set” to store the dictionary, and “dict” for segmented words in Python
  - For Tag Analysis, a precise word segmentation is unnecessary



## *Step 3: Word Tag Extraction and Statistical Computing*

- `str.split()` for all tags
- Discarding One-Char tags
- Discarding Stop-Word tags
- Select tags whose term frequencies are larger than a threshold (for example  $>2$ )
- Other statistical computing



## Step 4: Data Visualization and News Summarization



## *Data Visualization using Python*

- Necessity:
  - NumPy (Computing Package)
  - Scipy (Scientific Computing Package)
  - Pillow(Image)
  - Matplotlib (Diagram Package)
  - wordcloud (Word Cloud Package)
- Some packages also need some other required packages

*Installation  
Sequence*



原标题：【德媒】：【五大国】这次要被逼得联手了！——据德国《星期日世界报》20日报道，来自德国、法国、英国、俄罗斯和中国的外交官员正在协商一项新协议，希望借此挽救2015年签署的伊核协议，并说服特朗普解除对伊朗的制裁。这些外交官还将于25日在维也纳就此举行会议。不过路透社20日援引3名欧盟消息人士的话否认与会各方将讨论新协议。分析认为，鉴于欧盟自知力量有限，因此有意与中俄共同商讨新协议，但短期内，这个目标并不现实。——《星期日世界报》从欧盟高层人士获得的消息称，德国、法国、英国、俄罗斯和中国间的会谈定在下周周末，但美国不会出席，伊朗官员是否参加还不得而知。会谈的目的是商讨美国退出伊核协议后的下一步进程。——报道称，新协议和2015年的伊核协议相似，但新增限制伊朗弹道导弹和地区角色的条款，未来还可能增加对伊朗的财政援助内容。如果新协议能够达成，有助于说服特朗普解除对伊朗的制裁。——但3名曾参与阻止美国总统特朗普退出伊核协议谈判的欧盟消息人士20日晚些时候告诉路透社，上述消息并不正确，“本周五的维也纳会议将讨论伊核协议的实施问题和细节。”德国外交部目前尚未就有关消息予以回应。——虽然维也纳会议的具体议题尚不明确，但“为挽救伊核协议”，五国正组成“联合阵线”。——“德国之声”20日称，计划中的会议显示欧盟致力于确保伊核协议得以继续执行，即便这意味着他们要在脱离美国的情况下，与莫斯科、北京和德黑兰展开合作。——卡塔尔半岛电视台20日称，自5月8日特朗普宣布退出伊核协议以来，欧洲和德黑兰相互谨慎接近，双方声明遵守协议的“要求”，同时监测彼此的行为，以确保履行承诺。欧洲国家表示将尽力保持伊朗石油和投资流动，但同时也承认这并不容易。伊朗原子能机构负责人萨利希表示，如果欧洲国家未能保留协议，伊朗有多种选择，包括恢复提炼浓缩铀至纯度20%，并称欧盟只有几个星期的时间来履行其承诺。——而《星期日世界报》认为，之所以要寻找新途径，是因为欧洲官员知道，欧洲企业在美国的新制裁背景下难以在伊朗进行商业活动。欧盟希望伊朗知道，只要后者遵守伊核协议，欧盟就愿意为德黑兰注资。欧盟高级官员认为，布鲁塞尔就美国的制裁措施所采取的对策，对“伊朗经济的积极影响非常有限”，因此有必要与中俄缔结新的协议。——不过，中国社会科学院西亚非洲所副研究员王凤20日对《环球时报》记者表示，各方在短期内就伊核问题达成新的协议并不现实。因为研发弹道导弹一直是伊核计划的内容，很难要求伊朗停止研发弹道导弹以换取欧盟的金融支持，伊朗对欧盟的承诺并不放心。——广告——面对美国的强势，欧盟应该怎么办？美国《商业内幕》20日称，欧盟可以签署一个变动极小的协议，以绥靖特朗普，然后坐等他任期结束。

## Word Cloud (Main Code)

```
from numpy import *
from os import path
from matplotlib.pyplot import imread
# from scipy.misc import imread
import matplotlib.pyplot as plt; plt.rcParams()
from wordcloud import WordCloud, ImageColorGenerator
import time
def word_cloud_generate(word_tagging):
    # 获取当前文件路径
    # __file__ 为当前文件, 在ide中运行此行会报错, 可改为
    # d = path.dirname('.')
    d = path.dirname(__file__)
    # 设置背景图片
    alice_coloring = imread(path.join(d, "WordCloud3.png"))

    wc = WordCloud(background_color="white", # 背景颜色max_words=2000, #
词云显示的最大词数
                    mask=alice_coloring, # 设置背景图片
                    max_font_size=300, # 字体最大值
                    font_path="C:\\Windows\\Fonts\\simhei.ttf",
                    random_state=100)

    # 生成词云, 可以用generate输入全部文本(中文不好分词), 也可以我们计算好
    词频后使用generate_from_frequencies函数
    # wc.generate(text)
```

```
wc.generate_from_frequencies(word_tagging)
# txt_freq例子为[('词a', 100), ('词b', 90), ('词c', 80)]
# 从背景图片生成颜色值
image_colors = ImageColorGenerator(alice_coloring)
# 以下代码显示图片
plt.imshow(wc)
plt.axis("off")
# 绘制词云
plt.figure()
# recolor wordcloud and show
# we could also give color_func=image_colors directly in the
constructor
plt.imshow(wc.recolor(color_func=image_colors))
plt.axis("off")
# 绘制背景图片为颜色的图片
plt.figure()
plt.imshow(alice_coloring, cmap=plt.cm.gray)
plt.axis("off")
# plt.show()
# 保存图片

picurl="WordCloud" + str(time.time()) + ".png"
wc.to_file(path.join(d + '/static', picurl))
return picurl
```

