

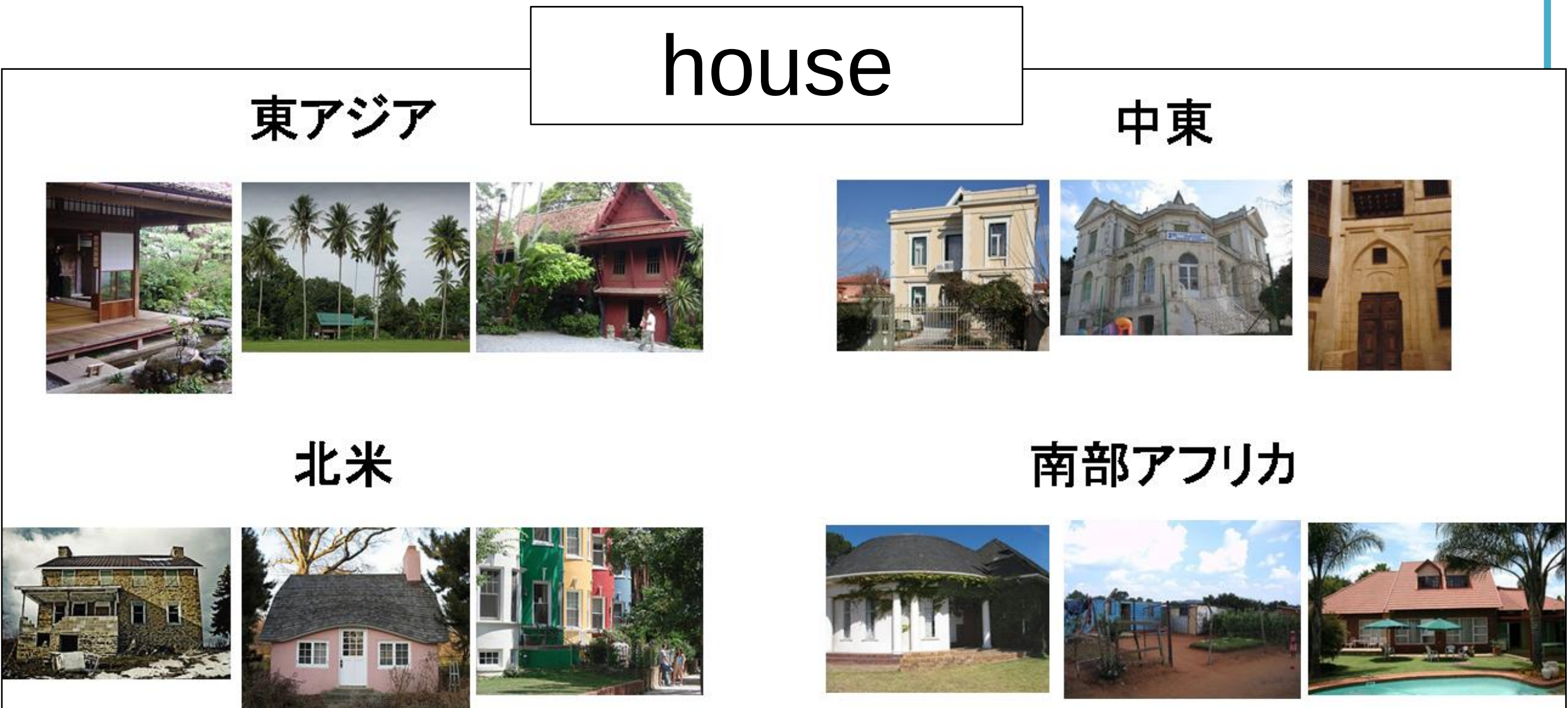
背景と目的

精度の高い一般物体認識のためには学習データの質が重要

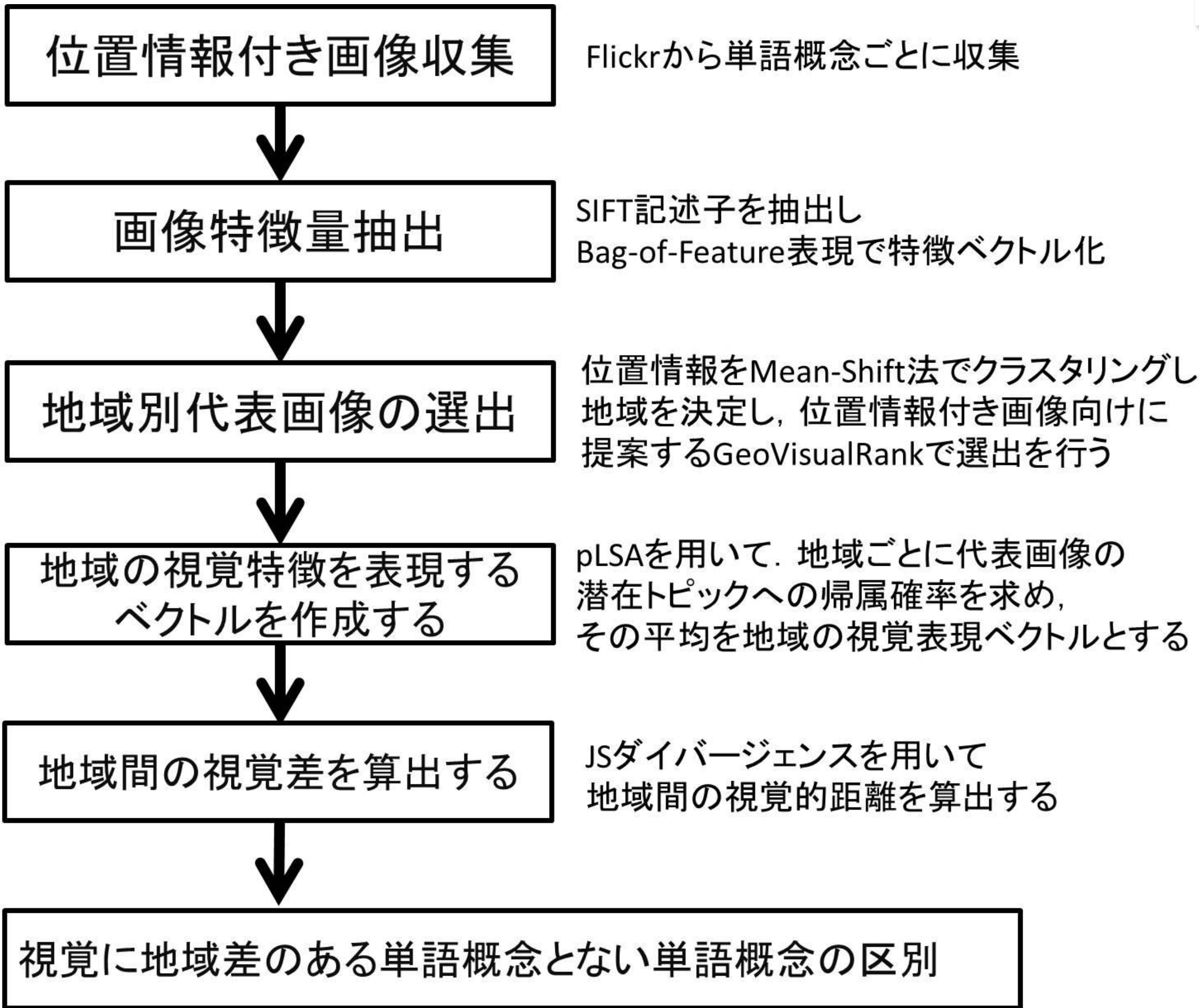
位置情報付き画像がWeb上で増加

◆大量の位置情報付き画像から単語概念に対応する「**視覚の地域性**」を定量化し、発見する

- 地域ごとの「代表的な視覚特徴」を自動抽出
- 地域別のデータセットを作成して、認識精度向上へ



手法概要



1. 画像収集

- 単語概念ごとに最大2000枚
- 同一ユーザからは最大20枚に制限

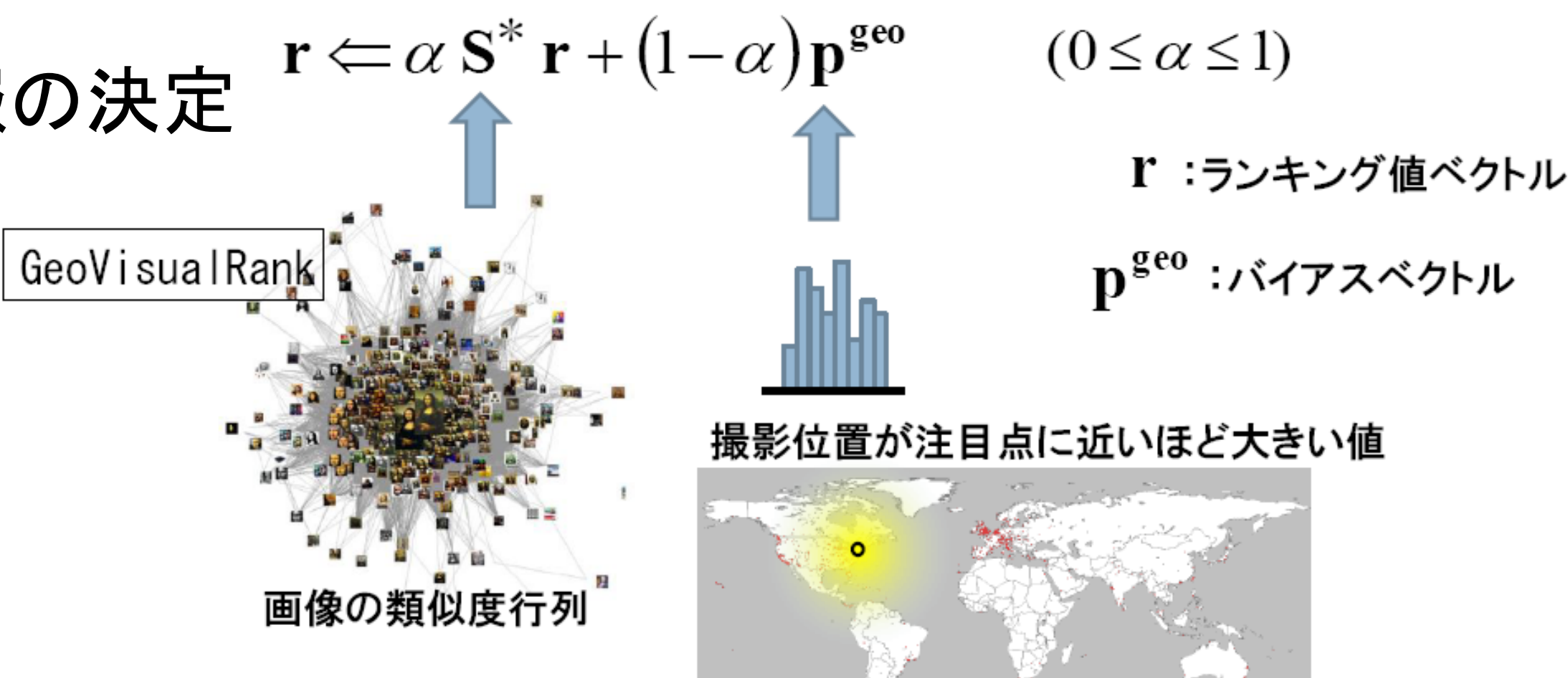
2. 画像特徴抽出

- SIFT
