

ครั้งที่ 5

การแบ่งกลุ่มข้อมูลอัตโนมัติ

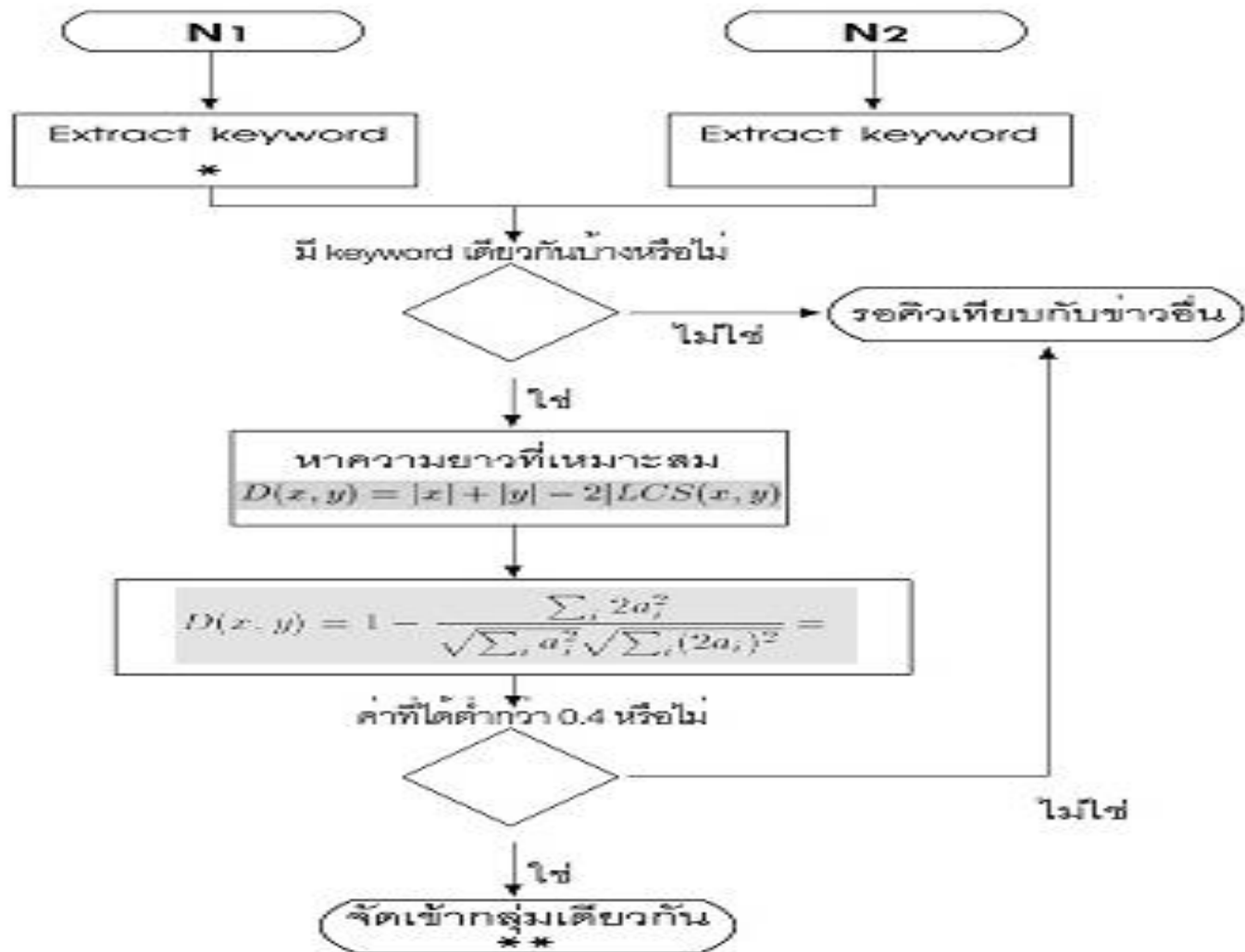
การแบ่งกลุ่มข้อมูลอัตโนมัติในระบบ **IR**

- การจัดกลุ่มคำสำคัญ
(Keyword clustering)

- การจัดกลุ่มของเอกสาร
(document clustering)



Clustering Architecture



Document Clustering Workbench screenshot

The screenshot displays the Carrot2 Workbench interface, which is used for document clustering. The main window is titled "clustering (100 documents from Yahoo Web (API), 19 clusters from Lingo)".

Left Panel (Search and Basic Settings):

- Source: Yahoo Web (API)
- Algorithm: Lingo
- Query (required): clustering
- Process button

Bottom Left Panel (Aduna Cluster Map Visualization):

This panel shows a hierarchical cluster map. The root node is "Clustering Data (19)", which branches into several sub-clusters, including "Clustering Algorithms (15)", "Definition of Clustering (5)", "Cluster/Analysis (5)", "Unsupervised Classification (5)", "K means Clustering Algo... (4)", and "Clustering Problem (5)".

Right Panel (Clusters and Documents):

The right panel is divided into two sections:

- Clusters:** A list of clusters with their respective document counts. The top-level cluster is "Clustering Data (19)". Other clusters include "Clustering Algorithms (15)", "New Clustering (10)", "Clustering Software (9)", "Cluster Analysis (5)", "Clustering Problem (5)", "Definition of Clustering (5)", "Clustering Resources (4)", "Linux Clustering (4)", "Clustering Technology (3)", "K-means Clustering Algorithm (3)", "Spectral Clustering (3)", "Storage Clustering (3)", "Application Server (2)", "Clustering Solution (2)", "Homework Help Math Clusters (2)", "SQL Server (2)", "Unsupervised Classification (2)", and "Other Topics (38)".
- Documents:** A list of documents with their titles and descriptions. The top document is "Cluster analysis - Wikipedia, the free encyclopedia", which describes clustering as a common technique for statistical data analysis. Other documents include "Clustering - Wikipedia, the free encyclopedia" and "Jboss.org: community driven".

Far Right Panel (Attributes and Settings):

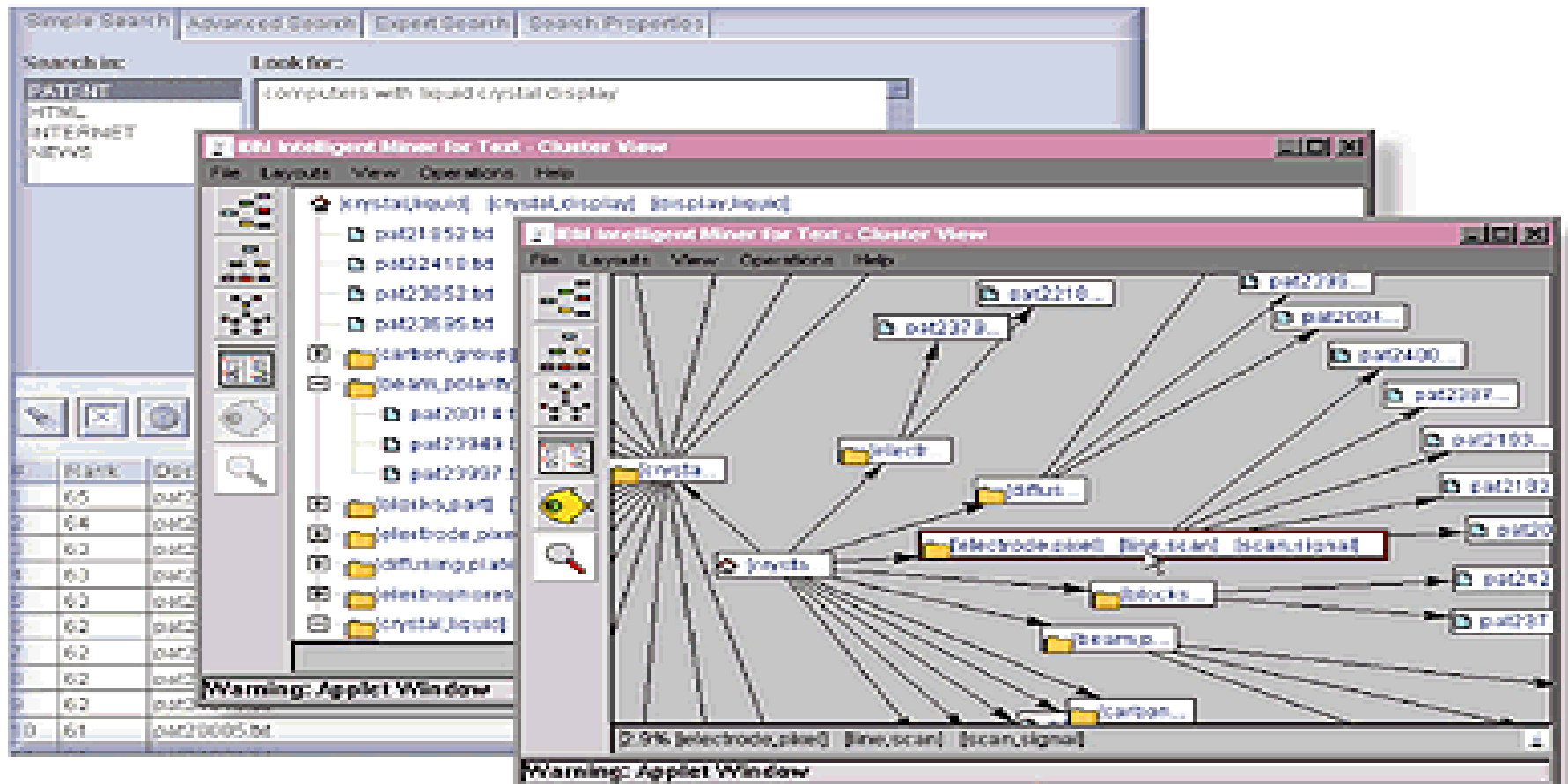
This panel contains settings for the clustering process:

- Clusters:** A slider for "Desired cluster count base" set to 25.
- Size-Score sorting ratio:** A slider set to 0.00.
- Label filtering:** Checkboxes for "Remove leading and trailing stop words", "Remove numeric labels", "Remove query words", and "Remove truncated phrases".
- Labels:** A section for "Attribute Info" and "Properties".

Attribute Info:

- Attribute:** Desired cluster count base
- Description:** Base factor used to calculate the number of clusters based on the number of documents on input. The larger the value, the more clusters will be created. The number of clusters created by the algorithm will be proportional to the cluster count base, but not in a linear way.
- Key:** LingoClusteringAlgorithm.desiredClusterCountBase
- Type:** Integer
- Default value:**

Hierarchical document clustering



R.M.Hayes ได้กล่าวไว้ว่า

- “โครงสร้างของการจัดหมวดหมู่ จะเป็นการรวมกลุ่มของไอเท็ม(**item**) ซึ่งอาจเป็นเอกสารหรือตัวแทนเอกสารเข้าด้วยกัน ซึ่งไอเท็มเหล่านี้จะถูกจัดให้รวมกลุ่มเปรียบเทียบให้เป็นหน่วยหนึ่ง ซึ่งไอเท็มเหล่านี้อาจจะสูญเสียคุณสมบัติเฉพาะไปบ้างไม่มากนัก”

โครงสร้างทางตรรกศาสตร์ (**Logical Organization**)

การแบ่งเอกสารโดยตรง ได้รับการ
พิสูจน์ว่ายากทางทฤษฎี ไม่
น่าเชื่อถือ

การแบ่งเอกสารโดยผ่านการคำนวณ
ของวิธีวัดความใกล้เคียงหรือความ
คล้ายคลึงกันของเอกสาร

วิธีวัดความคล้ายคลึง (*Measures of Association*)

ได้หลายวิธีดังนี้

Simple Coefficient

- เป็นรูปแบบที่ง่ายที่สุด คือ $|X \cap Y|$
- เป็นจำนวนของคำกรรชนีที่อยู่ทั้งในเอกสาร X และ Y วิธีนี้ไม่ได้นำเอาขนาดของ X และ Y มาพิจารณา