## Conclusions

本文与伊朗问题有关，可能跟武器和制裁有关，起决定力量的应该是美国、德国、中国和伊朗。欧洲与此消息关系较大。

官员 弹道导弹 星期日 承诺 伊朗 欧洲 消息 退出 协议 会议 表示 人士 美国 中国 德国 世界 黑兰 表示 人士 中国 并 为 裁 非 德 国



上海外国语大学  
SHANGHAI INTERNATIONAL STUDIES UNIVERSITY

machine learning approaches for data mining

# Data Mining Essentials

# Data Mining Essentials

## *Data Mining* 数据挖掘

- Data Mining is the power for producing high-quality journalism.
- Data Mining is an interdisciplinary subfield of computer science, and statistics.



# Data Mining Essentials

## *Social Demands*

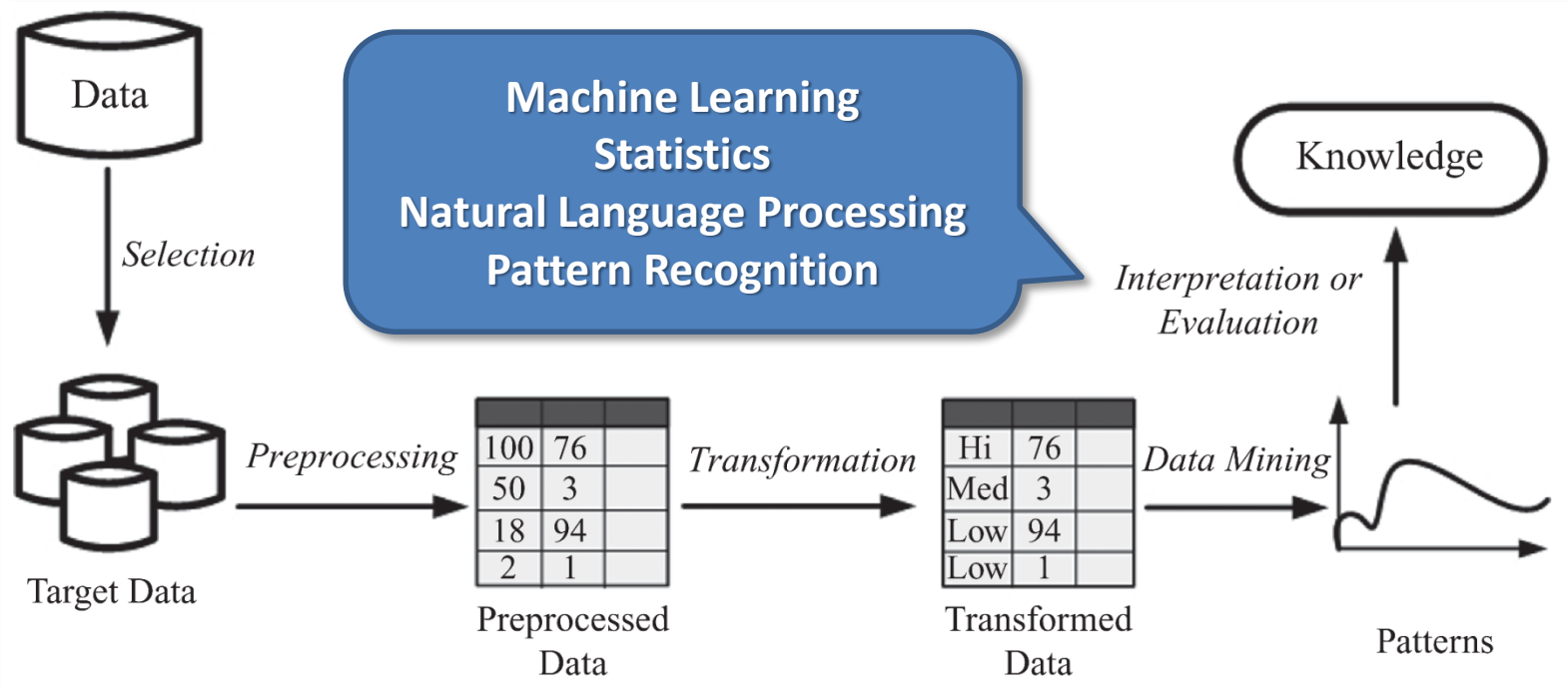
- Data production rate has increased dramatically (**Big Data**) and we are able to store much more data
  - E.g., purchase data, social media data, cell phone data
- Businesses and customers need useful or actionable knowledge to gain insight from raw data for various purposes
  - It's not just searching data or databases



The process of extracting useful patterns from raw data is known as **Knowledge Discovery in Databases (KDD)**

# Data Mining Essentials

## *KDD from Data Bases*



# Data Mining Essentials

## *Data* 数据

- Continuous Data 连续型数据
  - Regression
- Discrete Data 离散型数据
  - Classification



## Data Feature (1) 数字特征

Feature also called as Measurement, Attribute

- **Nominal 名词性**

- **Operations:**

- Mode (most common feature value), Equality Comparison

- E.g., {male, female}

- **Ordinal 序数性**

- Feature values have an intrinsic order to them, but the difference is not defined

- **Operations:**

- same as nominal, feature value rank

- E.g., {Low, medium, high}



## Data Feature (2) 数字特征

- Interval 间隔性

- Operations:

- Addition and subtractions are allowed whereas divisions and multiplications are not

