

空間自相關

Spatial Autocorrelation

授課教師：溫在弘

E-mail: wenthung@ntu.edu.tw

2025 期末考相關規定

- 準備階段 (5/19) : 課堂公告題目重點的大致方向
- 第一階段 (5/26) : 下午 2:20-5:30 (take-home or R305/電腦教室)
 - 於NTUCOOL公告題目以及必須使用的資料。
 - 繳交內容 : (1). 程式碼 : R Notebook的html檔案 ,
(2). 分析成果報告 (3分鐘以內的短影音 mp4 + 簡報檔 pptx)
 - 評量目的 : 在**有限時間內**完成初步分析的能力 , 評量對於分析方法和技術的熟悉程度
- 第二階段 (6/02) : 下午 5:30 之前上傳NTUCOOL
 - 繳交內容 : (1). **更新的**程式碼 : R Notebook的html檔案 ,
(2). 分析成果報告 (**更新的**3分鐘以內的短影音 mp4 + 書面報告 pdf)
 - 評量目的 : 對於問題思考的深度與面向討論的廣泛程度

學期成績評定 A + 的條件 (Beyond Expectation)

- 學期成績若已評定為A，以下條件之一可能調整成A +
 - 5/5 複習小考(紙筆測驗)的成績表現
 - 期末考的成果表現：
 - 第一階段：(至多)挑選一組
 - 第二階段：(至多)挑選一組

本週課程：空間自相關的觀念與計算

Moran's I coefficient

$$I = \frac{N}{W} \frac{\sum_i \sum_j w_{ij} (x_i - \bar{x})(x_j - \bar{x})}{\sum_i (x_i - \bar{x})^2}$$

N : no. of spatial units

$w_{i,j}$: a matrix of spatial weights

$$W = \sum_{i=1}^n \sum_{j=1}^n w_{i,j} \quad (\text{sum of all } w_{i,j})$$

教科書的重點研讀章節內容

TEXT_Spatial_Weights.pdf

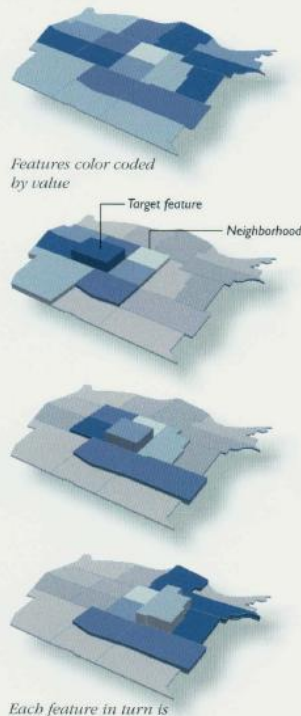
TEXT_Pattern.of.Feature.Value.pdf

Defining spatial neighborhoods and weights

Several of the methods discussed in chapters 3 and 4 show you how to analyze patterns and clusters of feature values. These methods look at both the difference between the values of features and the spatial relationship between the features (distance or other measure).

Specifically, the GIS compares the value of a feature (the “target”) to the values of neighboring features. It then moves to the next feature and does the same thing, and so on, for all the features in the study area. In order to do this, the GIS requires that you define the area surrounding each target feature within which feature values are compared—termed the “neighborhood”—and the nature of the spatial relationship between features. The GIS then assigns weights to each feature pair to specify whether the two features are in each other’s neighborhoods, and to represent the spatial relationship between the features.

You define the neighborhood based on the interaction between features. Features might influence each other—for example, the value



MEASURING THE SPATIAL PATTERN OF FEATURE VALUES

In addition to measuring the pattern formed by the locations of features, you can also measure patterns of attribute values associated with features, such as the pattern formed by median house values. These methods reveal whether similar values tend to occur near each other, or whether high and low values are interspersed.



Median house value by census tract.

The idea behind measuring patterns of feature values

Measuring the spatial pattern of feature values is based on the notion that things near each other are more alike than things far apart, an idea often attributed to geographer Waldo Tobler. The idea is consistent with our

Global and Local Measures

■ ***Global*** Measures

- ❑ A single value which applies to the entire data set
 - The same pattern occurs over the entire geographic area
 - An average for the entire area

■ ***Local*** Measures

- ❑ A value calculated for each observation unit
 - Different patterns or processes may occur in different parts of the region
 - A unique number for each location

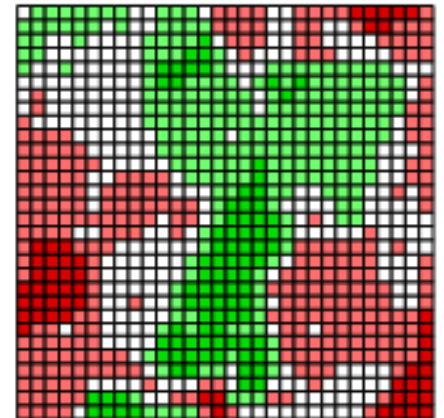
Global Analysis Methods 全域分析的方法

■ Point data *without* attributes

- Quantrat Analysis
- Nearest Neighbor Methods
 - K-order Nearest Neighborhood Analysis (NNA), G and F functions
- Ripley's K-function: $K(d)$ and $L(d)$

■ Point/Polygon data *with* attributes

- Definition of Neighborhoods or Spatial Structures
- Spatial Autocorrelation Index
 - **Moran's I** and Geary's C Ratio
- Spatial Concentration Index
 - **General G-statistic**

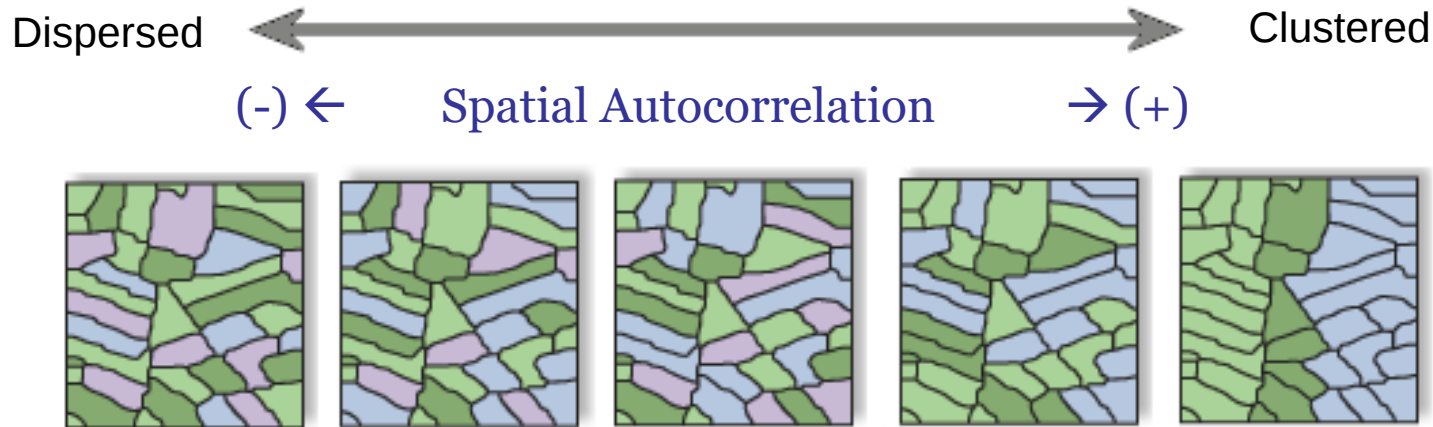


空間關連性與相依性

Spatial Relationship and Dependency

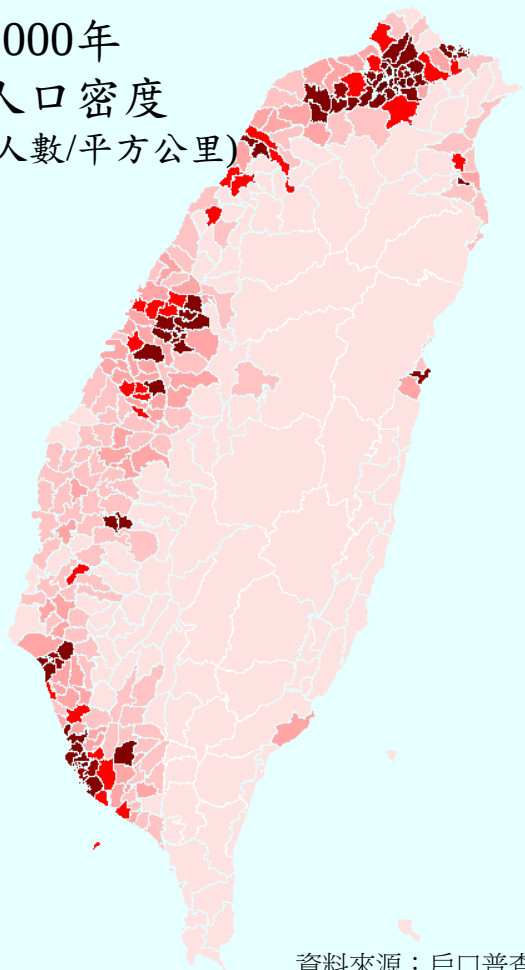
Tobler's *First Law of Geography* (1970) :

Everything is related to everything else,
but near things are more related than distant things.