สมมุติ $X = \{1, 2, 3\}$

$$Y = \{1, 4\}$$

$$X \cap Y = \{1\}$$

$$|X \cap Y| = 1$$

Dice's Coefficient

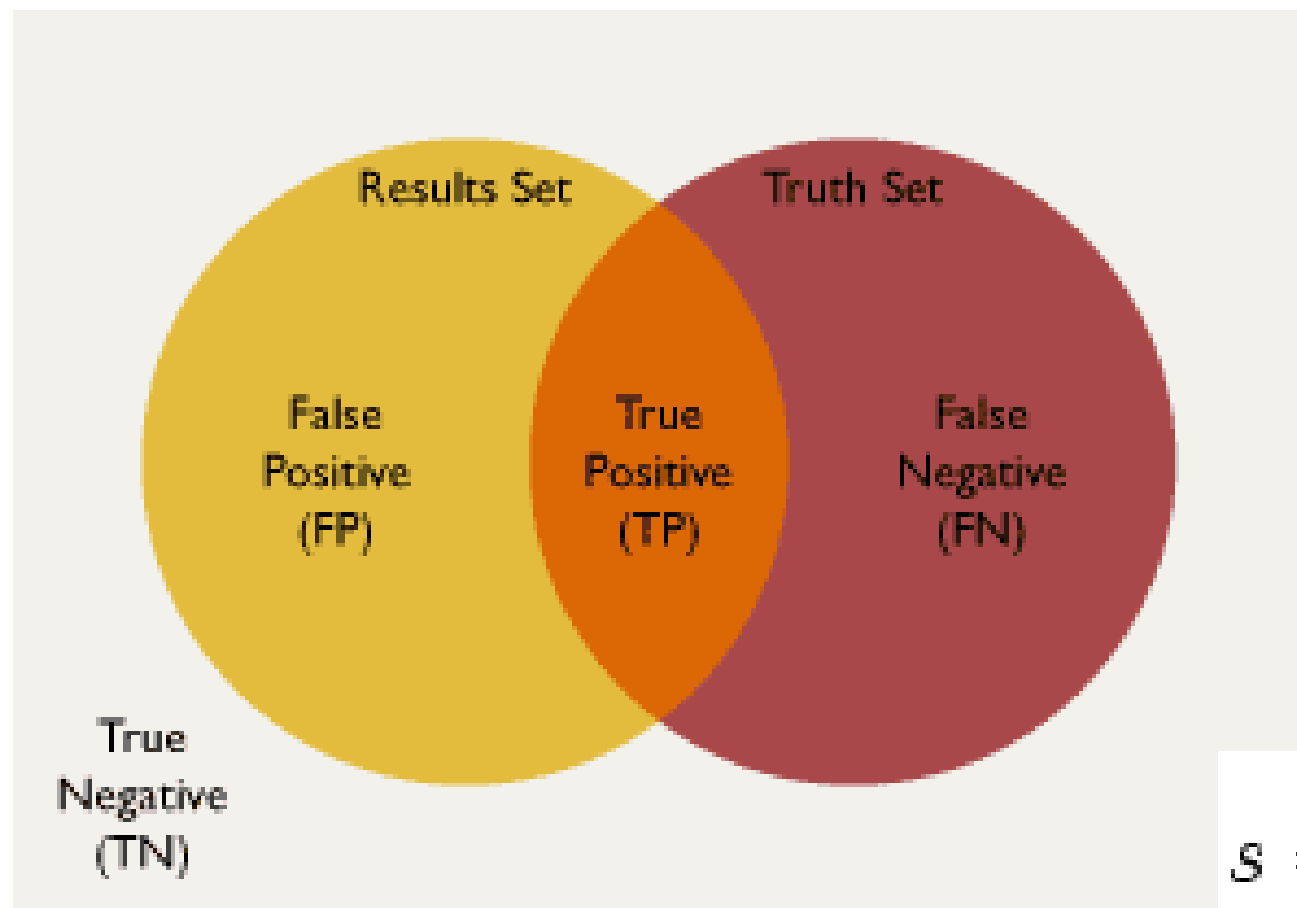
$$2 |X \cap Y| / |X| + |Y|$$

วิธีนี้มีการนำเอาขนาดของ X และ Y มาพิจารณา

$$\text{สมมุติ } X = \{1, 2, 3\} \quad Y = \{1, 4\}$$

$$X \cap Y = \{1\} \quad |X \cap Y| = 1$$

ดังนั้นค่าสัมประสิทธิ์มีค่าเท่ากับ $2*1 / 3+2 = 2/5$



$$s = \frac{2|X \cap Y|}{|X| + |Y|}$$

Jaccard's Coefficient

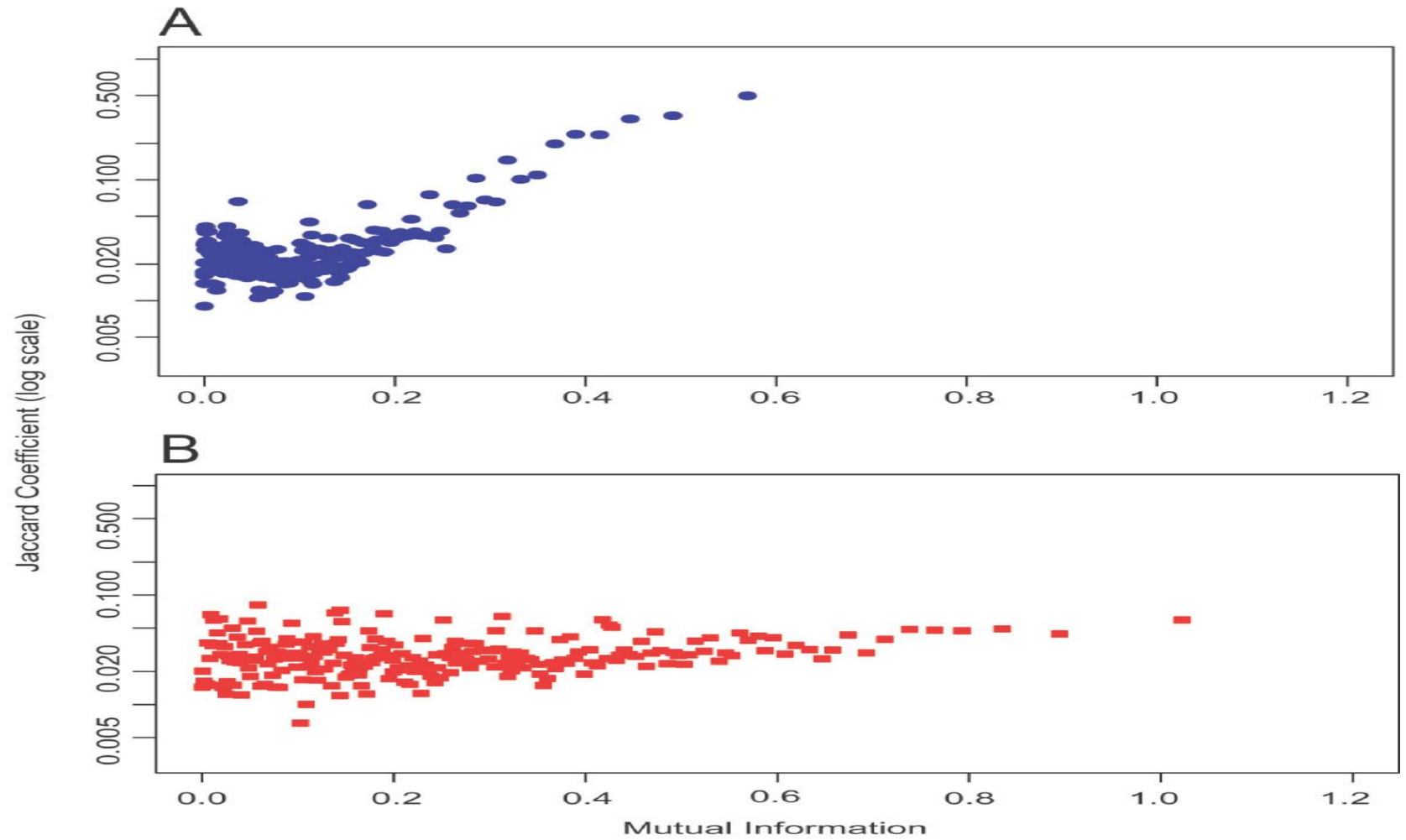
$$|X \cap Y| / |X \cup Y|$$

สมมุติ $X = \{1, 2, 3\}$ $Y = \{1, 4\}$

$$X \cap Y = \{1\} \quad |X \cap Y| = 1$$

$$X \cup Y = \{1, 2, 3, 4\} \quad |X \cup Y| = 4$$

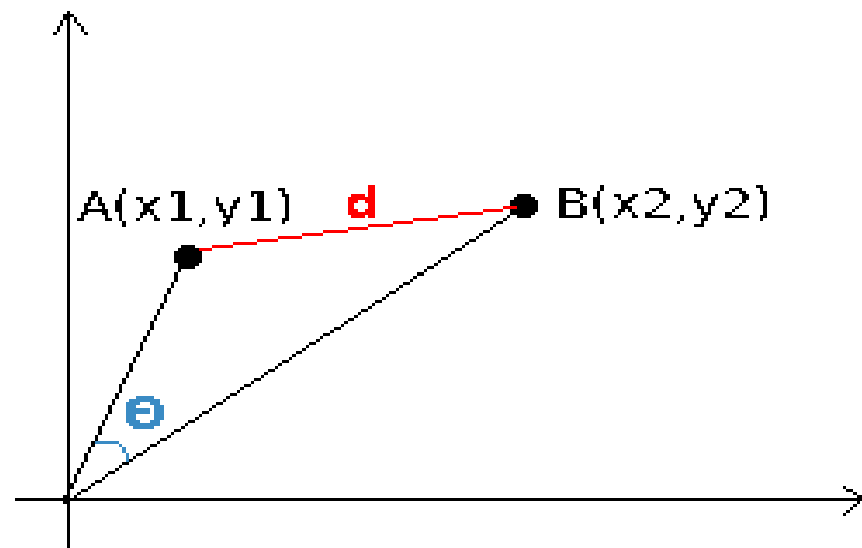
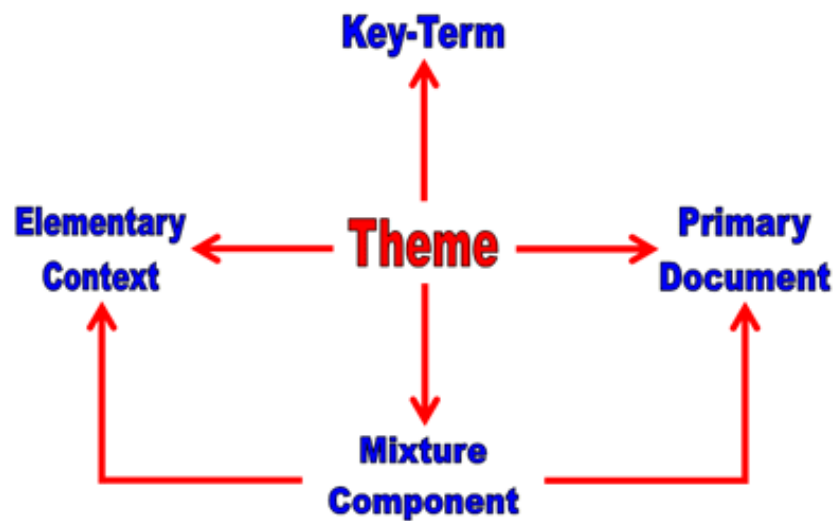
ดังนั้นค่าสัมประสิทธิ์มีค่าเท่ากับ $1 / 4$



Cosine Coefficient

- Salton ได้กำหนดตัวแทนเอกสารเป็นเวกเตอร์ของเลขฐานสอง ในลักษณะของ n มิติ ซึ่ง n เป็นจำนวนผลรวมของคำศัพท์นี้ ดังนั้น $(X \cap Y) / ||X|| ||Y||$ สามารถอธิบายได้เป็นมุม cosine ที่แยกเป็น 2 ไบนารีเวกเตอร์ X และ Y ดังนั้น สามารถเขียนตามเวกเตอร์จริงได้ดังนี้

$$|X \cap Y| / |X|^{1/2} * |Y|^{1/2}$$



Overlap coefficient

$$|X \cap Y| / \min(|X|, |Y|)$$

สมมุติ $X = \{1, 2, 3\}$ $Y = \{1, 4\}$

$$X \cap Y = \{1\} \quad |X \cap Y| = 1$$

ดังนั้นค่าสัมประสิทธิ์มีค่าเท่ากับ $1 / 2$

TABLE 6

Association Norms of German and French Students: Low Overlap Coefficient for "tranquil"

Rank	German Students		French Students	
1	forest	41	country	13
2	sleep	35	forest	11
3	church	20	nature	10
4	night	19	house	9
5	water	17	sea	9
6	read	14	room	8
7	cemetery	13	library	5
8	nature	9	landscape	5
9	bed	8	lonely person	5
10	lake	8	pensioner	5

Percentage of visual associations for the first 10 ranks.

Overlap coefficient = 0,155 (added relative frequencies of common associations, here calculated from all association percentages > 5 %); number of French students 210, number of German students 95.

Source: Institute for Consumer and Behavioral Research, 1992, Saarbrücken.

