



知的情報処理システム特論 I 第7回

二宮 崇

今日の講義の予定

- 生成モデルと識別モデル

- 確率的識別モデル

- 構造予測のための多クラスロジスティック回帰
- 条件付き確率場 (CRF)
- Linear-chain CRF

- 教科書

- 高村大也. 言語処理のための機械学習. コロナ社.
- Yusuke Miyao (2006) From Linguistic Theory to Syntactic Analysis: Corpus-Oriented Grammar Development and Feature Forest Model, Ph.D Thesis, University of Tokyo
- Cristopher M. Bishop “PATTERN RECOGNITION AND MACHINE LEARNING” Springer, 2006



Generative Model and Discriminative Model

生成モデルと識別モデル



生成モデルから識別モデルへ

- 生成モデル

- HMM

- 状態遷移による生成
 - 状態遷移の確率は一つ前の状態のみに依存
 - 任意の文 x と状態列 q に対する同時確率 $p(x, q)$ を計算できる



生成モデルの改良の先に

- 条件部を詳細化

- HMM: $p(q_j | q_{j-1}, q_{j-2})$

- 素性(特徴)という考え方

条件部をどんどんリッチにして線形補間をとればよいのでは？ → HMMの状態 x の確率 $p(x) = p(x | x \text{ 周辺の状況})$



生成モデルの問題点

- 生成モデルの問題点

- 独立性の仮定

- HMMは一つ前の状態にのみ依存する(マルコフ性)
 - 二つ前の状態と今の状態は何らかの関係があるかもしれないけど、独立した事象として扱われている。

- 例: 左辺("I have a")の方が右辺("I have a pen")より不自然なのに必ず確率が高くなってしまう。

$$p("I")p("have")p("a") > p("I")p("have") p("a")p("pen")$$



生成モデルの問題点

- 生成モデルでの予測

$$\hat{y} = \operatorname{argmax}_y p(y|x; \theta) = \operatorname{argmax}_y p(x, y; \theta)$$