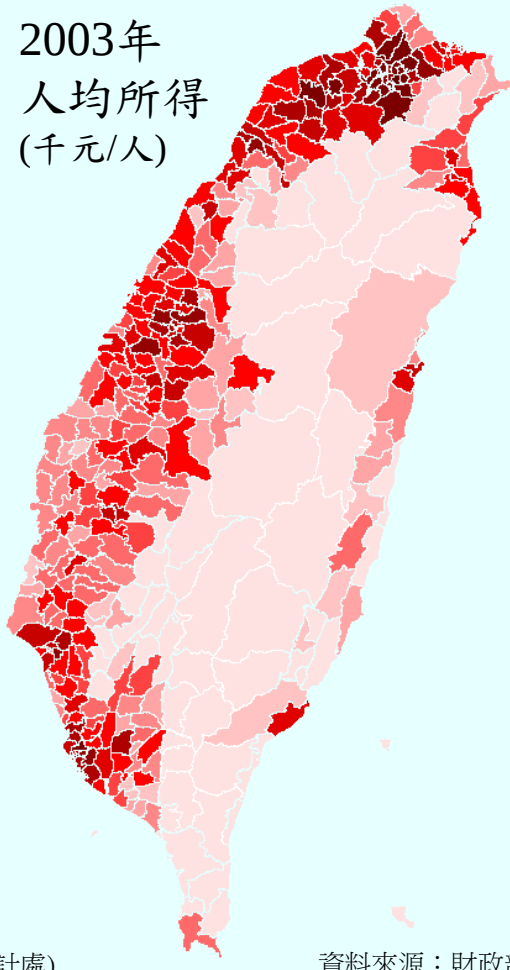
無所不在的空間自相關（空間相依的特性）

2000年
人口密度
(人數/平方公里)



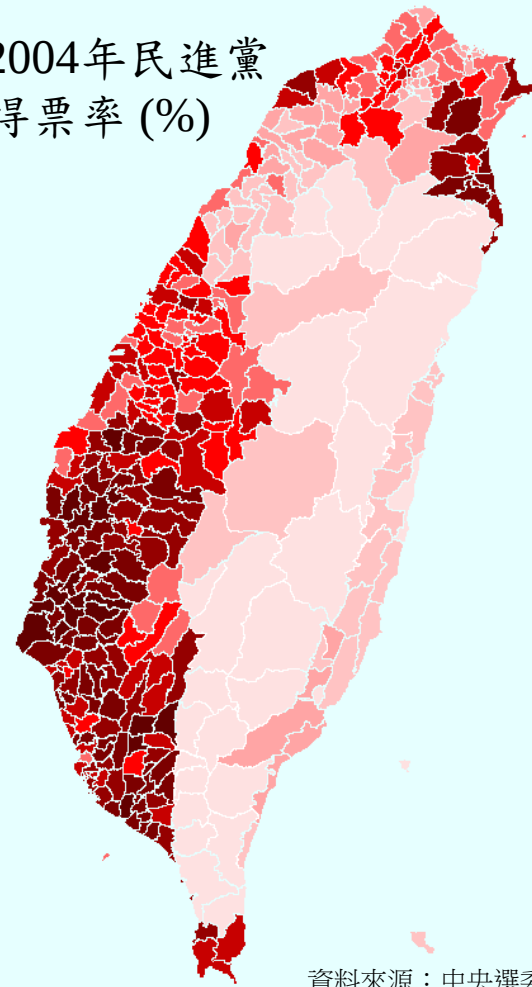
資料來源：戶口普查(主計處)

2003年
人均所得
(千元/人)



資料來源：財政部

2004年民進黨
得票率(%)



資料來源：中央選委會

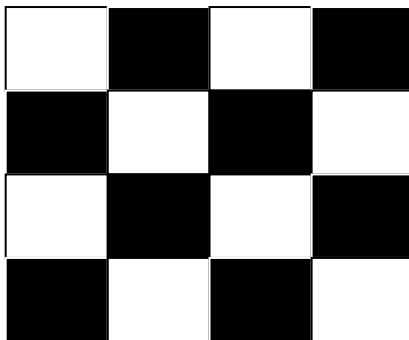
Definition of Spatial Autocorrelation

- Measurement of the *similarity of attributes* among spatial units within their *neighborhood*.

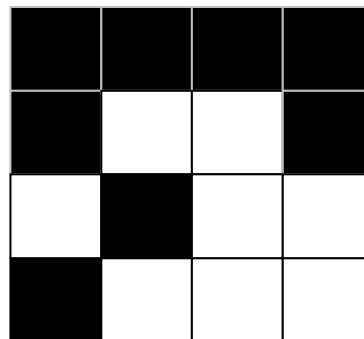
(-) ←

Spatial Autocorrelation

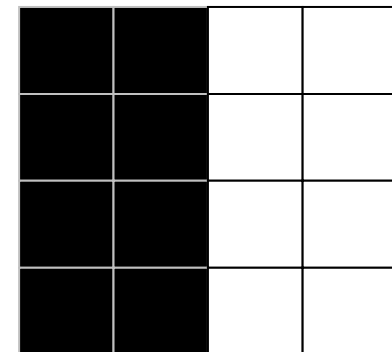
→ (+)



Dispersed



Random



Clustered

How we define the Neighborhood ?

1. Spatial adjacency
 - ❑ Physically contacted with each others
 2. Distances between the centroids
-

1. Spatial Adjacency

■ Rook's

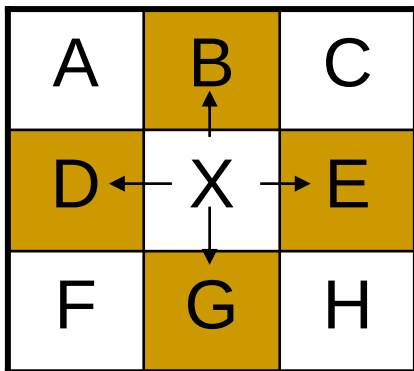
- Units that share common boundary with length greater than zero

■ Queen's

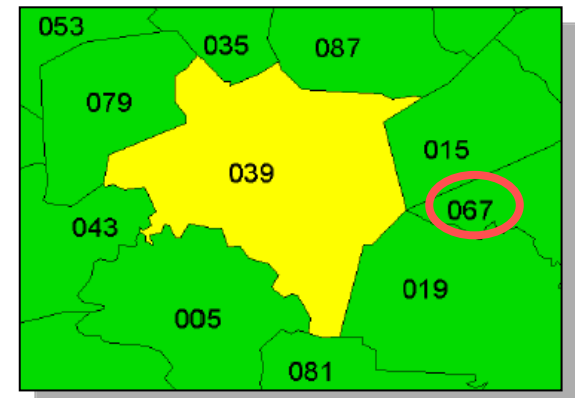
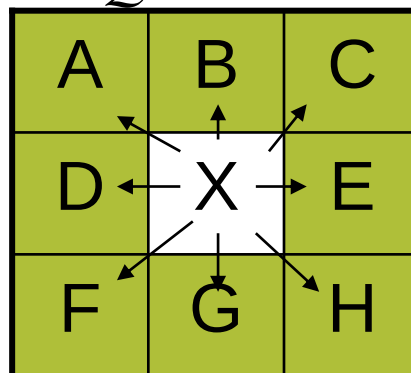
- Units that have common vertex are also included

(e.g. unit 067 in figure below)