Similarity Measures

Simple matching (coordination level match)

$$|Q \cap D|$$

Dice's Coefficient

$$2 \frac{|Q \cap D|}{|Q| + |D|}$$

Jaccard's Coefficient

$$\frac{|Q \cap D|}{|Q \cup D|}$$

Cosine Coefficient

$$\frac{|Q \cap D|}{|Q|^{\frac{1}{2}} \times |D|^{\frac{1}{2}}}$$

Overlap Coefficient

$$\frac{|Q \cap D|}{\min(|Q|, |D|)}$$

การวัดความคล้ายคลึงกันโดยใช้ **Vector Space**

กำหนดให้ D เป็นเอกสาร , Q เป็นข้อสอบถาม ,
 W เป็นน้ำหนักของแต่ละเทอม

$$D_i = w_{d_{i1}}, w_{d_{i2}}, \dots, w_{d_{iM}}$$

$$Q = w_{q1}, w_{q2}, \dots, w_{qt}$$

$w = 0$ if a term is absent

การวัดความคล้ายคลึงกันโดยใช้ **Vector Space**

การวัดความคล้ายคลึงกันกระทำโดยนำน้ำหนักเทอมไป normalize
โดยสูตร

$$sim(Q, D_i) = \frac{\sum_{j=1}^t w_{qj} * w_{d_y}}{\sqrt{\sum_{j=1}^t (w_{qj})^2 * \sum_{j=1}^t (w_{d_y})^2}}$$

คำนวณน้ำหนักของแต่ละเอกสาร

$$w_{ik} = tf_{ik} * \log(N / n_k)$$

T_k = term k in document D_i

tf_{ik} = frequency of term T_k in document D_i

idf_k = inverse document frequency of term T_k in C

N = total number of documents in the collection C

n_k = the number of documents in C that contain T_k

$$idf_k = \log\left(\frac{N}{n_k}\right)$$

ตัวอย่าง

กำหนดให้เวกเตอร์ Q มีค่าดังนี้

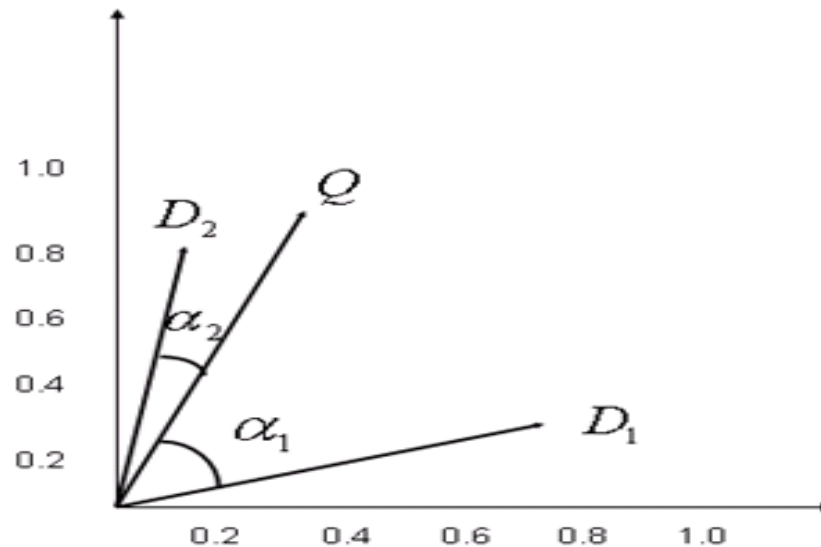
$$Q = (0.4, 0.8)$$

เอกสาร $D1$ มีค่าดังนี้

$$D1 = (0.8, 0.3)$$

เอกสาร $D2$ มีค่าดังนี้

$$D2 = (0.2, 0.7)$$



ตัวอย่าง

- เราสามารถหาความคล้ายคลึงจากเอกสารและข้อสอบถามได้โดยใช้สูตรคือ

$$\text{sim}(Q, D_i) = \frac{\sum_{j=1}^t w_{qj} * w_{d_y}}{\sqrt{\sum_{j=1}^t (w_{qj})^2 * \sum_{j=1}^t (w_{d_y})^2}}$$

ตัวอย่าง

$$sim(Q, D_i) = \frac{\sum_{j=1}^t w_{q_j} * w_{d_{ij}}}{\sqrt{\sum_{j=1}^t (w_{q_j})^2 * \sum_{j=1}^t (w_{d_{ij}})^2}}$$

$$Q = (0.4, 0.8) \quad D_2 = (0.2, 0.7)$$

$$\begin{aligned} sim(Q, D_2) &= \frac{(0.4 * 0.2) + (0.8 * 0.7)}{\sqrt{[(0.4)^2 + (0.8)^2] * [(0.2)^2 + (0.7)^2]}} \\ &= \frac{0.64}{\sqrt{0.42}} = 0.98 \end{aligned}$$

ตัวอย่าง

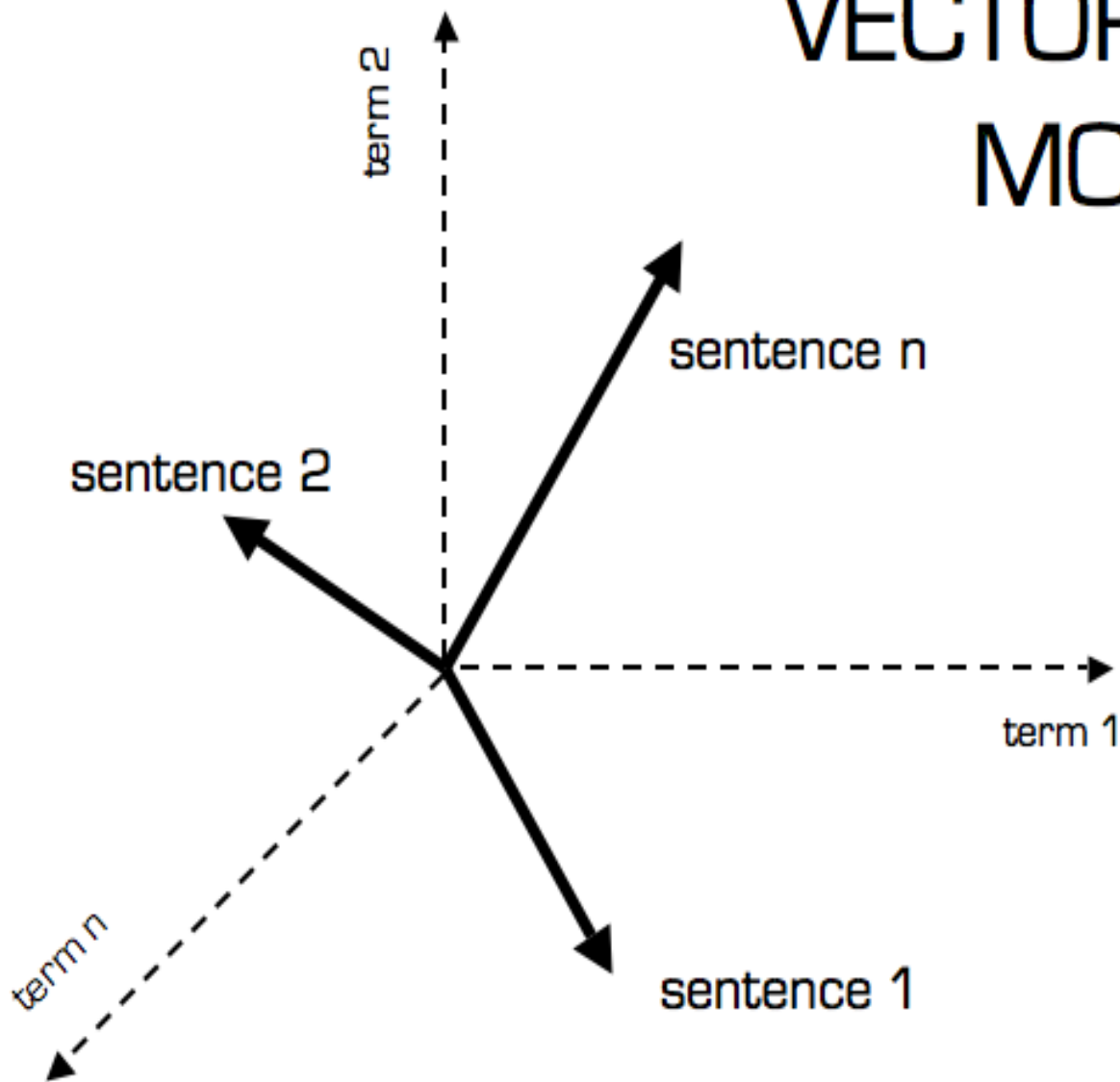
$$sim(Q, D_i) = \frac{\sum_{j=1}^t w_{qj} * w_{d_j}}{\sqrt{\sum_{j=1}^t (w_{qj})^2 * \sum_{j=1}^t (w_{d_j})^2}}$$

$$Q = (0.4, 0.8) \quad D1 = (0.8, 0.3)$$

แทนค่าในสูตรระหว่างข้อสอบถาม(Q) และเอกสาร D1 ได้ดังนี้

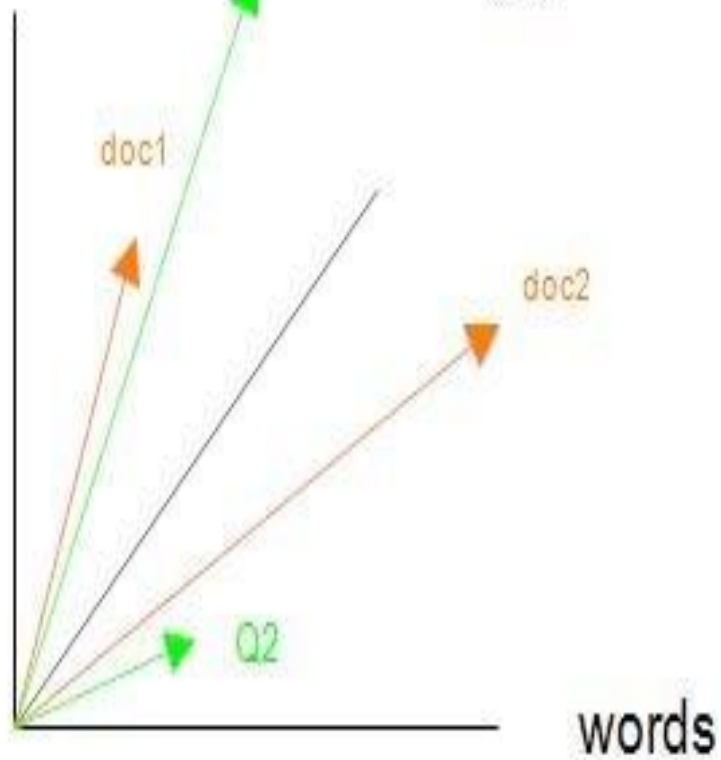
$$SIM(Q, D1) = \frac{0.56}{\sqrt{0.58}} = 0.74$$

VECTOR SPACE MODEL



all

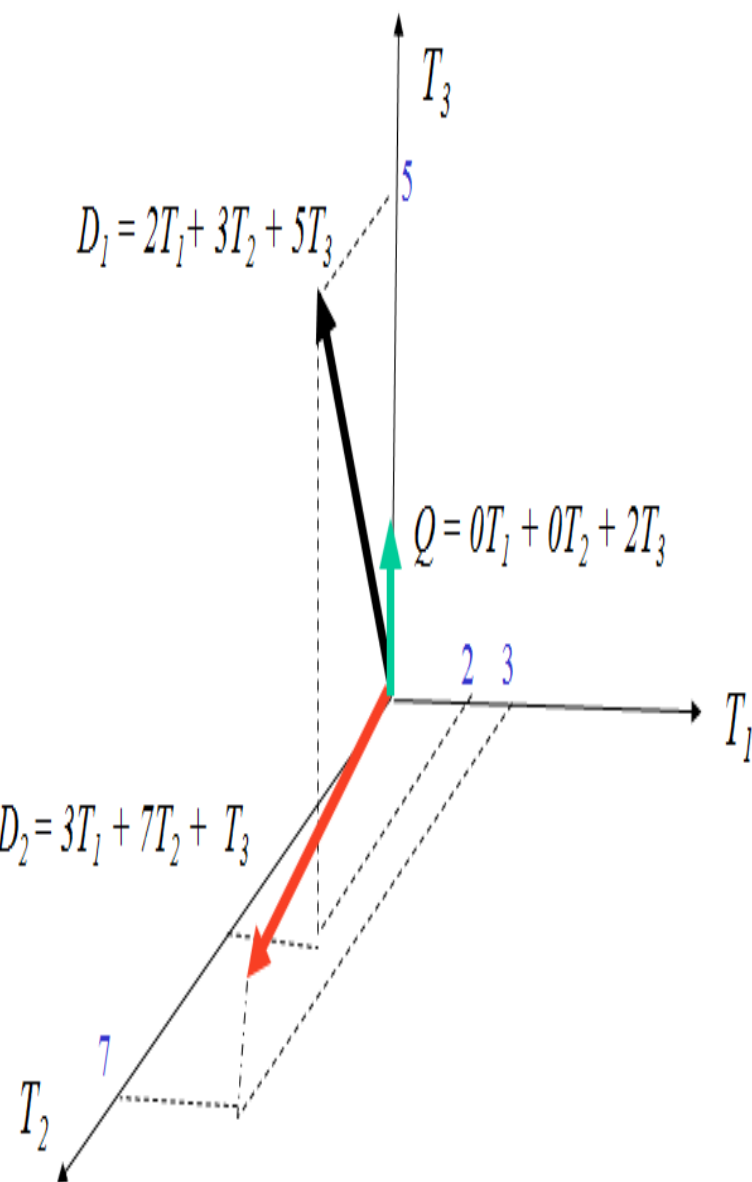
our



$$D_1 = 2T_1 + 3T_2 + 5T_3$$

$$D_2 = 3T_1 + 7T_2 + T_3$$

$$Q = 0T_1 + 0T_2 + 2T_3$$



ความไม่คล้ายคลึงกัน (**Dissimilarity**)

