



空間トピックモデルによる 土地利用のモニタリング 社会経済活動集積量

広島大学工学研究科 地球環境計画学研究室

塚井 誠人

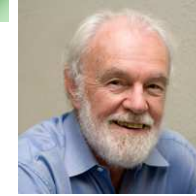
(塚野 裕太 : M2 & QUON ANN TRAN : D1)



Hedgehog

2018年6月22日
愛媛大学坊ちゃんセミナー

David Harveyの懸念@2004年夏の読書から



- ✓ 1935年、英国生まれ(存命)。ケンブリッジ大学より博士号取得。ブリストル大学講師、ジョンズ・ホプキンス大学准教授・教授、オックスフォード大学教授を経て、現在、ニューヨーク市立大学名誉教授。
- ✓ 『資本論』を中心とするマルクス主義を地理学に応用した批判的地理学の第一人者。社会的不公平に深い関心。地理学分野では、世界で最も多く論文が引用される学者。
wikipediaより

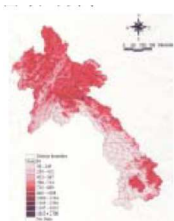


➢ 地理学は、地理的な共通性と相違に対する深い洞察が必要な学問。

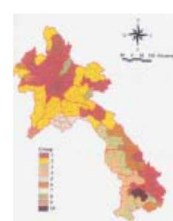
➢ 然るに、地理情報データに主成分分析などを当てはめて、**各地点に共通する特性にのみ着目するアプローチは、地理的な相違から目を逸らすもので、受け入れ難い**...

地理情報に主成分分析を...

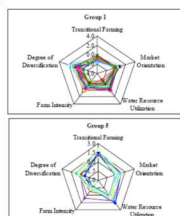
- 1972年当時では最新の研究分野だったはず：衛星画像に基づく土地被覆分類？
- google検索：N工営の国際事業部報告書がヒット



標高データ (DEM) や森林被覆図などラオス政府所有の既存の地理情報システム (GIS) グリッドデータ等をフルに活用し、全国の総合的な農業現況の把握に役立てた。



地理情報システム (GIS) の応用解析および、主成分分析、クラスター分析により、ラオス全国 141 の地域を農業開発の現況に基づき 10 のグループに分割。



抽出された農業開発に関わる5つの主成分に関して、レーダー・チャートを作成。レーダー・チャートにより、各グループの特性を把握。農業開発阻害要因と開発に必要な重点項目を特定。

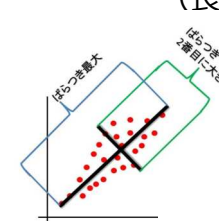
1. 主成分分析

2. クラスター分析

https://www.n-koei.co.jp/rd/thesis/pdf/200212/forum11_009g.pdf

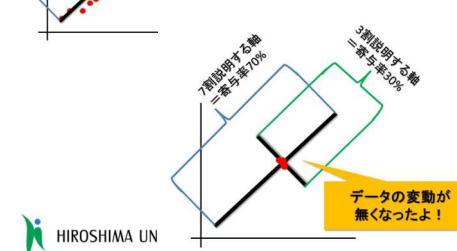
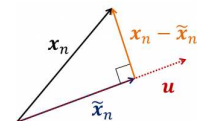
技術的なおさらい1：主成分分析

- J次元空間内のI個のベクトルxに関して、新たな軸uを発見する方法。xのu上への射影ベクトルの分散最大化問題として定式化。ただしuは単位ベクトル(長さ1)



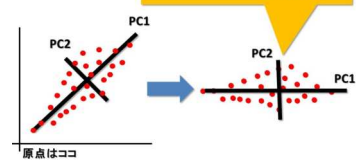
$$L(u) = \frac{1}{N} \sum_n (u^T x_n)^2 + \lambda(1 - u^T u)$$

Lを最大化 = 軸分散最大化



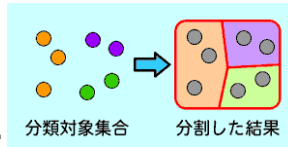
主成分得点

横軸にPC1
縦軸にPC2を置いた時の
X座標 = 第一主成分得点
Y座標 = 第二主成分得点

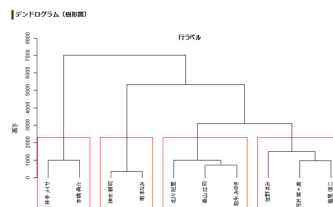


技術的なおさらい 2 - 1 : クラスタ分析

- J次元空間内のI個のサンプルから、サンプル間距離に基づいてk個のサンプルグループを構成。階層的な方法と非階層的な方法あり。



- 階層的な方法 : k=Iから始めてサンプルを併合し、最後はk=1になるデンドログラム（逆樹形図）／kは、図を見て選ぶ。



- サンプル間距離, グループ間距離の定義が複数

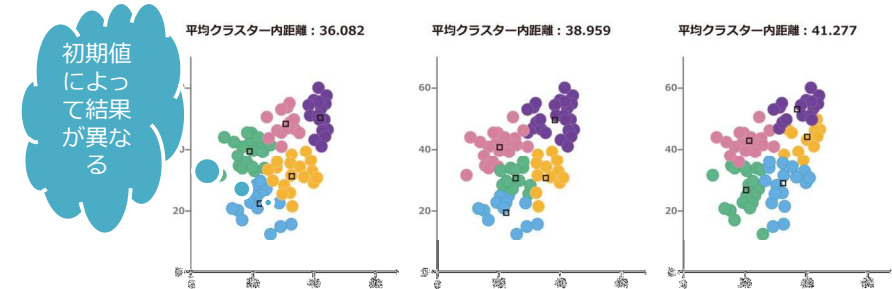
$$d(C_1, C_2) = \min_{x_1 \in C_1, x_2 \in C_2} d(x_1, x_2) \quad d(C_1, C_2) = \frac{1}{|C_1| |C_2|} \sum_{x_1 \in C_1} \sum_{x_2 \in C_2} d(x_1, x_2)$$

$$d(C_1, C_2) = \max_{x_1 \in C_1, x_2 \in C_2} d(x_1, x_2) \quad d(C_1, C_2) = E[C_1 \cup C_2] - E[C_1] - E[C_2]$$

4
59

技術的なおさらい 2 - 2 : クラスタ分析

- 非階層的な方法 : k-means法 / : 適当に選んだ代表点 (k個) と再近隣関係にある標本を併合する方法
- k, つまりグループ数を分析前に指定する方法
- もちろん, 初期値依存がある。



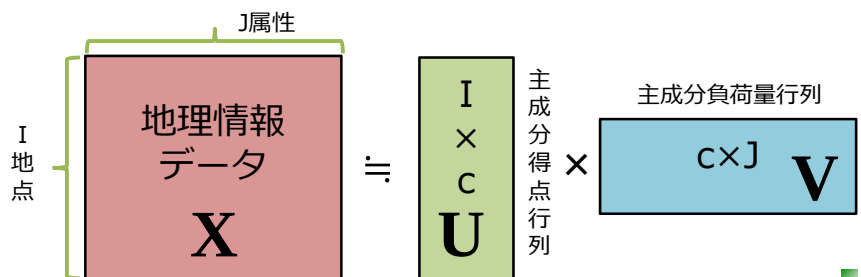
- 限界 : 各グループのサンプルは, 超球内に均等分布する前提あり。

5
59

N工営報告書のアルゴリズム 1

- 対象地域の地理情報データを入手
- 主成分分析から主成分軸を算出 : $c < J$
- サンプルの主成分負荷量, 主成分得点行列を作成

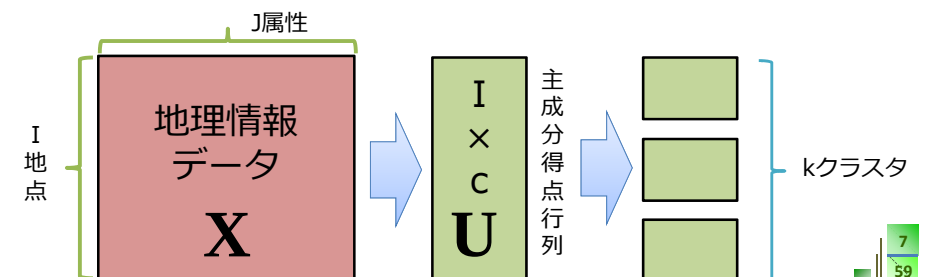
$$X \approx UV$$



6
59

N工営報告書のアルゴリズム 2

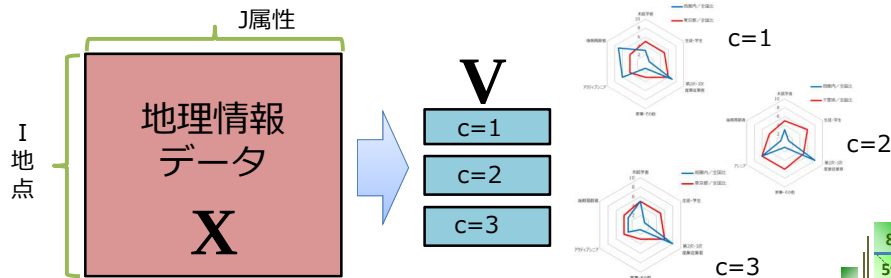
- 対象地域の地理情報データを入手
- 主成分分析から主成分軸を算出 : $c < J$
- サンプルの主成分負荷量, 主成分得点行列を作成
- 主成分得点行列から, クラスタ分析でサンプルをグループ分け = クラスタを得る : $k < I$