Rook's

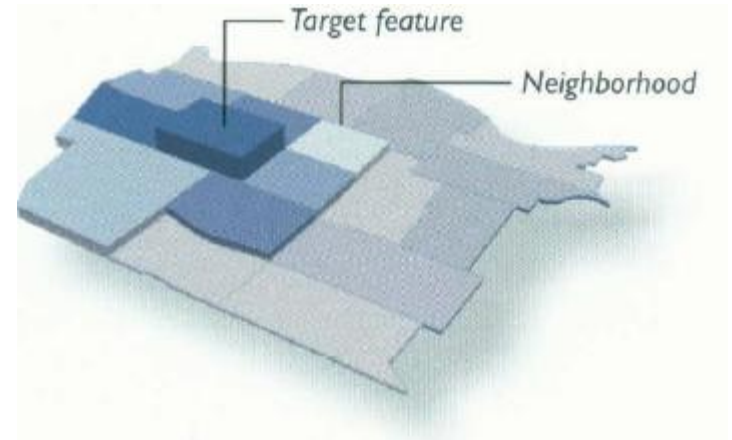


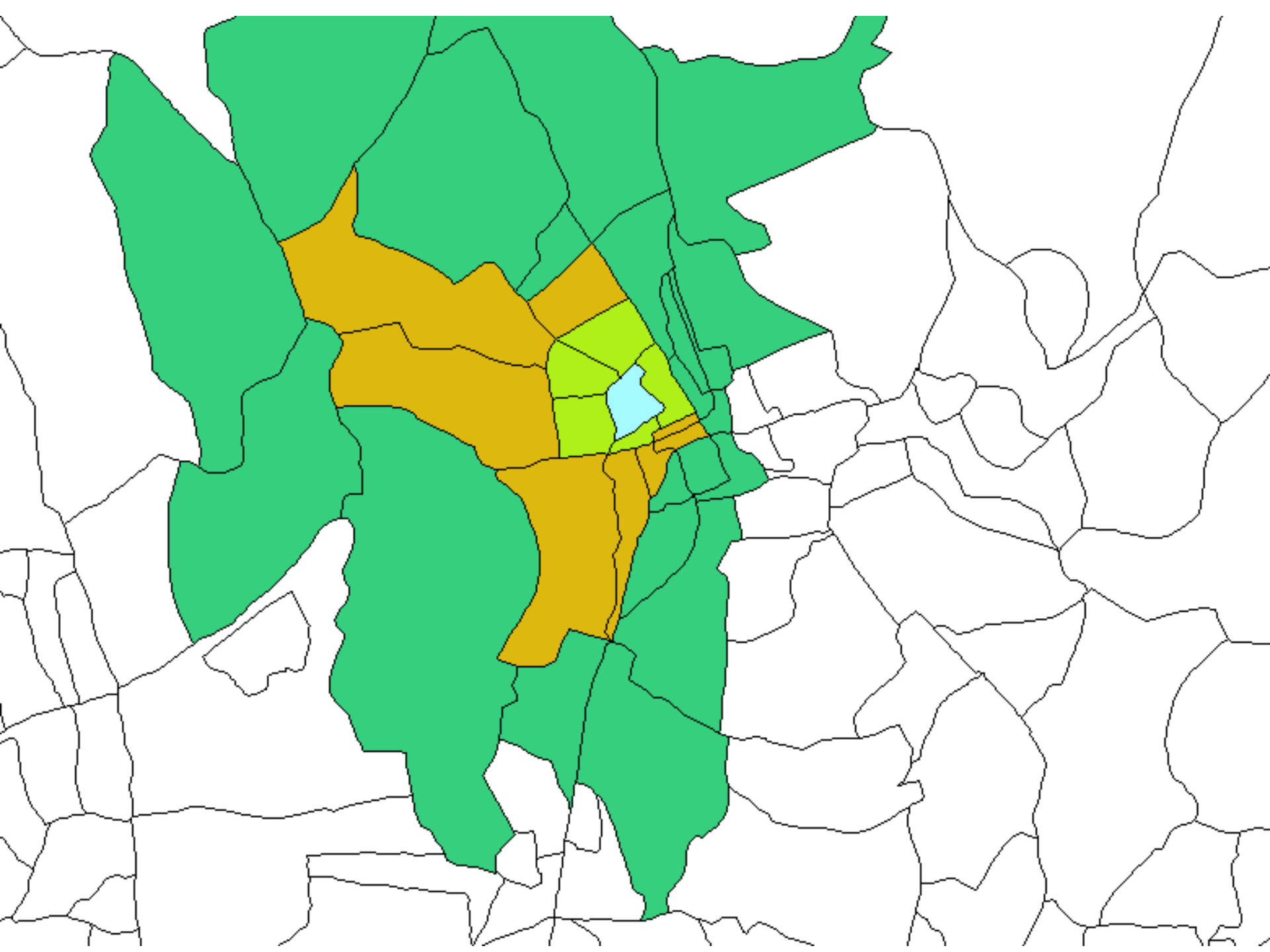
Queen's



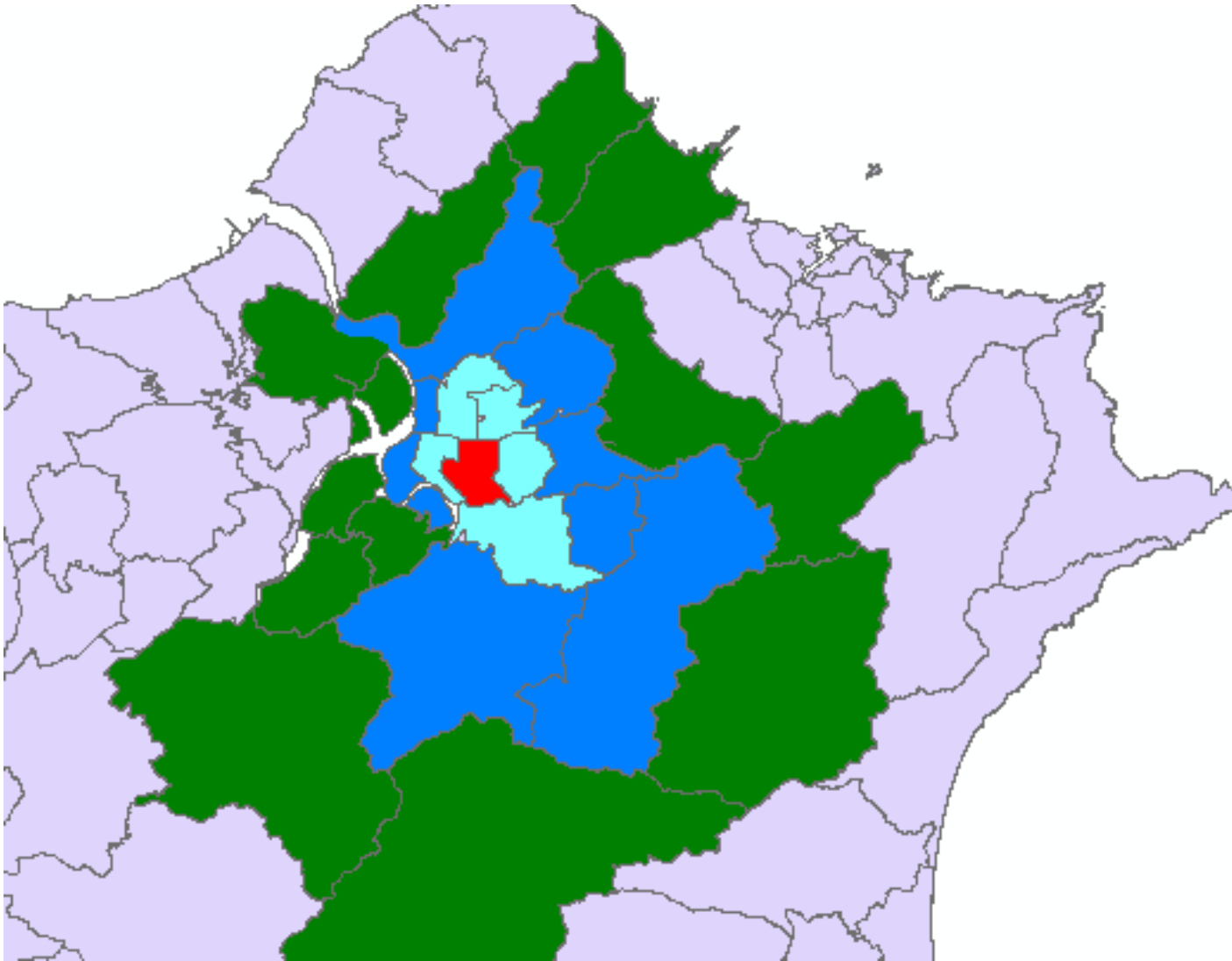
Order of Neighborhood

- 1st order
 - Immediate neighbor
 - Defined by Rook's or Queen's criteria
- 2nd order
- Higher order





台北市大安區的三階鄰近鄉鎮



Binary Connectivity Matrix

- Symmetrical $C_{ij} = C_{ji}$
- Values on diagonal are zeros
- Row sum $C_i = \sum C_{ij}$
 - The number of neighbors of unit i
- Same as connectivity matrix for network
- Not efficient for large numbers of objects
 - Redundant storage
 - Mostly zeros
 - Another way : **Sparse Matrix (稀疏矩陣)**

Sparse matrix

ID	1	2	3	4	5	6	7	8	9	10	11	12	13	14	Sum
長安里	文化里	中心里	中庸里	豐年里	溫泉里	大同里	清江里								7
中心里	泉源里	開明里	林泉里	中庸里	長名里	溫泉里									6
中庸里	中和里	開明里	文化里	智仁里	中心里	長名里									6
開明里	中和里	泉源里	智仁里	中心里	中庸里										5
文化里	智仁里	中庸里	豐年里	長名里	大同里										5
大同里	文化里	豐年里	長名里	清江里	中央里										5
中和里	泉源里	開明里	智仁里	中庸里											4
泉源里	中和里	開明里	林泉里	中心里											4
智仁里	中和里	開明里	文化里	中庸里											4
溫泉里	林泉里	中心里	長名里	清江里											4
清江里	長名里	溫泉里	大同里	中央里											4
林泉里	泉源里	中心里	溫泉里												3
豐年里	文化里	大同里													3
中央里	大同里	清江里													2

Binary Connectivity Matrix

ID	中和里	泉源里	開明里	文化里	智仁里	林泉里	中心里	中庸里	豐年里	長名里	溫泉里	大同里	清江里	中央里	Sum
中和里	0	1	1	0	1	0	0	1	0	0	0	0	0	0	4
泉源里	1	0	1	0	0	1	1	0	0	0	0	0	0	0	4
開明里	1	1	0	0	1	0	1	1	0	0	0	0	0	0	5
文化里	0	0	0	0	1	0	0	1	1	1	0	1	0	0	5
智仁里	1	0	1	1	0	0	0	1	0	0	0	0	0	0	4
林泉里	0	1	0	0	0	0	1	0	0	0	1	0	0	0	3
中心里	0	1	1	0	0	1	0	1	0	1	1	0	0	0	6
中庸里	1	0	1	1	1	0	1	0	0	1	0	0	0	0	6
豐年里	0	0	0	1	0	0	0	0	0	1	0	1	0	0	3
長安里	0	0	0	1	0	0	1	1	1	0	1	1	1	0	7
溫泉里	0	0	0	0	0	1	1	0	0	1	0	0	1	0	4
大同里	0	0	0	1	0	0	0	0	1	1	0	0	1	1	5
清江里	0	0	0	0	0	0	0	0	0	1	1	1	0	1	4
中央里	0	0	0	0	0	0	0	0	0	0	0	1	1	0	2

Sparse matrix 稀疏矩陣的儲存方式

Index

0	0	0	0	0	0	0
1	0	3	0	0	0	0
2	0	0	0	6	0	0
3	0	0	9	0	0	0
4	0	0	0	0	12	0

Sparse Matrix

5	6	4
1	1	3
2	3	6
3	2	9
4	4	12

Line #1: 這個矩陣是5X6矩陣，非零元素有4個

Line #2: 記錄其位置的列索引、行索引與儲存值

Stochastic Matrix

- Equally weighted for neighbors
 - $W_{ij} = C_{ij} / C_i$
- 又稱做 Row-standardized matrix (列標準化矩陣)
 - 考慮每一個相鄰的object 的影響量

Binary Connectivity Matrix

ID	中和里	泉源里	開明里	文化里	智仁里	林泉里	中心里	中庸里	豐年里	長名里	溫泉里	大同里	清江里	中央里	Sum
中和里	0	1	1	0	1	0	0	1	0	0	0	0	0	0	4
泉源里	1	0	1	0	0	1	1	0	0	0	0	0	0	0	4
開明里	1	1	0	0	1	0	1	1	0	0	0	0	0	0	5
文化里	0	0	0	0	1	0	0	1	1	1	0	1	0	0	5
智仁里	1	0	1	1	0	0	0	1	0	0	0	0	0	0	4
林泉里	0	1	0	0	0	0	1	0	0	0	1	0	0	0	3
中心里	0	1	1	0	0	1	0	1	0	1	1	0	0	0	6
中庸里	1	0	1	1	1	0	1	0	0	1	0	0	0	0	6
豐年里	0	0	0	1	0	0	0	0	0	1	0	1	0	0	3
長安里	0	0	0	1	0	0	1	1	1	0	1	1	1	0	7
溫泉里	0	0	0	0	0	1	1	0	0	1	0	0	1	0	4
大同里	0	0	0	1	0	0	0	0	1	1	0	0	1	1	5
清江里	0	0	0	0	0	0	0	0	0	1	1	1	0	1	4
中央里	0	0	0	0	0	0	0	0	0	0	0	1	1	0	2