- しかし、

$$p(x, y; \theta) = p(y|x; \theta)p(x; \theta)$$

であるから $p(x; \theta)$ が計算できてしまう分、冗長なモデルになっている。間接的に分類問題を解いている。

- 独立性を仮定した事象に分解しないと同時確率を計算することができない



識別モデル

- 識別モデル

直接

$$\hat{y} = \operatorname{argmax}_y p(y|x; \theta)$$

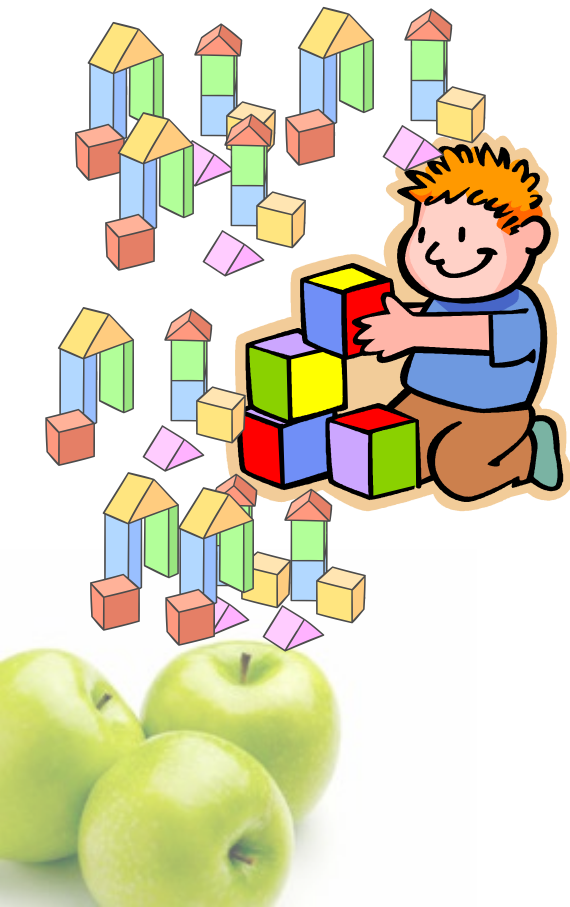
を解く

- 独立な事象を仮定しない
- 「条件部の確率」をモデルにいけない

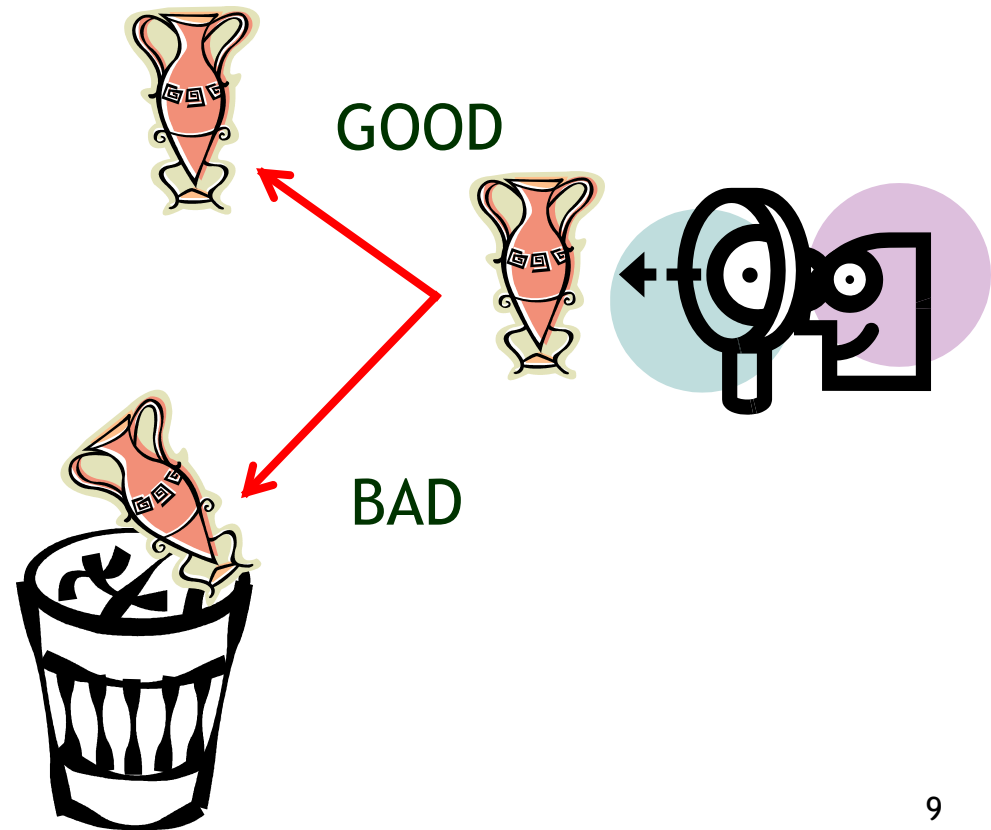


生成モデルと識別モデル (イメージ)

生成モデル



