

コンピュータビジョンにおける時系列パターン認識

川 嶋 宏 彰[†] 西 村 拓 一^{††}

本稿では、1980年代から盛んになってきたコンピュータビジョンにおける時系列パターン認識技術について概観する。特に非線形時間伸縮パターンの認識技術をパターンマッチングによる手法とモデルに基づく手法に分けて述べる。また、応用分野として表情認識および音響情報と視覚情報の統合に関して述べる。

Temporal Pattern Recognition in Computer Vision

KAWASHIMA, HIROAKI[†] and NISHIMURA, TAKUICHI^{††}

In this paper, we will survey researches on temporal pattern recognition which has been popular in 1980s in the field of computer vision. We will focus on dynamic time warping method categorizing it into pattern matching method and model based method. Furthermore, expression recognition and fusion of sound and image are introduced.

1. はじめに

1980年ごろから、カメラおよび計算機の能力向上に伴い、動画像を用いて対象の「動き」の理解を目指す研究が盛んになってきたこの動画像処理では、時間の変化とともに得られる連続した複数の画像系列、つまり時系列画像を扱う。時系列を扱うことで、対象の動きをも推定する必要が生じ問題が複雑となる一方、対象の切り出しや3次元形状復元に有利となる場合もある。2次元の時系列画像から3次元の物体の動きを推定する場合、まず、画像から変化部位を検出し、画像間の対応づけを行う。このとき、対象上の複数の点の対応づけが必要となるが、局所的な特徴だけでは、一意に決定できない事が多い。特に、物体の動きに伴って見え隠れする場合の対応づけは困難である。そこで、取り扱う世界および対象の形状および運動のモデルを導入する。これによって、特徴点の拘束条件を導入でき安定した動き推定が実現できる。例えば、剛体が運動している場合は、同一剛体の上の各点は共通した動きとなり、各点が融合したり分離しない。また、質量を持つため、速度や加速度の上限も存在する。これらの仮定により、物体の3次元運動を推定する研究が行われている¹⁰²⁾。

しかし、それまでの研究では、各時点での対象の位置、向き速度などを「計測する」ことが目的とされることが多かった。一方、計測した対象の位置などの時間的な変化パターンを認識することが必要となる場合もある。例えば、人のジェスチャを認識する場合では、各時点での人物の姿勢だけでなく、姿勢の時間的な変化を識別する必要がある。これは、長時間のジェスチャ動画像中から繰り返し出現する類似した動きを抽出し、対象人物の癖を把握するという課題の場合も同様である。

そこで、本稿では、動画像から得られる時系列データから、あるパターンを認識（時系列パターン認識）する研究のサーベイを行う。人間の動作認識についてのサーベイは、Pavlovicら¹⁰⁴⁾やAggarwalら¹⁰³⁾、Gavrila¹⁰⁵⁾が行っている。

コンピュータビジョンにおける時系列パターン認識の処理は次の3つのステップからなる。

- (1) 認識したい事象を的確に表す特徴量を取り出す特徴抽出
 - (2) 各カテゴリのサンプル集合を作成し、標準パターンやモデルを作る処理
 - (3) 新たに入力される時系列データを認識する処理
- 特徴抽出手法は、動画像の各フレームの画像から静的な特徴を抽出する手法や optical flow やそのエネルギー¹⁰⁶⁾、motion-history image¹⁰⁷⁾ など複数枚の画像を用い時間変化特徴を捉えた動的特徴を抽出する手法が用いられている。さらに、固有空間への射影、状態認識を行う場合もある。各カテゴリを認識するための事前知識は、DP マッチングのようなテンプレートマッチングの場合は標準パターン、HMM などの場合

[†] 京都大学大学院情報学研究科

Department of Intelligence Science and Technology,
Graduate School of Informatics, Kyoto University

^{††} 産業技術総合研究所情報技術研究部門

Information Technology Research Institute, National
Institute of Advanced Industrial Science and Technology (AIST)

はモデルと呼ばれる。認識手法では、多くの場合は特徴空間内での変動と時系列パターンの非線形時間伸縮を吸収することが要求される。これは、人物のジェスチャや動き、表情の認識問題などのように、時間的なスピードの変動や姿勢などの特徴量の変動が避けられない場合が多いからである。

ちなみに、出現順序が既知の複数の時系列パターン集合が存在する場合、一つの時系列パターンを認識できると、次に出現する時系列パターンを予測できる。また、パターンマッチングによる認識手法は、標準パターンをクエリ、認識する入力データをデータベースとすれば検索問題となる。つまり、認識問題は、検索や予測と関連するものである。ただし、検索や予測特有の問題に関しては本稿では取り扱わない。

以上の議論を踏まえ、時系列パターン認識を実現する手法を以下のように分類する。

- (1) 一般的なパターン認識による手法
- (2) 時系列パターン認識による手法
 - 軌道のマッチングに基づく手法
 - モデルに基づく手法

一般的なパターン認識では、動的特徴や一定時間区間の静的特徴を利用することで特徴抽出段階で時間情報を取り込んでいる。認識手法としてはテンプレートマッチング¹⁰⁷⁾やnearest neighbor¹⁰⁶⁾、ランダムサンプル手法の一つCONDENSATION法¹⁰⁸⁾などが初期のころに提案された。また、何らかの方法でパターン全体の時間伸縮率を求め、参照パターンと入力パターンの時間軸を一致させた後にゼロクロッシング、曲率などの特徴を抽出し比較する方法も研究されている。最近では、高次局所自己相関特徴¹⁰⁹⁾を時間方向に拡張した特徴を用いて異常行動の認識^{110),111)}やジェスチャ認識¹¹²⁾が提案されている。また、村瀬らは、一定時間区間の静的特徴のヒストグラムを用いて高速に検索するアクティブ探索法^{113),114)}を提案し、マルチメディアデータの検索¹¹⁵⁾を実現している。この検索手法は、前述のように複数個のクエリを登録しておく、時系列データを入力すれば、認識問題にも利用できる。

本稿では、時系列パターン認識による手法を、特徴空間内の軌道を示す標準パターンとのマッチングに基づく手法とHMMなどのモデルに基づく手法とに分類し、それぞれ2節、3節で研究動向を紹介する。また、4節にて2、3節で紹介できなかった応用研究について述べ、5節でまとめる。

2. 軌道のマッチングに基づく認識

本節では、軌道のマッチングに基づく認識手法として連続DPについて述べる。また、連続DPが得意とする検索や時系列データ中の類似区間検出について述べる。

2.1 連続 DP

動的計画法 (DP)¹¹⁶⁾は、1957年にR.Bellmanが定式化したもので、順序関係のない対象物についての「DP」であるため、順序の逆転も考えられていた。しかし、順序関係がその物理的な性質からの本質的な拘束となっている音声の時系列データを扱うに至って、順序の逆転の無い単調増加なパスのみがDPの適用において取られるようになった^{117),118)}。この拘束は、標準パターンと入力とのマッチングパスの傾きの制限である「傾斜制限」の名で呼ばれ、傾斜制限内で最小の累積距離に現時点の局所距離を累積する演算を漸的に繰り返すことで最適なマッチングパスと累積距離が求まる。その後、Darellらは、この手法を動画像を用いたジェスチャ認識に応用し、入力パターンの切り出しとマッチングをフレームワイズに同時に実現する(スポッティング認識可能な)手法を提案した¹⁾。しかし、入力フレームごとに過去にさかのぼってマッチングを繰り返すために計算量が多く、傾斜制限が0から ∞ と通常のジェスチャの時間伸縮より緩い制限となっていた。また、マッチングパスによって累積する局所距離の数が異なるにも関わらず、累積距離の正規化を行っていないために傾きの急なパスの累積距離が緩いものに比べて小さくなる問題点があった。