Stochastic Weighted Matrix

ID	中和里	泉源里	開明里	文化里	智仁里	林泉里	中心里	中庸里	豐年里	長名里	溫泉里	大同里	清江里	中央里
中和里	0.00	0.25	0.25	0.00	0.25	0.00	0.00	0.25	0.00	0.00	0.00	0.00	0.00	0.00
泉源里	0.25	0.00	0.25	0.00	0.00	0.25	0.25	0.00	0.00	0.00	0.00	0.00	0.00	0.00
開明里	0.20	0.20	0.00	0.00	0.20	0.00	0.20	0.20	0.00	0.00	0.00	0.00	0.00	0.00
文化里	0.00	0.00	0.00	0.00	0.20	0.00	0.00	0.20	0.20	0.20	0.00	0.20	0.00	0.00
智仁里	0.25	0.00	0.25	0.25	0.00	0.00	0.00	0.25	0.00	0.00	0.00	0.00	0.00	0.00
林泉里	0.00	0.33	0.00	0.00	0.00	0.00	0.33	0.00	0.00	0.00	0.33	0.00	0.00	0.00
中心里	0.00	0.17	0.17	0.00	0.00	0.17	0.00	0.17	0.00	0.17	0.17	0.00	0.00	0.00
中庸里	0.17	0.00	0.17	0.17	0.17	0.00	0.17	0.00	0.00	0.17	0.00	0.00	0.00	0.00
豐年里	0.00	0.00	0.00	0.33	0.00	0.00	0.00	0.00	0.00	0.33	0.00	0.33	0.00	0.00
長安里	0.00	0.00	0.00	0.14	0.00	0.00	0.14	0.14	0.14	0.00	0.14	0.14	0.14	0.00
溫泉里	0.00	0.00	0.00	0.00	0.00	0.25	0.25	0.00	0.00	0.25	0.00	0.00	0.25	0.00
大同里	0.00	0.00	0.00	0.20	0.00	0.00	0.00	0.00	0.20	0.20	0.00	0.00	0.20	0.20
清江里	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.25	0.25	0.25	0.00	0.25
中央里	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.50	0.50	0.00

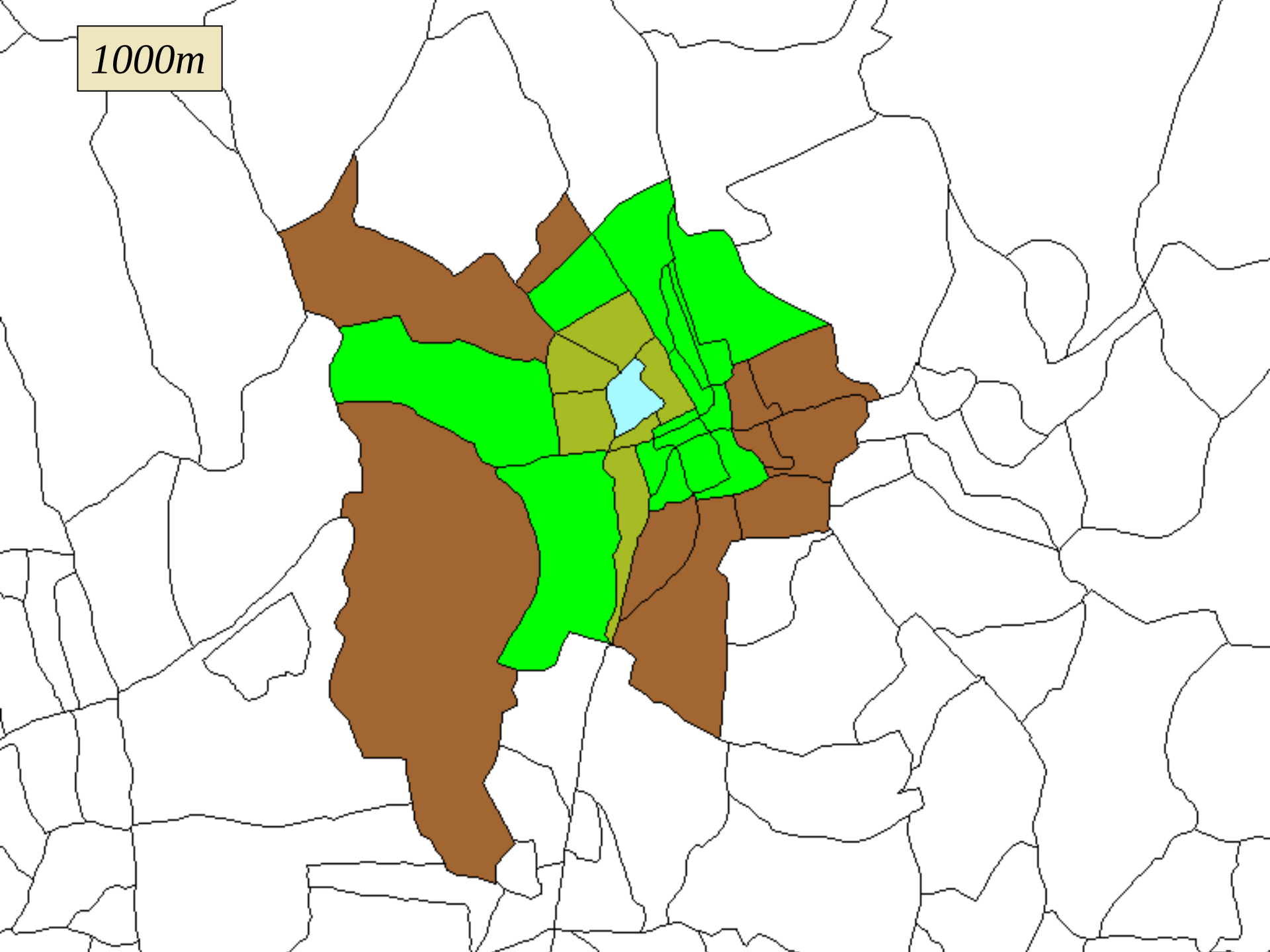
2. Distances

- Distance decay

- $W_{ij} = 1 / d_{ij}$

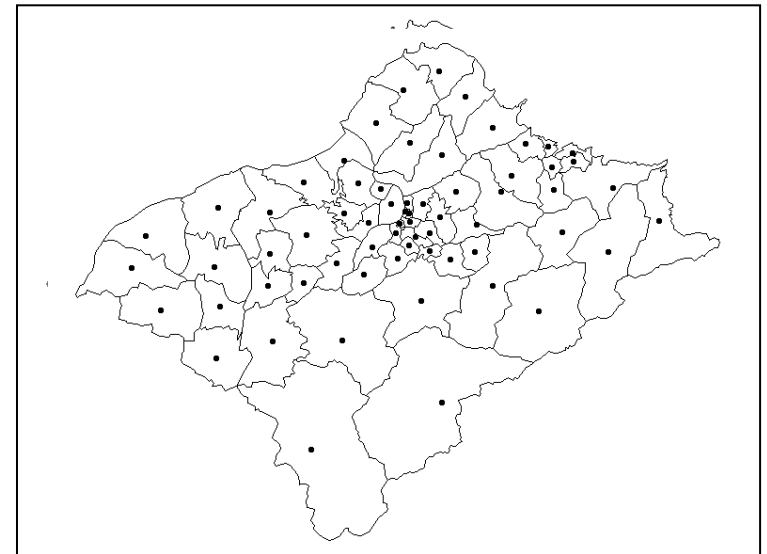
- $W_{ij} = 1 / d_{ij}^2$

1000m



Centroid Distances

- Distance between centroids
- Centroid : geometric center of the polygon
 - Affected by the shape of the polygon
 - May be located outside the polygon



Spatially Weighted Matrix

Using Centroid Distances



ID	中和里	泉源里	開明里	文化里	智仁里	林泉里	中心里	中庸里	豐年里	長安里	溫泉里	大同里	清江里	中央里
中和里	0	1359	1602	2229	1948	2334	1913	2142	3149	2530	2763	3169	3238	3375
泉源里	1359	0	2033	2966	2657	1737	2002	2602	3716	2824	2764	3618	3337	3758
開明里	1602	2033	0	974	684	1477	506	585	1692	930	1224	1633	1645	1812
文化里	2229	2966	974	0	309	2335	1309	607	952	932	1519	1073	1607	1316
智仁里	1948	2657	684	309	0	2103	1075	478	1204	903	1444	1265	1644	1496
林泉里	2334	1737	1477	2335	2103	0	1029	1747	2683	1678	1310	2461	1869	2508
中心里	1913	2002	506	1309	1075	1029	0	741	1790	829	856	1640	1381	1762
中庸里	2142	2602	585	607	478	1747	741	0	1115	450	966	1049	1209	1238
豐年里	3149	3716	1692	952	1204	2683	1790	1115	0	1005	1480	324	1203	522
長安里	2530	2824	930	932	903	1678	829	450	1005	0	595	816	760	935
溫泉里	2763	2764	1224	1519	1444	1310	856	966	1480	595	0	1208	581	1214
大同里	3169	3618	1633	1073	1265	2461	1640	1049	324	816	1208	0	882	247
清江里	3238	3337	1645	1607	1644	1869	1381	1209	1203	760	581	882	0	782
中央里	3375	3758	1812	1316	1496	2508	1762	1238	522	935	1214	247	782	0

Nearest Distances (較不常用)

- 兩個Polygon之間最短的點的距離
 - Will be **zero** for adjacent polygons

ID	中和里	泉源里	開明里	文化里	智仁里	林泉里	中心里	中庸里	豐年里	長安里	溫泉里	大同里	清江里	中央里
中和里	0	0	0	51	0	837	403	0	606	464	723	615	934	1075
泉源里	0	0	0	920	849	0	0	686	1236	731	653	1122	1095	1301
開明里	0	0	0	148	0	319	0	0	529	166	304	477	572	773
文化里	51	920	148	0	0	878	299	0	0	0	567	0	549	481
智仁里	0	849	0	0	0	879	322	0	405	260	600	415	753	870
林泉里	837	0	319	878	879	0	0	535	986	234	0	609	463	691
中心里	403	0	0	299	322	0	0	0	471	0	0	269	269	477
中庸里	0	686	0	0	0	535	0	0	89	0	209	89	407	558
豐年里	606	1236	529	0	405	986	471	89	0	0	583	0	547	67
長安里	464	731	166	0	260	234	0	0	0	0	0	0	0	169
溫泉里	723	653	304	567	600	0	0	209	583	0	0	70	0	179
大同里	615	1122	477	0	415	609	269	89	0	0	70	0	0	0
清江里	934	1095	572	549	753	463	269	407	547	0	0	0	0	0
中央里	1075	1301	773	481	870	691	477	558	67	169	179	0	0	0

Spatially Weighted Matrix

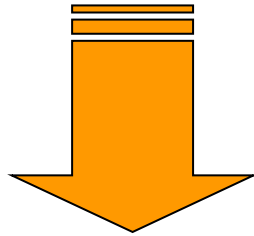
Using Nearest Distances



Spatial Weights Matrix Approaches

*Neighborhood
Definition*

- *Rook's Definition*
- *Queen's Definition*



*Spatial Weights
Matrix*

- *Binary Connective Matrix*
- *Stochastic or Row Standardized Weights Matrix*
- *Centroid Distances*
- *Nearest Distances*

教科書的研讀教材（列入考試範圍）

TEXT_Spatial_Weights.pdf

Defining spatial neighborhoods and weights

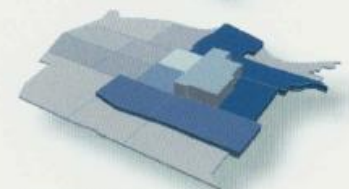
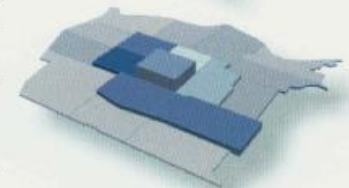
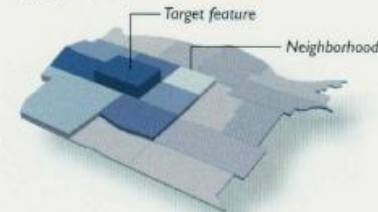
Several of the methods discussed in chapters 3 and 4 show you how to analyze patterns and clusters of feature values. These methods look at both the difference between the values of features and the spatial relationship between the features (distance or other measure).