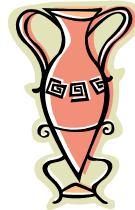
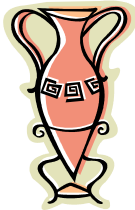
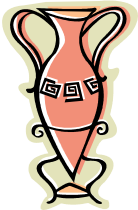
識別モデル



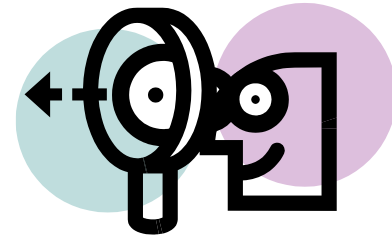
識別するための訓練

- 教師付学習

- 良い例と悪い例を与えて、どこに注目すればより良く識別出来るのか学習



good examples



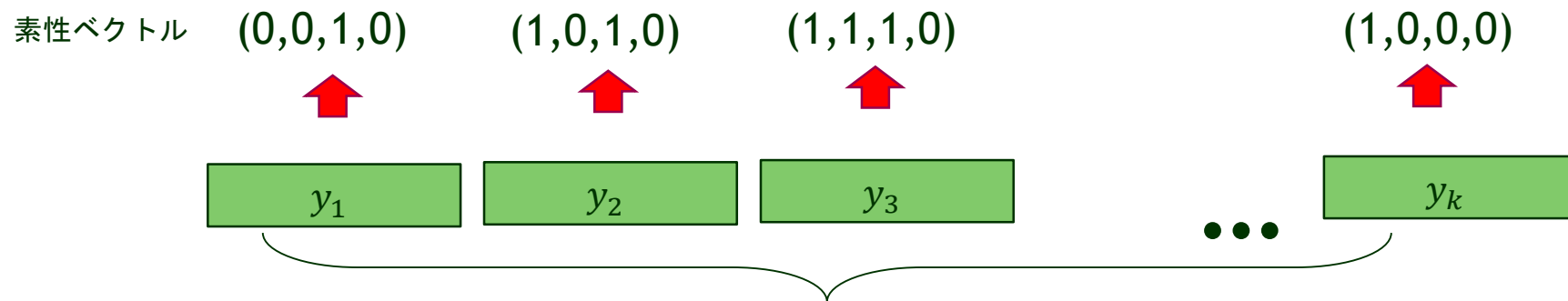
bad examples



識別モデル

- 品詞解析の例

- 入力 x = “A blue eye girl with white hair and skin walked”