7
59

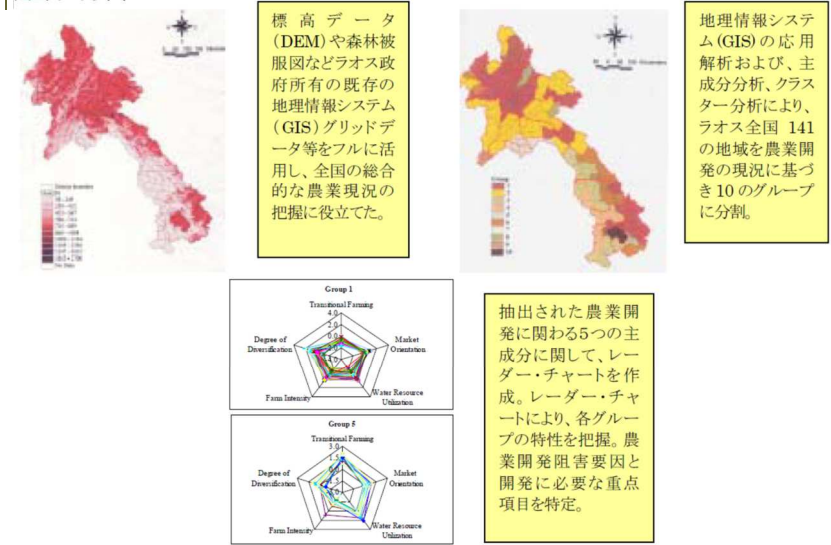
N工営報告書のアルゴリズム 3

1. 対象地域の地理情報データを入手
2. 主成分分析から主成分軸を算出: $c < J$
3. サンプルの主成分負荷量, 主成分得点行列を作成
4. 主成分得点行列から, クラスター分析でサンプルをグループ分け=クラスターを得る.
5. 主成分負荷量行列を主成分ごとに集計して, レーダーチャートに.



8
59

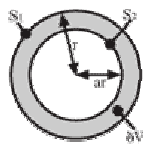
以上をまとめるとこのような結果に



9
59

数学的な問題

- 主成分分析は結局固有値問題: $X'X$ の固有値を求める問題は, 属性次元Jが大きくなっても, 高速で解ける方法がある.
- 一方, クラスター分析には2種類の呪いがかかる.



球面集中現象



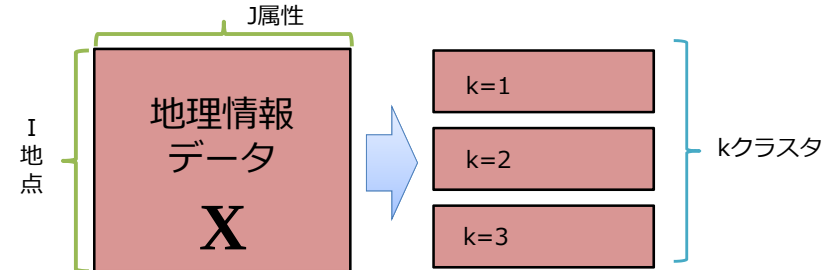
1. 属性次元の呪い: 半径がそれぞれ r と ar ($0 < a < 1$) の J 次元超球 S_1 と S_2 を考える. S_1 の体積 V に対する, 二つの超球の体積の差 δV (図のグレー部分) の比は $\delta V/V = 1 - a^J$. サンプルが超球内に均一分布しているとすると, V 内のサンプル数は超球の体積に比例. また, $\delta V/V$ は J の増加にともない1に近づく. つまり, J が大きな場合は S_1 中の対象は, ほとんど二つの超球の隙間に存在する. これは, 中心のサンプルから他のサンプルまでの距離が次元 J の増加に従って急速に大きくなることを表し, どのサンプル間の類似度も低くなることを意味する. クラスターリングは, 似ている対象をまとめる操作なので, 有意な解が得られなくなる.
2. サンプル数の呪い: 距離計算はサンプル数の二乗で増加.

10
59

このアルゴリズムの疑問点



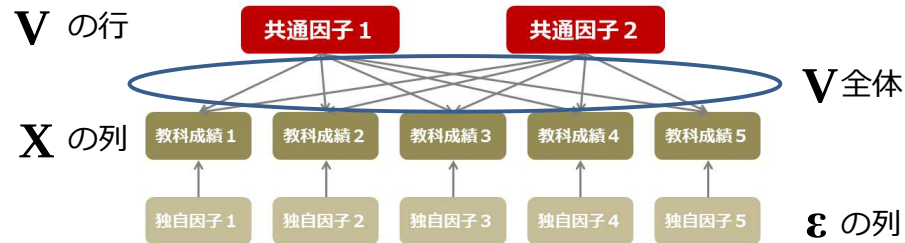
- アルゴリズムの概略: 1) 属性の主成分軸を求め, 2) クラスターを主成分行列から求める
- k と c の決定に関する手掛かりがない (分析者の経験値/主観あるのみ).
- k と c の役割の違いが不明. 結構似た意味かも...
- レーダーチャートは, X をクラスター分けした部分行列からも書ける. どちらが正しい?



11
59

代替法 1 : 因子分析 1

$$X = UV + \varepsilon$$

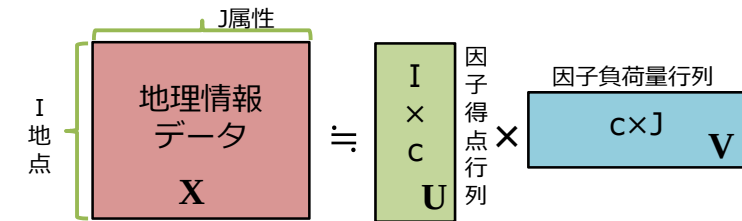


よく見る上の図に騙されそうになるが、
きちんと理解すると...

12
59

代替法 1 : 因子分析 2

$$X = UV + \varepsilon$$



- この問題もやはり固有値問題.
- 古典的には, まずVを推定.
- Vが解釈できるように因子回転をかけて, V'を得る
- 回転後のV'を所与としてUを推定.

誤差
行列
(独自因子)

13
59

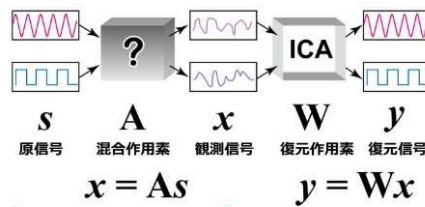
代替法 2 : 独立成分分析 (ICA) 1

$$X'(t) \equiv V'U'(t) \iff X \equiv UV$$

3. 独立成分分析とは何ぞや

まとめるとこんなイメージ

<https://www.slideshare.net/zansa/12121>
8-zansa13-for-web



未知の割合で混合した観測信号から

「独立・非ガウスの」という仮定のみに基づき

原信号を復元することができる!

- よく見る説明は音源分離 (Blind Source Separation)
- 記号が違う & 行列が転置されているので混乱するが, 整理するとわかる.

14
59

代替法 2 : 独立成分分析 (ICA) 2

$$X'(t) \equiv V'U'(t) \iff X \equiv UV$$

• 最尤法の導出

- 観測信号 $\mathcal{X} = \{x(1), \dots, x(T)\}$ に対する W の尤度

$$\mathcal{L}(W|\mathcal{X}) = \prod_{t=1}^T p(x(t)|W)$$

- 線形変換と確率密度関数

$$y = Wx \quad p(y) = \frac{1}{|\det W|} p(x)$$

- 分離信号 y の独立性を仮定

$$p(y) = \prod_{i=1}^I p(y_i) \quad p(y_i) \text{ はラプラス分布など}$$

- 以上から導かれる対数尤度を最大化する W を求める

$$\log \mathcal{L} = T \log |\det W| + \sum_{t=1}^T \sum_{i=1}^I \log p(y_i(t))$$

15
59

代替法 2 : 独立成分分析 (ICA) 3

$$X'(t) \equiv V'U'(t) \longleftrightarrow X \equiv UV$$

• 最尤法の導出

- サンプル数 T で割り, 最大化すべき目的関数を設定

$$\mathcal{J} = \log |\det \mathbf{W}| + \sum_{i=1}^I E\{\log p(y_i)\}$$

» 参考: $\mathbf{W}$ をユニタリ行列に限定すればFastICAと等価

- 勾配法により $\mathbf{W}$ を最適化

$$\mathbf{W} \leftarrow \mathbf{W} + \eta \cdot \frac{\partial \mathcal{J}}{\partial \mathbf{W}} \quad \eta \text{ は適切に設定されたステップサイズ}$$

$$\frac{\partial \mathcal{J}}{\partial \mathbf{W}} = (\mathbf{W}^T)^{-1} - E\{\Phi(\mathbf{y})\mathbf{x}^T\}$$

$$\Phi(\mathbf{y}) = \begin{bmatrix} \phi(y_1) \\ \vdots \\ \phi(y_I) \end{bmatrix} \quad \begin{array}{l} \phi(y_i) \text{ の具体的な形} \\ \text{ラプラス分布} \quad \phi(y_i) = \frac{y_i}{|y_i|} \\ \text{非線形関数 G} \quad \phi(y_i) = \frac{y_i}{\sqrt{y_i^2 + \alpha}} \end{array}$$

澤田宏: 独立成分分析入門, ~音の分離を題材として~, 部分空間法研究会チュートリアル, 2010

16
59

代替法 3 : 行列の特異値分解

- 実数行列に関する特異値分解定理

- 任意の $m \times n$ の行列 A は次の形に分解できる.

$$A = U \Sigma V^T$$

ここで,

- U : $m \times m$ の直交行列
- V : $n \times n$ の直交行列
- Σ : $m \times n$ の対角行列 $\Sigma = \text{diag}(\sigma_1, \sigma_2, \dots, \sigma_r)$
(ただし $\sigma_1 \geq \sigma_2 \geq \dots \geq \sigma_r > \sigma_{r+1} = \dots = \sigma_n = 0$)

- これらを使うと, A は以下のように表すことができる.