กำหนด P เป็นเซตของสิ่งของที่ถูกจัดกลุ่ม

D เป็นสัมประสิทธิ์ความไม่คล้ายคลึงกัน เป็นฟังก์ชันหนึ่งจาก

$P \times P$ ไปยังเลขจำนวนจริงที่ไม่เป็นค่าลบ

โดยทั่วไป D จะมีคุณสมบัติดังต่อไปนี้

$$D(X, Y) > 0 \text{ สำหรับทุกๆ } X, Y \text{ ของ } P$$

$$D(X, X) = 0 \text{ สำหรับทุกๆ } X \text{ ของ } P$$

$$D(X, Y) = D(Y, X) \text{ สำหรับทุกๆ } X, Y \text{ ของ } P$$

$$D(X, Y) < D(X, Z) + D(Y, Z)$$

สูตรความไม่คล้ายคลึงกัน

$$|X \Delta Y| / |X| + |Y| \quad \text{โดย} \quad |X \Delta Y| = |X \cup Y| - |X \cap Y|$$

ซึ่งสูตรนี้มาจากสูตรความคล้ายคลึงของ **Dice's Coefficient**

$$2 |X \cap Y| / |X| + |Y|$$

สามารถหาความไม่คล้ายคลึงกันได้ดังนี้

$$\begin{aligned} \text{Dissimilarity} &= 1 - (2 |X \cap Y| / |X| + |Y|) \\ &= (|X| + |Y| - 2 |X \cap Y|) / (|X| + |Y|) \\ &= (|X \cup Y| - |X \cap Y|) / (|X| + |Y|) \\ &= |X \Delta Y| / |X| + |Y| \end{aligned}$$

$$|X \Delta Y| / |X| + |Y| \quad \text{โดย} \quad |X \Delta Y| = |X \cup Y| - |X \cap Y|$$

$$\text{สมมุติ } X = \{1, 2, 3\} \quad Y = \{1, 4\}$$

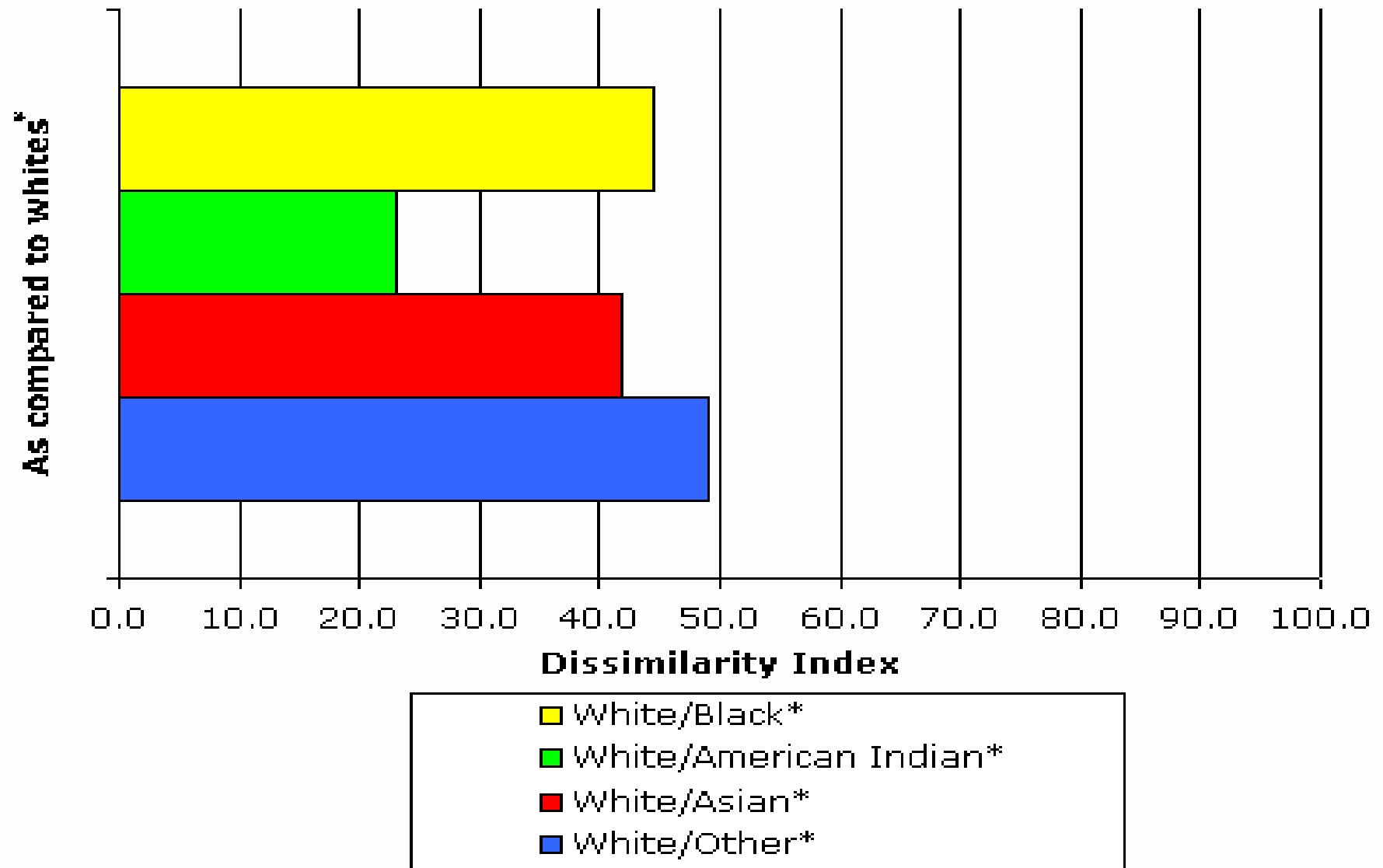
$$X \cap Y = \{1\} \quad |X \cap Y| = 1$$

$$X \cup Y = \{1, 2, 3, 4\} \quad |X \cup Y| = 4$$

ดังนั้นค่าสัมประสิทธิ์ความไม่คล้ายคลึงกัน

$$\text{มีค่าเท่ากับ } (4 - 1) / (3 + 2) = 3/5$$

Dissimilarity Indices for Selected Multiracial Groups



Jaccard

- ใช้เลขฐานสองแทนค่าบรรชีของเอกสาร มีผลให้บรรชีตัวที่ i จะแทนด้วยเลข **0** หรือ **1** ในตำแหน่งที่ i ตามลำดับ ซึ่งค่า 1 นั้นจะแสดงถึงค่าสำคัญอยู่ในเอกสาร สำหรับเลข 0 หมายถึงค่าสำคัญไม่ได้อยู่ในเอกสารนั้นในตำแหน่งที่ระบุ

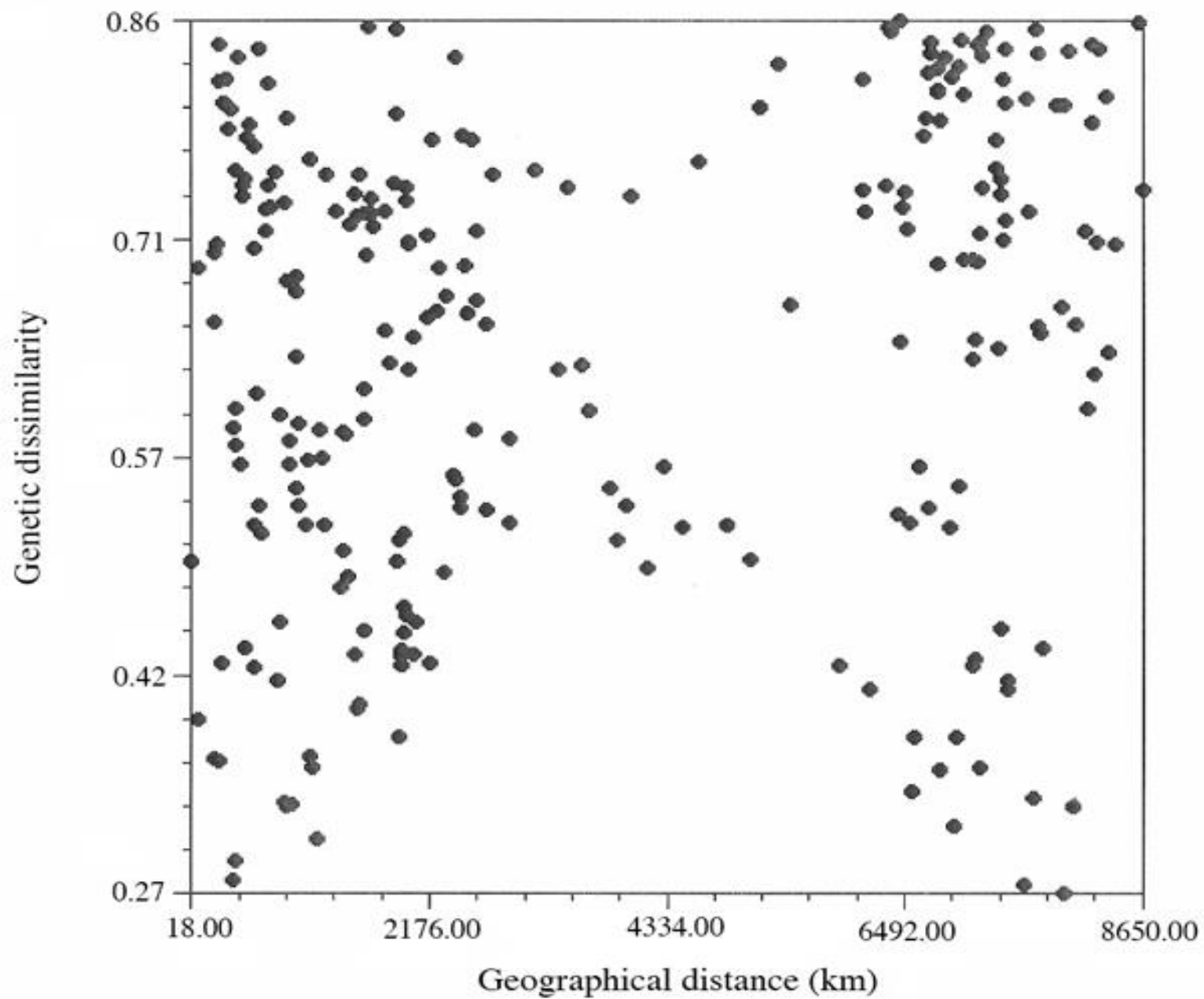
ถ้า $|X| = \sum X_i$ โดยที่ $i \in 1..N$