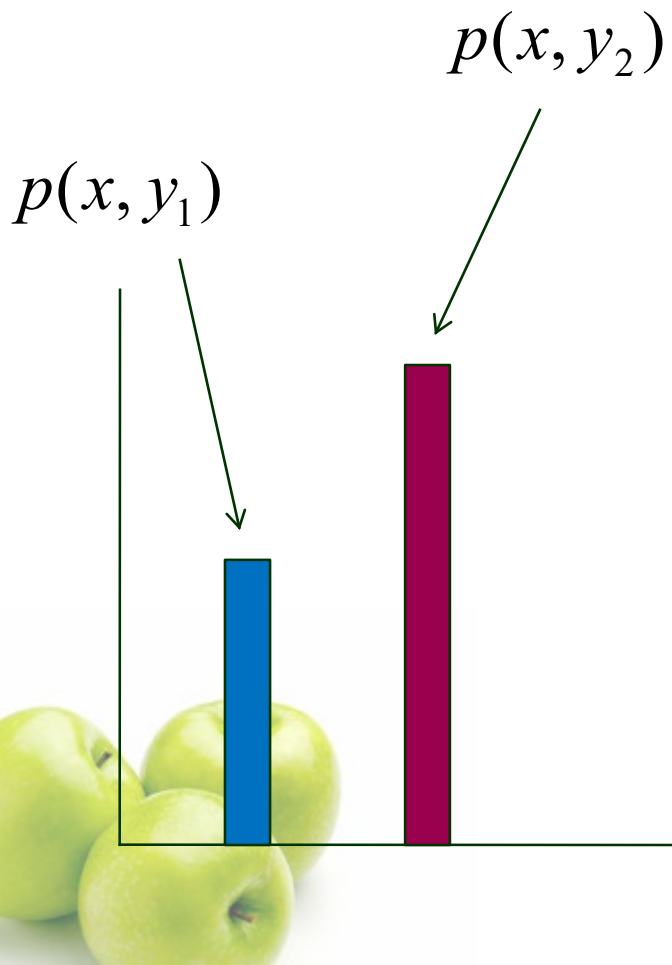
x に対する全ての可能な品詞列集合

$p(y_3|x)$ は $\{y_1, y_2, y_3, \dots, y_k\}$ から y_3 を選択する確率

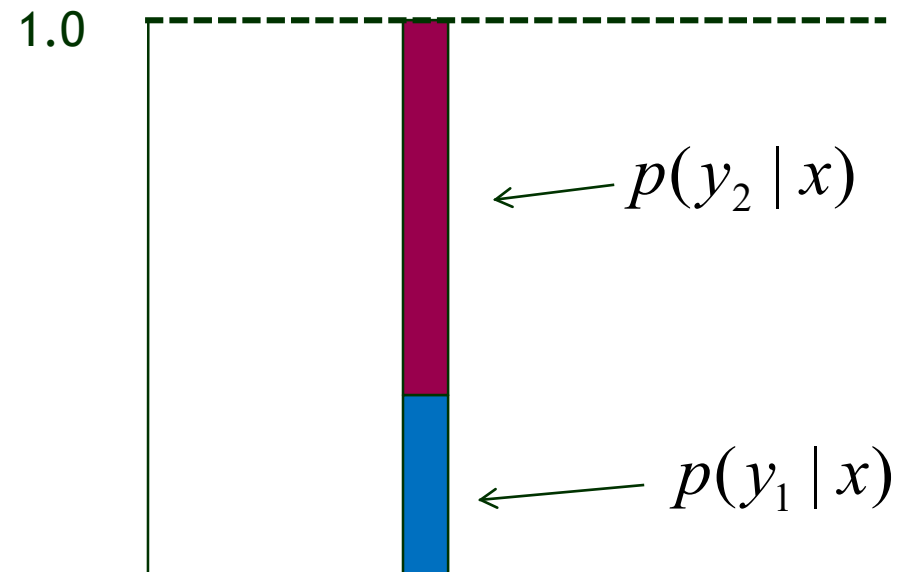


生成モデルと識別モデル

生成モデル



識別モデル



確率的識別モデル



問題設定

- 入力: x (任意の入力)
- 出力: $y \in Y(x)$
 - $Y(x)$: 入力 x に依存して列挙できる有限の無向グラフ集合 ($Y(x) \subset G$)
- 訓練データ
 - $(x_i, y_i)_{i=1, \dots, n}$
 - 例
 - x は競馬情報、 y は1位の馬 (y はラベル)
 - x は文、 y は x に対する正解の品詞列 (y はラベル列)
 - x は文、 y は x に対する正解の構文木 (y は木グラフ)
- 推論
 - ある未知の入力 x に対する出力 y の予測



素性関数

- 入力と出力から特徴を抽出する素性関数(feature function)を複数定義

$$\phi_j(x, y) \quad j = 1, \dots, D$$

- 注意
 - 人手で設計するか、素性関数を学習する仕組み(深層学習など)が必要
 - x だけでなく y からも特徴を抽出できる (y から特徴を抽出しなければならない)

- 素性ベクトル(feature vector)

$$\phi(x, y) = (\phi_1(x, y), \phi_2(x, y), \dots, \phi_D(x, y))$$



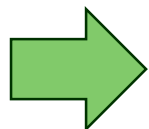
全体の流れ(1/2)

- パラメータ推定

- 各素性 ϕ_j に対する重み w_j を学習

訓練データ

入力	出力
x_1	y_1
x_2	y_2
...	...
x_n	y_n



素性ベクトル
$\phi(x_1, y_1) = (\phi_1(x_1, y_1), \phi_2(x_1, y_1), \dots, \phi_D(x_1, y_1))$
$\phi(x_2, y_2) = (\phi_1(x_2, y_2), \phi_2(x_2, y_2), \dots, \phi_D(x_2, y_2))$
...
$\phi(x_n, y_n) = (\phi_1(x_n, y_n), \phi_2(x_n, y_n), \dots, \phi_D(x_n, y_n))$