$$A = \sigma_1 u_1 v_1^T + \sigma_2 u_2 v_2^T + \dots + \sigma_r u_r v_r^T$$

$$A = \begin{bmatrix} \times & \times & \times \\ \times & \times & \times \\ \times & \times & \times \end{bmatrix} \begin{bmatrix} \sigma_1 \\ \sigma_2 \\ \sigma_r \end{bmatrix} \begin{bmatrix} \times & \times & \times \\ \times & \times & \times \\ \times & \times & \times \end{bmatrix} \begin{bmatrix} v_1 \\ v_2 \\ v_r \end{bmatrix}$$

<https://www.slideshare.net/HA19801/20141204journal-club>

HIROSH

17
59

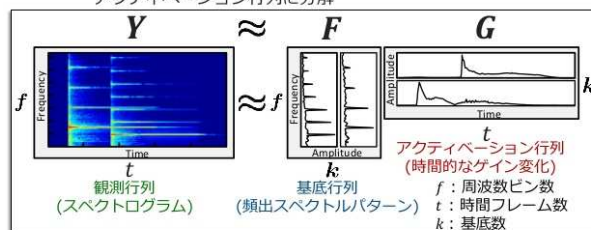
代替法 4 : 非負値行列因子分解 (NMF) 1

$$X' \equiv V'U' \longleftrightarrow X \equiv UV$$

非負値行列因子分解 (NMF)

- NMF [Lee & Seung, 1999]

- 非負値行列を非負値行列の積に低ランク近似
- 画像処理, 自動採譜など応用先は様々
- 音源分離の場合, 音源のスペクトログラムを基底行列とアクティベーション行列に分解



- ここでも音源分離が例に

- 記号が違ふ & 行列が転置されているので混乱するが, 整理するとわかる

<https://www.slideshare.net/DaichiKitamura/ss-66384298>

18
59

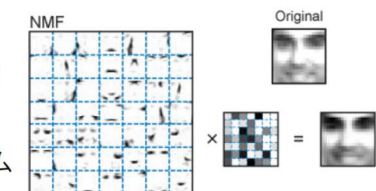
代替法 4 : 非負値行列因子分解 (NMF) 2

NMFの生まれた背景

- 画像処理分野で生まれた技術 [Lee1999]
 - 顔画像から顔パーツを抽出するのが目的
 - 概念自体は90年代前半に登場 [Paatero1994]

VQ: vector quantization

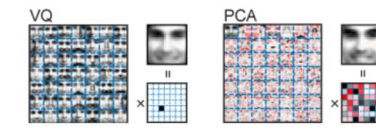
- 音のスペクトルを画像と見なして適用(後述)
→ 音声分離・自動採譜等 [Smaragdis2003] 以て極めて多数...



- 効率的な反復アルゴリズム [Lee2000]

$$H_{\omega,k} \leftarrow H_{\omega,k} \frac{[YU^T]_{\omega,k}}{[HUU^T]_{\omega,k}}$$

$$U_{k,t} \leftarrow U_{k,t} \frac{[H^T Y]_{k,t}}{[H^T H U]_{k,t}}$$



亀岡弘和: 非負値行列因子分解—チュートリアル, 音楽情報科学研究会資料

19
59

代替法 4 : 非負値行列因子分解 (NMF) 3

NMFの定式化

- T 個の観測データ(非負値ベクトル) $y_1, \dots, y_T$
- K 個の非負値基底ベクトル $h_1, \dots, h_K$ の非負結合で
どの観測データも良く表現できる基底セットを求めたい

$$y_t \simeq \sum_{k=1}^K h_k U_{k,t}, \quad U_{k,t} \geq 0 \quad (t = 1, \dots, T)$$

$$\begin{aligned} Y &:= (y_1, \dots, y_T) \\ H &:= (h_1, \dots, h_K) \\ U &:= (U_{k,t})_{K \times T} \end{aligned} \quad \left. \begin{aligned} Y &\simeq HU \\ H_{\omega,k} &\geq 0, U_{k,t} \geq 0 \end{aligned} \right\}$$

Y と HU の乖離度を表す規準

$$\begin{aligned} &\text{minimize} \quad \mathcal{D}(Y|HU) \\ &\text{subject to} \quad H_{\omega,k} \geq 0, U_{k,t} \geq 0 \end{aligned}$$

代替法 4 : 非負値行列因子分解 (NMF) 4

非負値であることのメリット

- **データ行列の非負性**
 - 実世界には非負値データが多い
(例) パワースペクトル, 画素値, 度数,

OD表も,
モバ空も

- **基底行列の非負性**
 - 「非負値データの構成要素もまた非負値データであるべき(でないとう物理的に意味をなさない!)」
という考え方
(例) 負の値をもったパワースペクトルなんて解釈のしようがない

$$X' = V'U'$$

基底とは、 V' の列 = V の行のこと

- **係数行列の非負性**
 - 構成要素の混ざり方は「足し算」のみ
 - 係数行列をスパースに誘導 → 基底の情報量をアップ

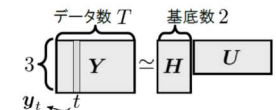
代替法 4 : 非負値行列因子分解 (NMF) 5

非負値で基底を張るとは…?

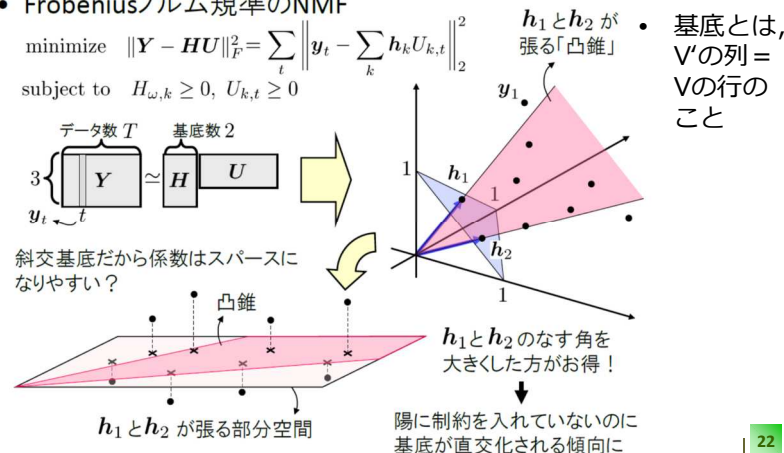
$$X' = V'U'$$

- Frobeniusノルム規準のNMF

$$\begin{aligned} &\text{minimize} \quad \|Y - HU\|_F^2 = \sum_t \|y_t - \sum_k h_k U_{k,t}\|_2^2 \\ &\text{subject to} \quad H_{\omega,k} \geq 0, U_{k,t} \geq 0 \end{aligned}$$



斜交基底だから係数はスパースになりやすい?



代替法 4 : 非負値行列因子分解 (NMF) 6

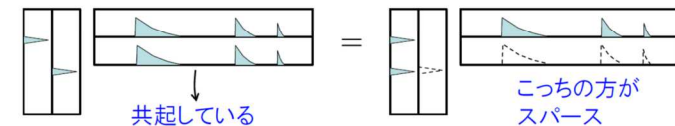
負荷量行列がスパースとはどんな状態? $X' = V'U'$

$$Y \simeq H U$$

これがスパースになるよう誘導される

- H がどうなれば U はスパースになる?

- 共起するパーツがあればひとまとめにした方がスパースに!

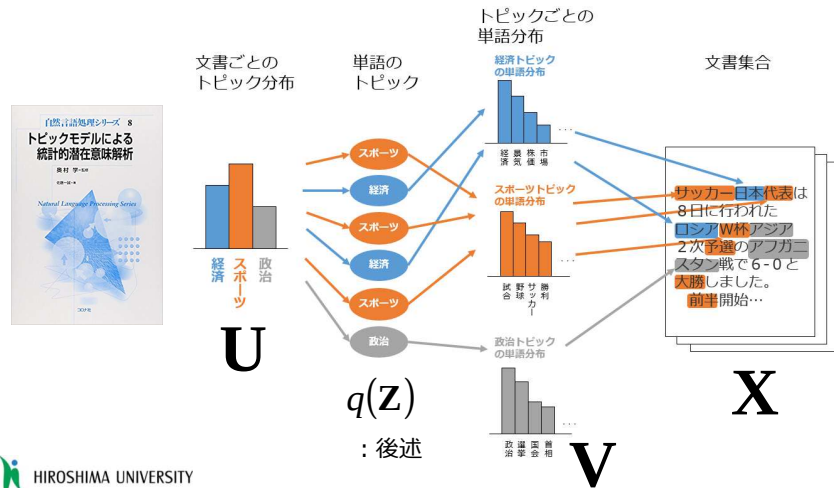


- 共起頻度が高いパーツのペアをひとまとめにした方がスパースに!
(上のひとまとめ操作では $\triangleleft$ を3つ分減らせている)

頻出するひとまとまりのパーツが
各基底ベクトルとして得られる傾向になる

トピックモデルへ

$$X \equiv UV$$



文書データの作成 1

- One-hot表現の作成

「今日私は愛媛大学にきた」

形態素解析

私 は きた
に た 愛媛大学

単語
順序

	昨日	今日	明日	私	君	は	が	松山	愛媛大学	に	くる	帰る
1		1										
2				1								
3						1						
4									1			
5										1		
6											1	

文書データの作成 2

- 単語順序のOne-hot表現を文書単位で集計すると、Bag-of-Wordsに

「今日私は愛媛大学にきた」

形態素解析

私 は きた
に た 愛媛大学

文書

	昨日	今日	明日	私	君	は	が	松山	愛媛大学	に	くる	帰る
A		1		1		1			1	1	1	

全文書に現れる全語彙 (= 属性)