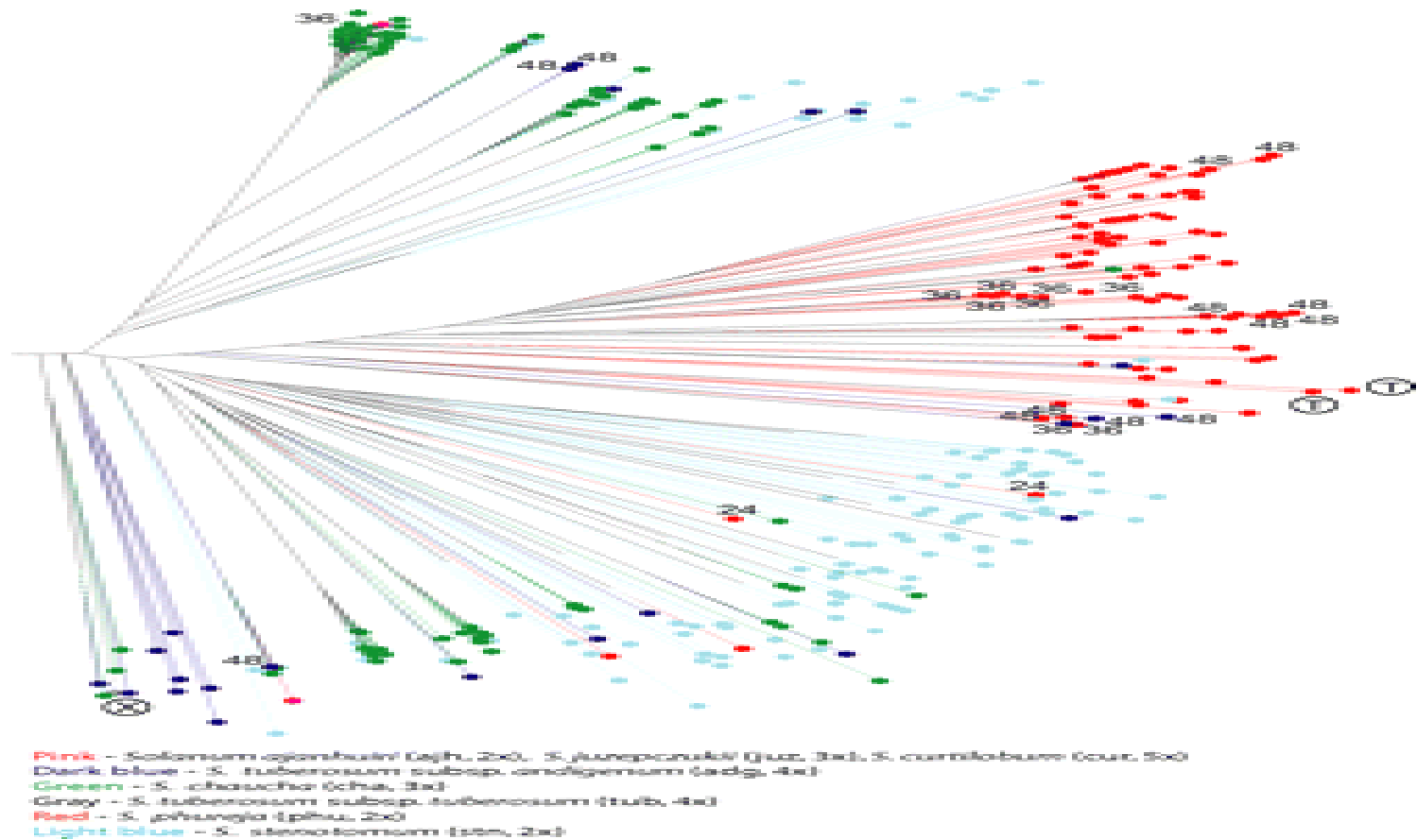
และ N เป็นจำนวนของค่าบรรชีทั้งหมดในเอกสาร

ดังนั้น $|X \cap Y| = \sum X_i Y_i$ โดยที่ $i \in 1..N$

ความไม่คล้ายคลึงกันมีค่าเท่ากับ

$$|X \Delta Y| / |X| + |Y| = (\sum X_i (1-Y_i) + \sum Y_i (1-X_i)) / (\sum X_i + \sum Y_i)$$





การปรับปรุงประสิทธิภาพระบบค้นคืนสารสนเทศ

ใช้ความคาดหวังเพื่อใช้ในการวัดความเกี่ยวข้องได้

สามารถกระจายความน่าจะเป็น $P(X_i)$ และ $P(X_j)$

$$P(x_i, x_j) = \sum_{x_i, x_j} P(x_i, x_j) \quad \log \frac{P(x_i, x_j)}{P(x_i)P(x_j)}$$

Jardine และ Sibson

- วัดความไม่คล้ายคลึงกันระหว่าง สองออปเจ็กต์ โดยแบ่งการกระจายความน่าจะเป็นเป็นสองจุดคือ 1 และ 0 ดังนั้น $P_1(1), p_1(0), P_2(1), P_2(0)$ เป็นการกระจายความน่าจะเป็นของออปเจ็กต์ทั้งสอง ซึ่งการวัดความไม่คล้ายคลึงกันของ Jardine และ Sibson นี้เรียกว่ารัศมีสารสนเทศ (Information Radius)

มีการคำนวณดังนี้

$$\begin{aligned} & w_{P_1}^P(1) \log \frac{P_1(1)}{w_{P_1}^P(1) + w_{P_2}^P(1)} + w_{P_2}^P(1) \log \frac{P_2(1)}{w_{P_1}^P(1) + w_{P_2}^P(1)} + \\ & w_{P_1}^P(0) \log \frac{P_1(0)}{w_{P_1}^P(0) + w_{P_2}^P(0)} + w_{P_2}^P(0) \log \frac{P_2(0)}{w_{P_1}^P(0) + w_{P_2}^P(0)} \end{aligned}$$

วิธีจัดแบ่งกลุ่ม (Classification Methods)

ออฟเจกต์

อาจเป็นเอกสาร(documents)

คำสรรชนี(keyword)

ตัวอักษรที่เขียนด้วยมือคน

(hand written character)หรือ

ชนิดของพืช(species)

ส่วนคำอธิบาย(Description)

เป็นโครงสร้างของข้อมูลภายใต้ชื่อต่างๆ
อาจเป็นคุณสมบัติในหลายๆสถานะ

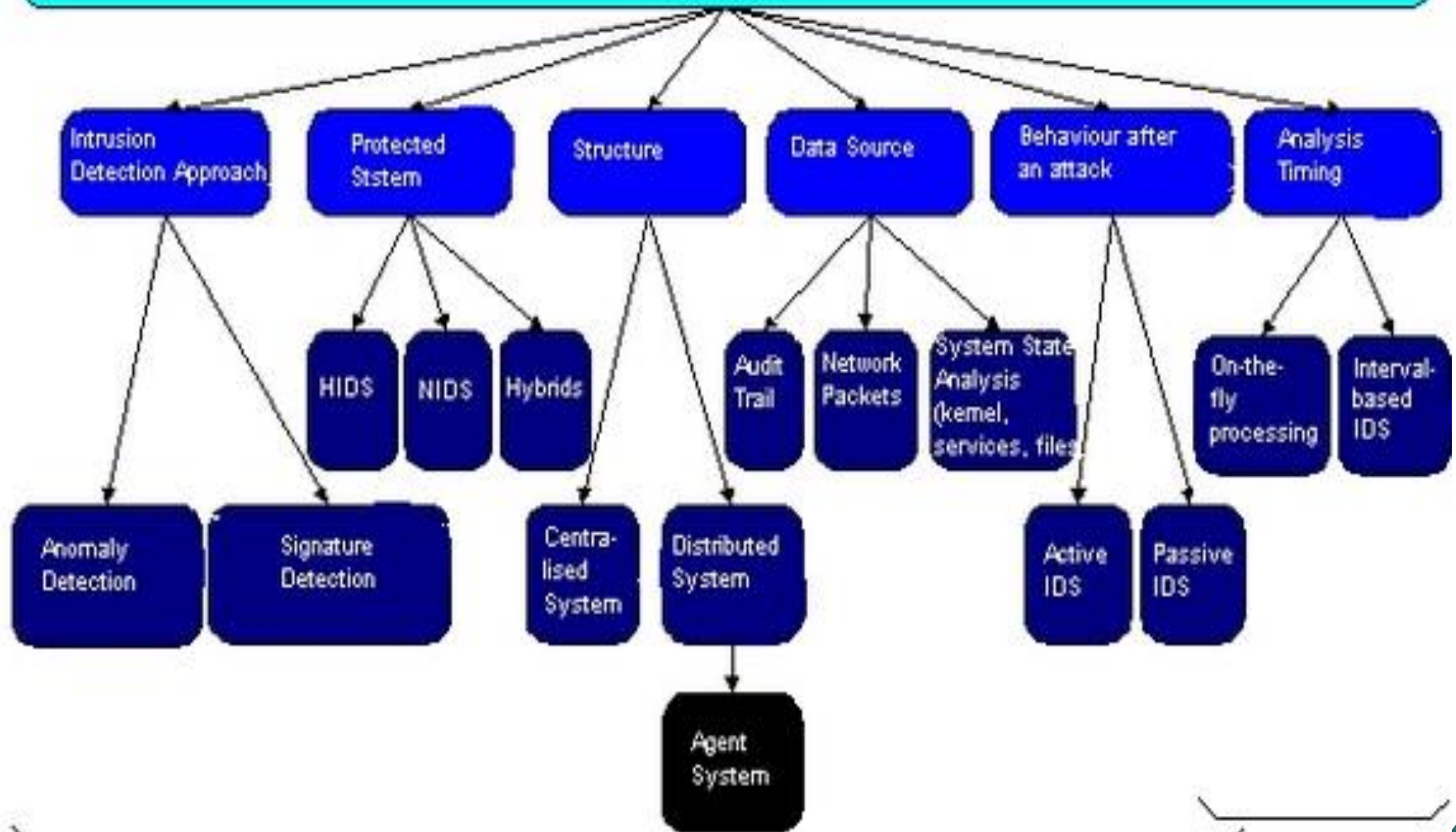
- สี

- สถานะของเลขฐานสอง เช่น คำสรรชนี
(Keyword)

- จำนวน เช่น มาตรการวัดความแข็ง หรือ
น้ำหนักของสรรชนี เป็นต้น

- การกระจายความน่าจะเป็น
(Probability Distributions)

Intrusion Detection System

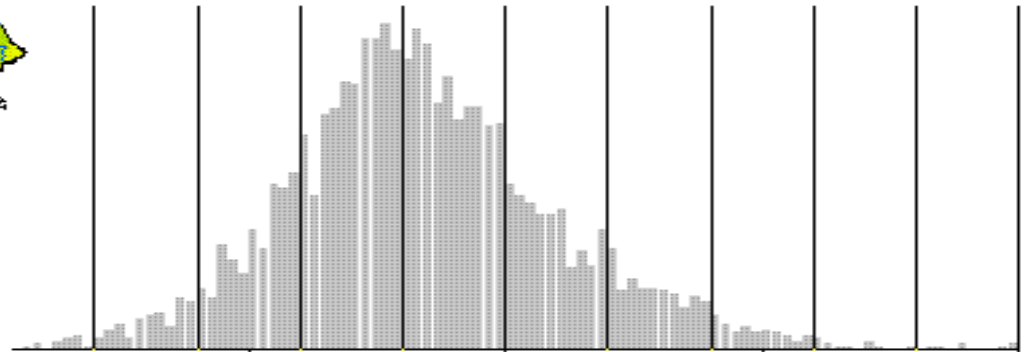
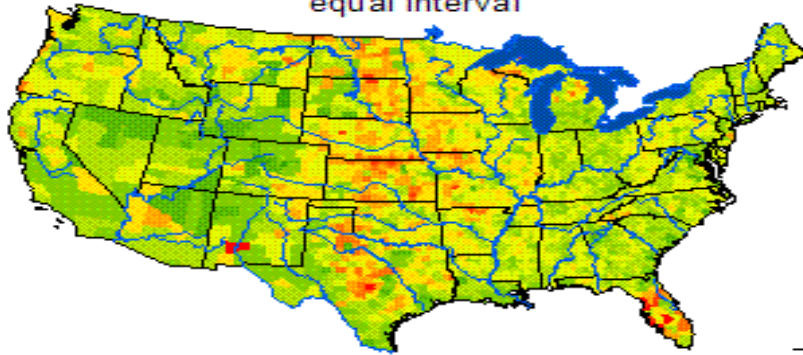


COMPARISON OF CLASSIFICATION METHODS

Percentage of US Population Age 65 or Older, by County

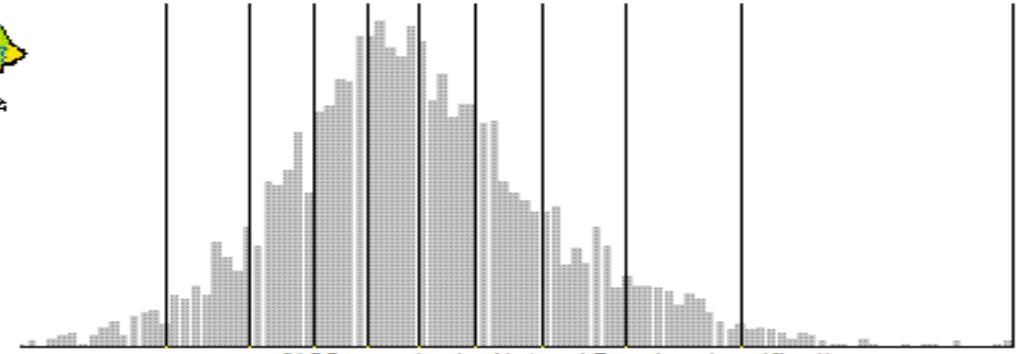
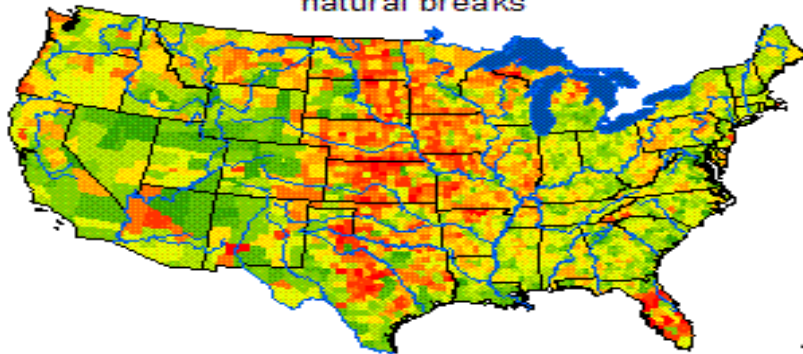
Distributions of Counties