- E.g., 3:08 PM, calendar dates

- Ratio 比例性

- Operations:

- divisions and multiplications are allowed

- E.g., Height, weight, money quantities



## Data Quality 数据质量

- **Noise 噪声数据**
  - Noise is the distortion of the data
- **Outliers 异常值**
  - Outliers are data points that are considerably different from other data points in the dataset
- **Missing Values 缺失值**
  - Missing feature values in data instances
  - **Solution:**
    - Remove instances that have missing values
    - Estimate missing values, and
    - Ignore missing values when running data mining algorithm
- **Duplicate data 重复数据**



- *Data Preprocessing (1)*

## 数据预处理

- **Aggregation** 聚合

- It is performed when multiple features need to be combined into a single one or when the scale of the features change
- Example: image width , image height -> image area (width x height)

- **Discretization** 离散化

- From continues values to discrete values
- Example: money spent -> {low, normal, high}



- **Data Preprocessing (2) 数据预处理**
- **Feature Selection 特征选择**
  - Choose relevant features
- **Feature Extraction 特征提取**
  - Creating new features from original features
  - Often, more complicated than aggregation
- **Sampling 取样**
  - Random Sampling
  - Sampling with or without replacement
  - Stratified Sampling: useful when having class imbalance
  - Social Network Sampling

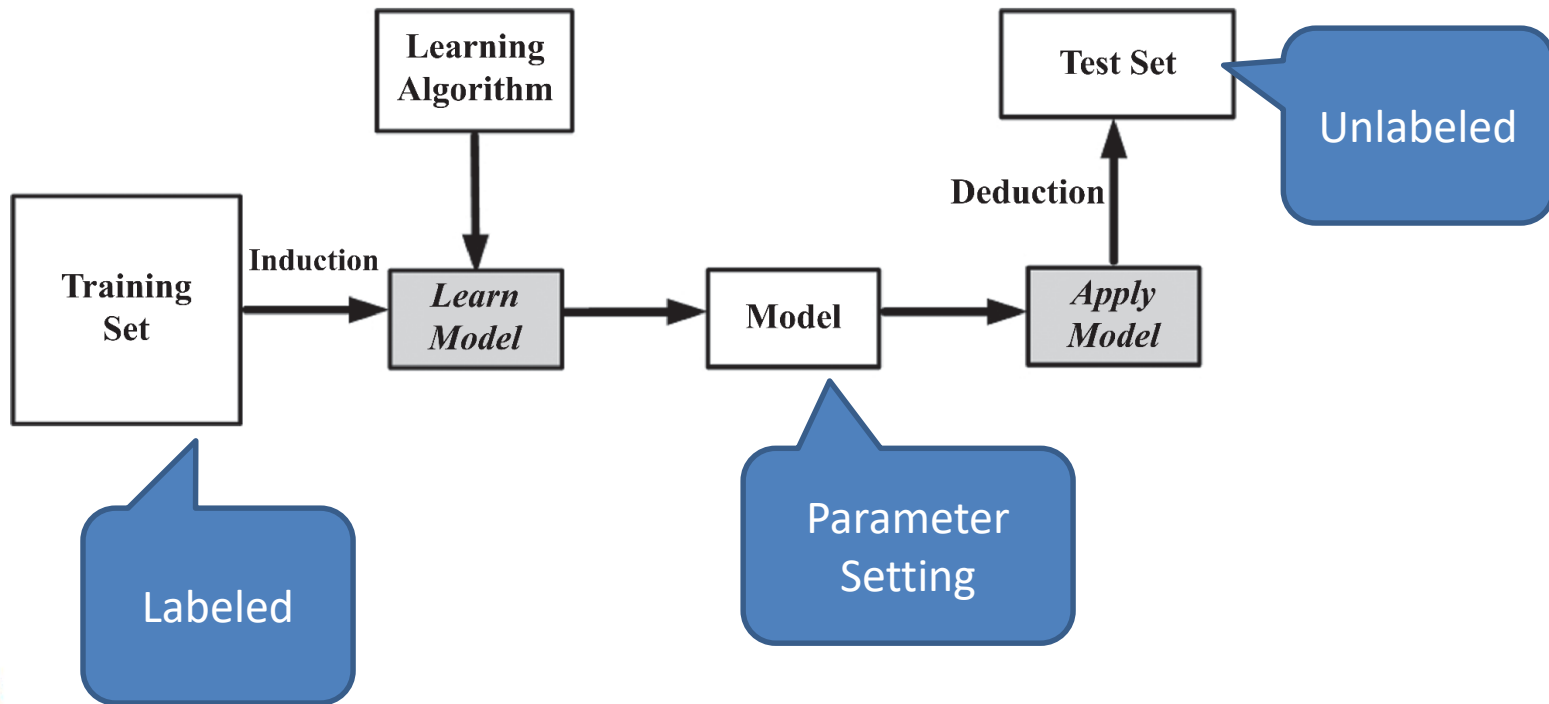
## *Machine Learning* 机器学习

- Supervised Learning
  - Classification
  - Regression
- Unsupervised Learning
  - Clustering
  - Dimensional Reduction