- Bag-of-Wordsでは、文書内の語の順序は無視する。
- Bag-of-Wordsの全語彙集合からは、低頻度語を除く処理を加えることがある。

行列分解法としてのトピックモデル

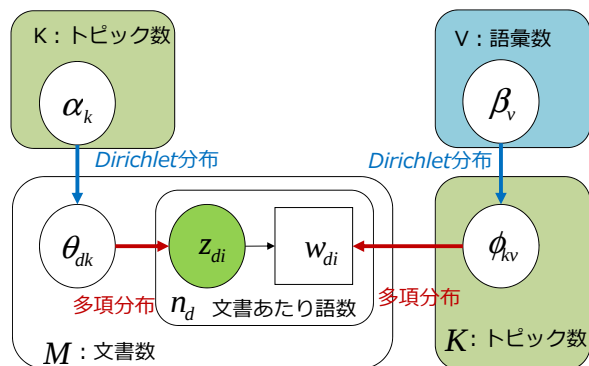
トピックモデルのはたらき

J個の属性をJより少ないk個のトピックにまとめる
階層ベイズモデル：斜交成分の抽出
確率構造の明示：尤度基準によるモデル探索Vからトピックを命名
Uから文書ごとのトピックを推定

$$\begin{matrix} \text{I 文書} \\ \left\{ \begin{array}{c} \text{文書データ BOW} \\ X \end{array} \right\} \end{matrix} \equiv \begin{matrix} \begin{array}{c} I \times k \\ U \end{array} \\ \text{文書トピック行列} \end{matrix} \times \begin{matrix} \begin{array}{c} k \times J \\ V \end{array} \\ \text{トピック語彙行列} \end{matrix}$$

トピックモデルの概要

- 観測された語: w_{di}
- 語とトピックの関係: z_{di}
- トピックと文書の関係: θ_{dk}
- トピックの分布(集中度): α_k
- 語彙とトピックの関係: ϕ_{kv}
- 語彙の分布(集中度): β_v



28
59

トピックモデルの定式化

多項分布:

$$p(\{n_k\}_{k=1}^K | \pi) = \frac{n!}{\prod_{k=1}^K n_k!} \prod_{k=1}^K \pi_k^{n_k}$$

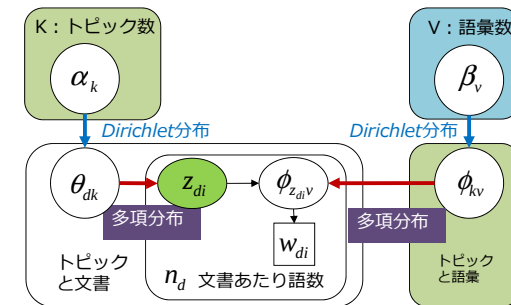
Dirichlet分布:

$$p(\pi | \alpha) = \frac{\Gamma(\sum_{k=1}^K \alpha_k)}{\prod_{k=1}^K \Gamma(\alpha_k)} \prod_{k=1}^K \pi_k^{\alpha_k - 1}$$

- $z_{di} \sim \text{Multi}(\theta_d), k=1, \dots, K$
- $w_{di} \sim \text{Multi}(\phi_{z_{di}, v}), v=1, \dots, V$
- $\theta_{dk} \sim \text{Dir}(\alpha), k=1, \dots, K$
- $\phi_{kv} \sim \text{Dir}(\beta), v=1, \dots, V$

$$p(w, z, \theta, \phi | \alpha, \beta) = p(w | z, \phi) p(z | \theta) p(\theta | \alpha) p(\phi | \beta)$$

同時分布の因数分解近似



29
59

周辺化変分ベイズ法: CVB0の導出

周辺化変分下限の導出

$$p(w | \alpha, \beta) = \sum_z p(w, z | \alpha, \beta) \quad \text{潜在変数} z \text{の導入}$$

$$= \int \sum_z p(w, z, \theta, \phi | \alpha, \beta) d\theta d\phi$$

$$= \int \sum_z p(w | \phi, z) p(z | \theta) p(\theta | \alpha) p(\phi | \beta) d\theta d\phi$$

$$\log p(w | \alpha, \beta) \quad \text{事後分布} p \text{と近似事後分布} q \text{の距離} KL \text{の最小化}$$

$$= F[q(z, \theta, \phi)] + KL[q(z, \theta, \phi) \| p(z, \theta, \phi | w, \alpha, \beta)]$$

⇒変分下限Fの最大化

$$\log p(w | \alpha, \beta) = \sum_z \log p(w, z | \alpha, \beta)$$

$$= \sum_z \log q(z) \frac{p(w, z | \alpha, \beta)}{q(z)}$$

$$\geq \sum_z q(z) \log \frac{p(w, z | \alpha, \beta)}{q(z)}$$

$$= F_{CVB}[q(z)]$$

周辺化による θ, ϕ の消去 (qの因数分解近似より)

周辺化変分下限更新式の導出

$$z = (z_{di}, z^{di}) \quad \text{逐次更新のための分解}$$

$$q^*(z_{di} = k \& w_{di} = v) \propto \exp \left(\sum_z q(z^{di} = k) \times \log p(z_{di} = k, w_{di} = v | w^{di}, z^{di}, \alpha, \beta) \right)$$

$$= \exp E_{q(z^{di})} \left[\log p(z_{di} = k, w_{di} = v | w^{di}, z^{di}, \alpha, \beta) \right]$$

$$= \exp E_{q(z^{di})} \left[\log p(w_{di} = v | z_{di} = k, w^{di}, z^{di}, \alpha, \beta) + \log p(z_{di} = k | z^{di}, \alpha) \right]$$

$$= \exp E_{q(z^{di})} \left[\log \frac{n_{kv}^{di} + \beta_v}{n_k^{di} + \sum_v \beta_v} + \log \frac{n_{dk}^{di} + \alpha_k}{n_d^{di} + \sum_k \alpha_k} \right]$$

$$\propto \frac{\exp E_{q(z^{di})} [\log(n_{kv}^{di} + \beta_v)]}{\exp E_{q(z^{di})} [\log(n_k^{di} + \sum_v \beta_v)]} \times \exp E_{q(z^{di})} [\log(n_{dk}^{di} + \alpha_k)]$$

$$\approx \frac{E[n_{kv}^{di}] + \beta_v}{E[n_k^{di}] + \sum_v \beta_v} \cdot E[n_{dk}^{di}] + \alpha_k \quad \text{テイラー展開 \& 0次近似}$$

30
59

計算法: Collapsed Variational Bayes 0

- 周辺化変分ベイズ法: 決定論アルゴリズム
- MCMC法と異なり, 初期乱数の後は不動点探索を行う

- 0-1. Give the number of K
- 0-2. Initialize parameters: α, β
- 0-3. Randomly Initialize: $q(z_{di})$
- 0-4. Calculate $E[N_{kv}], E[N_{dk}]$ by $q(z_{di})$

1. $q(z_{di}) = f(E[N_{kv}], E[N_{dk}], \alpha, \beta)$
2. $E[N_{kv}], E[N_{dk}]$ Update by $q(z_{di})$
3. α, β Update by $E[N_{kv}], E[N_{dk}]$
4. Check the convergence in α, β

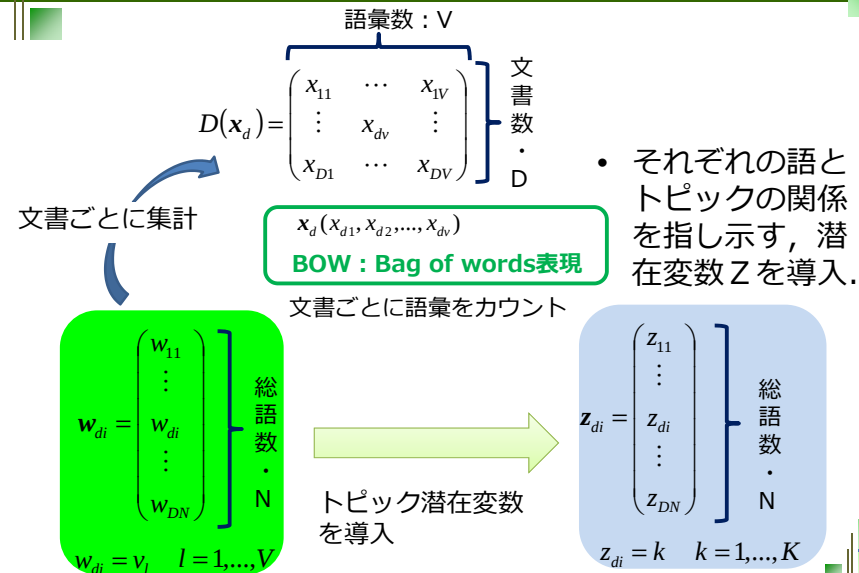
$$\mathbf{V} \quad \phi_{kv} = \frac{E[N_{kv}]}{N_k} \quad E[N_{kv}] \quad \text{トピック別の語彙分布}$$

$$\mathbf{U} \quad \theta_{dk} = \frac{E[N_{dk}]}{N_d} \quad E[N_{dk}] \quad \text{文書別のトピック分布}$$

Converged

31
59

CVB0アルゴリズムの入出力 1



CVB0アルゴリズムの入出力 2

- 語とトピックの関係を示す潜在変数 Z を, トピック帰属確率に分解する.

$$q(z_{di}) = \begin{pmatrix} q_{z11} & \cdots & q_{z1K} \\ \vdots & & \vdots \\ q_{zN1} & \cdots & q_{zNK} \end{pmatrix} \quad \left. \begin{array}{l} \text{総語数} \\ \cdot \\ N \end{array} \right\}$$

$$q(z_{di}) = \sum_{k=1}^K q_k(z_{di}) = 1 \quad \text{for all } z_{di}$$

トピック数 : K

- それぞれの語とトピックの関係を指し示す, 潜在変数 Z を導入.

$$z_{di} = \begin{pmatrix} z_{11} \\ \vdots \\ z_{di} \\ \vdots \\ z_{DN} \end{pmatrix} \quad \left. \begin{array}{l} \text{総語数} \\ \cdot \\ N \end{array} \right\}$$