一方、岡はDarellらより前にこれらの問題点を解決した「連続DP」^{119),120)}というスポッティング認識可能な手法を提案し、音声認識やジェスチャ認識にて活用されるようになった^{120)~122)}。この連続DPは、傾斜制限が1/2から2で、かつ過去にさかのぼってマッチングしなくても入力フレームに対する演算のみで正規化された累積距離を出力するという特長がある。その後、いくつかの連続DPの改良手法^{123),124)}が提案されている。また、オンライン教示¹²⁵⁾や複数方向からのカメラ映像の特徴を用いて人物の回転にロバストな認識手法¹²⁶⁾も提案されている。画像中の部分画像を検出可能な2次元連続DP¹²⁷⁾は、動画中の変形物体のトラッキング認識へ適用されている¹²⁸⁾。

2.2 検索への適用

十分な量の学習データが用意できる場合には、HMMのモデルは連続DPの標準パターンよりカテゴリーの特性を精度よく表現する。このためHMMの方が高い認識率となるものの、学習データが少数の場合はモデルパラメータの推定精度が上がらない。それに対して、連続DPでは通常学習パターンの平均と分散(学習データが1個の場合は学習データそのもの)を標準パターンにする。この1個のデータを標準パターンにできるという特性があるため、標準パターンをクエリ、入力をデータベースとして検索問題に適用できる¹²⁹⁾。また、ハミングの音高変化パターンをノイズの多い音響信号の時間-音高データベースから検索する手法¹³⁰⁾や視野画像列による自己位置推定¹³¹⁾も提案されている。

また、入力のみ切り出しだけでなく標準パターンの切り出しをも同時に行うことで、類似した部分系列を検出する RIFCDP¹³²⁾ も提案され話題境界検出¹³³⁾ に応用された。その後、この RIFCDP の計算量とメモリ量を削減するために局所距離にかかる重みを過去に遡るに従って指数関数的に減少させて累積距離を計算する重み減衰型 RIFCDP¹³⁴⁾ や連続 DP の効率的なシフトにより同機能を実現させる手法¹³⁵⁾ が提案されている。

一方、アクティブ検索においても類似区間を検出する手法¹³⁶⁾ が提案されており、最近では杉山が複数の時系列へ拡張している¹³⁷⁾。

これらの類似区間検出により大規模な時系列データの圧縮、構造化、学習サンプルの一次抽出、複数時系列間の出現関係理解、予測などが可能となる。

3. モデルに基づく認識

軌道のマッチングに基づいて認識を行う場合、わずかなデータでも高速な学習が実現できる反面、時系列パターンの分布を表現できない。そこで、1990 年代ごろからは、時系列パターンの変化に関して、あらかじめ状態遷移に基づくモデルを持つことで認識を行う、モデルベースの手法が多く用いられるようになっていく。

時系列パターンの認識やモデル化に関しては、ビジョンに限らず、音声や言語、機械学習、制御、統計、生物学 (DNA 解析) といった多くの分野での興味の対象であり、ビジョンの分野における時系列パターンのモデルも、音声や言語などの分野で行われてきたものを、ビジョン特有の特徴量や表現を導入・拡張した上で利用している場合が多い。特に、hidden Markov model (HMM)^{6)~8)} は、現在の音声認識の標準的な認識手法として用いられており、ビジョンの研究でも最も多く用いられている時系列パターンのモデルである。一方で、ジェスチャや表情などの人の動きを認識・解析する上では、その速度や持続時間といった力学的な特徴が重要になる場合も多い。そこで、制御のシステム方程式や微分方程式といった力学的なモデルが、1990 年代後半からしばしば用いられている。

どのモデルを用いるかは、認識したい対象の時間的振る舞いによって使い分ける必要があり、本稿では以下のような分類をする*。

- (1) 離散事象系 (discrete event system)
- (2) 力学系 (dynamical system)
- (3) 両者の混在系 (hybrid (dynamical) system)

離散事象系を対象とする場合は、離散状態の集合をまず定義しておき、離散状態の順序構造をモデル化する。

* 系 (システム) は実際の対象 (実体) を指し、モデルは数式などによるシステムの記述方法を指すが、系 (システム) という用語はモデルとしての意味で用いられることも多い。

る。これには、有限状態機械に基づく方法⁹⁾ や、その確率的モデルの 1 つである HMM¹⁰⁾ があり、さらに複雑な構造を扱うために確率化した文脈自由文法もしばしば用いられる¹¹⁾。離散事象系のモデルは、「ランプカードを切る」、「持ち上げる」といった比較的抽象度の高い行動の間の関係 (例えば文脈を踏まえた行動解析など) を表現するのに適しているが、HMM に関しては「足を上げる」際の角度の変化といった、信号レベルに近い変化のモデル化にも多く用いられる。

力学系を対象とする場合は、連続状態空間 (実ベクトル空間など) を仮定し、微分方程式や差分方程式によって状態変化を記述する。ビジョンの分野では、制御理論の線形動的システムや、自己帰帰モデル、パーティクルフィルタが多く用いられている。HMM などの離散事象系が連続的な信号変化 (特徴量変化) の情報を失うのに対して、力学系モデルは連続的な変化をそのままダイナミクスとして表現できるため、認識だけでなく生成にもしばしば用いられる¹²⁾。ただし、線形な力学系では複雑な構造を扱うことが困難であるため、離散事象系と力学系を統合したハイブリッドシステムが、特に 1990 年代後半から、ビジョンや制御などの多くの分野で提案されている。

以下では、離散事象系、力学系、ハイブリッドシステムの 3 つの分類に従って、3.1 節、3.2 節、3.3 節においてそれぞれの代表的な手法を紹介する。

3.1 離散事象系に基づく時系列パターン認識

3.1.1 有限状態機械を用いる手法

有限状態機械 (finite state machine, FSM) (もしくは有限オートマトン) では、HMM と異なり、状態が直接観測できるものとする。認識したい事象が離散的な状態の系列としてあらかじめ記述できる場合 (事象集合や要素間の順序、分岐構造などが既知、シナリオが与えられている場合)、もしくはその構造が学習可能な場合にしばしば用いられる。

Bobick らの提案したジェスチャ認識手法⁹⁾ では、学習時には、手やマウスのジェスチャから得られた特徴空間中のサンプル点集合から、個々のクラスの典型的な曲線 (prototype curve) を見つける。次に、この prototype curve を離散化し、得られたベクトル集合をクラスタリングすることで状態の集合を決定する。このとき、状態は複数のクラス間で共通して用いられる場合もある。すると、各クラスのジェスチャは、この状態集合中の状態の系列としてそれぞれ表現できるようになる。認識時には、得られた特徴点の系列を状態系列にした上で、学習時に得られている各クラスの典型的な状態系列と、動的計画法によってマッチングを行うことでクラス識別をする。