equal interval



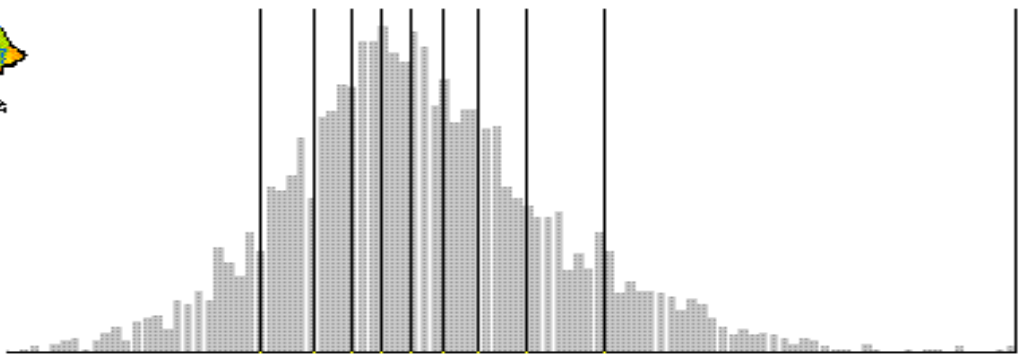
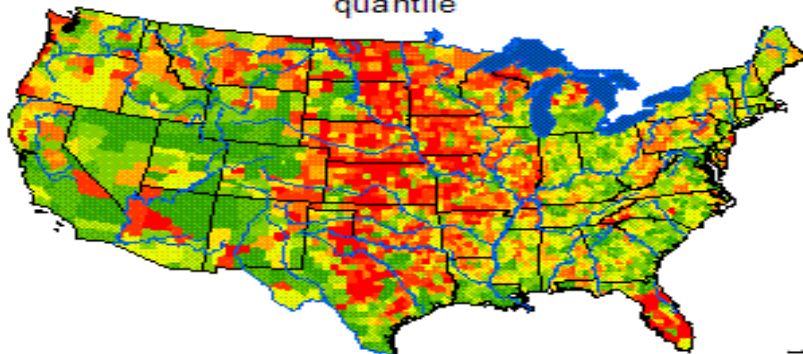
%65up -- equal data intervals

natural breaks



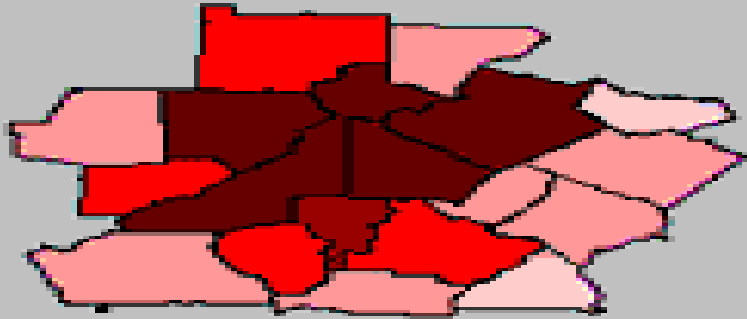
%65up -- Jenks Natural Breaks classification

quantile

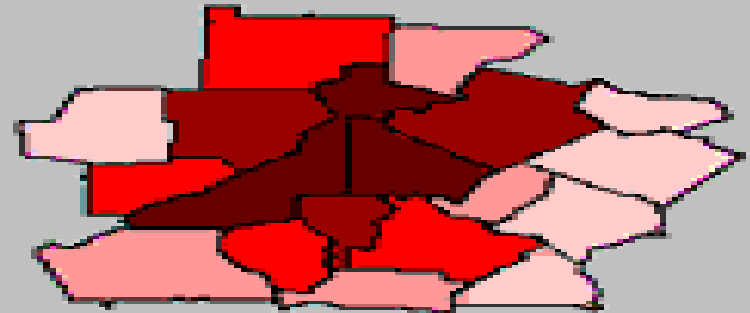


%65up -- equal number of counties in each category

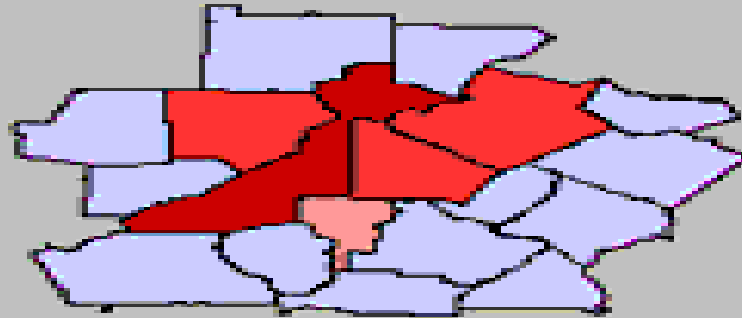
 Natural Breaks



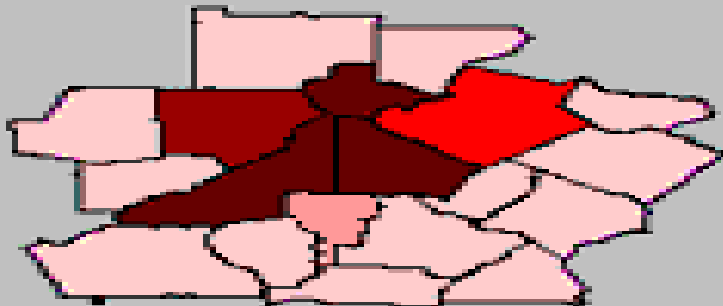
 Equal Area



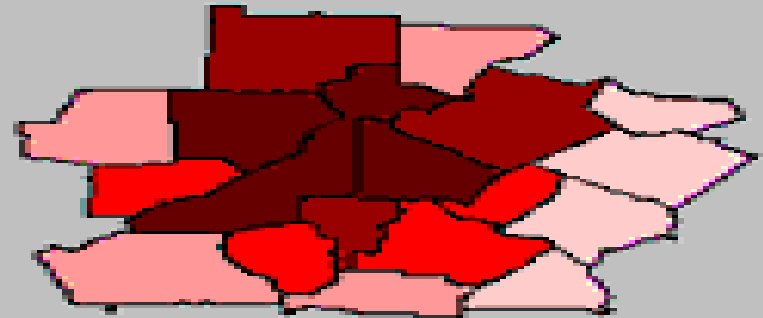
 Standard Deviation



 Equal Interval



 Quantile



are grouped into five classes using different classification methods.

Sparck Jones

ได้แบ่งวิธีการจัดกลุ่มตามความสัมพันธ์ได้
หลายลักษณะดังนี้

ความสัมพันธ์ระหว่างคุณสมบัติและคลาส

- **Monothetic**

กลุ่มของ **Individual** ที่มี
คุณสมบัติเหมือนกัน

- **Polythetic**

คุณสมบัติในสปีชีส์จะมีการจะมี
คุณสมบัติที่เหมือนกัน 3 ข้อที่
เหมือนกัน เป็นต้น กล่าวคือไม่
เหมือนกันหมด

ความสัมพันธ์ระหว่างคุณสมบัติและคลาส

รูปแสดงความแตกต่างระหว่าง **Monothetic** และ **Polythetic**

	A	B	C	D	E	F	G	H
1	+	+	+					
2	+	+		+				
3	+		+	+				
4		+	+	+				
5					+	+	+	
6					+	+	+	
7					+	+		+
8					+	+		+

Polythetic

Monothetic



ความสัมพันธ์ระหว่างออปเจกต์และคลาส

Exclusive Class

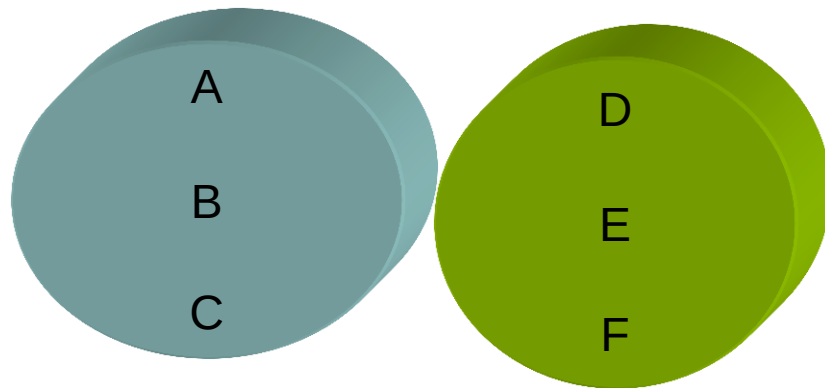
กลุ่มของ individual ที่เป็น
สมาชิกเฉพาะ ไม่มี

**Overlapping
Individuals**

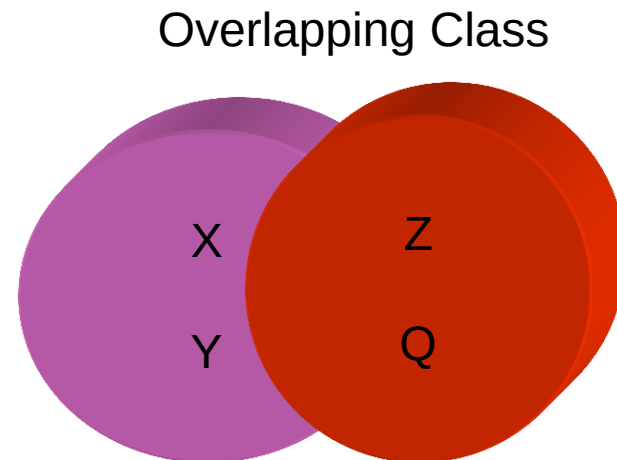
Overlapping Class

กลุ่มของ Individuals ที่อาจมี
สมาชิกหนึ่งของกลุ่มนี้ ไปเป็น
สมาชิกของกลุ่มอื่น

ความสัมพันธ์ระหว่างออปเจกต์และคลาส



Exclusive Class



ความสัมพันธ์ระหว่างคลาสกับคลาส

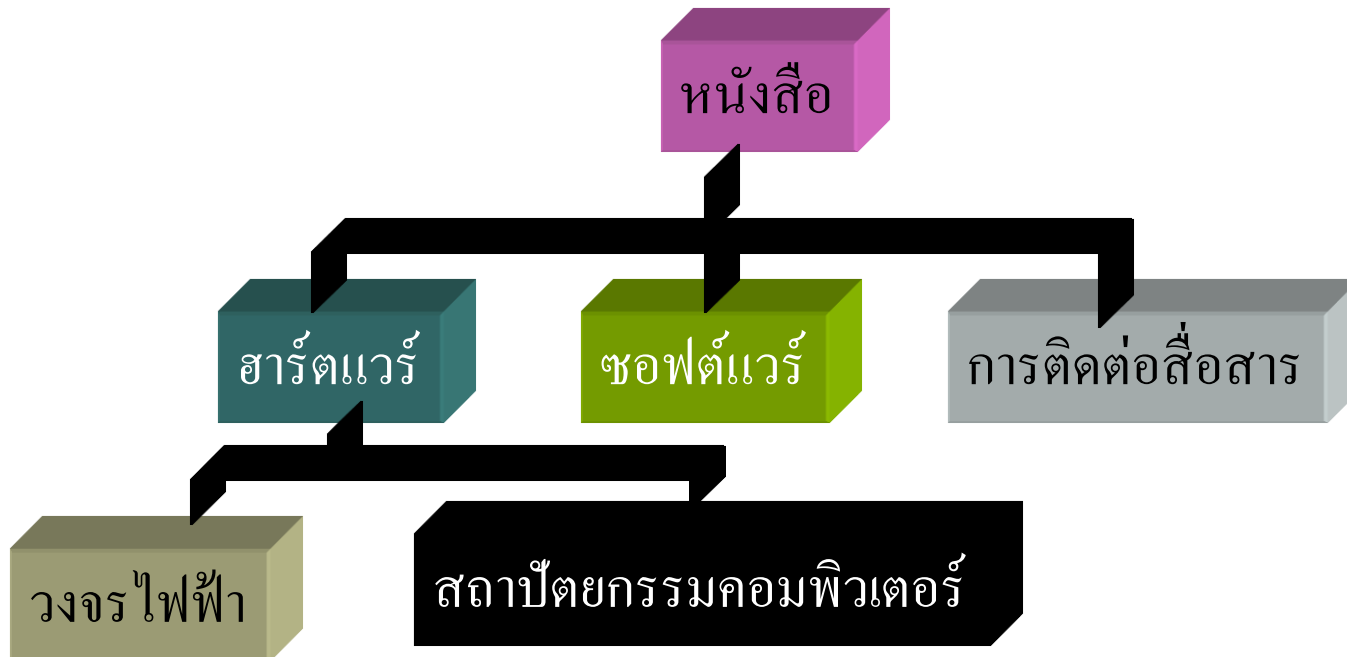
Ordered Classification

เป็นการจัดแบ่งกลุ่มแบบมีการจัดลำดับของคลาส เช่น การจัดแบ่งกลุ่มเป็นลำดับชั้น (Hierarchical) ซึ่งคลาสที่อยู่ในระดับล่างจะอยู่ภายใต้ขอบข่ายของคลาสที่อยู่ระดับบน