$z_{di} = k \quad k = 1, \dots, K$

33

59

CVB0アルゴリズムの入出力 3

- 推定パラメータが, 語数に関する加法条件を満たしている

N : 全文書の総語数 N_d : 文書 d (1 to D) の語数

N_v : 語彙 v (1 to V) の語数

V : 全文書の語彙数 V_k : トピック k (1 to K) の語彙数

N_{dk} : 文書 d のトピック k の語数 N_{kv} : トピック k の語彙 v の語数

$$\sum_{d=1}^D N_d = N \quad \sum_{k=1}^K N_{dk} = N_d \quad \text{for all } d \quad \sum_{k=1}^K N_{kv} = N_v \quad \text{for all } v$$

$$\sum_{v=1}^V N_v = N \quad \sum_{v=1}^V N_{kv} = \sum_{d=1}^D N_{dk} = N_k \quad \text{for all } k : \text{トピック } k \text{ の語数} \\ \text{= トピック } k \text{ の大きさ}$$

34

59

各手法の特徴まとめ

	データ	基底順序	基底回転	基底直交	基底数探索	解の正値性	係数の疎性	解の唯一性	解の初期値依存	興味は?
PCA 主成分分析	連続	有	無	直交	累積寄与	無	無	有	無	V?
FA 因子分析	連続	無	有	直交 →斜交	累積寄与?	無	無	無	無	U?
ICA 独立成分分析	連続	無	無	斜交	尤度	無	無	無	無	V, U
SVD 行列特異値分解	連続	有	無	直交	×	無	無	無	無	
NMF 非負値行列因子分解	連続	無	無	斜交	残差基準	有	有	無	有	
トピックモデル	離散	無	無	斜交	尤度	有	有	無	有	

35

59

地理情報へのトピックモデルの適用

GISの進歩に伴い、
地理情報データ整備が進む。

土地利用や人口動態など、
多種類の情報が**詳細な地点単位**
で入手できる。

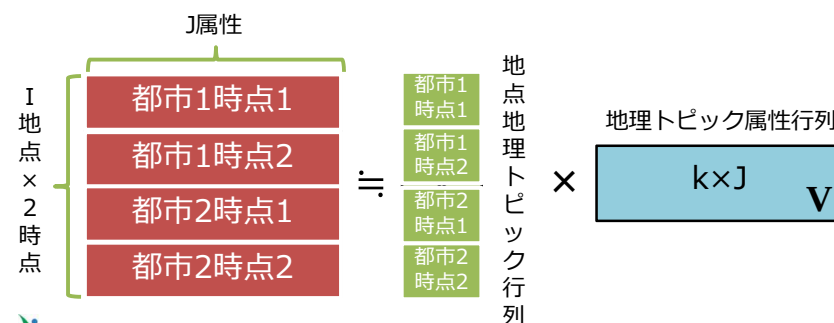
そのままでは扱いづらい、膨大な情報量
の適切な**圧縮手法**が求められるように。



出典) esriジャパン,
<https://www.esri.com/n/gis-guide/gis-datamodel/gis-datamodel/>

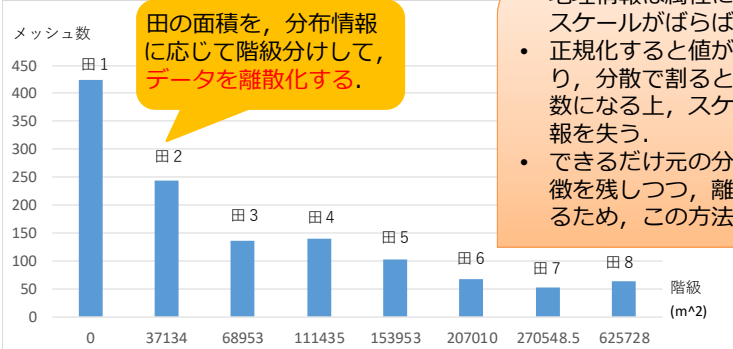
入力データの構成

- 2都市2時点のデータをプール
- 共通する = 都市や時点によって変わらない地理トピック (V) を抽出
- トピック重み (U) の違いが、年次別の土地利用特性の違いに現れるようにデータ作成



データ加工：離散化

- 属性ごとに8段階に自然階級分類によって**階級分け**。
- 階級別に**該当/非該当をダミー変数**で与える。
- 下図において田1とは、1メッシュ内の田の面積が**もっとも小さい**階級を示す。
- データ欠損や秘匿のメッシュは、NA階級。



田の面積を、分布情報
に応じて階級分けして、
データを離散化する。

- 地理情報は属性によって
スケールがばらばら。
- 正規化すると値が負になり、分散で割ると妙な実数になる上、スケール情報を失う。
- できるだけ元の分布の特徴を残しつつ、離散化するため、この方法を採用。

データの概要

- 対象地域：福岡都市圏、北九州都市圏
- データソース：国勢調査地域メッシュ統計、経済センサス、事業所企業統計調査、国土数値情報土地利用
- メッシュの種類：3次メッシュ
- メッシュ数：2113×2時点=4226メッシュ
- 属性数：36変数×9階級→324属性：データの**離散化**

土地	田 他農用地
世帯	1人世帯数, 2人世帯数, 3人世帯数, 4人世帯数, 5人以上世帯数, 1戸建て世帯数, 共同住宅世帯数, 核家族世帯数
就業者	第1次産業就業者総数 第2次産業就業者総数 第3次産業就業者総数
事業所・従業者	第2次産業事業所数, 第2次産業従業者数, 卸売・小売業事業所数 金融・保険業事業所数 不動産・物品賃貸業事業所数 宿泊業・飲食サービス業事業所数, 教育・学習支援事業所数, 医療・福祉事業所数, 卸売・小売業従業者数, 金融・保険従業者数, 不動産業従業者数, 宿泊業・飲食サービス業従業者数, 教育・学習支援従業者数, 医療・福祉従業者数, 公務事業所数, 公務従業者数
人口	0~14歳人口総数, 15~64歳人口総数, 65歳以上人口総数, 居住期間10年未満総数, 居住期間出走時から総数

トピック数の決定方法

● コサイン類似度 (V行列/U行列)

$$\cos(\mathbf{e}_i, \mathbf{e}_{i'}) = \frac{\mathbf{e}_i \cdot \mathbf{e}_{i'}}{\|\mathbf{e}_i\| \|\mathbf{e}_{i'}\|}$$

- トピック比較の場合, i' は別のトピック・属性ベクトル
- 同地点の時点間比較の場合, i' は同地点異時点の地点・地理トピックベクトル

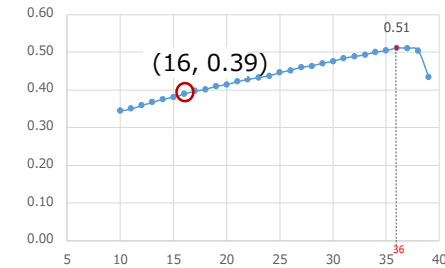
● 対数尤度 (U・V行列)

$$\log L(u, v|x) = \sum_{d=1}^D \sum_{i=1}^{N_d} \sum_{k=1}^K \log u_{dk} v_{k,j=x}$$

トピック間類似度が低く、対数尤度の高いトピック数を選択

尤度・トピック類似度

- 尤度はトピック数36で極大
- $k=36$ で、得られたトピック内容の重複を確認
- k を減らしながらトピック類似度を確認したところ, $k=16$ で、トピック類似度が全て0.8を下回ることを確認→16トピックを採用