日常生活で現れるより自然なジェスチャを表現するために、同じく Bobick らは、各状態の持続長分布や動きの大きさの分布を持たせたモデルを提案している¹³⁾。ジェスチャにおける動きの休止の入れ方などに

基づいて、二相性 (bi-phasic) と三相性 (tri-phasic) のジェスチャを認識している。なお、このとき持続長をモデル化した HMM¹⁴⁾ を参考にしているが、これについては 3.1.2 節で述べる。

Wada らは、非決定性有限オートマトンを用いることで、複数人物の入退室の認識を行った¹⁵⁾。このとき、各視点のカメラについてそれぞれオートマトンを設け、認識時には互いの状態に対して抑制をかけることで、動的に仮説を統合する機構を導入している。

3.1.2 Hidden Markov Model を用いる手法

ビジョンにおける HMM の利用としては、手書き文字の認識が 1980 年代中ごろより行われているが、特にカメラで撮影した動画像を入力として時系列パターン認識を行ったものとしては、Yamato らのデンスのスイングの認識^{10),16)} が初期のものであり、さらに Starner と Pentland による手話認識¹⁷⁾ がよく知られている。その後、HMM はジェスチャや表情、行動認識などの様々な目的に用いられ、様々な拡張が提案されているが、ここでは特に

- 空間的変動のモデル化 (a)
- 時間的変動のモデル化 (b)
- 複雑な状態遷移の構造のモデル化 (c,d)
- 複雑な依存関係のモデル化 (e,f)
- HMM の集合 (クラス) の学習法 (g)

のそれぞれに焦点を当てながら分類を行う (括弧内はパラグラフ番号)。

HMM は状態遷移によって時間方向の伸縮に対応し、各状態での出力分布によって空間的な変動に対応するが、以下で述べるように、目的に応じてこれらにさらに拡張が行われる。特に時間方向の変動に関しては、状態の持続長を明示的に持つマルコフモデルとして segment model¹⁸⁾ が提案されている。これらは HMM と区別されることもあるが、本稿では HMM と共にまとめて紹介をする。

HMM の状態遷移のトポロジー (状態の接続関係に関する構造) として、構造や状態数の学習方法についての検討がなされている。また、階層的な構造を持った HMM がいくつか提案されており、これらについても紹介をする。

複数のプロセス間の相互関係 (モダリティ間の関係や左右の手の動きの関係など) や、過去への依存関係を柔軟に表現するために、状態変数や観測間の様々な依存関係 (因果関係など) をモデル化した HMM について紹介を行う。これらを統一的に記述する dynamic Bayesian network については、3.1.3 節でまとめて紹介を行う。

なお、各状態が出力だけでなく入力も持つモデルとして input-output HMM があり、ジェスチャ認識への応用が検討されているが^{19),20)}、ほとんどの場合には、HMM を生成モデルとして用いる。したがって、以下ではデータを HMM へ「入力する」という言い

方をした場合は、その HMM の観測データとして用いることを意味するものとする。

(a) 空間的変動をモデル化するための HMM

HMM では出力確率の分布によって状態と特徴量の対応をモデル化するため、どのような分布を用いるかが、特徴選択と共に重要となる。ビジョンにおける初期の認識手法としては、音声認識と同様に、離散 HMM (discrete HMM) を用いた手法^{10),16),17)} や連続分布 HMM (continuous density HMM) による手法^{21),22)} が用いられ、現在も多く of 認識システムに利用されている。一方で、ジェスチャなどの認識においては音声以上に特徴量の変動が大きくなるという問題があり、これを解決するために、出力確率に対してパラメタを設け、系統的な変動をモデル化した parametric HMM が提案されている^{23),24)}。最近では、出力確率をパラメトリックな分布としてではなく事例として持つという考え方による手法²⁵⁾ があり、これは非常に多くの出力記号を持つ離散 HMM とも考えることができる。

(b) 持続長をモデル化するための HMM

HMM はもともと離散事象系の確率モデルであり、状態遷移の時間的なタイミングは表現しない (オートマトンでは基本的には入力が得られたときに状態遷移を行う)。特に、通常の認識システムでは固定長の周期でサンプリングした特徴ベクトル系列を入力として用いるため、各サンプリング時刻で状態遷移を考えることになる。ある離散時刻における HMM の状態を i とし、次の時刻で再び同じ状態に遷移する確率を a_{ii} とすれば、持続時間 t (t サンプリング) だけ状態 i が持続する確率は、 t 回セルフループする確率として計算できるため

$$d_i(t) = a_{ii}^t (1 - a_{ii}) \quad (1)$$

となる⁷⁾。 $d_i(t)$ は持続時間が長くなるほど確率が低くなる指数分布となるため、特定の持続時間の確率が高い場合などは、適切にモデル化できない。そこで、音声の分野では、各状態の持続長分布を明示的に持たせた time duration HMM^{6),7),14)} が提案されている。これをさらに持続時間を持ったマルコフモデルとして一般化した segment model¹⁸⁾ もしくは hidden semi-Markov model が提案されている。また、マルチグラムモデル²⁶⁾ と呼ばれるモデルが、入力系列を可変長の系列に分節化する際に用いられることがある。time-duration HMM, segment model およびマルチグラムモデルの関係も含めて、Murphy によって統一的に述べられている²⁷⁾。Segment model の中には、自己回帰モデルや動的システムと組み合わせたものも含まれているが、これらは 3.3 節にて改めて述べる。

ビジョンの分野では、3.1.1 節で述べた文献¹³⁾ 以外にも、time duration HMM に基づく時系列パターン認識がしばしば検討されている^{28),29)}。

(c) HMM のトポロジーの学習法

HMM の状態間の接続関係は経験的に決められる

ことが多く、特に left-to-right モデルがよく用いられる。ただし、認識対象によっては最適なトポロジーが自明でない場合があり、自然言語の分野では、Brants によって状態のトポロジーを学習する手法が提案されている。これには、状態のマージに基づく手法^{30),31)}と状態の分割に基づく手法³²⁾がある。ビジョンの分野では、これらの手法とは独立に、HMM の状態のクラスタリングを EM アルゴリズムに基づいて行う方法³³⁾や、エントロピーの最小化によって HMM の状態数やそのトポロジーを学習する手法³⁴⁾が提案されている。

(d) 階層的な構造を扱うための HMM

複雑な事象は、簡単な事象の組み合わせとして表現できることが多く、HMM に階層的な構造を導入した hierarchical HMM が Fine らによって提案されている³⁵⁾。これは、HMM の各状態が出力確率ではなく、下位の HMM を持てるようになっていて、そして下位の HMM の状態もさらにその下位の HMM を持てるという再帰構造になっている。いったん下位の HMM のいずれかの状態に遷移した後に、その HMM の最終状態に遷移すると、再び元の上位の状態に戻る。したがって再帰的に下位の状態を活性化させることができ、後述する確率文脈自由文法 (SCFG) の簡略化されたモデルとなっている。また、SCFG よりも少ない計算量で学習することができる。Hierarchical HMM の拡張³⁶⁾やビジョンへの応用³⁷⁾は Bui らによって積極的に行われており、室内での行動解析などに用いられている。