学習

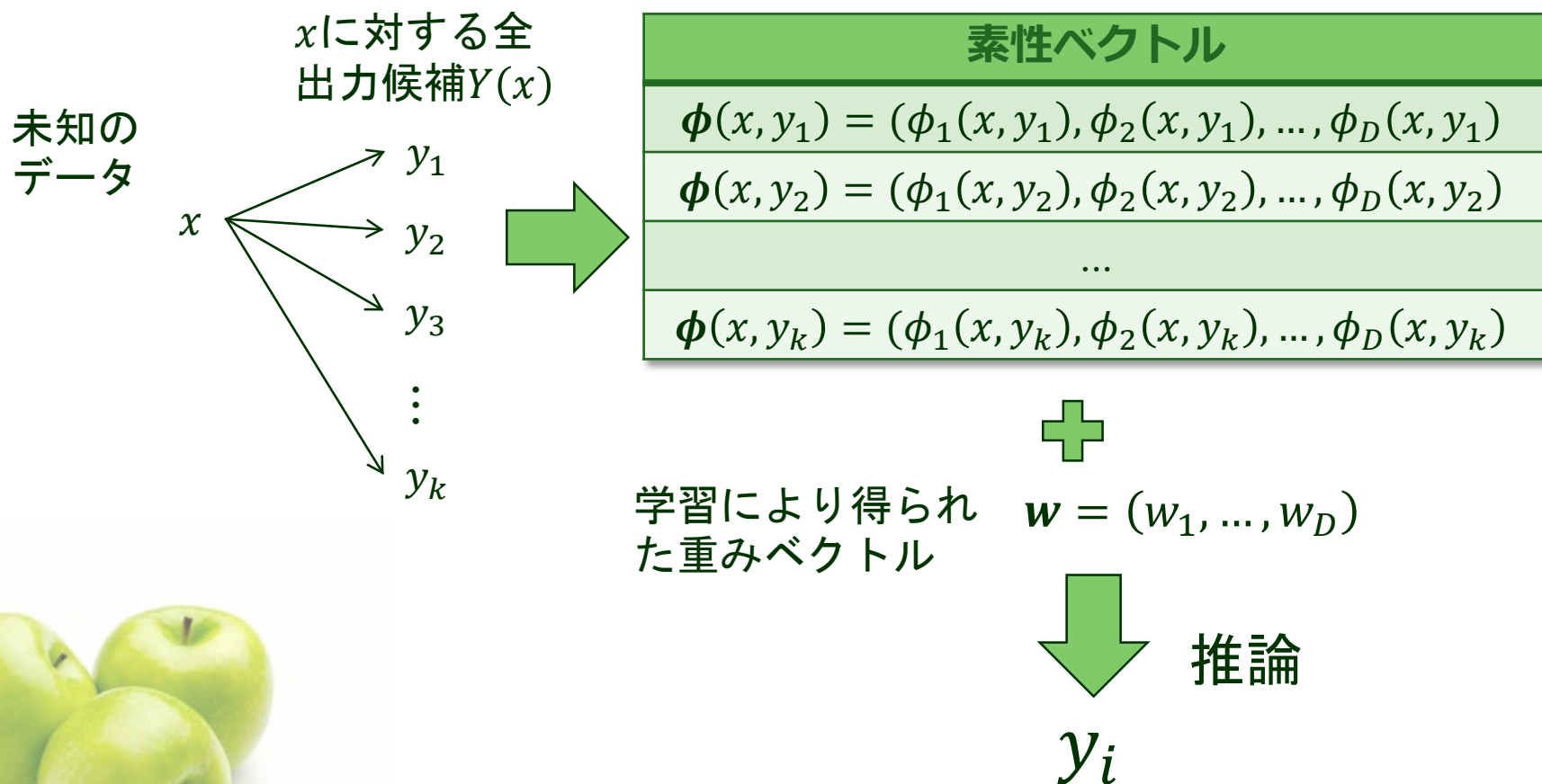
$$\mathbf{w} = (w_1, w_2, \dots, w_D)$$



全体の流れ(2/2)

● 推論

- 未知のデータ x に対する出力 y の推定



確率的識別モデル

- 構造予測のための多クラスロジスティック回帰

$$p(y|x) = \frac{1}{Z(x)} \exp(\underset{\substack{\text{重みベクトル} \\ \downarrow}}{\mathbf{w}^T} \underset{\substack{\text{素性ベクトル} \\ \nwarrow}}{\boldsymbol{\phi}(x, y)})$$

分配関数
(Partition function)

ただし、

$$Z(x) = \sum_{y' \in Y(x)} \exp(\mathbf{w}^T \boldsymbol{\phi}(x, y'))$$



パラメータ推定

- 勾配法を用いてパラメータを求める
 - 訓練データに対する損失関数(負対数尤度)の勾配を求めれば良い
- 訓練データに対する損失関数(負対数尤度)

$$\begin{aligned} & -\log \prod_{i=1}^n p(y_i | x_i) \\ &= -\sum_{i=1}^n \log p(y_i | x_i) \\ &= -\sum_{i=1}^n \log \left\{ \frac{1}{Z(x_i)} \exp(\mathbf{w}^T \boldsymbol{\phi}(x_i, y_i)) \right\} \\ &= \sum_{i=1}^n \log Z(x_i) - \sum_{i=1}^n \mathbf{w}^T \boldsymbol{\phi}(x_i, y_i) \end{aligned}$$

$$\begin{aligned} \log ab &= \log a + \log b \\ \log_e \exp(x) &= \log_e e^x = x \end{aligned}$$

パラメータ推定

- 損失関数

$$L(\mathbf{w}) = \sum_{i=1}^n \log Z(x_i) - \sum_{i=1}^n \mathbf{w}^T \boldsymbol{\phi}(x_i, y_i)$$

- 勾配

$$\begin{aligned} \frac{\partial L(\mathbf{w})}{\partial w_j} &= - \sum_{i=1}^n \phi_j(x_i, y_i) + \sum_{i=1}^n \frac{1}{Z(x_i)} \frac{\partial Z(x_i)}{\partial w_j} \\ &= - \sum_{i=1}^n \phi_j(x_i, y_i) + \sum_{i=1}^n \frac{1}{Z(x_i)} \sum_{y \in Y(x_i)} \phi_j(x_i, y) \exp(\mathbf{w}^T \boldsymbol{\phi}(x_i, y)) \\ &= - \sum_{i=1}^n \phi_j(x_i, y_i) + \sum_{i=1}^n \sum_{y \in Y(x_i)} p(y|x_i) \phi_j(x_i, y) \\ \nabla L(\mathbf{w}) &= \left(\frac{\partial L(\mathbf{w})}{\partial w_1}, \dots, \frac{\partial L(\mathbf{w})}{\partial w_D} \right) \end{aligned}$$