トピック類似度：補足

● $k=16$ のトピック類似度

0.007	0.047	0.085	0.007	0.678	0.013	0	0.006	0.013	0	0.711	0.024	0.054	0.065	0.04															
0.018	0.025	0.004	0.001	0.052	0.722	0.095	0.005	0.031	0.003	0.019	0.007	0.015	0.073																
0.374	0.105	0.01	0.284	0.036	0.07	0.098	0.111	0.019	0.482	0.132	0.181	0.033																	
0.052	0.017	0.101	0.018	0.038	0.057	0.044	0.02	0.157	0.116	0.139	0.058																		
0.000	0.001		0.002	0.11	0.16	0.06	0.592	0.314	0.014	0.055	0.718	0.268	0.007																
0.000	0.000	0.276		0.003	0.001	0.001	0.001	0.004	0.005	0.011	0.014	0.008																	
0.000	0.000	0.002	0.000	0.128	0.347	0.163	0.346	0.01	0.385	0.085	0.345	0.011																	
0.002	0.002	0.004	0.000	0.000	0.262	0.166	0.187	0.003	0.027	0.109	0.103	0.04																	
0.000	0.000	0.000	0.000	0.742	0.000	0.106	0.24	0.003	0.085	0.049	0.119	0.01																	
0.002	0.002	0.000	0.000	0.000	1.000	0.000		0.414	0.006	0.048	0.655	0.281	0.01																
0.002	0.002	0.000	0.000	0.000	0.998	0.000	0.998	0.01	0.114	0.268	0.277	0.001																	
0.001	0.005	0.106	0.003	0.000	0.735	0.000	0.717	0.715	0.012	0.016	0.021	0.01																	
0.000	0.000	0.018	0.072	0.007	0.000	0.005	0.000	0.000	0.056	0.122	0.019																		
0.000	0.000	0.003	0.004	0.015	0.000	0.029	0.000	0.000	0.000	0.000	0.204																		
0.000	0.000	0.050	0.317	0.000	0.000	0.000	0.000	0.001	0.256	0.013		0.251	0.04																
0.000	0.000	0.000	0.000	0.586	0.000	0.756	0.000	0.000	0.000	0.011	0.040	0.000		0.048															
0.000	0.000	0.169	0.341	0.000	0.003	0.000	0.000	0.000	0.052	0.094	0.008	0.183	0.000																
0.000	0.000	0.004	0.007	0.001	0.001	0.001	0.000	0.000	0.020	0.015	0.014	0.005	0.003	0.007															
0.000	0.000	0.000	0.007	0.001	0.000	0.002	0.000	0.001	0.002	0.008	0.001	0.002	0.001	0.288															
0.000	0.000	0.000	0.006	0.001	0.000	0.001	0.000	0.000	0.001	0.000	0.000	0.000	0.000	0.566															
0.000	0.000	0.000	0.002	0.282	0.000	0.403	0.000	0.003	0.000	0.007	0.026	0.001	0.468	0.001	0.077	0.286	0.271												
0.000	0.000	0.000	0.013	0.005	0.000	0.009	0.000	0.000	0.001	0.005	0.007	0.001	0.009	0.001	0.174	0.398	0.355	0.709											
0.000	0.000	0.000	0.003	0.305	0.000	0.422	0.000	0.000	0.000	0.007	0.034	0.000	0.400	0.001	0.064	0.265	0.317	0.447	0.423										
0.000	0.000	0.001	0.001	0.106	0.000	0.152	0.000	0.000	0.000	0.047	0.215	0.010	0.173	0.003	0.113	0.108	0.016	0.147	0.058	0.119									
0.000	0.000	0.000	0.003	0.008	0.000	0.015	0.000	0.000	0.000	0.066	0.099	0.037	0.023	0.047	0.126	0.095	0.029	0.045	0.072	0.039	0.204								
0.000	0.000	0.001	0.013	0.010	0.000	0.018	0.000	0.000	0.001	0.009	0.021	0.005	0.020	0.001	0.289	0.418	0.267	0.586	0.664	0.239	0.120	0.114							
0.000	0.000	0.002	0.009	0.001	0.000	0.002	0.000	0.000	0.006	0.000	0.002	0.000	0.002	0.001	0.117	0.329	0.349	0.471	0.572	0.466	0.021	0.041	0.434						
0.000	0.000	0.002	0.054	0.001	0.000	0.001	0.000	0.000	0.005	0.007	0.031	0.053	0.008	0.131	0.170	0.055	0.017	0.010	0.043	0.011	0.055	0.151	0.065	0.027					
0.000	0.001	0.006	0.000	0.050	0.001	0.101	0.000	0.001	0.030	0.000	0.000	0.000	0.004	0.005	0.002	0.001	0.001	0.034	0.001	0.032	0.000	0.000	0.000	0.005	0.001				
0.000	0.000	0.016	0.064	0.007	0.001	0.014	0.000	0.000	0.005	0.033	0.042	0.023	0.021	0.060	0.431	0.178	0.106	0.062	0.259	0.075	0.085	0.103	0.290	0.163	0.154	0.001			
0.001	0.002	0.066	0.019	0.066	0.017	0.001	0.001	0.024	0.309	0.017	0.002	0.003	0.001	0.072	0.097	0.089	0.074	0.029	0.164	0.040	0.002	0.028	0.159	0.116	0.058	0.033	0.287		
0.000	0.000	0.006	0.021	0.028	0.000	0.058	0.000	0.002	0.001	0.026	0.074	0.002	0.068	0.002	0.274	0.296	0.149	0.212	0.402	0.166	0.218	0.136	0.478	0.218	0.089	0.001	0.377	0.212	
0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	
0.001	0.004	0.151	0.005	0.006	0.028	0.001	0.002	0.001	0.095	0.001	0.000	0.001	0.001	0.079	0.033	0.003	0.002	0.002	0.005	0.002	0.001	0.002	0.005	0.012	0.009	0.066	0.015	0.469	0.006

● $k=32$ のトピック類似度

V行列：トピックの命名 1

寄与の大きな属性の並びからトピックを命名

高集積	準高集積
他農用地0	金融業・保険業従業者数1
四人世帯数7	卸売業・小売業事務所数1
田0	卸売業・小売業従業者数2
一戸建世帯数7	不動産業事務所数2
第一次産業従業者総数7	第二次産業従業者数2
共同住宅世帯数7	医療・福祉事務所数0
65歳以上人口総数7	公務従業者数0
三人世帯数7	三人世帯数5
核家族世帯数7	不動産業従業者数2
居住期間出生時から総数7	他農用地0
15.64歳人口総数7	第二次産業従業者総数7
三人世帯数7	15.64歳人口総数5
0.14歳人口総数7	教育・学習支援業従業者数2
卸売業・小売業従業者数3	一戸建世帯数5
不動産業事務所数3	第一次産業従業者総数5
第二次産業従業者数3	五人以上世帯数4
三人世帯数6	宿泊業・飲食サービス業従業者数2
不動産業従業者数3	金融業・保険業事務所数2
15.64歳人口総数6	共同住宅世帯数5
五人以上世帯数7	教育・学習支援業事務所数3

- トピックごとに総重みが1になるように基準化
- 表は、20位まで or 重みが0.01以上のみ作成

V行列：トピックの命名2

居住（密度高）		居住（密度中）		居住（密度中低）	
五人以上世帯数3	0.036	第一次産業就業者総数3	0.039	15.64歳人口総数2	0.060
第一次産業就業者総数4	0.035	一戸建世帯数3	0.039	三人世帯数2	0.058
一戸建世帯数4	0.034	金融業、保険業事務所数1	0.036	一戸建世帯数2	0.055
居住期間10年未満総数2	0.033	不動産従業者数1	0.035	共同住宅世帯数2	0.053
共同住宅世帯数4	0.032	15.64歳人口総数3	0.033	第一次産業就業者総数2	0.053
15.64歳人口総数4	0.030	五人以上世帯数2	0.033	居住期間出生時から総数2	0.049
第三次産業就業者総数2	0.030	共同住宅世帯数3	0.032	居住期間10年未満総数1	0.040
三人世帯数4	0.029	卸売業、小売業事務所数1	0.029	0.14歳人口総数2	0.039
二人世帯数4	0.024	金融業、保険業従業者数1	0.028	第三次産業就業者総数1	0.037
居住期間出生時から総数4	0.023	第二次産業就業者総数4	0.028	65歳以上人口総数2	0.037
第2次産業従業者数2	0.022	三人世帯数3	0.028	二人世帯数2	0.036
0.14歳人口総数4	0.021	第三次産業就業者総数1	0.027	第二次産業就業者総数3	0.035
核家族世帯数4	0.021	居住期間出生時から総数3	0.026	核家族世帯数2	0.027
四人世帯数4	0.021	65歳以上人口総数3	0.025	四人世帯数3	0.026
教育、学習支援業事務所数2	0.021	二人世帯数3	0.024	五人以上世帯数2	0.025
0.14歳人口総数3	0.020	医療、福祉事務所数0	0.023	核家族世帯数3	0.024
第二次産業就業者総数5	0.019	公務従業者数0	0.023	不動産従業者数1	0.021
宿泊業、飲食サービス業事務所数2	0.019	居住期間10年未満総数1	0.023	四人世帯数2	0.021
四人世帯数5	0.018	核家族世帯数3	0.021	金融業、保険業事務所数1	0.020
他農用地0	0.018	教育、学習支援業従業者数1	0.021	五人以上世帯数1	0.017

V行列：トピックの命名3

低密度居住		低密度居住+農業		低密度居住+低密度商業	
第二次産業就業者総数1	0.071	0.14歳人口総数1	0.060	教育、学習支援業事務所数1	0.099
核家族世帯数1	0.067	五人以上世帯数1	0.058	宿泊業、飲食サービス業事務所数1	0.088
65歳以上人口総数1	0.065	居住期間出生時から総数1	0.055	医療、福祉従業者数1	0.087
15.64歳人口総数1	0.065	共同住宅世帯数1	0.054	教育、学習支援業従業者数1	0.087
第一次産業就業者総数1	0.064	三人世帯数1	0.053	卸売業、小売業事務所数0	0.053
三人世帯数1	0.063	居住期間10年未満総数1	0.049	第2次産業従業者数1	0.053
二人世帯数1	0.063	15.64歳人口総数1	0.048	金融業、保険業従業者数0	0.053
居住期間出生時から総数1	0.063	第一次産業就業者総数1	0.045	卸売業、小売業従業者数1	0.053
四人世帯数1	0.062	四人世帯数2	0.045	公務従業者数0	0.039
五人以上世帯数1	0.061	第二次産業就業者総数2	0.041	医療、福祉事務所数0	0.039
一戸建世帯数1	0.061	一戸建世帯数1	0.040	不動産従業者数0	0.034
共同住宅世帯数1	0.059	核家族世帯数2	0.034	金融業、保険業事務所数0	0.034
0.14歳人口総数1	0.058	第三次産業就業者総数0	0.027	第三次産業就業者総数1	0.030
居住期間10年未満総数1	0.055	65歳以上人口総数1	0.026	他農用地1	0.026
第三次産業就業者総数0	0.048	二人世帯数1	0.025	公務事務所数1	0.026
一人世帯数1	0.025	公務事務所数1	0.023	居住期間10年未満総数1	0.025
一人世帯数2	0.014	田7	0.019	65歳以上人口総数2	0.023
		一人世帯数4	0.018	第2次産業事務所数1	0.021
		二人世帯数2	0.017	一人世帯数3	0.015
		一人世帯数3	0.016		