Specifically, the GIS compares the value of a feature (the “target”) to the values of neighboring features. It then moves to the next feature and does the same thing, and so on, for all the features in the study area. In order to do this, the GIS requires that you define the area surrounding each target feature within which feature values are compared—termed the “neighborhood”—and the nature of the spatial relationship between features. The GIS then assigns weights to each feature pair to specify whether the two features are in each other’s neighborhoods, and to represent the spatial relationship between the features.

You define the neighborhood based on the interaction between features. Features might influence each other—for example, the value



Features color coded by value



Each feature in turn is

Measuring Spatial Autocorrelation

- Spatial weighting W_{ij}
 - Contiguity [binary or row-standardized]
 - Common Border
 - Distance [centroids or nearest]
 - Distance band
 - K^{th} -nearest neighbors

1. Index of Spatial Autocorrelation: Moran's I

$$I = \frac{N}{W} \frac{\sum_i \sum_j w_{ij} (x_i - \bar{x})(x_j - \bar{x})}{\sum_i (x_i - \bar{x})^2}$$

where

N is the number of cases

$\bar{x}$ is the mean of the variable

x_i is the variable value at a particular location i

w_{ij} is a **spatial weight indexing** location of i relative to j

$$W = \sum_{i=1}^n \sum_{j=1}^n w_{i,j} \quad (\text{sum of all } w_{i,j})$$

- Applied to a **continuous variable** for polygons or points

這個公式是怎麼想出來的？

皮爾森相關係數 Pearson's correlation coefficient (r):

共變異數 $\text{cov}(X,Y)$ 的觀念

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

$$\frac{\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y}) / n}{\sqrt{\sum_{i=1}^n (x_i - \bar{x})^2 / n} \sqrt{\sum_{i=1}^n (y_i - \bar{y})^2 / n}}$$

這個公式是怎麼想出來的？[續]

$$\frac{\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y}) / n}{\sqrt{\sum_{i=1}^n (x_i - \bar{x})^2 / n} \sqrt{\sum_{i=1}^n (y_i - \bar{y})^2 / n}}$$

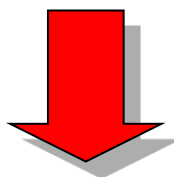
Neighboring values of variable y replace those of x

把 x_i 替換成 y_i 的鄰居 ($c_{ij} \cdot y_i$)

$$\frac{\sum_{i=1}^n \left[\sum_{j=1}^n c_{ij} (y_i - \bar{y}) \right] (y_j - \bar{y}) / \sum_{i=1}^n \sum_{j=1}^n c_{ij}}{\sqrt{\sum_{i=1}^n (y_i - \bar{y})^2 / n} \sqrt{\sum_{i=1}^n (y_i - \bar{y})^2 / n}}$$

這個公式是怎麼想出來的？[續]

$$\frac{\sum_{i=1}^n \sum_{j=1}^n c_{ij} (y_i - \bar{y})(y_j - \bar{y}) / \sum_{i=1}^n \sum_{j=1}^n c_{ij}}{\sqrt{\sum_{i=1}^n (y_i - \bar{y})^2 / n} \sqrt{\sum_{i=1}^n (y_i - \bar{y})^2 / n}}$$



$$I = \frac{N}{W} \frac{\sum_i \sum_j w_{ij} (x_i - \bar{x})(x_j - \bar{x})}{\sum_i (x_i - \bar{x})^2}$$

Spatial Autocorrelation: Moran's I Statistic

$$I = \frac{n \sum_{i=1}^n \sum_{j=1}^n w_{ij} (y_i - \bar{y})(y_j - \bar{y})}{\left(\sum_{i=1}^n (y_i - \bar{y})^2 \right) \left(\sum_{i \neq j} \sum w_{ij} \right)}$$

Product of the deviation from the mean for all pairs of adjacent regions ($w_{ij}=1$)

Essentially a measure of variance across the regions

Sum of the weights (count of all adjacent pairs)

- ▶ n = number of regions
- ▶ w_{ij} = measure of spatial proximity between region i and j

Moran's I Interpretations

- Similar to correlation coefficient, range between ± 1.0
 - 0 indicates no spatial autocorrelation, approximate technically it is $-1/(n-1)$
 - Highly auto-correlated, if I is closed to 1 or -1
 - Sign of values indicate negative/positive autocorrelation
- Can be used as index for dispersion/random/cluster patterns
 - 0: random
 - Positive : more toward clustering
 - Negative: more toward dispersion/uniform

Significance Tests for Moran's I

- Z-score: $Z = (I - E(I)) / S_{\text{Err}}(I)$
 - I: Moran's I of sample
 - $E(I)$: Expected value of I ; $E(I) = -1/(n-1)$
 - S_{Err} : Standard error
 - Depend on if free or non-free sampling is used
- $\alpha = 0.05$, Critical Z value = ± 1.96
 - will be ± 1.645 for $\alpha = 0.1$
- 檢定是否為達到顯著差異
 - H_0 : 無差異 (隨機分佈)
 - At $p < 0.05$, Reject H_0 if $|Z| > 1.96$

$$E_N(I) = E_R(I) = \frac{-1}{n-1}$$

$$VAR_N(I) = \frac{(n^2 S_1 - n S_2 + 3W^2)}{W^2(n^2 - 1)} - [E_N(I)]^2$$

$$VAR_R(I) = \frac{n[(n^2 - 3n + 3)S_1 - nS_2 + 3W^2]}{(n-1)(n-2)(n-3)W^2} - \frac{k[(n^2 - n)S_1 - nS_2 + 3W^2]}{(n-1)(n-2)(n-3)W^2} - [E_R(I)]^2,$$

where

$$W = \sum_{i=1}^n \sum_{j=1}^n w_{ij}$$

$$S_1 = \frac{\sum_{i=1}^n \sum_{j=1}^n (w_{ij} + w_{ji})^2}{2}$$

$$S_2 = \sum_{i=1}^n (w_{i.} + w_{.i})^2$$

$$k = \frac{\sum_{i=1}^n (x_i - \bar{x})^4}{\left(\sum_{i=1}^n (x_i - \bar{x})^2 \right)^2}.$$

free sampling
(with replacement.)
(normality)

non-free sampling
(without replacement)
(randomization)

Randomization vs. normality sampling

Randomization sampling assumes that the observed spatial pattern of your data represents one of many possible spatial arrangements—the number of features having a particular value is always going to be the same (based on the observed number of each), but the arrangement can change. Suppose you have the number of cases per census tract of a disease. Some tracts have several cases of the disease; many don't have any. Your null hypothesis would be that the disease strikes randomly. If you could take the case count values for your study area and scatter them on a map of the census tracts, making sure every census tract got a value, you'd create a random pattern. The randomization null hypothesis postulates that if you could perform this operation infinite times, most of the time you would produce a pattern that was not markedly different from the observed pattern. If your significance test indicates that you should reject the null hypothesis at the specified confidence level, then you know that the observed arrangement of values would significantly differ from this randomly produced pattern.

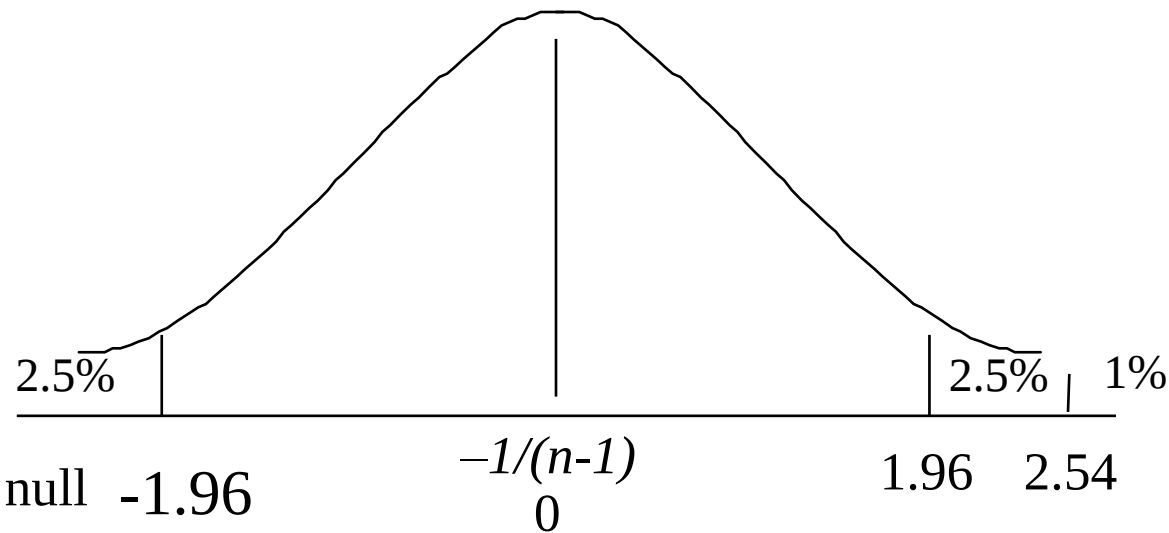
Randomization vs. normality sampling (cont'd)

Normalization sampling, in contrast, assumes that the number of cases associated with any particular census tract could be derived from an infinitely large, normally distributed population of values (through some random sampling process). Rather than scattering the observed values on the map of census tracts, you'd pick values from this hypothetical normal distribution and scatter those values on the map to create the random pattern.