Oliver らによる layered HMM³⁸⁾では、hierarchical HMM とは異なる考え方で階層化を行っている。これは、下位の HMM 組 (クラス数を K とする) の尤度 (K 次元ベクトル、もしくは最大の尤度) を、上位の HMM への入力とするものであり、オフィス環境におけるマルチモダリティを統合した行動認識システムへ応用されている。

(e) 2 次以上のマルコフ性をモデル化する HMM

HMM は通常単純マルコフ性を仮定して、離散時刻 t の状態 s_t が、直前の時刻 $t-1$ の状態 s_{t-1} によってのみ決まるものとしている。ただし、状態によっては直前だけでなくさらに過去の状態に依存する場合がある。そこで、2 次以上のマルコフ性を導入することが考えられる⁸⁾。ただし、どれだけ過去に依存するかが状態によって異なる場合がある。そこで、機械学習の分野では可変長 N グラムモデルと呼ばれるモデルが提案されており³⁹⁾、マルコフ性の次数を状態に応じて可変となるようなマルコフモデルを学習できる。これに基づいて、ビジョンの分野では Galata らによる variable-length HMM⁴⁰⁾が提案されている。

(f) 複数のプロセス間の関係を扱うための HMM

両手の動きからなるジェスチャは、右手と左手の時系列データの組からなると考えることができる。この

ような複数の時系列データを認識するために、複数の HMM を結合した coupled HMM と呼ばれるモデルが Brand によって提案されている⁴¹⁾。2 つの HMM を結合する場合を例にとれば、一方の HMM のある時刻における状態に対する、他方の HMM の次の時刻の状態を確率として持つことで、2 つの HMM の状態の共起関係をモデル化することができる。

(g) 複数の HMM のクラスタリング

一般的な HMM による時系列パターン認識では、クラスの数だけ HMM を用意しておき、入力された時系列に対して最も高い尤度を持つ HMM を決定することでクラス識別する。しかし、そもそも何種類の HMM が必要であるかが未知である場合も多く、さらに個々の HMM を学習するためには、あらかじめクラスごとに分節化された学習データが必要になるという問題がある。これらの問題を、複数の HMM の (モデルベースの) クラスタリングによって解決しようとする方法が検討されている^{42),43)}。

3.1.3 Dynamic Bayesian Network

ベイジアンネットの観点から、HMM における状態をいくつかの因子が組み合わさったものとして表現できるようにしたものを、dynamic Bayesian network (DBN、もしくは dynamic belief network) と呼ぶ。Ghahramani は、HMM と離散時間カルマンフィルタとの共通性について DBN の観点からの説明している。特に、HMM の前向きアルゴリズムとカルマンフィルタリング、HMM の後向きアルゴリズムとカルマンスムージング、さらに両者の学習方法について、DBN の確率推論・学習の観点から統一的に述べている⁴⁴⁾。また、Ghahramani 自身の提案している factorial HMM⁴⁵⁾についても DBN の観点から解説がなされている。

Murphy は、前節で述べた様々な HMM について DBN の観点からまとめると共に、DBN の確率推論の近似方法などについても検討している⁴⁶⁾。DBN はこの数年で急速に広まっており、特に行動解析・認識において複雑な文脈を表現する場合や、複数の時系列データの統合などの、構造を持った時系列パターンの認識に用いられることが多い^{47)~49)}。

3.1.4 文脈自由文法に基づく手法

有限状態機械 (FSM) と HMM の関係と同様に、文脈自由文法の確率モデルとして確率文脈自由文法 (stochastic context-free grammar, SCFG) があり、構造的な事象を、HMM よりも直接的に扱うことが可能である。Ivanov と Bobick は、SCFG をジェスチャ認識および駐車場における状況認識に適用している⁵⁰⁾。これは、プリミティブとなる離散事象 (手を上げる、下げるといった単純な動きなど) を入力された動画像から検出していく (例えばジェスチャ認識では HMM を利用)。あらかじめこれら事象間の関係 (文脈) を SCFG によってモデル化しておくことで、確

率推論によってあいまい性を解消しながら安定した認識が可能となる。

Moore と Essa は, Ivanov らの手法に対して誤り検出やリカバリーなどの拡張を行い, カードゲームにおける各プレイヤーの行動 (behavior) や戦略の認識を行っている¹¹⁾。

3.2 力学系に基づく時系列パターン認識

力学系モデルは時間の経過と共に自律的に状態を遷移させていくことのできるモデルであり, 時系列パターンの認識や表現に用いられる線形な力学系モデルとしてはカルマンフィルタなど (線形動的システム) が, 非線形な力学系モデルとしては, 拡張カルマンフィルタ, リカレントニューラルネットなど (非線形動的システム) などがよく用いられる^{*}。離散時間の動的システムでは, あるサンプリング時刻の状態 (実ベクトル) が, 直前のサンプリング時刻の状態にのみ依存するという単純マルコフ性を仮定する。また, 状態に対して変換を行うことで, 観測データ (実ベクトル) が得られる。このように, 単純マルコフ性を仮定した状態遷移, および状態からの観測によって時系列をモデル化する点で, 動的システムと HMM は類似点が多い (3.1.3 節や文献⁵¹⁾ を参照のこと)。

しかし, 動的システムが時間の経過にしたがって自律的に状態を遷移させる (離散時間モデルはその近似である) のに対し, HMM は状態の順序のみをモデル化し, 入力 that 得られるたびに状態遷移を考える (離散事象系である) という点で, 両者の時間や状態の考え方は全く異なる。そのため, 動的システムを時系列パターン認識へ用いる際も, 状態 (ジェスチャのポーズや離散的事象) の順序というよりは, その変化の度合いや速度, 力のかかり方といった物理的特性そのもの (メトリックを持った情報) が重要になる場合に用いられる。また, 動画像や関節の動きといった時系列パターンの生成にしばしば応用されている^{12), 52)}。

3.2.1 線形動的システム

線形動的システムとしては, 状態に単一のガウス分布を仮定するカルマンフィルタがよく知られている⁵³⁾。一方, 複数人物のトラッキングなどではこの仮定が成り立たない場合も多い。そこで, 粒子 (パーティクル) の集まりによって状態の非ガウス分布を表現する方法として, Blake らの CONDENSATION⁵⁴⁾ を始めとするパーティクルフィルタがある。

状態にガウス分布を仮定した線形動的システムを用いて, 時系列パターンの認識を行う初期の研究としては Bregler⁵⁵⁾ のモデルがあるが, これは HMM との混在系の中で用いており, 3.3 節で改めて述べる。線形

動的システムを単体で用いるものとして, Rao は, ロバスト推定とカルマンフィルタを組み合わせることによって, 遮蔽物があるような入力画像に対して, あらかじめ学習された画像系列を想起できる手法を提案している⁵⁶⁾。これに近いモデルとして, 与えられた動画像から線形動的システムを同定し, 同定されたシステムから動画像の生成や認識を行う方法 (Dynamic texture) を提案している¹²⁾。Rittscher と Blake は, 自己回帰モデルに基づいてエアロビクスにおける人の動きをモデル化し, システム行列の固有値などを解析している⁵⁷⁾。Bis-sacco らは, 人の walking, running, dancing, jumping などの動きをそれぞれ線形動的システムとしてモデル化し, あらかじめシステム間の距離を定義しておくことで, 認識時には, 登録されたシステムと入力データから同定されたシステムとの距離による nearest neighbor によって認識を行っている⁵⁸⁾。