Unordered Classification

เป็นการแบ่งกลุ่มแบบไม่มีลำดับ เช่น การสร้าง **Thesaurus** แบบอัตโนมัติ ซึ่งเป็นพจนานุกรมข้อมูลที่ทำให้วิธีการค้นหาข้อมูลมีประสิทธิภาพ

ความสัมพันธ์ระหว่างคลาสกับคลาส



Ordered Classification

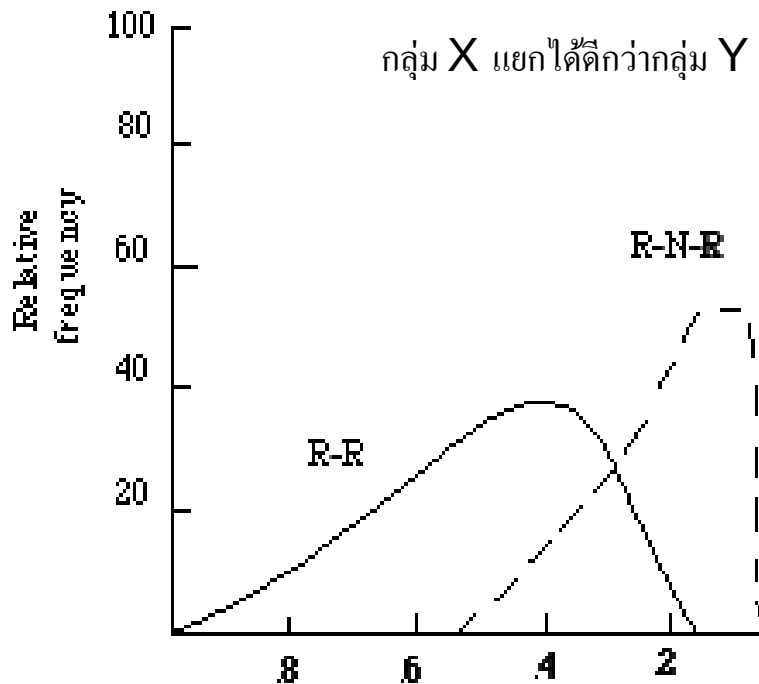
ข้อสมมุติฐานของคลัสเตอร์ (Cluster Hypothesis)

- “เอกสารที่เกี่ยวข้องกันอย่างใกล้ชิดมีแนวโน้มที่จะมีการเกี่ยวข้องกับคำถามที่ร้องขอเหมือนกัน “
- “**closely associated documents tend to be relevant to the same requests “**

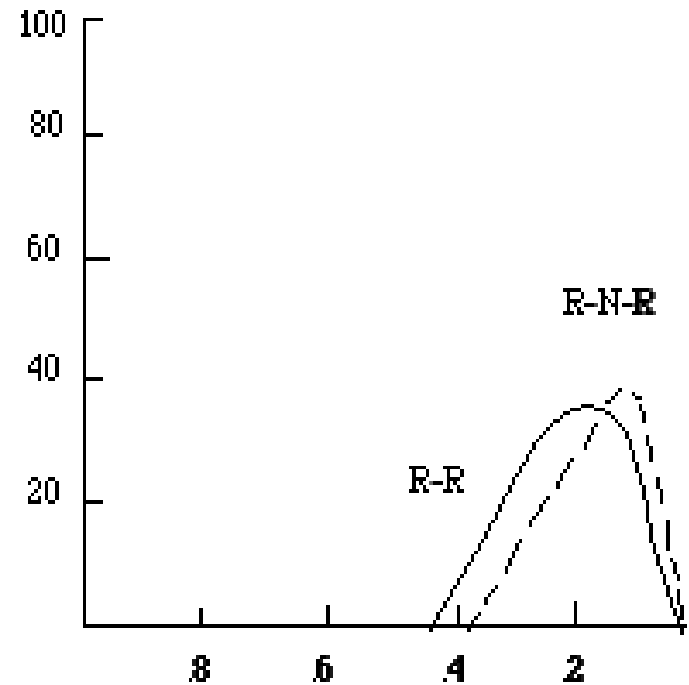
รูป ■ แสดงถึงผลรวมของกลุ่มข้อเรียกร้อง

relevant-relevant(R-R) และ relevant-non-relevant(R-N-R)

Collection X



Collection Y



เอกสารที่เกี่ยวข้อง(relevant) จะคล้ายคลึงกันและกัน มากกว่าจะคล้ายกับเอกสารที่ไม่เกี่ยวข้อง(Non-relevant)

ทำไมต้องจัดกลุ่มเอกสาร (**document clustering**)

- ช่วยให้การดึงข้อมูลได้ประสิทธิภาพมากกว่าการค้นหาตามลำดับ เพราะวิธีการค้นหาตามลำดับ จะไม่สนใจความสัมพันธ์ระหว่างเอกสาร แต่คลัสเตอร์สามารถจัดโครงสร้างของกลุ่มโดยให้เอกสารที่คล้ายคลึงกันมากอยู่ในคลาสเดียวกัน ทำให้สามารถเพิ่มความเร็วในการดึงข้อมูล ทำให้มีประสิทธิภาพมากขึ้น

clustering

- กระบวนการที่รวมออปเจกต์ที่มีลักษณะคล้ายกัน หรือเหมือนกันไว้ในกลุ่มเดียวกัน หรือกำหนดกฎเกณฑ์ขึ้นมาภายในคลัสเตอร์ ซึ่งออปเจกต์ที่อยู่ภายในกลุ่มจะมีลักษณะตรงตามที่กำหนดเอาไว้
นั่นเอง

กระบวนการคลัสเตอร์ริง(**Clustering algorithm**)

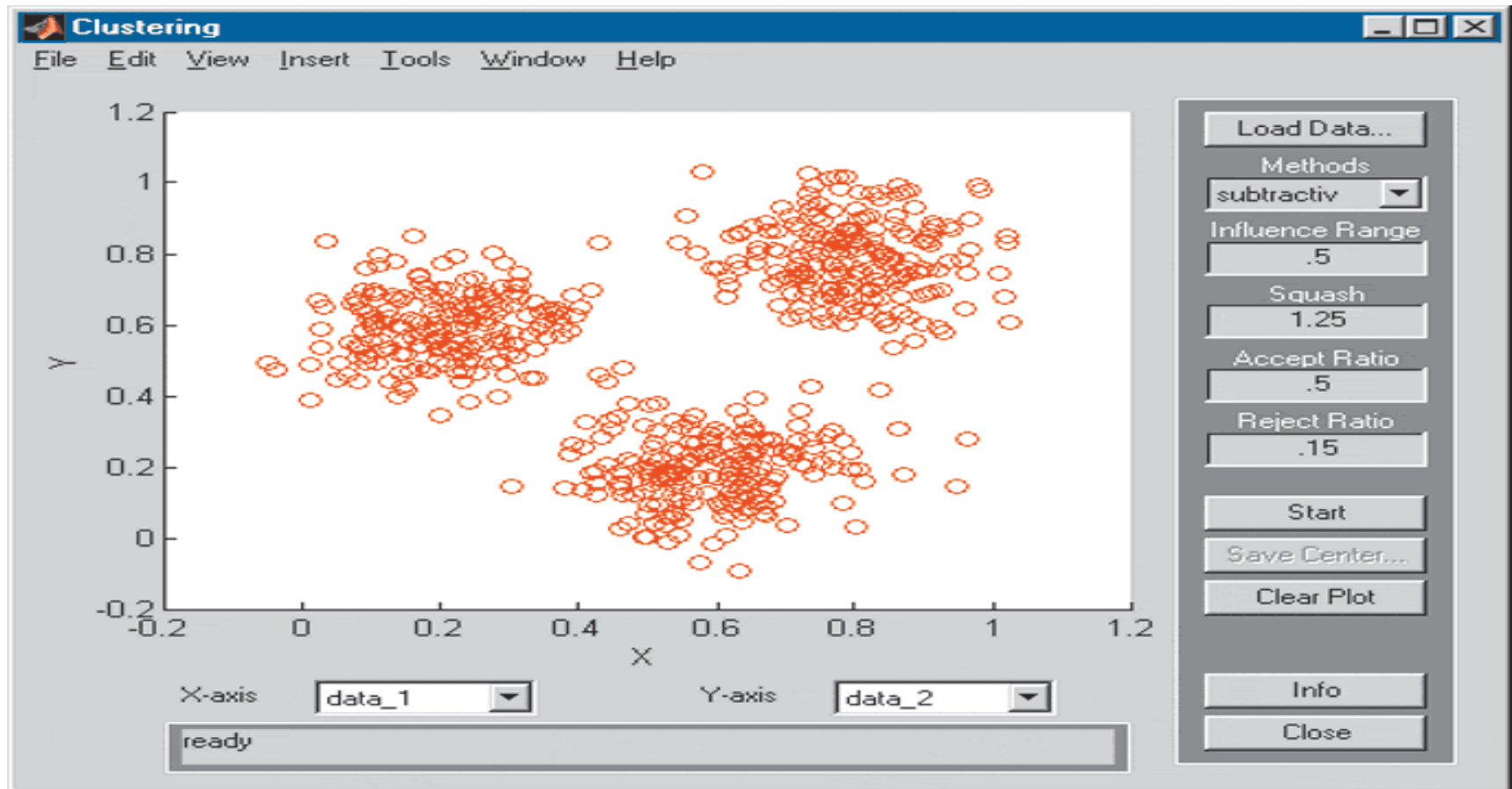
สามารถนำไปใช้ได้หลายทาง

- การจัดกลุ่มลูกค้าที่มีลักษณะของการซื้อขายที่เหมือนกันไว้ในฐานข้อมูลเดียวกัน
- ทางชีววิทยาสามารถนำลักษณะเด่นของพืชและสัตว์มาจัดกลุ่มไว้ในกลุ่มเดียวกัน
- ทางด้านบรรณารักษ์ จะนำไปใช้ในการจัดหนังสือและแยกประเภทของหนังสือ
- การประกันภัยสามารถนำไปใช้ในการแบ่งแยกกลุ่มของกรรมกรมประกันภัยของลูกค้า
- การจัดทำผังเมือง ใช้แบ่งที่ตั้งที่อาศัยตามสภาพภูมิศาสตร์
- การเฝ้าสังเกตการณ์แผ่นดินไหว ใช้ในการเฝ้าสังเกตว่าส่วนใดเป็นพื้นที่อันตรายที่เสี่ยงต่อการเกิดแผ่นดินไหว