識別モデルのいいところ

- **独立性を仮定していない**
 - (戦略として) 思いつく限りいろんな素性をいれる
 - 訓練データに対してより良い予測ができる
 - 逆にoverfittingする可能性がある
 - c.f. ガウス事前分布によるMAP推定でoverfittingを緩和



素性に関する注意

- 素性の組み合わせ

- 素性同士の共起情報が別素性として自動的に組み込まれるわけではない
 - 一つ前の単語が”the”で、今見ている単語が”flies”。一つ前の品詞が動詞で、今見ている品詞が冠詞など。
 - SVMでは多項式カーネルを用いるとこれらの組み合わせを自動的に計算したことになる
- 素性の組み合わせを手で指示する。自動的に組み合わせる研究もある。
c.f. 素性選択



系列解析のための確率的識別モデル

LINEAR-CHAIN CRF



系列解析のための確率的識別モデル

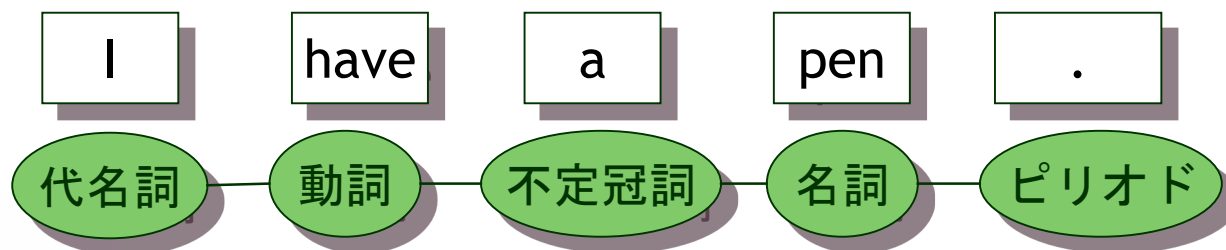
Linear-chain CRF

- CRF (Conditional Random Fields, 条件付確率場)

- 構造予測のための確率的識別モデル $p(y|x)$
- 出力 y が無向グラフ
- 素性を各ノードの近傍で定義

- Linear-chain CRF

- グラフのノードが線形につながったCRF (系列解析のためのCRF)
- 例えば、品詞解析の場合、入力 x が文で、出力 y が単語-品詞列となる。単語と品詞のペアがノードになる。素性を各ノードの近傍で定義する。



x = "I have a pen."

y = "I-代名詞 have-動詞 a-不定冠詞 pen-名詞 .-ピリオド"

に対し $p(y|x)$ を求める



Linear-chain CRF: 素性テンプレート

- 素性テンプレート

- どのような素性を採用するか表記したもの
- 訓練データに対し素性テンプレートを適用し、実際に出現した値だけを素性として採用

- 素性テンプレートの例

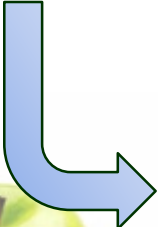
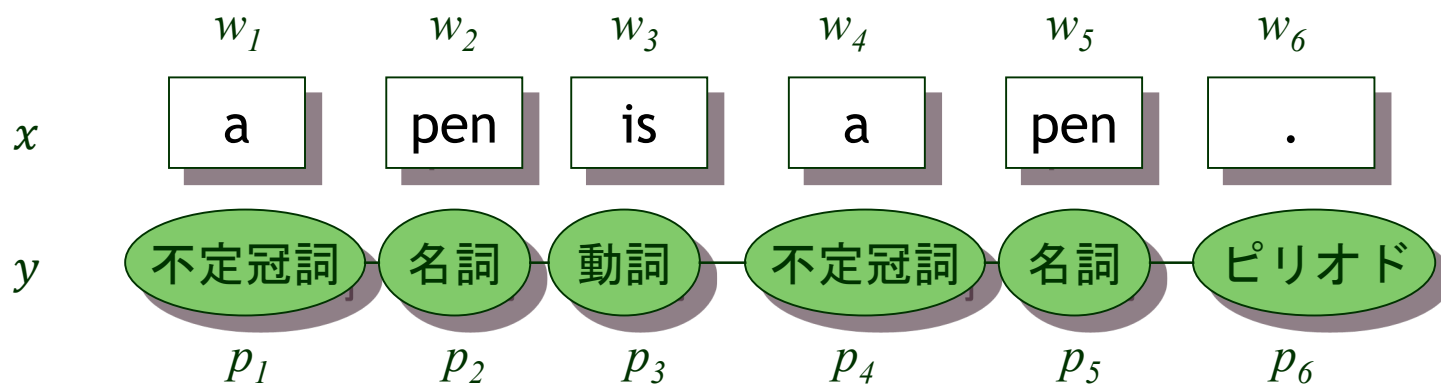
- i 番目のノードを c_i とする
- i 番目の品詞を p_i とする
- i 番目の単語を w_i とする
- HMMと同様の素性の場合
 - 各ノード c_i に対し、 $p_{i-1}p_i$ [最大で品詞数 $\times$ 品詞数の次元ができる]
 - 各ノード c_i に対し、 w_ip_i [最大で単語数 $\times$ 品詞数の次元ができる]
- 他にもいろんな組合せの素性を追加できる
 - 各ノード c_i に対し、 $w_{i-1}w_ip_i$
 - 各ノード c_i に対し、 $p_{i-2}p_{i-1}p_i$