Randomization vs. normality sampling (比較)

- The **normalization** null hypothesis not only assumes that your data is a sample, but also that the sample was obtained randomly and that the population from which the sample was obtained has a normal distribution of values. Every time you make an assumption about the data or the sample, you're potentially inducing error into the test.
- **Randomization makes fewer assumptions than normalization, so it's safer to use**, unless you know for sure your data matches the assumptions of normalization.

Test Statistic for Normal Frequency Distribution



Null Hypothesis: no spatial autocorrelation

*Moran's $I = 0$ * *technically* $-1/(n-1)$

Alternative Hypothesis: spatial clustering exists

*Moran's $I > 0$ (單尾檢定)

Reject *Null Hypothesis* if Z test statistic > 1.96 (or < -1.96)

---less than a 2.5% chance that, in the population, there is no spatial clustering

---97.5% confident that spatial clustering exists

Monte-Carlo Significance Test

■ Permutation test (排列検定)

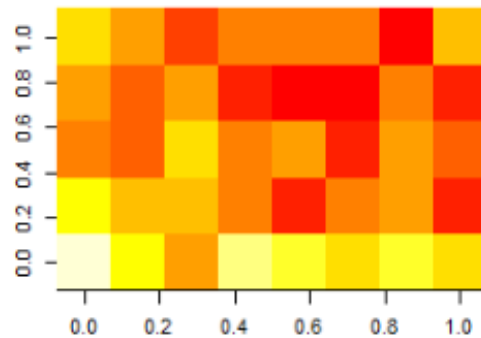
- The null hypothesis is that the data were determined and then assigned to their spatial locations at random.
- The alternative is that the assignment to each location depended on the assignment at that location's neighbors.
- The permutation test does not randomize over the possible sets of data values--it considers them given--
but **conditional on the data observed**, considers all possible ways of reassigning them to the locations.

Monte-Carlo Significance Test (cont'd)

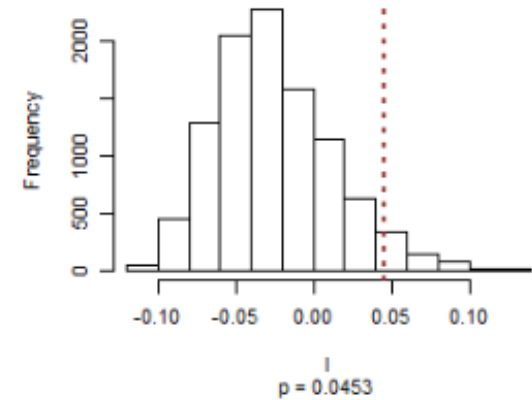
- **Permutation test (排列検定)**
- Such a reassignment is a *permutation*. For n data points, there are $n! = n \times (n-1) \times (n-2) \times \dots \times (2) \times (1)$ permutations. For n much larger than 10 or so, that's too many to generate.
- There usually is no simple analytical expression for the full permutation distribution.
- Accordingly, we typically resort to sampling from the set of all permutations at random, giving them all equal weight. The distribution of the autocorrelation statistic in a sufficiently large sample (usually involving at least 500 permutations) approximates the true distribution.

Examples

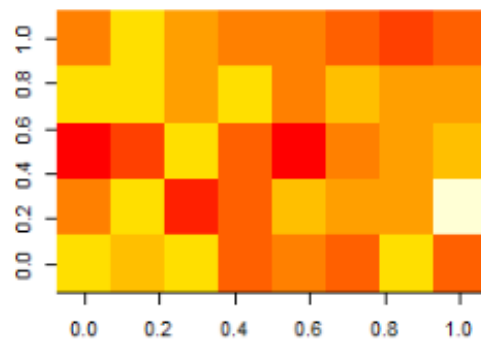
Autocorrelated Data



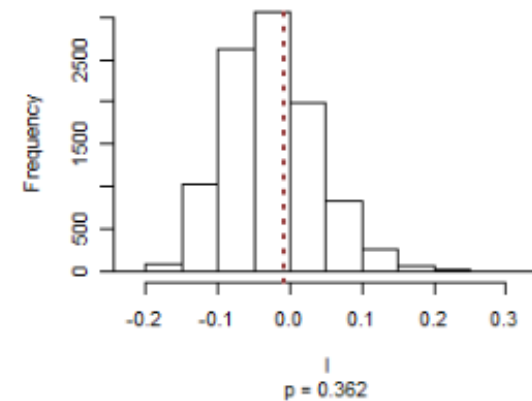
Null Distribution of Moran's I



Uncorrelated Data



Null Distribution of Moran's I



Output in R: Moran's I statistic

```
> M<-moran.test(Popn, listw=TWN_nb_w, zero.policy=T); M
```

Moran I test under randomisation

data: Popn
weights: TWN_nb_w

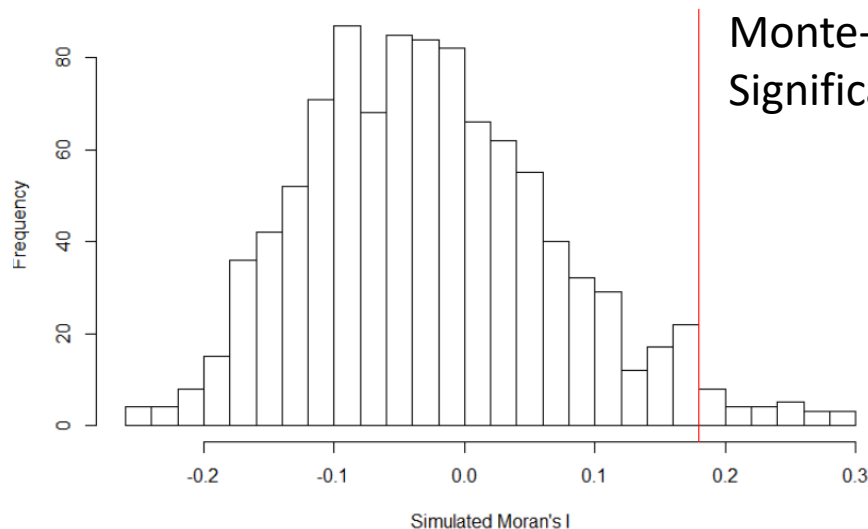
Moran I statistic standard deviate = 2.1678, p-value = 0.01509

alternative hypothesis: greater

sample estimates:

Moran I statistic	Expectation
0.181359104	-0.025000000

Variance
0.009062094



Moran.test()

`moran.test {spdep}`

R Documentation

Moran's I test for spatial autocorrelation

Description

Moran's test for spatial autocorrelation using a spatial weights matrix in weights list form. The assumptions underlying the test are sensitive to the form of the graph of neighbour relationships and other factors, and results may be checked against those of `moran.mc` permutations.

Usage

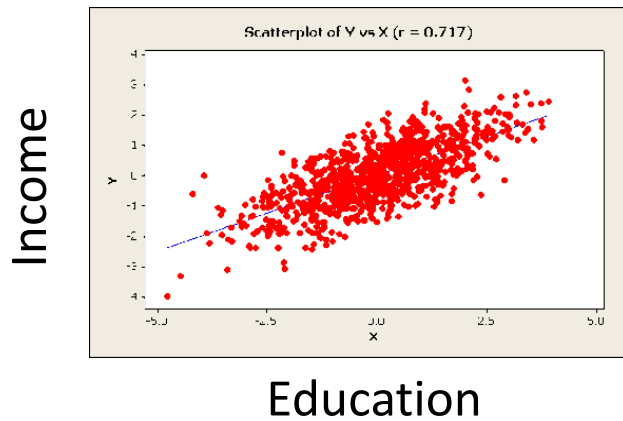
```
moran.test(x, listw, randomisation=TRUE, zero.policy=NULL,  
  alternative="greater", rank = FALSE, na.action=na.fail, spChk=NULL, a
```

<code>randomisation</code>	variance of I calculated under the assumption of randomisation, if FALSE normality
----------------------------	--