# Data Mining Essentials

## Supervised Machine Learning 有监督学习

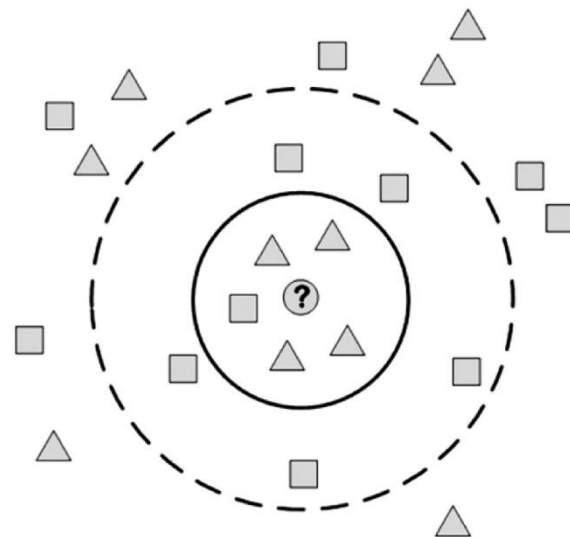


# Data Mining Essentials

## Classification 分类

Prediction Result with Labeled Discrete Value

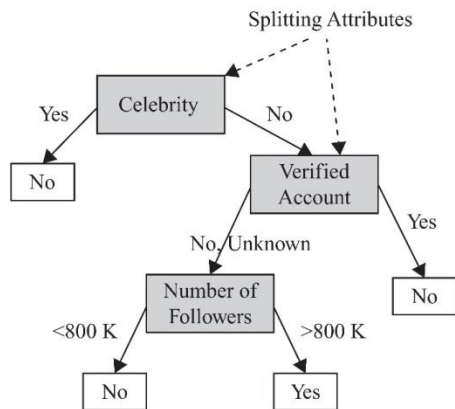
- KNN(K-Nearest Neighbors) K临近原则
- Linear Classifier 线性分类器
- Neural Networks 神经网络
- Support Vector Machine 支撑向量机
- Decision Tree 决策树



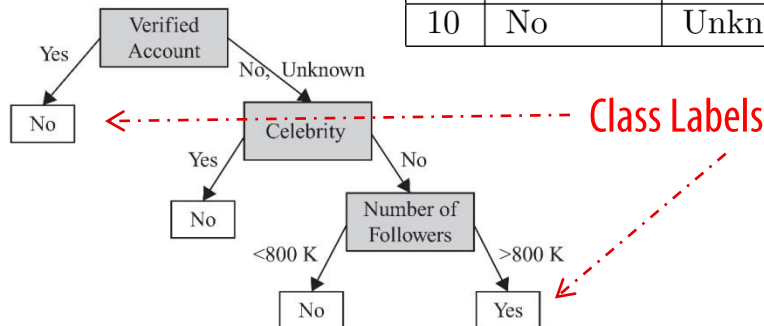
# Data Mining Essentials

Multiple **decision trees** can be learned from the same dataset

| ID | Celebrity | Verified Account | # Followers | Influential? |
|----|-----------|------------------|-------------|--------------|
| 1  | Yes       | No               | 1.25M       | No           |
| 2  | No        | Yes              | 1M          | No           |
| 3  | No        | Yes              | 600K        | No           |
| 4  | Yes       | Unknown          | 2.2M        | No           |
| 5  | No        | No               | 850K        | Yes          |
| 6  | No        | Yes              | 750K        | No           |
| 7  | No        | No               | 900K        | Yes          |
| 8  | No        | No               | 700K        | No           |
| 9  | Yes       | Yes              | 1.2M        | No           |
| 10 | No        | Unknown          | 950K        | Yes          |



(a) Learned Decision Tree 1

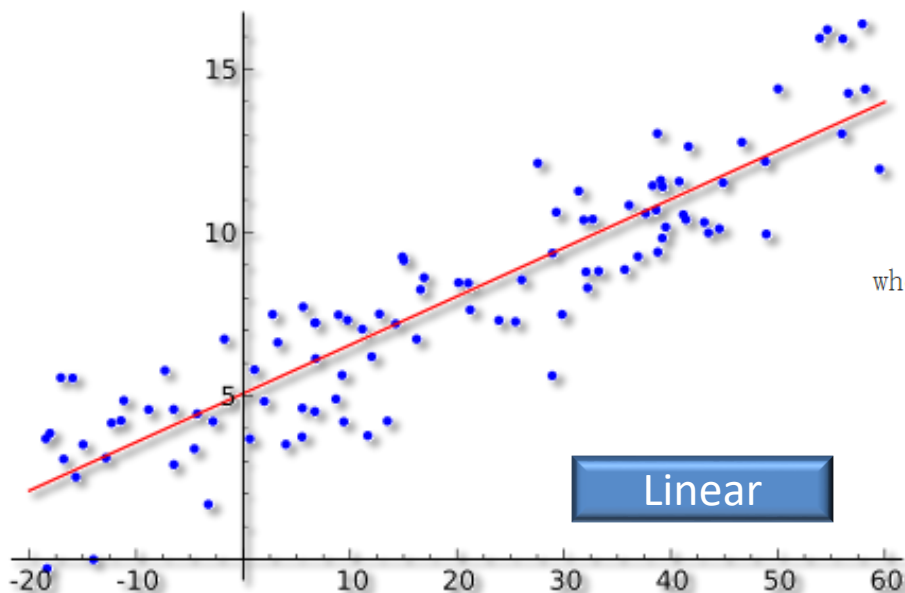


(b) Learned Decision Tree 2

# Data Mining Essentials

## Regression (1) 回归

Prediction Result with Unlabeled Continuous Value



Eg. Linear least squares 线性最小二乘法

$$\mathbf{X}\boldsymbol{\beta} = \mathbf{y},$$

where

$$\mathbf{X} = \begin{bmatrix} X_{11} & X_{12} & \cdots & X_{1n} \\ X_{21} & X_{22} & \cdots & X_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ X_{m1} & X_{m2} & \cdots & X_{mn} \end{bmatrix}, \quad \boldsymbol{\beta} = \begin{bmatrix} \beta_1 \\ \beta_2 \\ \vdots \\ \beta_n \end{bmatrix}, \quad \mathbf{y} = \begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_m \end{bmatrix}.$$