- Bag-of-Features表現
- 500次元

3. 地域別画像の選出

- クラスタリングによる位置情報の決定
- GeoVisualRankによる評価
- 位置情報のバイアスペクトルによる補正

- クラスタ中心に近いほど
ランキング値をより大きく



4. 表現ベクトル作成

- pLSAにより潜在トピックへの帰属ベクトルへ変換

$$P(z_k | R) = \frac{1}{N} \sum_{i=1}^N P(z_k | d_i^{\text{TOP}})$$

R : 地域
 z_k : 潜在トピック
 d_k^{TOP} : ランキング上位*i*番目の画像

5. 距離計算

- JSダイバージェンス
- 閾値以下なら関係性あり

$$D_{KL}(P \parallel Q) = \sum_i P(i) \log \frac{P(i)}{Q(i)}$$
$$D_{JS}(P \parallel Q) = \frac{D_{KL}(P \parallel (P/2 + Q/2))}{2} + \frac{D_{KL}(Q \parallel (Q/2 + P/2))}{2}$$

6. 視覚の地域性の推定

- 単語概念内で他地域との視覚的距離を計算

➢単語概念Cの地域性の
指標値VisualRegionality
を計算

$$VR(C) = \max_{R_i \in R_C} \left\{ \frac{1}{|R_C| - 1} \sum_{j=1}^{|R_C|} D_R(R_i \parallel R_j) \right\}$$