Concept of Moran Scatter Plots

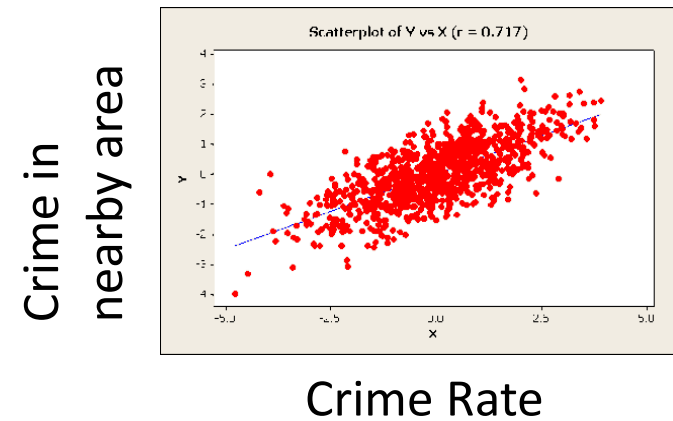
Scatter Plot

Two variables



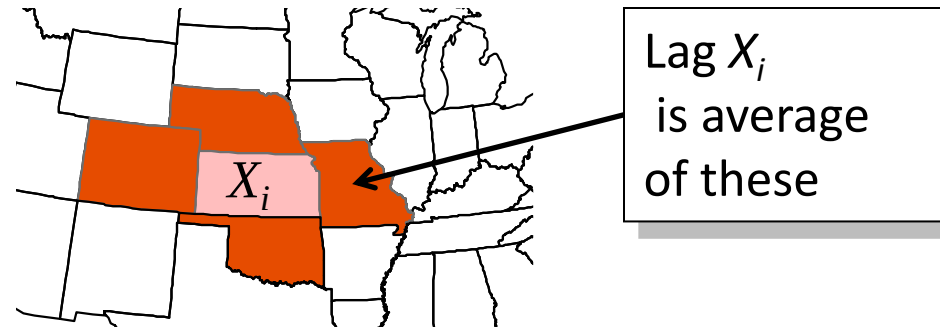
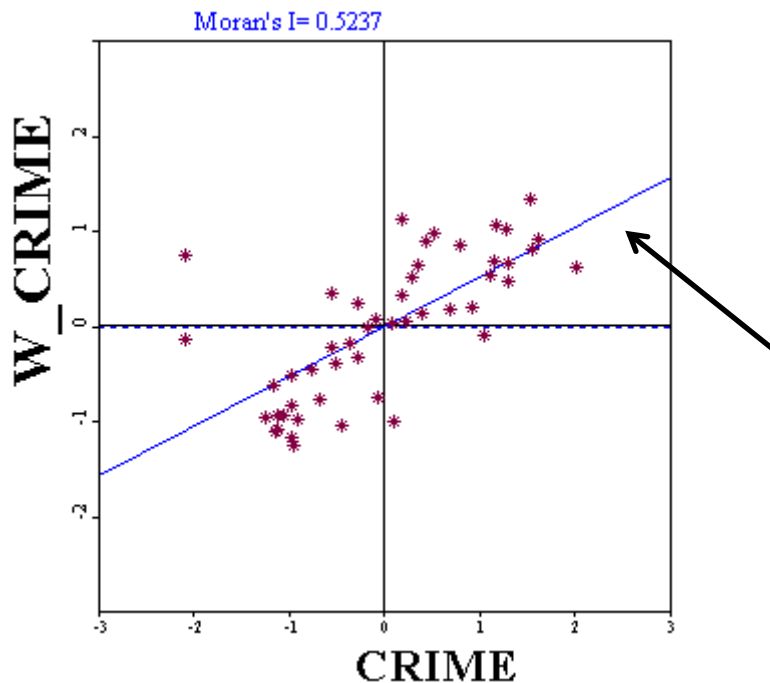
Moran Scatter Plot

Only one variable



Moran Scatter Plots

Moran's I can be interpreted as the correlation between variable, X , and the “spatial lag” of X formed by averaging all the values of X for the neighboring polygons.

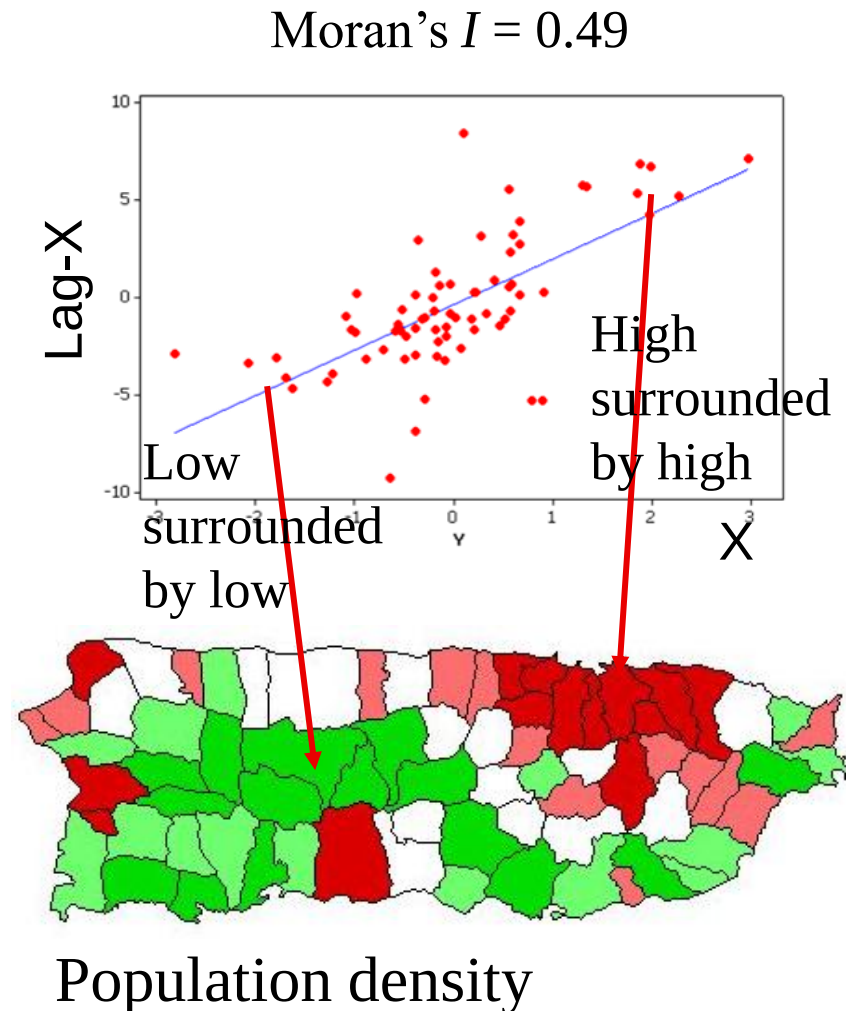


Least squares “best fit” line to the points.

The slope of this *regression line* is Moran's I

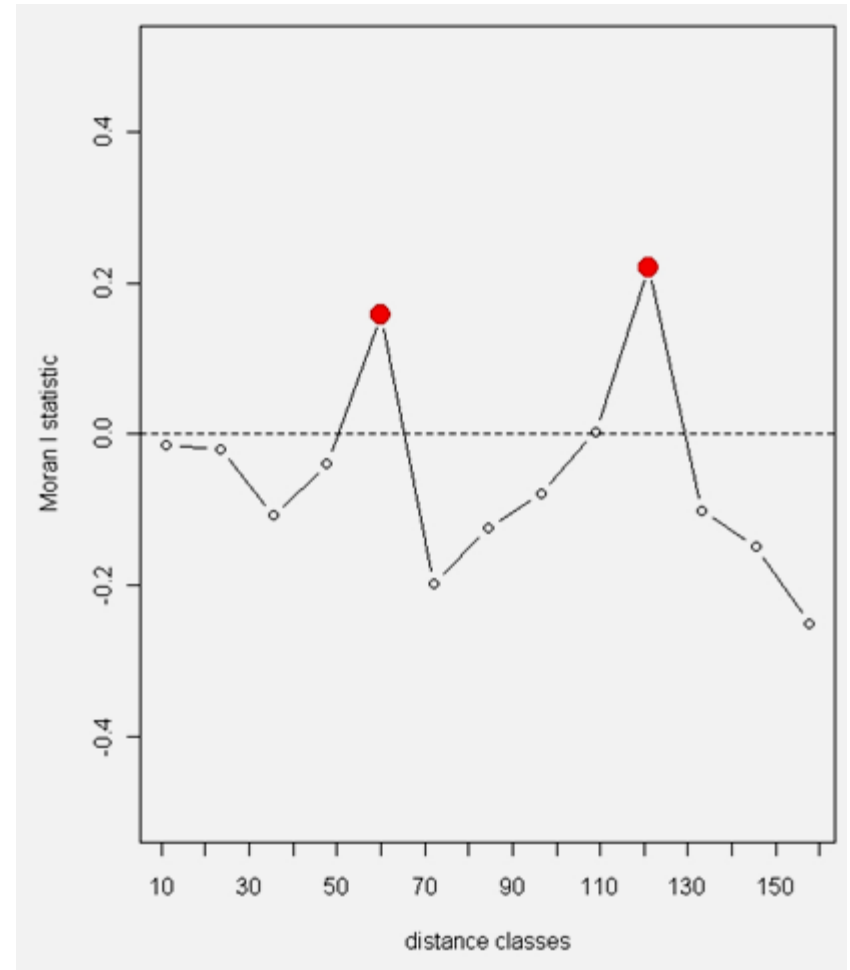
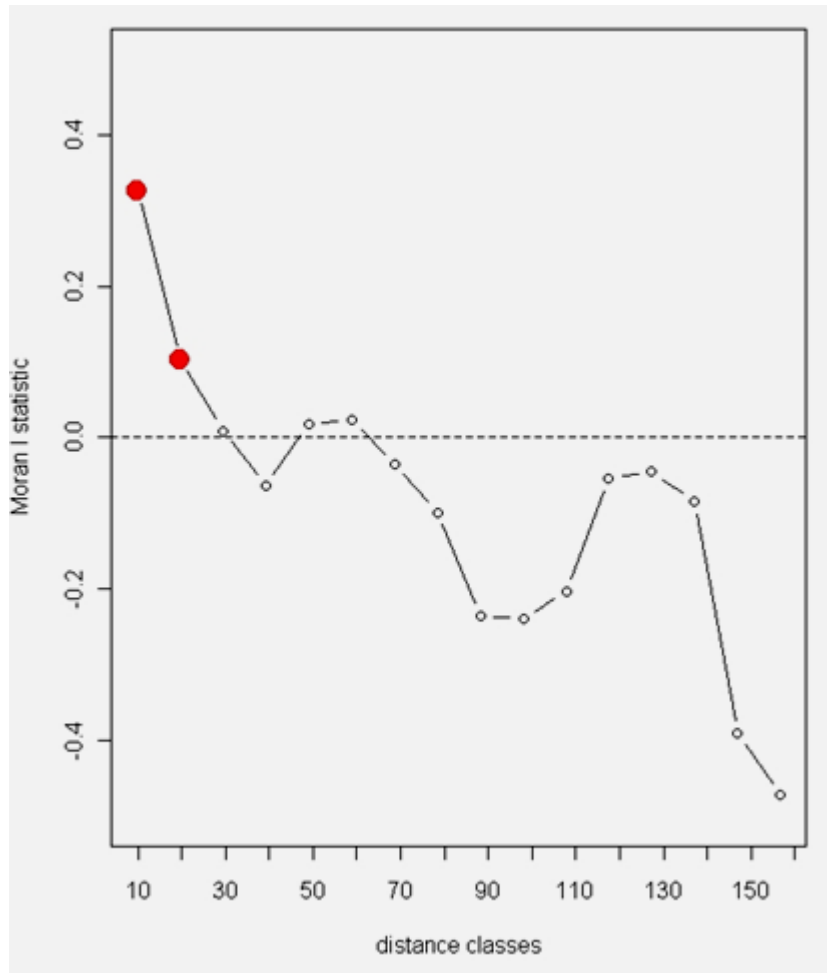
Moran Scatter Plot: example

- Scatter plot of **X** vs. **Lag-X**
- The slope of the regression is Moran's I



Moran Correlograms

Correlogram: plot distance on X-axis against correlation coefficient on Y-axis



Getis-Ord General G-statistic

- Moran's I 無法區別
“hot spots” or “cold spots”
- ***Spatial Concentration*** method
- Definition

$$G(d) = \frac{\sum \sum w_{ij}(d) x_i x_j}{\sum \sum x_i x_j}$$

d : neighborhood distance
 w_{ij} : 1 if it is within d , 0 otherwise

- Calculation of G must begin by identifying a neighborhood distance within which cluster is expected to occur

公式的詳細說明

The General G statistic of overall spatial association is given as:

$$G = \frac{\sum_{i=1}^n \sum_{j=1}^n w_{i,j} x_i x_j}{\sum_{i=1}^n \sum_{j=1}^n x_i x_j}, \quad \forall j \neq i$$

where x_i and x_j are attribute values for features i and j , and $w_{i,j}$ is the spatial weight between feature i and j . n is the number of features in the dataset and $\forall j \neq i$ indicates that features i and j cannot be the same feature.