# Data Mining Essentials

## Regression (2) 回归

### Nonlinear Regression 非线性回归计算

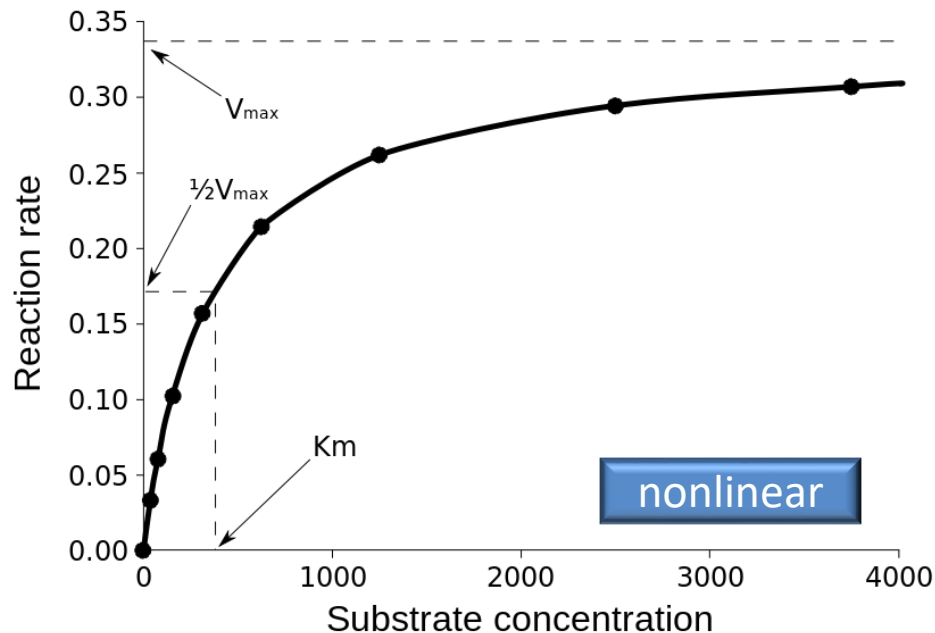
- Linearization 线性化方法

- Transformation 变形法

$$y = ae^{bx}U \Rightarrow \ln(y) = \ln(a) + bx + u$$

- Segmentation 分割法

split up into classes or segments and *linear regression* can be performed per segment



# Data Mining Essentials

## *Unsupervised Machine Learning*

### 无监督学习

machine learning task of inferring a function to describe hidden structure from unlabeled data



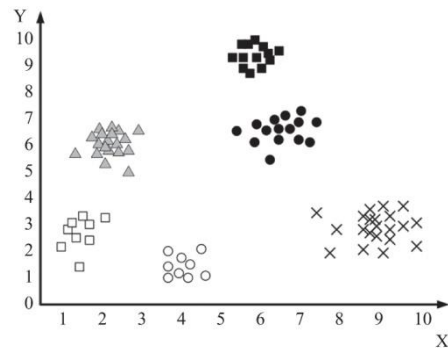
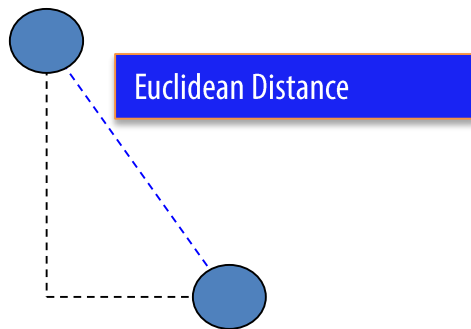
# Data Mining Essentials

## Clustering 聚类

- **Clustering Goal:** Group together similar items
- Clustering algorithms group together **similar items**
  - The algorithm does not have examples showing how the samples should be grouped together (unlabeled data)

### Similarity Computing (1) 相似度计算

- The most popular (dis)similarity measure for continuous features are **Euclidean Distance** and **Pearson Linear Correlation**



$$d(X, Y) = \sqrt{(x_1 - y_1)^2 + (x_2 - y_2)^2 + \cdots + (x_n - y_n)^2} = \sqrt{\sum_{i=1}^n (x_i - y_i)^2}$$

# Data Mining Essentials

## Similarity Computing (2) 相似度计算

*X and Y are n Dimensional Vectors*

$$X = (x_1, x_2, \dots, x_n)$$

$$Y = (y_1, y_2, \dots, y_n)$$

| Measure Name            | Formula                                          | Description                                                                    |
|-------------------------|--------------------------------------------------|--------------------------------------------------------------------------------|
| Mahalanobis             | $d(X, Y) = \sqrt{(X - Y)^T \Sigma^{-1} (X - Y)}$ | X, Y are features vectors and $\Sigma$ is the covariance matrix of the dataset |
| Manhattan ( $L_1$ norm) | $d(X, Y) = \sum_i  x_i - y_i $                   | X, Y are features vectors                                                      |
| $L_p$ -norm             | $d(X, Y) = (\sum_i  x_i - y_i ^n)^{\frac{1}{n}}$ | X, Y are features vectors                                                      |

Once a distance measure is selected, instances are grouped using it.