実験・結果

実行結果 公開用URL

<http://mm.cs.uec.ac.jp/kawakuh/visualregionality/>

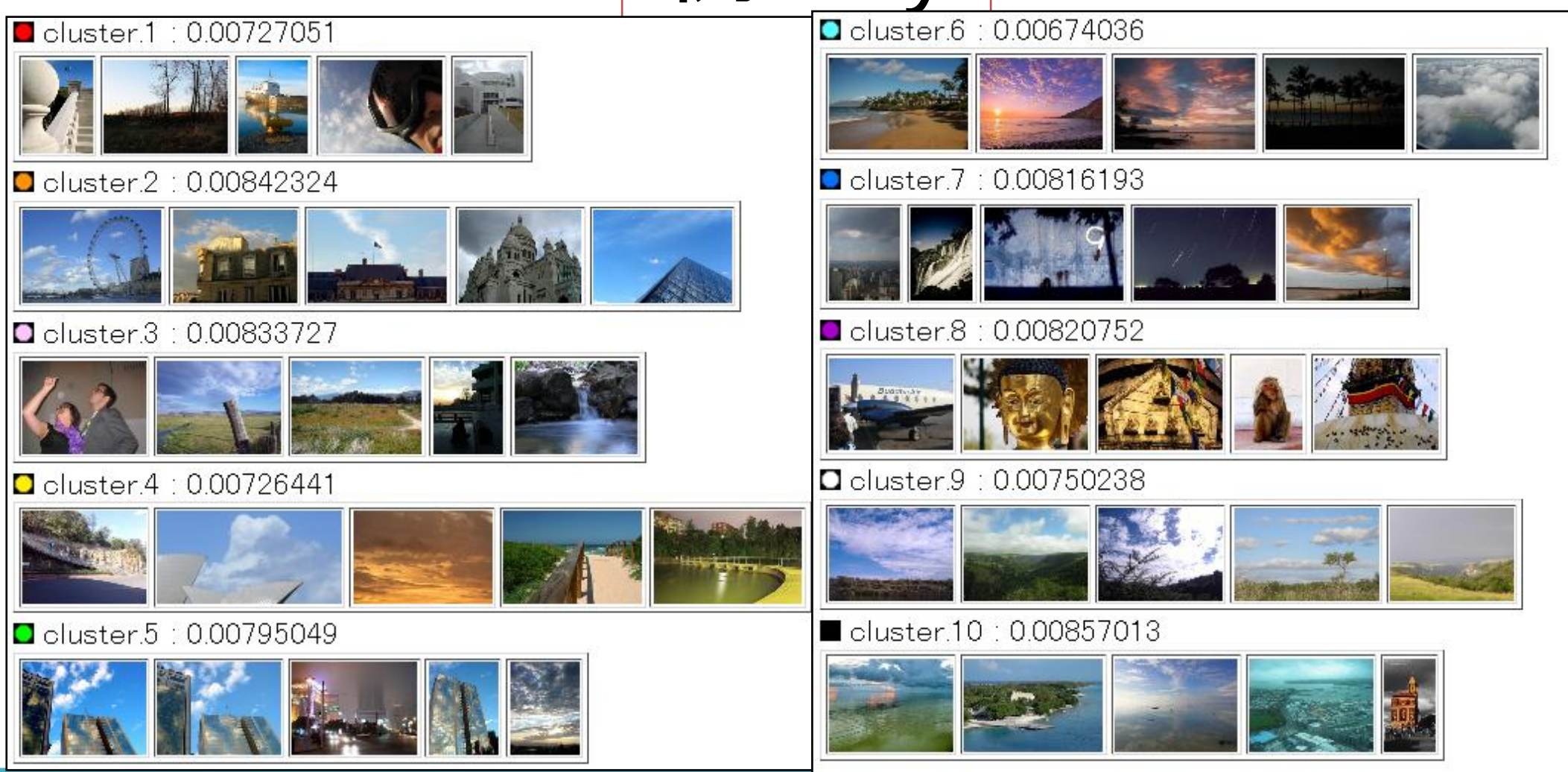
□ 名詞250語、形容詞100語

パラメータ	用いた値のリスト
Mean-Shift法での半径パラメータ(km)	50, 100, 250, 500, 750, 1000
GeoVisualRankでのバイアスパラメータα	0.7, 0.75, 0.8, 0.85, 0.9, 0.95
ランキング上位何枚を代表画像とするか	5, 10, 20, 50

例: doragon



例: sky



例: building

まとめ

- ◆ 単語概念の視覚の地域性を定量化する手法を提案
- ◆ ほぼ同一の複数の画像で学習した場合、影響が生じる

- ◆ Flickrの投稿IDや位置情報、撮影日時から重複・類似画像を省く必要