V行列：トピックの命名4

低密度商業		低密度工業		低密度工業+農業	
宿泊業、飲食サービス業従業者数1	0.175	教育、学習支援業従業者数0	0.097	不動産事務所数0	0.100
不動産事務所数1	0.171	宿泊業、飲食サービス業事務所数0	0.097	宿泊業、飲食サービス業従業者数0	0.100
公務従業者数0	0.068	第2次産業従業者数1	0.091	公務事務所数1	0.072
医療、福祉事務所数0	0.067	卸売業、小売業従業者数1	0.086	第2次産業事務所数1	0.070
他農用地0	0.066	医療、福祉従業者数0	0.080	医療、福祉事務所数0	0.058
公務事務所数1	0.058	教育、学習支援業事務所数0	0.080	公務従業者数0	0.058
金融業、保険業従業者数0	0.054	不動産従業者数0	0.066	卸売業、小売業事務所数0	0.055
卸売業、小売業事務所数0	0.053	金融業、保険業事務所数0	0.065	金融業、保険業従業者数0	0.055
田0	0.048	金融業、保険業従業者数0	0.059	金融業、保険業事務所数0	0.050
卸売業、小売業従業者数1	0.041	卸売業、小売業事務所数0	0.059	不動産従業者数0	0.050
第2次産業従業者数1	0.034	医療、福祉事務所数0	0.056	教育、学習支援業従業者数0	0.045
不動産従業者数1	0.029	公務従業者数0	0.056	宿泊業、飲食サービス業事務所数0	0.045
金融業、保険業事務所数1	0.029	第2次産業事務所数1	0.029	医療、福祉従業者数0	0.045
田1	0.024	田2	0.018	教育、学習支援業事務所数0	0.044
第2次産業事務所数2	0.022	他農用地0	0.018	第2次産業従業者数0	0.033
第三次産業就業者総数1	0.017			卸売業、小売業従業者数0	0.033
一人世帯数1	0.015			田1	0.017
				他農用地1	0.015
				田3	0.010
				他農用地0	0.010

V行列：トピックの命名5

未利用1		未利用2	
公務事務所数0	0.072	核家族世帯数0	0.071
第2次産業事務所数0	0.072	0.14歳人口総数0	0.070
不動産従業者数0	0.061	一人世帯数0	0.069
金融業、保険業事務所数0	0.061	共同住宅世帯数0	0.066
医療、福祉従業者数0	0.059	一戸建世帯数0	0.065
教育、学習支援業事務所数0	0.058	二人世帯数0	0.063
卸売業、小売業従業者数0	0.057	第三次産業就業者総数0	0.062
第2次産業従業者数0	0.057	四人世帯数0	0.058
金融業、保険業従業者数0	0.057	居住期間10年未満総数0	0.057
卸売業、小売業事務所数0	0.056	五人以上世帯数0	0.055
公務従業者数0	0.055	居住期間出生時から総数0	0.054
医療、福祉事務所数0	0.055	第二次産業就業者総数0	0.053
教育、学習支援業従業者数0	0.051	第一次産業就業者総数0	0.052
宿泊業、飲食サービス業事務所数0	0.050	65歳以上人口総数0	0.050
宿泊業、飲食サービス業従業者数0	0.041	三人世帯数0	0.049
田0	0.041	15.64歳人口総数0	0.045
不動産事務所数0	0.040	他農用地0	0.032
他農用地0	0.033	田0	0.023

V行列：トピックの命名6

秘匿1		秘匿2		秘匿3	
公務従業者数NA	0.058	第2次産業事務所数NA	0.030	第三次産業就業者総数NA	0.061
医療、福祉従業者数NA	0.058	第2次産業従業者数NA	0.030	第二次産業就業者総数NA	0.061
教育、学習支援従業者数NA	0.058	卸売業、小売業事務所数NA	0.030	第一次産業就業者総数NA	0.061
宿泊業、飲食サービス業従業者数NA	0.058	金融業、保険業事務所数NA	0.030	核家族世帯数NA	0.061
不動産業従業者数NA	0.058	不動産業事務所数NA	0.030	共同住宅世帯数NA	0.061
金融業、保険業従業者数NA	0.057	宿泊業、飲食サービス業事務所数NA	0.030	一戸建世帯数NA	0.061
卸売業、小売業従業者数NA	0.057	教育、学習支援業事務所数NA	0.030	居住期間出生時から総数NA	0.061
公務事務所数NA	0.057	医療、福祉事務所数NA	0.030	居住期間10年未満総数NA	0.061
医療、福祉事務所数NA	0.057	公務事務所数NA	0.030	五人以上世帯数NA	0.061
教育、学習支援業事務所数NA	0.057	卸売業、小売業従業者数NA	0.030	四人世帯数NA	0.061
宿泊業、飲食サービス業事務所数NA	0.057	金融業、保険業従業者数NA	0.030	三世帯数NA	0.061
不動産業事務所数NA	0.057	不動産業従業者数NA	0.030	二世帯数NA	0.061
金融業、保険業事務所数NA	0.057	宿泊業、飲食サービス業従業者数NA	0.030	一世帯数NA	0.061
卸売業、小売業事務所数NA	0.057	教育、学習支援業従業者数NA	0.030	65歳以上人口総数NA	0.061
第2次産業従業者数NA	0.057	医療、福祉従業者数NA	0.030	15.64歳人口総数NA	0.061
第2次産業事務所数NA	0.057	公務従業者数NA	0.030	0.14歳人口総数NA	0.061
田1	0.021	他農用地0	0.030		
他農用地1	0.015	田0	0.029		
		0.14歳人口総数NA	0.029		
		15.64歳人口総数NA	0.029		

土地利用数値化指標1

● 優勢トピック（U行列）

$$\hat{k}_i = \left\{ k \mid k \text{は, } \max(u_{i1}, \dots, u_{ik}, \dots, u_{iK}) \text{ を満たすトピック} \right\}$$

- トピックモデルは、本来ソフトクラスタリングモデル=どれか一つのトピックに分類されない。
- 地図を描くための方便

● 都市化推移指標（U行列）

$$C_i = \sum_{k=1}^K w_k (u_{ik} - u_{i'k})$$

- i' は同地点異時点の地点・地理トピックベクトル
- w_k は、集積程度に応じたトピック k の重み：高集積=大きな値→低集積=小さな値。

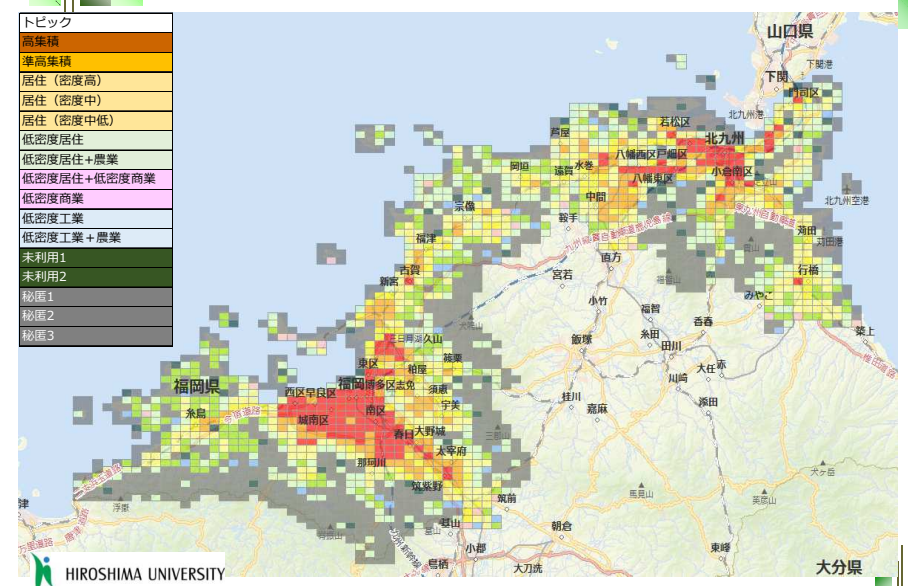
土地利用数値化指標2

● 建物開発率（U行列）

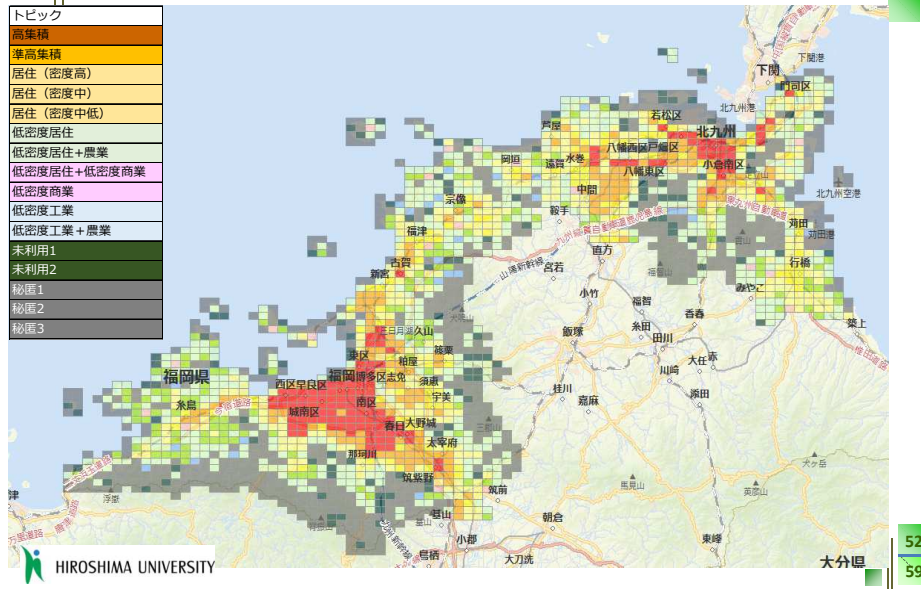
$$B_k = \frac{\sum_{i \in R} p_{ki}^{t+1} b_i^{t+1}}{\sum_{i \in R} p_{ki}^t b_i^t} \quad p_{ki}^t = \frac{u_{ik}^t}{\sum_k u_{ik}^t}$$