## Pearson Linear Correlation 皮尔逊线性相关

### Correlation Coefficient 相关系数

$$\rho_{X,Y} = \frac{\text{cov}(X, Y)}{\sigma_X \sigma_Y}$$

Where,  $\text{cov}$  is the covariance

$\sigma$  is the standard deviation

$$\text{cov}(X, Y) = E[(X - \mu_X)(Y - \mu_Y)]$$

$$\rho_{X,Y} = \frac{E[XY] - E[X]E[Y]}{\sqrt{E[X^2] - [E[X]]^2} \sqrt{E[Y^2] - [E[Y]]^2}}$$

Relations between  
Variance and Covariance

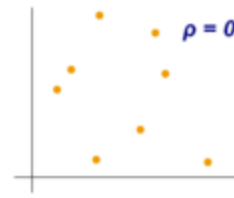
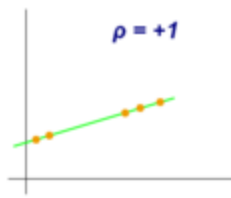
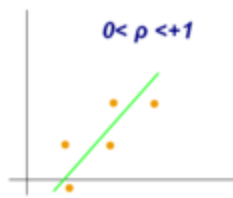
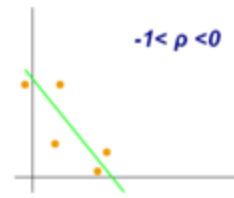
$$\mu_X = E[X]$$

$$\mu_Y = E[Y]$$

$$\sigma_X^2 = E[(X - E[X])^2] = E[X^2] - [E[X]]^2$$

$$\sigma_Y^2 = E[(Y - E[Y])^2] = E[Y^2] - [E[Y]]^2$$

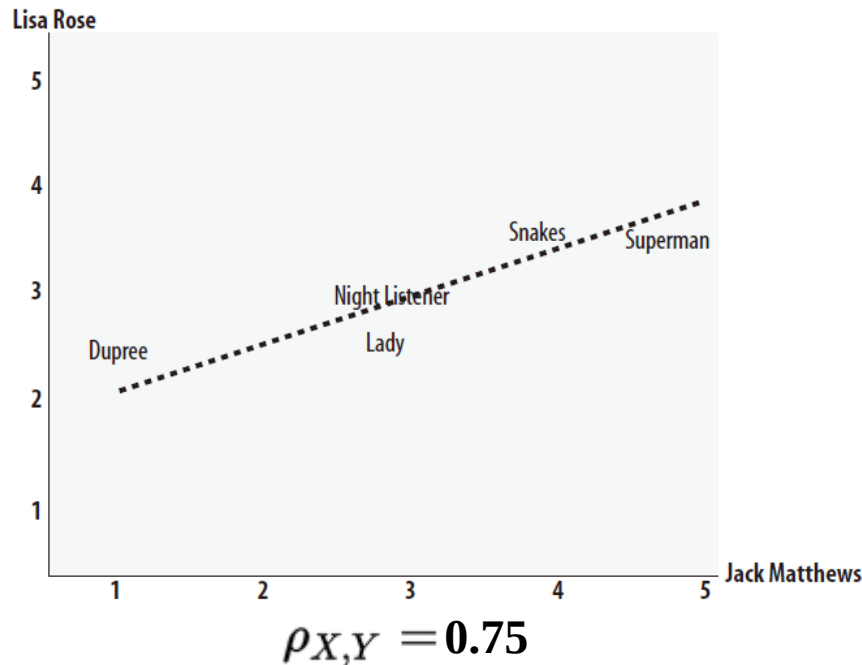
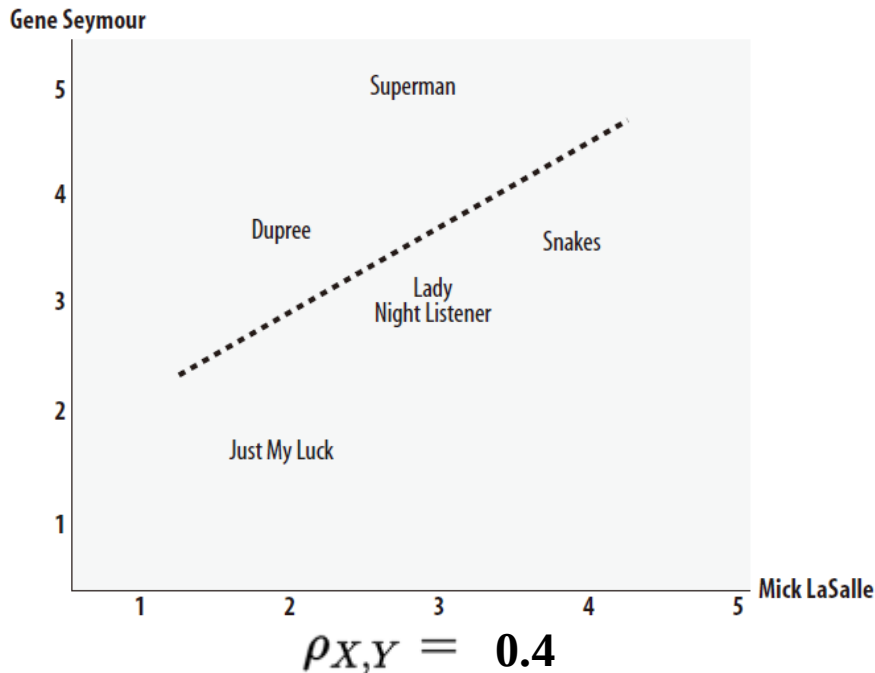
$$E[(X - \mu_X)(Y - \mu_Y)] = E[(X - E[X])(Y - E[Y])] = E[XY] - E[X]E[Y]$$



# Data Mining Essentials

## Film Ranking Correlation

*Superman* was rated 3 by Mick LaSalle and 5 by Gene Seymour, so it is placed at (3,5) on the chart.