Linear-chain CRF: 素性テンプレートと素性

● 素性テンプレートの例

- HMMと同様の素性の場合
 - 各ノード c_i に対し、 $p_{i-1}p_i$
 - 各ノード c_i に対し、 $w_i p_i$

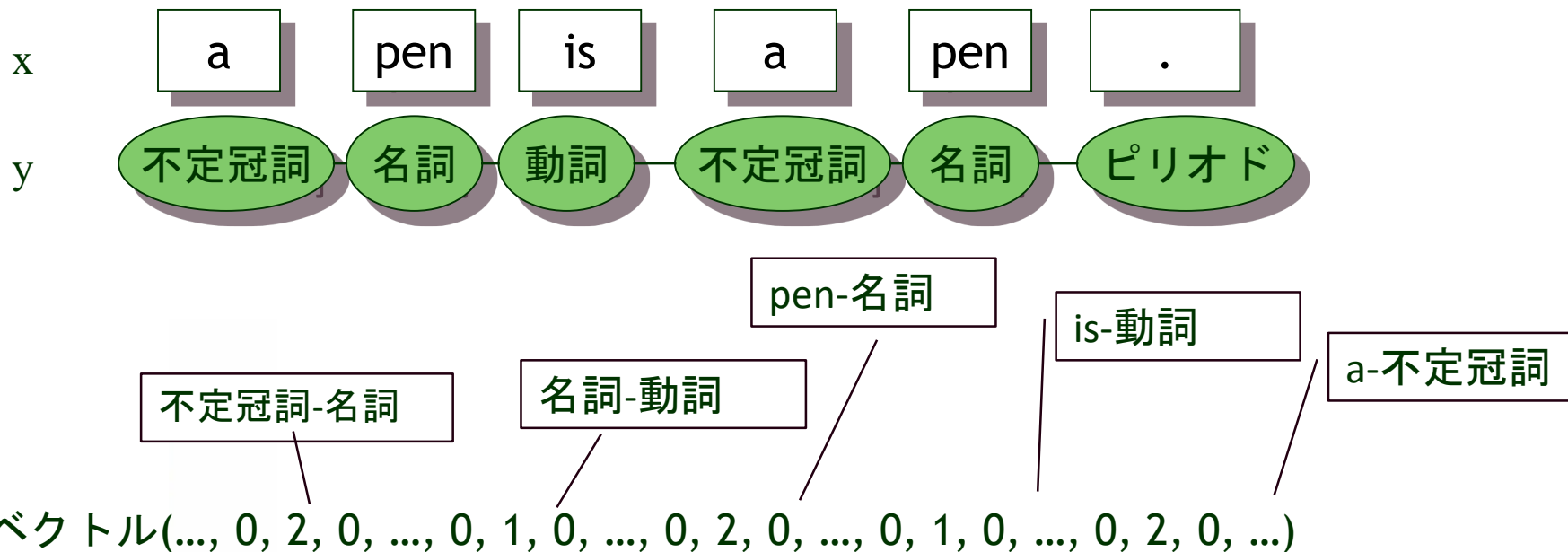


$p_{i-1}p_i$ から生成される素性	$w_i p_i$ から生成される素性
不定冠詞-名詞	a-不定冠詞
名詞-動詞	pen-名詞
動詞-不定冠詞	is-動詞
名詞-ピリオド	.-ピリオド