計算範例

<i>a</i> 20	<i>b</i> 20	<i>c</i> 30
<i>d</i> 20	<i>e</i> 10	<i>f</i> 30
<i>g</i> 20	<i>h</i> 20	<i>i</i> 20

$$G(d) = \frac{\sum_{i=1}^n \sum_{j=1}^n w_{ij}(d) x_i x_j}{\sum_{i=1}^n \sum_{j=1}^n x_i x_j}$$

計算範例 step 1

$$G(d) = \frac{\sum_{i=1}^n \sum_{j=1}^n w_{ij}(d) x_i x_j}{\sum_{i=1}^n \sum_{j=1}^n x_i x_j}$$

	a	b	c	d	e	f	g	h	i
a	400	400	600	400	200	600	400	400	400
b	400	400	600	400	200	600	400	400	400
c	600	600	900	600	300	900	600	600	600
d	400	400	600	400	200	600	400	400	400
e	200	200	300	200	100	300	200	200	200
f	600	600	900	600	300	900	600	600	600
g	400	400	600	400	200	600	400	400	400
h	400	400	600	400	200	600	400	400	400
i	400	400	600	400	200	600	400	400	400
	=3800	=3800	=5700	=3800	=1900	=5700	=3800	=3800	=3800

$$\sum x_i x_j = 36100$$

$$36100 - (\text{對角線總和}) = 36100 - 4300 = 31800$$

計算範例 step 2

$$G(d) = \frac{\sum_{i=1}^n \sum_{j=1}^n w_{ij}(d) x_i x_j}{\sum_{i=1}^n \sum_{j=1}^n x_i x_j}$$

<i>a</i> 20	<i>b</i> 20	<i>c</i> 30
<i>d</i> 20	<i>e</i> 10	<i>f</i> 30
<i>g</i> 20	<i>h</i> 20	<i>i</i> 20

w_{ij}^a	a	b	c	d	e	f	g	h	i
a	0	1	0	1	0	0	0	0	0
b	1	0	1	0	1	0	0	0	0
c	0	1	0	0	0	1	0	0	0
d	1	0	0	0	1	0	1	0	0
e	0	1	0	1	0	1	0	1	0
f	0	0	1	0	1	0	0	0	1
g	0	0	0	1	0	0	0	1	0
h	0	0	0	0	1	0	1	0	1
i	0	0	0	0	0	1	0	1	0

計算範例 step 3

$$G(d) = \frac{\sum_{i=1}^n \sum_{j=1}^n w_{ij}(d) x_i x_j}{\sum_{i=1}^n \sum_{j=1}^n x_i x_j}$$

0	400	0	400	0	0	0	0	0
400	0	600	0	200	0	0	0	0
0	600	0	0	0	900	0	0	0
400	0	0	0	200	0	400	0	0
0	200	0	200	0	300	0	200	0
0	0	900	0	300	0	0	0	600
0	0	0	400	0	0	0	400	0
0	0	0	0	200	0	400	0	400
0	0	0	0	0	600	0	400	0
=800	=1200	=1500	=1000	=900	=1800	=800	=1000	=1000

$$\sum_{i=1}^n \sum_{j=1}^n w_{ij}(d) x_i x_j = 10000$$

計算範例 Final

$$G(d) = \frac{\sum_{i=1}^n \sum_{j=1}^n w_{ij}(d) x_i x_j}{\sum_{i=1}^n \sum_{j=1}^n x_i x_j} = \frac{10000}{31800} = 0.314465$$

<i>a</i> 20	<i>b</i> 20	<i>c</i> 30
<i>d</i> 20	<i>e</i> 10	<i>f</i> 30
<i>g</i> 20	<i>h</i> 20	<i>i</i> 20

與spdep套件計算結果相互驗證

```
library(sf)
library(spdep)

# 建立sf資料與屬性
data = c(20,20,20,20,10,30,20,20,30)
sfc = st_sfc(st_polygon(list(rbind(c(0,0), c(1,0), c(1,1), c(0,0)))))
grid = st_as_sf(st_make_grid(sfc, cellsize = .4))
grid$data = data

# 建立鄰近矩陣
w_list = nb2listw(poly2nb(grid,queen=F),style="B")

globalG.test(grid$data, listw=w_list )
```

```
> globalG.test(grid$data, listw=w_list )
```

Getis-Ord global G statistic

data: grid\$data

weights: w_list

standard deviate = -1.1237, p-value = 0.8694

alternative hypothesis: greater

sample estimates:

Global G statistic	Expectation	Variance
0.3144654088	0.3333333333	0.0002819107

Getis-Ord General G-Statistic

- General G-statistic **can distinguish between hot/cold spots.**
It identifies **spatial concentrations.**
 - G is relatively large if high values cluster together
 - G is relatively low if low values cluster together
- G statistic is interpreted relative to its expected value
 - $G > E(G) \rightarrow$ potential ***“hot spot”***
 - $G < E(G) \rightarrow$ potential ***“cold spot”***
 - $G = E(G) \rightarrow$ no spatial association
- 所謂之larger/smaller不能單從值的大小判斷，
需以 **Z test statistic** 來檢定差異的統計顯著性。

Significance Test for Getis-Ord General G

- Statistical Significance Test

$$Z = \frac{G - E(G)}{S_{Err}(G)}$$

- Expected G : $E(G) = \frac{W}{n(n-1)}$; where $W = \sum_i \sum_j w_{ij}(d)$,

- Standard Error will depend on the sampling method (free / non-free)

Getis-Ord General G-statistic in R

```
> G<-globalG.test(Popn, listw=TWN_ran1_wb); G
```

```
Getis-Ord global G statistic
```

```
data: Popn
```

```
weights: TWN_ran1_wb
```

```
standard deviate = 3.4804, p-value = 0.0002504
```

```
alternative hypothesis: greater
```

```
sample estimates:
```

Global G statistic	Expectation	Variance
0.6216802233	0.5402439024	0.0005474983

教科書的研讀教材（列入考試範圍）

TEXT_Pattern.of.Feature.Value.pdf

MEASURING THE SPATIAL PATTERN OF FEATURE VALUES

In addition to measuring the pattern formed by the locations of features, you can also measure patterns of attribute values associated with features, such as the pattern formed by median house values. These methods reveal whether similar values tend to occur near each other, or whether high and low values are interspersed.



Median house value by census tract.

The idea behind measuring patterns of feature values

Measuring the spatial pattern of feature values is based on the notion that things near each other are more alike than things far apart, an idea often attributed to geographer Waldo Tobler. The idea is consistent with our

實習：介紹 R package: spdep

Spatial Dependence: Weighting Schemes, Statistics and Models

spdep: Spatial Dependence: Weighting Schemes, Statistics and Models

A collection of functions to create spatial weights matrix objects from polygon 'contiguities', from point patterns by distance and tessellations, for summarizing these objects, and for permitting their use in spatial data analysis, including regional aggregation by minimum spanning tree; a collection of tests for spatial 'autocorrelation', including global 'Morans I', 'APLE', 'Gearys C', 'Hubert/Mantel' general cross product statistic, Empirical Bayes estimates and 'Assunção/Reis' Index, 'Getis/Ord' G and multicoloured join count statistics, local 'Moran's I' and 'Getis/Ord' G, 'saddlepoint' approximations and exact tests for global and local 'Moran's I'; and functions for estimating spatial simultaneous 'autoregressive' ('SAR') lag and error models, impact measures for lag models, weighted and 'unweighted' 'SAR' and 'CAR' spatial regression models, semi-parametric and Moran 'eigenvector' spatial filtering, 'GM SAR' error models, and generalized spatial two stage least squares models.

spdep 重要函數

■ Spatial Neighbors

- Contiguity: QUEEN vs. ROOK `poly2nb(); nb2mat()`
- K-nearest Neighbors (KNN) `knn2nb(); knearneigh(coords, k=2)`
- Distance-based `dnearneigh()`

■ From Spatial Neighbors to ListW (Weighting matrix)

- `nb2listw()`

■ Spatial Autocorrelation

- Mapping the attribute `tmap::tm_shape()`
- Moran's I Statistic `moran.test()`
- Monte-Carlo simulation `moran.mc()`
- Moran correlogram `sp.correlogram()`
- Moran Scatter Plot `moran.plot()`
- Getis-Ord General G Statistic `globalG.test()`

■ 台灣鄉鎮市區人口密度的空間型態分析

- 計算以下統計量與繪製圖表，說明其參數設定，並解釋其意義。
包括：Moran's I coefficient, Monte-Carlo simulation, Moran scatter plot, Correlogram, and General G statistic.
- 利用以下三種不同的空間鄰近定義，計算Moran's I coefficient，
比較其數值的差異，並討論可能的原因。

Spatial Neighbors: Contiguity; K-nearest Neighbors (KNN); Distance-based

【9】空間自相關：MORAN'S I

 授課投影片

 授課程式碼

 助教投影片

 LAB9

 助教課影片

<https://wenlab501.github.io/GEOG2017/>

作業：實作

資料：Popn_TWN2.shp

- 共三題 (以Contiguity定義鄰近)
 - 繪製各鄉鎮的鄰居數的直方圖。
 - 找出台灣本島最多鄰居的鄉鎮是哪一個? (TOWN_ID)
 - 繪製台灣各鄉鎮的1st-order鄰居人口密度的面量圖。