状態の非ガウス分布を導入した手法として, Blake らは, 3.3 節で述べるハイブリッドシステムの観点から, ジャグリングにおける玉の動きの力学的変化 (切り替わり方) を CONDENSATION に基づいてモデル化している⁵⁹⁾。一方, Dornaika らは, 単一のパーティクルフィルタを用いて顔の追跡と表情認識を同時に行っている⁶⁰⁾。

線形動的システムの学習法としては, 制御と同様にシステム同定を用いる方法^{12), 58)} や, HMM と同様に EM アルゴリズムを用いる方法^{61), 62)}, 学習するパラメタを状態とみなしてカルマンフィルタによって推定する方法などがある⁵⁶⁾。

3.2.2 リカレントニューラルネット

リカレントニューラルネットは, 非線形力学系として時系列パターン認識に用いられることがあり^{63)~68)}, 脳のモデル化という観点からも, 力学系モデルの上に離散事象系を構築することを目的として多くの興味深い研究が行われている。ただし, 実際の応用には多くのハイパーパラメタを設定しておく必要がある。そのため, モデルの見通しのよさに欠け, ビジョンの分野 (特に認識精度を競う研究) における事例は比較的に少数に留まっている。

3.3 ハイブリッドシステム~力学・離散事象混在系

力学系と離散事象系が混在した系をハイブリッドシステムと呼び, そのモデル化や解析方法について, 1990 年前後から計算機科学や制御などの分野で盛んに研究行われるようになってきた^{69)~71)}。ビジョンの分野ではこれとは独立に, 人の構造的な行為や運動などを記述するためのモデルとして, 力学系と離散事象系を組み合わせたシステムがいくつか提案されている。例えば, Siskind らは, 缶などの物体を操作するときの力学的特徴を推論に組み込むことで, 動的なシーンの解釈を与える方法などについて検討を行っている^{72), 73)}。

線形動的システムと HMM を直接組み合わせる方

^{*} 力学系と動的システムは共に dynamical system であるが, 制御理論などで特に線形な系を対象とする場合は動的システムが, カオスなどの非線形の場合は力学系という言葉が用いられることが多いようである。なお, それぞれの論文の使用法に合わせ, ここでは具体的なモデルのことをシステムと呼ぶ場合がある。

法として、カルマンフィルタの出力（例えば分散）を、HMM への入力とする方法もある⁷⁴⁾。ただし以下では、両者がより密接に結びついたシステムに注目し、複数の線形動的システム（カルマンフィルタや自己回帰モデル、パーティクルフィルタ等）の間の遷移を HMM によってモデル化したものを取り上げる。

3.3.1 Switching linear dynamical system

線形動的システムと HMM を組み合わせたものとして、特に Bregler⁵⁵⁾ のモデルがよく知られており、画像上における人の足の動きを、複数の線形動的システムの切り替わりとして表現することで、skipping や hopping の認識などを行っている。つまり、複雑な足の動きを、単純な動き（足を前に出すなど）のフェーズの組み合わせとして考え、HMM の各状態が、それぞれ異なるフェーズ（この場合は 2 次の線形動的システム）に対応付けられている。

Dynamic Bayesian network の観点から両者を統合したものとして、HMM の状態に応じて状態空間自体を切り替える⁷⁵⁾ や、状態空間は共有する switching linear dynamical system (SLDS)^{76),77)} があり、Murphy の文献⁷⁸⁾ の中でこれらの違いについて述べられている。

パーティクルフィルタと HMM を組み合わせたものとしては、前節で述べた Blake らの方法⁵⁹⁾ がある。

複数の HMM を用いて coupled HMM⁴¹⁾ を構成したのと同様に、複数の SLDS において状態の共起関係をモデル化したものとして、coupled SLDS⁷⁹⁾ がある。また、HMM を parametric HMM²³⁾ へ拡張したのと同様に、SLDS に対して大域的な変動を表現するようなパラメータを導入したものとして、parametric SLDS⁸⁰⁾ がある。

3.3.2 持続時間を表現するハイブリッドシステム

前節で述べた、SLDS を始めとするモデルでは、観測サンプルが得られるごとに、HMM の離散的な状態と線形動的システムの連続的な状態の遷移同時に考える。すると、3.1.2 節で述べたように、1 つの離散状態が持続する長さは幾何分布に従うことになり、離散事象系にとっては不自然なモデルとなっている。そこで、離散事象系をサンプリング時刻に同期した遷移から切り離し、状態の持続長を明示的に分布としてモデル化したものとして、segment model¹⁸⁾ や文献⁸¹⁾ がある。

3.3.3 運動生成への応用

Bregler のモデルを始めとするこれらのシステムは生成モデルであり、時系列パターンの認識だけでなく、生成にも用いられることがある。あらかじめモーションキャプチャデータによって学習されたモデルから、人の複雑なダンスなどの動きを生成できるようにしたシステムとしては Motion texture⁵²⁾ がよく知られている。

3.3.4 ハイブリッドシステムの学習法

線形動的システムと HMM を組み合わせたハイブリッドシステムの学習では、HMM の状態数（線形動的システムの個数に一致）および、それぞれの線形動的システムのパラメータを推定する必要がある。歩行などの足や手の周期的動きの認識では、離散事象系の状態数が既知とする場合が多い。この場合、EM アルゴリズムによって反復的な学習を行うことによって、各線形動的システムのパラメータと、その間の遷移確率が推定される。また、行列演算を用いた学習も検討されている⁸²⁾。

一方、状態数が未知とする場合は、いったん一定の範囲から推定した線形動的システムに、続く時系列データを入力（フィッティング）していき、フィッティング精度（尤度もしくはその近似として予測誤差）が閾値を超える場合に新たな線形動的システムを追加するといった、欲張り法によるアルゴリズムがある^{52),83)}。ただし、ノイズが多いデータに対してはこの閾値の決め方が困難である。そこで、観測データ系列からボトムアップに線形動的システムの集合やその個数を学習する方法が提案されている^{81),84)}。これは、複数の線形動的システムのうち性質の最も近いベアを順に併合する階層的クラスタリングであり、各システムの固有値が複素平面の単位円内に収まるように制約をかけることで、収束性の線形動的システムを、少ないデータから安定して学習可能である。状態数（線形動的システムの数）の決定は、通常の階層型クラスタリングと同様に、併合時の尤度（フィッティング精度）のカーブが急激に変化する時点をその候補としている。

4. 応用分野

4.1 表情認識

表情認識に関する研究は非常に多く、特にオブティカルフローに基づく動的な特徴量を抽出して、テンプレートマッチングや PCA、判別分析、SVM などの（静的）パターン認識手法を適用する方法がしばしば用いられる^{85)~90)}。一方で、HMM などの時系列パターン認識手法を用いる手法も多く提案されている^{91)~94)}。Cohen らは、静的なパターン認識と動的なパターン認識 (HMM) の比較を行っており、特定人物の表情を認識する場合や、表情を識別するための適当な時間区間の分節化が十分ではない場合には、静的なパターン認識手法に対して HMM が優位であるとコメントしている⁹⁵⁾。また最近の研究では、3.2 節で述べたパーティクルフィルタを用いた表情認識が提案されている⁶⁰⁾。