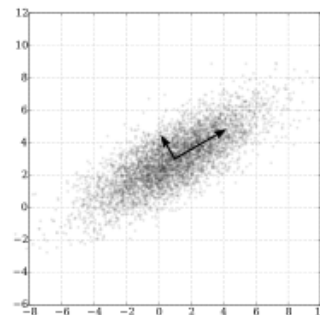
**Conclusion: Films recommended to Lisa, also can be recommended to Jack.**

# Data Mining Essentials

## *Dimensional Reduction* 降维

### Principal Component Analysis (PCA) 主成份分析

1. PCA is a statistical procedure **converts** a set of observations of possibly correlated variables **into** a set of values of linearly uncorrelated variables called principal components.
2. The number of principal components is less than or equal to the number of original variables.
3. This transformation is defined in such a way that the first principal component has **the largest possible variance**, and each succeeding component in turn has the highest variance possible under the constraint that it is orthogonal to the preceding components.





上海外国语大学  
SHANGHAI INTERNATIONAL STUDIES UNIVERSITY

# Reference

## *Books and Chapters (1)*

<https://item.jd.com/11983227.html>

Chapter 1-2

Machine Learning Package Installation

Machine Learning Theory Foundations



## *Books and Chapters (2)*

<https://item.jd.com/11803260.html>

Chapter 5

Data Mining Essentials

Online Reference:

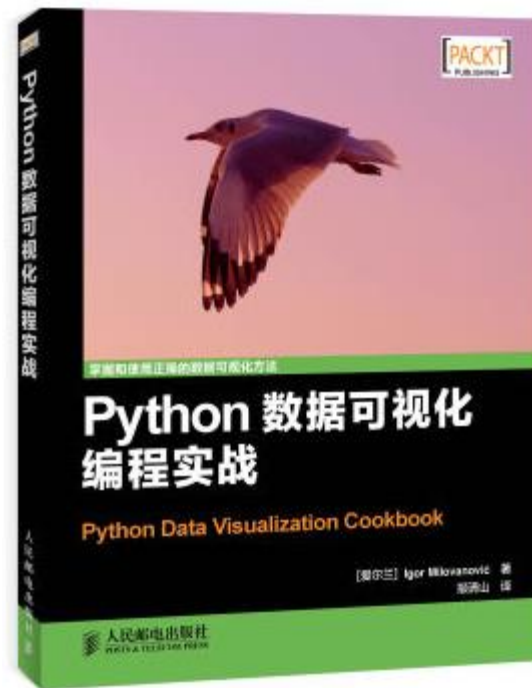
<http://www.public.asu.edu/~huanliu/>



## *Books and Chapters (3)*

<https://item.jd.com/11676691.html>

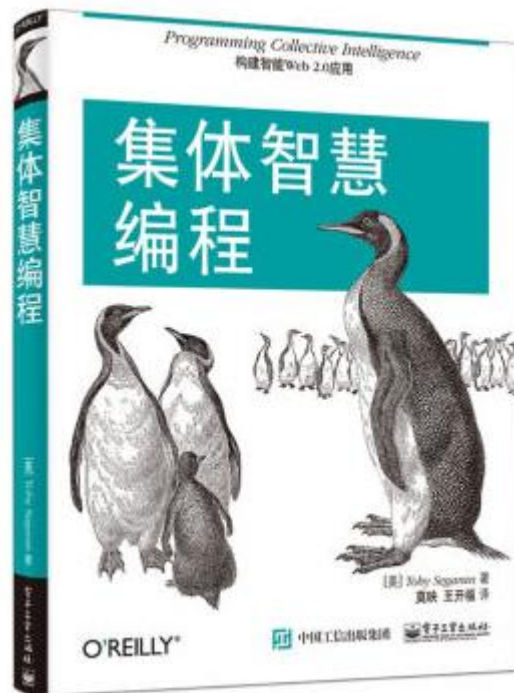
Python Data Visualization



## *Books and Chapters (4)*

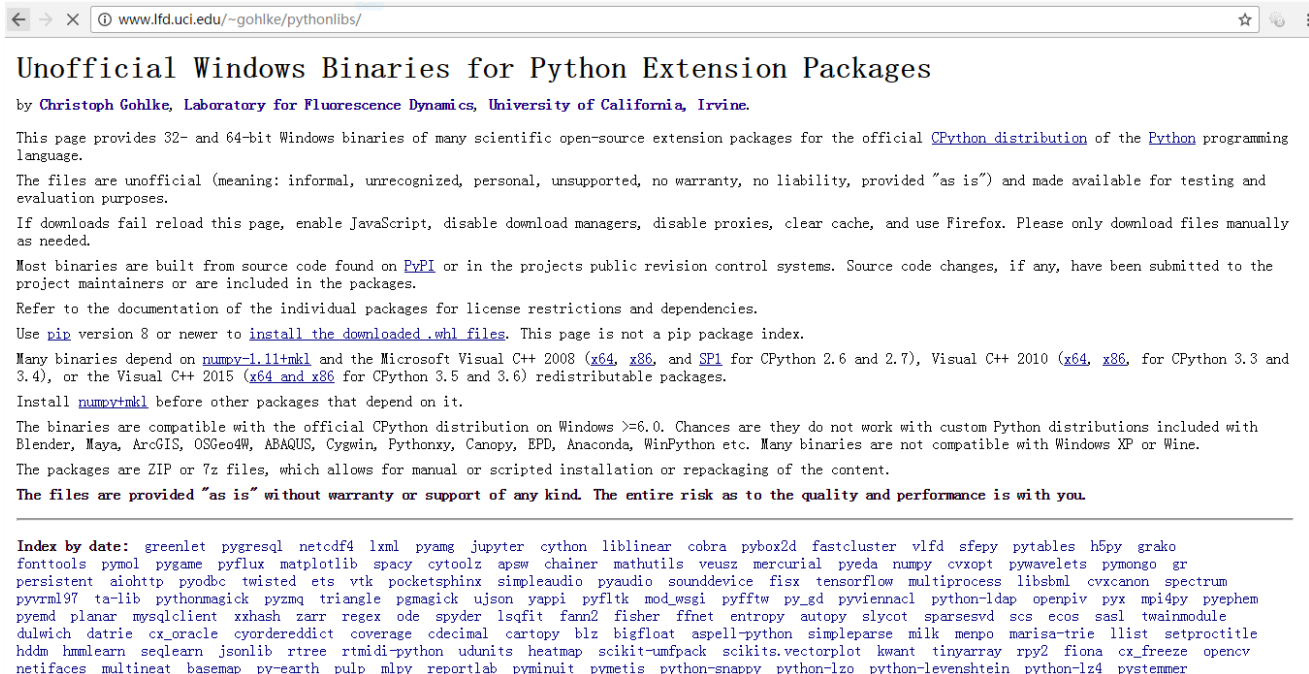
<https://item.jd.com/11667512.html>

Programming Collective Intelligence



## Python Extension Packages

<http://www.lfd.uci.edu/~gohlke/pythonlibs/>



← → × ⓘ www.lfd.uci.edu/~gohlke/pythonlibs/ ☆

### Unofficial Windows Binaries for Python Extension Packages

by [Christoph Gohlke](#), [Laboratory for Fluorescence Dynamics](#), [University of California, Irvine](#).

This page provides 32- and 64-bit Windows binaries of many scientific open-source extension packages for the official [CPython distribution](#) of the [Python](#) programming language.

The files are unofficial (meaning: informal, unrecognized, personal, unsupported, no warranty, no liability, provided "as is") and made available for testing and evaluation purposes.

If downloads fail reload this page, enable JavaScript, disable download managers, disable proxies, clear cache, and use Firefox. Please only download files manually as needed.

Most binaries are built from source code found on [PyPI](#) or in the projects public revision control systems. Source code changes, if any, have been submitted to the project maintainers or are included in the packages.

Refer to the documentation of the individual packages for license restrictions and dependencies.

Use [pip](#) version 8 or newer to [install the downloaded .whl files](#). This page is not a pip package index.

Many binaries depend on [numpy-1.11+mk1](#) and the Microsoft Visual C++ 2008 ([x64](#), [x86](#), and [SP1](#) for CPython 2.6 and 2.7), Visual C++ 2010 ([x64](#), [x86](#), for CPython 3.3 and 3.4), or the Visual C++ 2015 ([x64 and x86](#) for CPython 3.5 and 3.6) redistributable packages.

Install [numpy+mk1](#) before other packages that depend on it.

The binaries are compatible with the official CPython distribution on Windows >=6.0. Chances are they do not work with custom Python distributions included with Blender, Maya, ArcGIS, QGIS, ABAQUS, Cygwin, Pythonxy, Canopy, EPD, Anaconda, WinPython etc. Many binaries are not compatible with Windows XP or Wine.

The packages are ZIP or 7z files, which allows for manual or scripted installation or repackaging of the content.

**The files are provided "as is" without warranty or support of any kind. The entire risk as to the quality and performance is with you.**

---

**Index by date:** greenlet pygresql netcdf4 lxml pyang jupyter cython liblinear cobra pybox2d fastcluster vlfd sfepy pytables h5py grako fonttools pymol pygame pyflux matplotlib spacy cytoolz apsw chainer mathutils veusz mercurial pyeda numpy cvxopt pywavelets pymongo gr persistent aiohttp pyodbc twisted ets vtk pocketsphinx simpleaudio pyaudio sounddevice fixx tensorflow multiprocessing libsbml cvxcanon spectrum pyvrml97 ta-lib pythonmagick pyzmq triangle pgmagick ujson yappi pyfltk mod\_wsgi pyfftw py\_gd pyviennacl python-ldap openpiv pyx mpi4py pyephem pyemd planar mysqlclient xxhash zarr regex ode spyder lsqfit fann2 fisher ffnets entropy autopy slycot sparsesvd scs ecos sasl twainmodule dulwich datetrie cx\_oracle cyordereddict coverage cdecimal cartopy blz bigfloat aspell-python simpleparse milk mempo marisa-trie llist setproctitle hddm hmmlearn seqlearn jsonlib rtree rtmidi-python udunits heatmap scikit-umfpack scikits.vectorplot kwant tinyarray rpy2 fiona cx\_freeze opencv netifaces multineat basemap py-earth pulp mlpy reportlab pyminuit pymetis python-snappy python-lzo python-levenshtein python-lz4 systemer



## *Data Visualization in Python*

- <http://it.sohu.com/20151119/n427117609.shtml>
- <http://www.oschina.net/translate/python-data-visualization-libraries>



## *Using WordCloud*

- <http://blog.csdn.net/tanzuozecheng/article/details/50789226>
- [https://www.oschina.net/code/snippet\\_2294527\\_56155](https://www.oschina.net/code/snippet_2294527_56155)

## *Chinese Display*

- <http://blog.csdn.net/u012705410/article/details/47379957>



## *Provided Repositories for Social Mining*

- <http://socialcomputing.asu.edu>
- <http://snap.Stanford.edu>
- <https://github.com/caesar0301/awesome-public-datasets>





# The End of Lecture 9

Thank You

<http://www.wangting.ac.cn>