Linear-chain CRF: 全体の素性ベクトルと局所的素性ベクトルの関係

- 各素性関数の値は1文中に各素性が発火(出現)した回数

$$\phi_j(x, y) = \sum_c \phi_j(c) \quad (\forall j) \quad \Leftrightarrow \quad \Phi(x, y) = \sum_c \Phi(c)$$



Linear-chain CRF: 全体のスコアと局所的なスコアの関係

- 全体のスコア=各ノード c に対するスコアの積

$$p(y|x) = \frac{1}{Z} \exp(\mathbf{w}^T \boldsymbol{\phi}(x, y)) = \frac{1}{Z} \exp\left(\mathbf{w}^T \sum_c \boldsymbol{\phi}(c)\right) = \frac{1}{Z} \exp\left(\sum_c \mathbf{w}^T \boldsymbol{\phi}(c)\right)$$
$$= \frac{1}{Z} \prod_c \exp(\mathbf{w}^T \boldsymbol{\phi}(c))$$

文全体の素性ベクトル: (2,0,3,1,1)

$$e^{w_1 \cdot 2 + w_2 \cdot 0 + w_3 \cdot 3 + w_4 \cdot 1 + w_5 \cdot 1}$$

掛算

$$e^{w_1 \cdot 1 + w_2 \cdot 0 + w_3 \cdot 1 + w_4 \cdot 1 + w_5 \cdot 0}$$

$$e^{w_1 \cdot 0 + w_2 \cdot 0 + w_3 \cdot 1 + w_4 \cdot 0 + w_5 \cdot 0}$$

$$e^{w_1 \cdot 1 + w_2 \cdot 0 + w_3 \cdot 1 + w_4 \cdot 0 + w_5 \cdot 1}$$

(1,0,1,1,0)

w_1

p_1

(0,0,1,0,0)

w_2

p_2

(1,0,1,0,1)

w_3

p_3

$$e^{a+b} = e^a e^b$$



Linear-chain CRF: スコアの分解とHMMとの対応

- 単語-品詞列の確率

$$p(y|x) = \frac{1}{Z(x)} \exp(\mathbf{w}^T \boldsymbol{\phi}(x, y)) = \frac{1}{Z(x)} \prod_c \exp(\mathbf{w}^T \boldsymbol{\phi}(c))$$

品詞列集合全体のスコアの和 (前向きアルゴリズムで計算)

HMMの遷移確率と出力確率に対応

- 解析: ビタビアルゴリズムの適用

- 遷移確率と出力確率 $\Leftrightarrow$ ノードのスコア $\exp(\mathbf{w}^T \boldsymbol{\phi}(c))$

- 学習: 前向き後向きアルゴリズムの適用

- 遷移回数と出力回数 $\Leftrightarrow$ 素性値(素性の発火回数)



1次のLinear-chain CRF

- ラベル集合を Q 、入力を x 、出力を $q_1 q_2 \dots q_T$ とする。
- Linear-chain CRF

$$p(q_1 \dots q_T | x) = \frac{1}{Z(x)} \prod_{t=1}^T \exp\{\mathbf{w}^T \boldsymbol{\phi}(q_{t-1}, q_t, x, t)\}$$
$$Z(x) = \sum_{q'_1 \in Q, \dots, q'_T \in Q} \prod_{t=1}^T \exp\{\mathbf{w}^T \boldsymbol{\phi}(q'_{t-1}, q'_t, x, t)\}$$

ただし、 $Z(x)$ は分配関数、 w は ϕ に対する重み、 $q_0 = DUMMY$ (ダミーのラベル)とし、素性関数 ϕ の引数は次のように定義される。

ϕ (時刻 $t - 1$ のラベル, 時刻 t のラベル, 入力, 時刻)

- 入力(第三引数)と時刻(第四引数)を与えることによって、入力の情報を素性として用いることができる
 - 例えば、品詞解析の場合、ラベルが品詞であり、入力が単語列となり、時刻 t やその前後における単語を素性として用いることができる。

まとめ

- 生成モデルと識別モデル
- 確率的識別モデル
 - 構造予測のための多クラスロジスティック回帰
 - 条件付き確率場 (CRF)
 - Linear-chain CRF
- 今までの講義資料

<http://aiweb.cs.ehime-u.ac.jp/~ninomiya/iips/>