一方で、happy や sad といった感情的なカテゴリではなく、社交的な笑いと自然な笑いといった微細な表情の認識では、人間は顔のパーツ（目や口など）の動きのタイミングやその持続時間を用いているという報告がある⁹⁶⁾。そこで、ハイブリッドシステムによ

てあらかじめ各パーツの動きを分節化しておき、その動きの間のタイミングによって、自発的と意図的な笑いの判別を検討した研究がある⁹⁷⁾。また、表情というよりは顔の動きから人の状態(覚醒、眠りなど)を認識するために、dynamic Bayesian network などが用いられている⁹⁸⁾。

4.2 音響情報と視覚情報の統合

音声と動画を統合することで、雑音環境下でも頑健な音声認識を実現させる audio-visual speech recognition(AVSR)の研究が、音声やビジョンなどの研究者によって行われている。特に OpenCV などでも用いられている、coupled HMM による方法⁹⁹⁾や、dynamic Bayesian network を利用した手法¹⁰⁰⁾などがよく知られている。これらは音声と動画の間の共起関係を、隣り合う時刻の離散状態間の同時確率によってモデル化している。一方、音声から唇や顔の動画を生成する研究もあり¹⁰¹⁾、音声で学習された HMM を用いて動画の特徴ベクトル系列を学習することで、各状態において音声と動画の対応をモデル化している。

5. ま と め

本稿では、主に非線形時間伸縮機能を持つ時系列パターン認識手法について説明した。まず、軌道に基づく認識手法として連続 DP について概説し検索問題への適用について述べた。次に、音声認識の分野と同様に盛んに研究されているモデルに基づく認識手法について離散事象系、力学系、両者の混在系の3種に分けて最新の研究動向について説明した。さらに、応用分野として表情認識、音響情報と視覚情報の統合に関する研究を紹介した。

時系列パターンの認識は、近年のセンサ技術および計算機の発展によりユーザインタフェースとしてだけでなく、ユビキタス情報環境システム、ウェアラブル機器、Web 空間における大量の時系列データのハンドリングにおいて重要になっており、データの構造化、検索、予測などの問題と同時に議論していく必要があるだろう。

謝 辞

2 節に関しては会津大学岡隆一教授のご協力を頂いた。

参 考 文 献

- 1) Darrell, T. and Pentland, A.: Space-Time Gestures, *Proc. IJCAI'93 Looking at People Workshop* (1993).
- 2) Niyogi, S. and Adelson, E.: Analyzing and Recognizing Walking Figures in XYT, *Proc. IEEE Conference on Computer Vision and Pattern Recognition*, pp.469-474 (1994).

- 3) Campbell, L. and Bobick, A.: Recognition of human body motion using phase space constraints, *Proc. IEEE Int. Conference on Computer Vision*, pp.624-630 (1995).
- 4) Ohno, H. and Yamamoto, M.: Gesture Recognition Using Character Recognition Techniques on Two-Dimensional Eigenspace, *Proc. IEEE Int. Conference on Computer Vision*, pp.151-156 (1999).
- 5) Chu, S., Keogh, E., Hart, D. and Pazzani, M.: Iterative Deepening Dynamic Time Warping for Time Series, *The Second SIAM Int. Conference on Data Mining* (2002).
- 6) Rabiner, L.R.: A tutorial on hidden Markov models and selected applications in speech recognition, *Proc. IEEE*, pp.257-286 (1989).
- 7) Huang, H.D., Ariki, Y. and Jack, M.A.: *Hidden Markov Models for Speech Recognition*, Edinburgh Univ. (1990).
- 8) 中川聖一: 音声認識研究の動向, 電子情報通信学会論文誌, Vol. J83-D-II, No.2, pp.433-457 (2000).
- 9) Bobick, A.F. and Wilson, A.D.: A State-Based Approach to the Representation and Recognition of Gesture, *IEEE Trans. Pattern Analysis and Machine Intelligence*, Vol.19, No.12, pp.1325-1337 (1997).
- 10) Yamato, J., Ohya, J. and Ishii, K.: Recognizing Human Action in Time-Sequential Images Using Hidden Markov Model, *Proc. Int. Conference on Computer Vision and Pattern Recognition*, pp.379-385 (1992).
- 11) Moore, D. and Essa, I.: Recognizing multitasked activities from video using stochastic context-free grammar, *Proc. National Conference on Artificial intelligence*, pp.770-776 (2002).
- 12) Doretto, G., Chiuso, A., Wu, Y.N. and Soatto, S.: Dynamic Textures, *Int. Journal of Computer Vision*, Vol.51, No.2, pp.91-109 (2003).
- 13) Wilson, A. D., Bobick, A. F. and Cassell, J.: Temporal Classification of Natural Gesture with Application to Video Coding, *Proc. IEEE Conference on Computer Vision and Pattern Recognition*, pp.948-954 (1997).
- 14) Levinson, S. E.: Continuously variable duration hidden Markov models for automatic speech recognition, *Computer Speech and Language*, Vol.1, pp.29-45 (1986).
- 15) Wada, T. and Matsuyama, T.: Multiobject Behavior Recognition by Event Driven Selective Attention Method, *IEEE Trans. Pattern Analysis and Machine Intelligence*, Vol.22, No.8, pp.873-887 (2000).