- p は、時点 t 、メッシュ i の地理トピック k のシェア
- b は、時点 t 、メッシュ i の建物用地面積

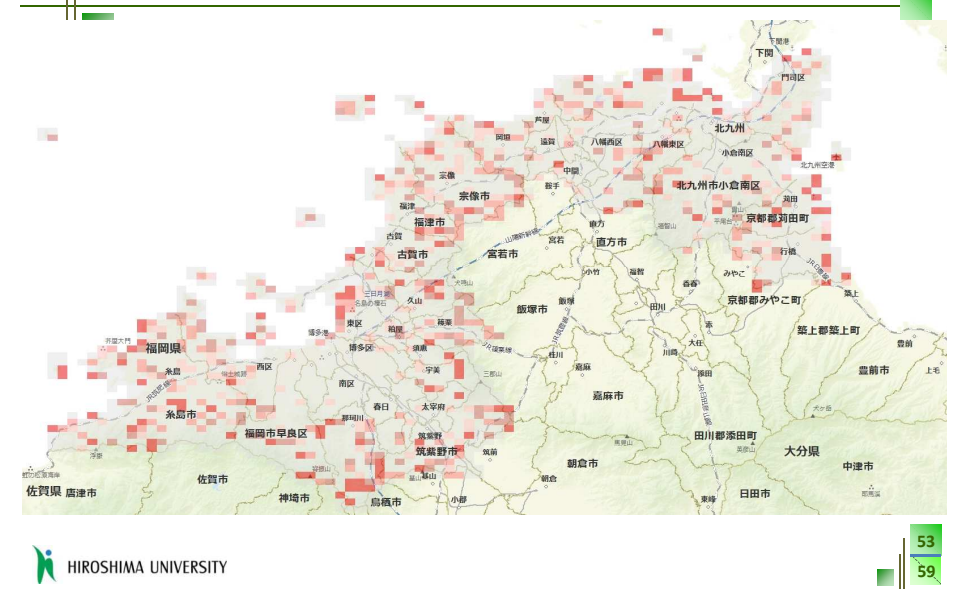
U行列：2000年の優勢トピック



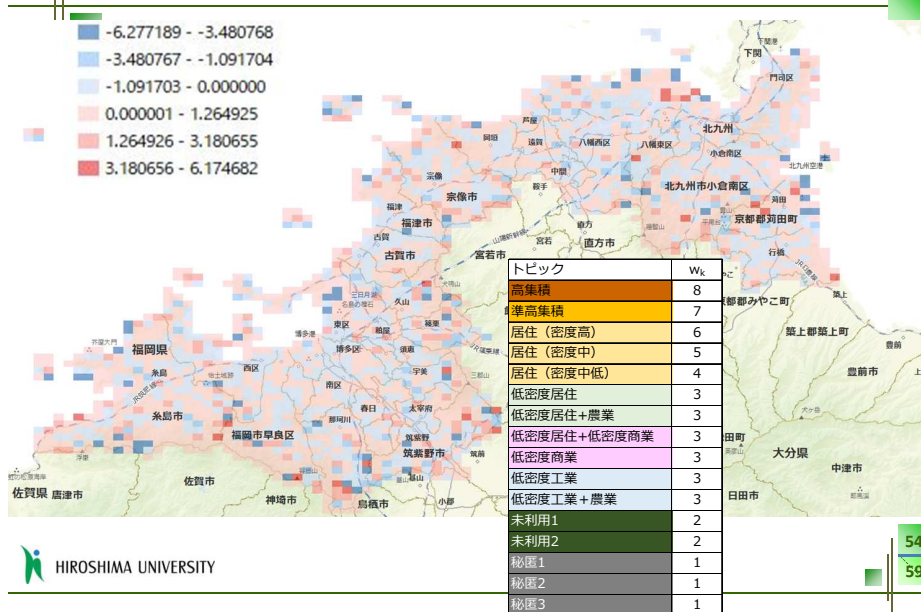
U行列：2010年の優勢トピック



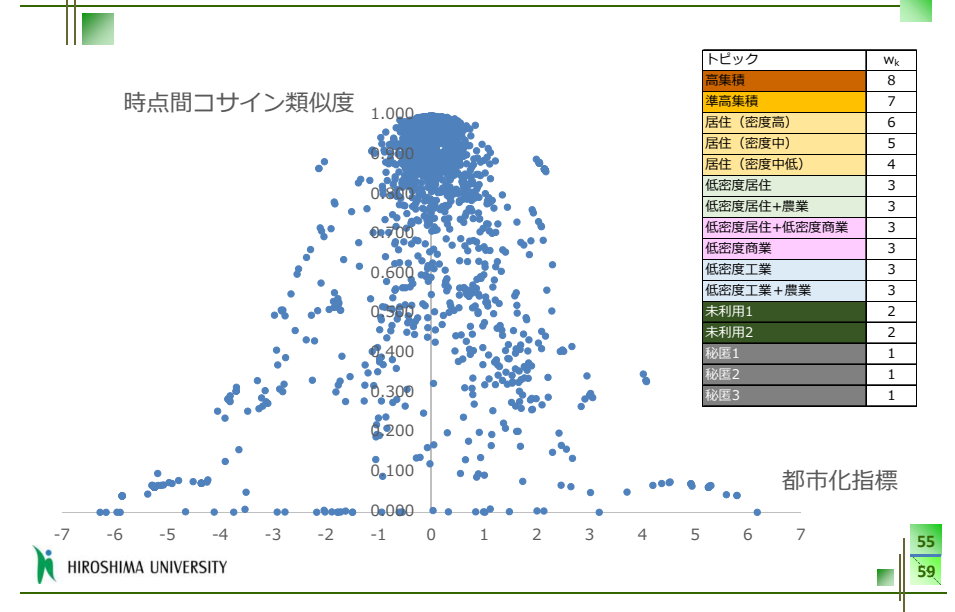
地点別トピックのコサイン類似度



都市化推移指標



都市化指標と時点間コサイン類似度



各都市圏のトピック構成比 (%)

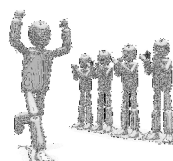
トピック	福岡都市圏			北九州市圏		
	2000年	2010年	差分	2000年	2010年	差分
高集積	6.6	7.2	0.6	5.4	4.7	-0.7
準高集積	6.6	6.9	0.3	7.5	7.9	0.4
居住（密度高）	4	4.2	0.2	5.6	6.5	0.9
居住（密度中）	4.9	5.2	0.3	6.6	6.6	0
居住（密度中低）	4.2	4	-0.2	6.9	7.1	0.2
低密度居住	8.9	11.7	2.8	10.2	12.1	1.9
低密度居住+農業	6.7	5.6	-1.1	5.3	3.8	-1.5
低密度居住+低密度商業	4.9	4.7	-0.2	5.3	5.4	0.1
低密度商業	3.9	4.4	0.5	5.1	5.7	0.6
低密度工業	4.8	4.9	0.1	5.3	5.3	0
低密度工業+農業	6	5.6	-0.4	6.9	6	-0.9
未利用1	4.5	5	0.5	4.4	4.9	0.5
未利用2	1	3.8	2.8	1.7	2.7	1
秘境1	5.8	5.4	-0.4	3.6	3.3	-0.3
秘境2	18.8	17.2	-1.6	14.2	13.3	-0.9
秘境3	8.4	4.2	-4.2	6	4.7	-1.3
	100	100		100	100	

建物用地の伸び (km²)

トピック	福岡都市圏			北九州市圏		
	2000年	2010年	伸び率	2000年	2010年	伸び率
高集積	69.8	87.2	1.25	31.7	31.5	0.99
準高集積	54.5	72.1	1.32	38.9	48.5	1.25
居住（密度高）	24.1	36.2	1.50	22.6	33.9	1.50
居住（密度中）	22.9	33.7	1.47	19.8	25.5	1.29
居住（密度中低）	1.4	15.6	11.14	14.3	20.2	1.41
低密度居住	7.1	13.6	1.92	7.1	11.2	1.58
低密度居住+農業	9.4	11.6	1.23	5.9	6.2	1.05
低密度居住+低密度商業	14.4	16.7	1.16	9.8	12.1	1.23
低密度商業	11.4	17.8	1.56	12.9	17.2	1.33
低密度工業	6.7	9.3	1.39	6.6	10.2	1.55
低密度工業+農業	7	8.1	1.16	5.9	6.6	1.12
未利用1	2.5	3.2	1.28	2.7	3.3	1.22
未利用2	0.4	1.2	3.00	1.3	2.4	1.85
秘境1	1.2	1.2	1.00	1.1	0.8	0.73
秘境2	0.8	0.6	0.75	1.1	3.8	3.45
秘境3	2	0.6	0.30	3.3	5.4	1.64
	235.6	328.7	1.40	185	238.8	1.29

地理トピックモデルの成果

- 地理情報+主成分分析に関する長年の違和感が解消した。
- 独立成分分析への苛立ち（分解の面白さ ⇔ 成分と負荷量の負値）も解消した。
- トピック分析（50周年幹事の宿題）に答えつつ、作成したアルゴリズムの転用に成功した。
- 教科書を首ったけで読み解いてアルゴリズムを自作する、しんどい作業が報われた。
- 連続データの離散化による**属性数インフレ**という暴挙を試みたが、見込み通り計算ができた。
- 土地利用→社会経済活動集積量の推移を可視化できて満足だった。
- 都市計画方面への応用を考えられそう。



サンプル分類問題の“解呪法”

- データの属性に基づいてサンプル分類するとき、全ての属性を同じ重みで評価することに疑問があった。
- トピックモデルによる“解呪”の仕組み
 - Dirichlet分布=単体の模式図によると、この分布でパラメータをうまく選べば、いくつかの属性の寄与を、ほぼ完全に無視できる。
 - つまり分類グループを構成する上で重要な=都合の良い属性だけを眺めて特徴を把握している=V行列。
 - 隠れ変数qの集計から合成されるU行列は、特徴と整合的なソフトクラスタリングを与える。
- 「おっさん」でも、数学の勉強は大事。