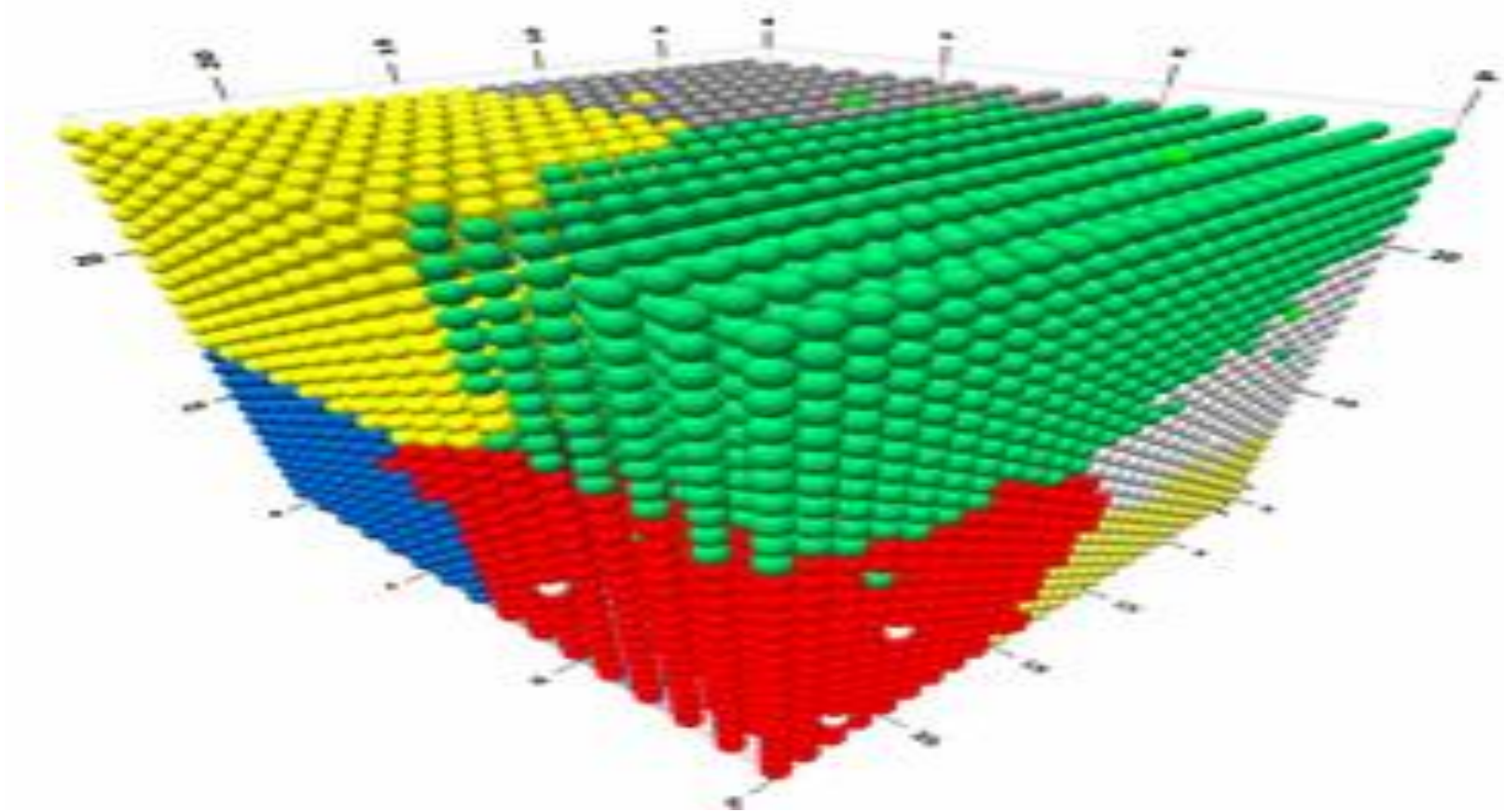
Customer Clustering



The Fuzzy Clustering GUI uses numerical data to develop classification

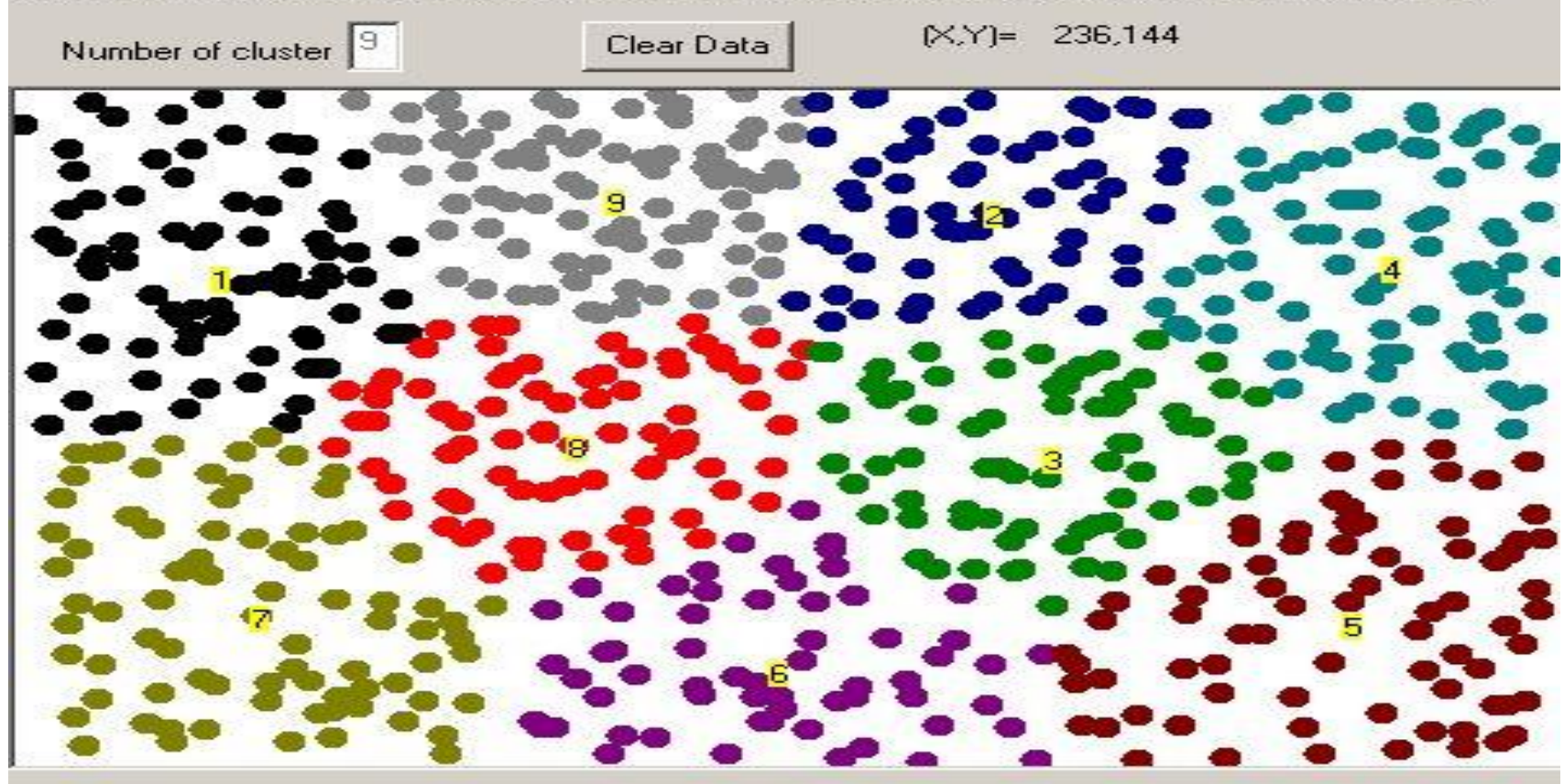


Visual Clustering

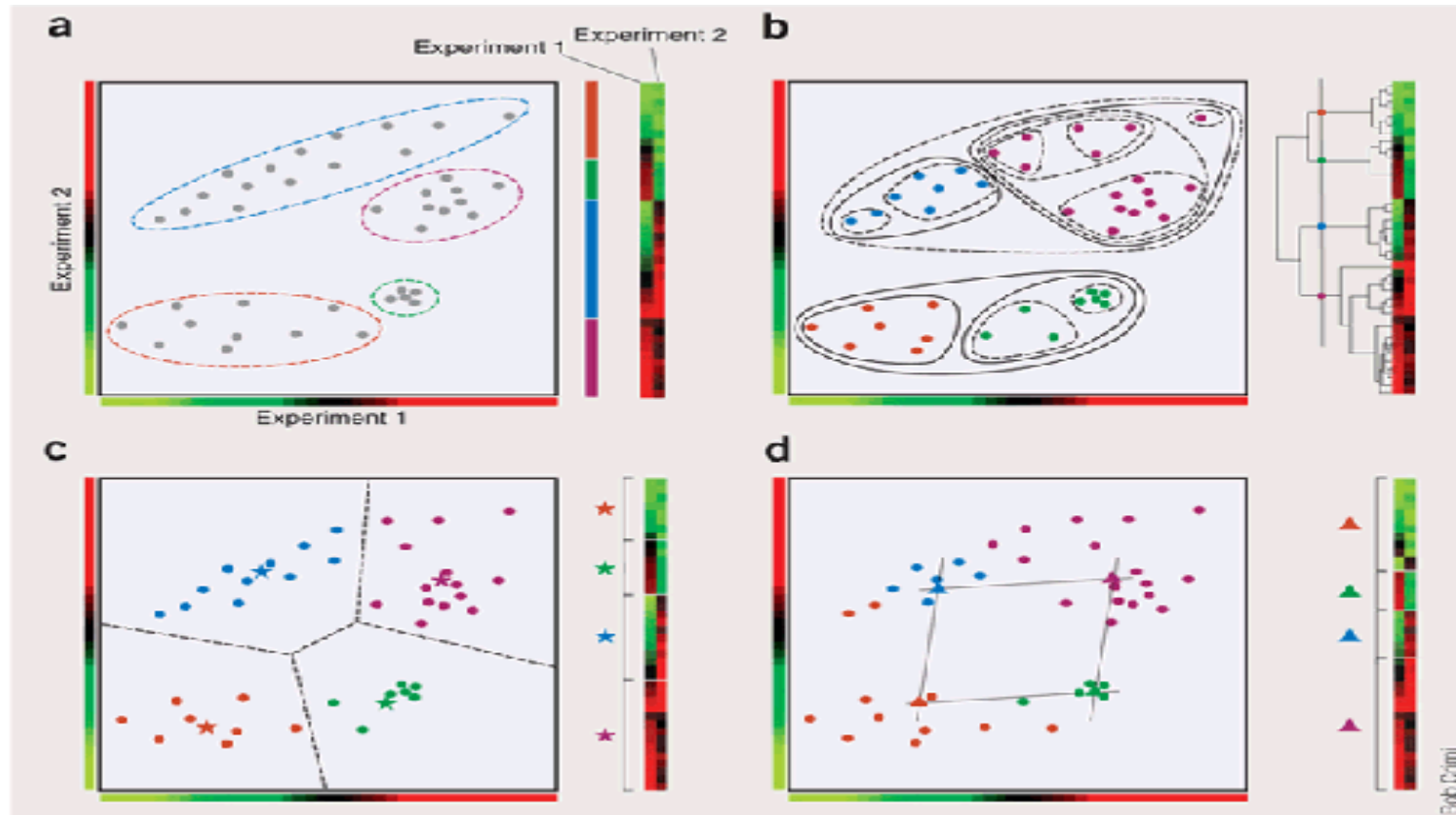


K-Mean Clustering

Click data in the picture box below. The program will automatically cluster the data by color code



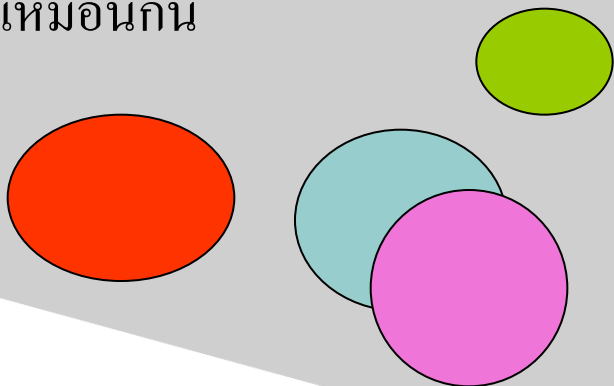
applying cluster analysis to your own gene expression data sets.



การใช้วิธีคลัสเตอร์ในระบบค้นคืนสารสนเทศ

Exclusive Clustering

เป็นการจัดกลุ่มที่ข้อมูลที่อยู่ในกลุ่มเดียวกันจะมีลักษณะเหมือนกัน



Overlapping Clustering

เป็นการจัดกลุ่มที่ข้อมูลภายในกลุ่มแบ่งออกให้อยู่ในรูปของเซตย่อยๆ แต่ละจุดอาจเป็นส่วนหนึ่งของคลัสเตอร์มากกว่าหนึ่งก็ได้และจำนวนสมาชิกภายในเซตต่างๆมีค่าแตกต่างกันได้

การใช้วิธีคลัสเตอร์ในระบบค้นคืนสารสนเทศ

Hierarchical Clustering

เป็นการจัดกลุ่มที่มีการรวม
คุณสมบัติพื้นฐานของ

Exclusive Clustering และ
Overlapping Clustering

ไว้ด้วยกันแต่มีการกำหนดเงื่อนไข
หรือกฎที่ใช้ในการกำหนดข้อมูล
ในแต่ละคลัสเตอร์

Probabilistic Clustering

เป็นการจัดกลุ่มโดยใช้วิธีการทางสถิติ

การเลือกวิธีที่เหมาะสมในการคลัสเตอร์นั้นมีหลักเกณฑ์ **2** ประการ

ความมั่นคงของวิธีการ

- (1) เมื่อมี ออกเป็เจ็ก อื่นๆเพิ่มเข้ามา คลัสเตอร์จะไม่ค่อยเปลี่ยนแปลง คือเจริญเติบโตแบบคงที่
- (2) วิธีการนั้นจะต้องมั่นคง กล่าวคือความผิดพลาดเล็กน้อยๆ ก็จะก่อให้เกิดการเปลี่ยนแปลงอย่างเล็กน้อยๆในคลัสเตอร์
- (3) วิธีการต้องเป็นอิสระไม่ขึ้นอยู่กับลำดับเริ่มแรกของออกเป็เจ็กต์

ประสิทธิภาพ

ด้านความเร็ว และ ความต้องการเนื้อที่เก็บความจำ ซึ่งเป็นความสามารถในการดึงข้อมูลที่ต้องการ และมีข้อมูลที่ไม่ต้องการมาน้อยที่สุด