- 16) 大和淳司, 大谷 淳, 石井健一郎: 隠れマルコフモデルを用いた動画像からの人物の行動認識, 電子情報通信学会論文誌, Vol.J76-D-II, No.12, pp.2556-2563 (1993).
- 17) Starner, T. and Pentland, A.: Visual Recognition of American Sign Language Using Hidden Markov Models, *Proc. Int. Workshop on Automatic Face and Gesture Recognition*, pp. 189-194 (1995).
- 18) Ostendorf, M., Digalakis, V. and Kimball, O.A.: From HMMs to Segment Models: A Unified View of Stochastic Modeling for Speech Recognition, *IEEE Trans. Speech and Audio Process*, Vol.4, No.5, pp.360-378 (1996).
- 19) Marcel, S., Bernier, O., Viallet, J.-E. and Collobert, D.: Hand Gestur  Recognition Using Input-Output Hidden Markov Models, *Proc. IEEE Int. Conference on Automatic Face and Gesture Recognition*, pp.456-461 (2000).
- 20) Just, A., Bernier, O. and Marcel, S.: HMM and IOHMM for the Recognition of Mono- and Bi-Manual 3D Hand Gestures, *British Machine Vision Conference* (2004).
- 21) Wilson, A. and Bobick, A.: Learning Visual Behavior for Gesture Analysis, *IEEE International Symposium on Computer Vision*, p.5A Motion II (1995).
- 22) Hamdan, R., Heitz, F. and Thoraval, L.: Gesture Localization and Recognition using Probabilistic Visual Learning, *Proc. IEEE Conference on Computer Vision and Pattern Recognition*, pp.98-103 (1999).
- 23) Wilson, A.D. and Bobick, A.F.: Nonlinear PHMMs for the Interpretation of Parameterized Gesture, *Proc. IEEE Conference on Computer Vision and Pattern Recognition*, pp.879-884 (1998).
- 24) Wilson, A.D. and Bobick, A.F.: Parametric hidden Markov models for gesture recognition, *IEEE Trans. Pattern Analysis and Machine Intelligence*, Vol.21, No.9, pp.884-900 (1999).
- 25) Elgammal, A., Sht, V., Yacoob, Y. and Davis, L.S.: Learning Dynamics for Exemplar-based Gesture Recognition, *Proc. IEEE Conference on Computer Vision and Pattern Recognition*, pp.571-578 (2003).
- 26) Deligne, S. and Bimbot, F.: Inference of Variable-Length Linguistic and Acoustic Units by Multigrams, *Speech Communication*, Vol.23, pp.223-241 (1997).
- 27) Murphy, K.P.: Hidden semi-Markov models (HSMMS), *Informal Notes* (2002).
- 28) Hongeng, S. and Nevatia, R.: Large-Scale Event Detection Using Semi-Hidden Markov Models, *Proc. IEEE Int. Conference on Computer Vision*, pp.1455-1462 (2003).
- 29) Duong, T.V., Bui, H.H., Phung, D.Q. and Venkatesh, S.: Activity Recognition and Abnormality Detection with the Switching Hidden Semi-Markov Model, *Proc. IEEE Conference on Computer Vision and Pattern Recognition*, pp.838-845 (2005).
- 30) Brants, T.: Estimating HMM topologies, *Logic and Computation* (1995).
- 31) Brants, T.: Better Language Models with Model Merging, *Proc. the Conference on Empirical Methods in Natural Language Processing* (1996).
- 32) Brants, T.: Estimating Markov Model Structure, *Proc. Int. Conference on Spoken Language Processing*, pp.893-896 (1996).
- 33) Gong, S., Walter, M. and Psarrou, A.: Recognition of Temporal Structures: Learning Prior and Propagating Observation Augmented Densities via Hidden Markov States, *Proc. IEEE Int. Conference on Computer Vision*, pp.157-162 (1999).
- 34) Brand, M. and Kettner, V.M.: Discovery and Segmentation of Activities in Video, *IEEE Trans. Pattern Analysis and Machine Intelligence*, Vol.22, No.8, pp.844-851 (2000).
- 35) Fine, S., Singer, Y. and Tishby, N.: The Hierarchical Hidden Markov Model: Analysis and Applications, *Machine Learning*, Vol.32, No.1, pp.41-62 (1998).
- 36) Bui, H.H., Phung, D.Q. and Venkatesh, S.: Hierarchical Hidden Markov Models with General State Hierarchy, *Proc. National Conference on Artificial Intelligence*, pp.324-329 (2004).
- 37) Nguyen, N.T., Phung, D.Q., Venkatesh, S. and Bui, H.: Learning and Detecting Activities from Movement Trajectories Using the Hierarchical Hidden Markov Model, *Proc. IEEE Conference on Computer Vision and Pattern Recognition*, pp.955-960 (2005).
- 38) Oliver, N., Garg, A. and Horvitz, E.: Layered representations for learning and inferring office activity from multiple sensory channels, *Computer Vision and Image Understanding*, Vol.96, No.2, pp.163-180 (2004).
- 39) Ron, D., Singer, Y. and Tishby, N.: The power of amnesia: Learning Probabilistic Automata with Variable Memory Length, *Machine Learning*, Vol.25, No.2-3, pp.117-149 (1996).
- 40) Galata, A., Johnson, N. and Hogg, D.: Learning Variable-Length Markov Models of Behavior, *Computer Vision and Image Understanding*, Vol.81, No.3, pp.398-413 (2001).

- 41) Brand, M., Oliver, N. and Pentland, A.: Coupled Hidden Markov Models for complex action recognition, *Proc. IEEE Conference on Computer Vision and Pattern Recognition*, pp.994-999 (1997).
- 42) Juang, B.H. and Rabiner, L.R.: A probabilistic distance measure for hidden Markov models, *AT & T Technical Journal*, Vol.64, No.2, pp.391-408 (1985).
- 43) Alon, J., Sclaroff, S., Kollios, G. and Pavlovic, V.: Discovering clusters in motion time-series data, *Proc. IEEE Conference on Computer Vision and Pattern Recognition*, pp.375-381 (2003).
- 44) Ghahramani, Z.: Learning Dynamic Bayesian Networks, *Lecture Notes in Artificial Intelligence. Springer-Verlag.*, Vol.1387, pp.168-197 (1998).
- 45) Ghahramani, Z. and Jordan, M. I.: Factorial Hidden Markov Models, *Machine Learning*, Vol.29, No.2-3, pp.245-273 (1997).
- 46) Murphy, K.P.: Dynamic Bayesian Networks: Representation, Inference and Learning, *PhD Thesis, UC Berkeley, Computer Science Division* (2002).
- 47) Ayers, D. and Chellappa, R.: Scenario recognition from video using a hierarchy of dynamic belief networks, *Proc. Int. Conference on Pattern Recognition*, pp.835-838 (2000).
- 48) Mittal, A., Cheong, L.F. and Sing, L.T.: Dynamic Bayesian framework for extracting temporal structure in video, *Proc. IEEE Conference on Computer Vision and Pattern Recognition*, pp.110-115 (2001).
- 49) Xiang, T. and Gong, S.: Beyond Tracking: Modelling Activity and Understanding Behaviour, *International Journal of Computer Vision*, Vol.67, No.1, pp.21-51 (2006).
- 50) Ivanov, Y.A. and Bobick, A.F.: Recognition of visual activities and interactions by stochastic parsing, *IEEE Trans. Pattern Analysis and Machine Intelligence*, Vol.22, No.8, pp.852-872 (2000).
- 51) Minka, T.P.: From Hidden Markov Models to Linear Dynamical Systems, *MIT Media Lab Vision and Modeling TR531* (1999).
- 52) Li, Y., Wang, T. and Shum, H.-Y.: Motion Texture: A Two-Level Statistical Model for Character Motion Synthesis, *Proc. SIGGRAPH*, pp.465-472 (2002).
- 53) Anderson, B. D.O. and Moor, J.B.: *Optimal Filtering*, Prentice-Hall (1979).
- 54) Isard, M. and Blake, A.: Condensation - conditional density propagation for visual tracking, *Int. Journal of Computer Vision*, Vol.29, No.1, pp.5-28 (1998).
- 55) Bregler, C.: Learning and Recognizing Human Dynamics in Video Sequences, *Proc. Int. Conference on Computer Vision and Pattern Recognition*, pp.568-574 (1997).
- 56) Rao, R.: Dynamic Appearance-Based Recognition, *Proc. Int. Conference on Computer Vision and Pattern Recognition*, pp.540-546 (1997).
- 57) Rittscher, J. and Blake, A.: Classification of Human Body Motion, *Proc. IEEE Int. Conference on Computer Vision*, pp.634-639 (1999).
- 58) Bissacco, A., Chiuso, A., Ma, Y. and Soatto, S.: Recognition of Human Gaits, *Proc. IEEE Conference on Computer Vision and Pattern Recognition*, pp.52-57 (2001).
- 59) Blake, B. N.A., Isard, M. and Rittscher, J.: Learning and Classification of Complex Dynamics, *IEEE Trans. Pattern Analysis and Machine Intelligence*, Vol.22, No.9, pp.1016-1034 (2000).
- 60) Dornaika, F. and Dovoine, F.: Simultaneous Facial Action Tracking and Expression Recognition Using a Particle Filter, *Proc. IEEE Int. Conference on Computer Vision*, pp.1733-1738 (2005).
- 61) Ghahramani, Z. and Hinton, G.E.: Parameter Estimation for Linear Dynamical Systems, *Technical Report CRG-TR-96-2* (1996).
- 62) Roweis, S. and Ghahramani, Z.: A Unifying Review of Linear Gaussian Models, *Neural Computation*, Vol.11, No.2, pp.305-345 (1999).
- 63) 二見亮弘, 星宮 望: 時系列パターン認識の神経回路モデル, *電子情報通信学会論文誌, Vol.J71-D-II, No.10*, pp.2181-2190 (1988).
- 64) Elman, J.L.: Finding Structure in Time, *Cognitive Science*, Vol.14, pp.179-211 (1990).
- 65) Robinson, A.J.: An Application of Recurrent Nets to Phone Probability Estimation, *IEEE Trans. Neural Networks*, Vol.5, No.2, pp.298-305 (1994).
- 66) Dorffner, G.: Neural Networks for Time Series Processing, *Neural Network World*, Vol.6, No.4, pp.447-468 (1996).
- 67) 森田昌彦, 村上 聡: 非準調神経回路網による時系列パターンの認識, *電子情報通信学会論文誌, Vol.J81-D-II, No.7*, pp.1679-1688 (1998).
- 68) 内山 徹, 高橋治久: リカレントニューラル予測モデルを用いた不特定話者単語音声認識, *電子情報通信学会論文誌, Vol.J83-D-II, No.2*, pp.776-783 (2000).
- 69) Gollu, A. and Varaiya, P.: Hybrid Dynamical Systems, *Proc. Conference on Decision and*

- Control*, pp.2708–2712 (1989).
- 70) Maler, O., Manna, Z., de Bakker, A. P. J.W., Huizing, C., de Roeper, W. P. and Rozenberg, G.: *From Timed to Hybrid Systems, Real-Time: Theory in Practice*, pp. 447–484, Springer-Verlag (1991).
 - 71) R.Alur, e.a.: The Algorithmic Analysis of Hybrid Systems, *Theoretical Computer Science*, Vol.138, No.1, pp.3–34 (1995).
 - 72) Mann, R., Jepson, A. and Siskind, J.M.: The Computational Perception of Scene Dynamics, *Computer Vision and Image Understanding*, Vol.65, No.2, pp.113–128 (1997).
 - 73) Siskind, J.M.: Reconstructing force-dynamic models from video sequences, *Artificial Intelligence*, Vol.151, No.1-2, pp.91–154 (2002).
 - 74) Dockstader, S.L., Imenkov, N.S. and Tekalp, A.M.: Markov-Based Failure Prediction for Human Motion Analysis, *Proc. IEEE Int. Conference on Computer Vision*, pp. 1283–1288 (2003).
 - 75) Ghahramani, Z. and Hinton, G.E.: Switching State-Space Models, *Dept. of Computer Science* (1996).
 - 76) Pavlovic, V., Rehg, J.M., Cham, T. and Murphy, K.P.: A Dynamic Bayesian Network Approach to Figure Tracking Using Learned Dynamic Models, *Proc. IEEE Int. Conference on Computer Vision*, pp.94–101 (1999).
 - 77) Pavlovic, V., Rehg, J.M. and MacCormick, J.: Impact of Dynamic Model Learning on Classification of Human Motion, *Proc. IEEE Conference on Computer Vision and Pattern Recognition*, pp.788–795 (2000).
 - 78) Murphy, K.P.: Switching Kalman Filter, *Technical report, U. C. Berkeley* (1998).
 - 79) Jeong, M. H., Kuno, Y. and Shimada, N.: Two-Hand Gesture Recognition using Coupled Switching Linear Model, *Proc. Int. Conference on Pattern Recognition*, pp.529–532 (2002).
 - 80) Oh, S.M., Rehg, J.M., Balch, T. and Delaert, F.: Learning and Inference in Parametric Switching Linear Dynamical Systems, *Proc. IEEE Int. Conference on Computer Vision*, pp. 1161–1168 (2005).
 - 81) Kawashima, H. and Matsuyama, T.: Multi-phase Learning for an Interval-based Hybrid Dynamical System, *IEICE Trans. Fundamentals*, Vol.E88-A, No.11, pp.3022–3035 (2005).
 - 82) Bissacco, A.: Modeling and Learning Contact Dynamics in Human Motion, *Proc. IEEE Conference on Computer Vision and Pattern Recognition*, pp.421–428 (2005).
 - 83) Lu, C., Liu, H. and Ferrier, N.J.: Multidimensional Motion Segmentation and Identification, *Proc. IEEE Conference on Computer Vision and Pattern Recognition*, pp.2629–2636 (2000).
 - 84) Kawashima, H. and Matsuyama, T.: Hierarchical Clustering of Dynamical Systems based on Eigenvalue Constraints, *3rd International Conference on Advances in Pattern Recognition (S. Singh et al. (Eds.): ICAPR 2005, LNCS 3686)*, pp.229–238 (2005).
 - 85) Essa, I.A. and Pentland, A.P.: Facial Expression Recognition using a Dynamic Model and Motion Energy, *Proc. IEEE Int. Conference on Computer Vision*, pp.360–367 (1995).
 - 86) Yacoob, Y. and Davis, L.S.: Recognizing Human Facial Expressions from Long Image Sequences Using Optical-Flow, *IEEE Trans. Pattern Analysis and Machine Intelligence*, Vol.18, No.6, pp.636–642 (1996).
 - 87) Essa, I.A. and Pentland, A.P.: Coding, Analysis, Interpretation, and Recognition of Facial Expressions, *IEEE Trans. Pattern Analysis and Machine Intelligence*, Vol.19, No.7, pp. 757–763 (1997).
 - 88) Kimura, S. and Yachida, M.: Facial Expression Recognition and Its Degree Estimation, *Proc. IEEE Conference on Computer Vision and Pattern Recognition*, pp.295–300 (1997).
 - 89) Donato, G., Bartlett, M.S., Hager, J.C., Ekman, P. and Sejnowski, T.: Classifying facial actions, *IEEE Trans. Pattern Analysis and Machine Intelligence*, Vol.21, No.10, pp.974–989 (1999).
 - 90) Bartlett, M. S., Littlewort, G., Frank, M., Lainscsek, C., Fasel, I. and Movellan, J.: Recognizing Facial Expression: Machine Learning and Application to Spontaneous Behavior, *Proc. IEEE Conference on Computer Vision and Pattern Recognition*, pp.568–573 (2005).
 - 91) Otsuka, T. and Ohya, J.: Recognizing Abruptly Changing Facial Expressions from Time-Sequential Face Images, *Proc. IEEE Conference on Computer Vision and Pattern Recognition*, pp.808–813 (1998).
 - 92) Hoey, J. and Little, J.J.: Representation and Recognition of Complex Human Motion, *Proc. IEEE Conference on Computer Vision and Pattern Recognition*, pp.752–759 (2000).
 - 93) Chang, Y., Hu, C. and Turk, M.: Probabilistic Expression Analysis on Manifolds, *Proc. IEEE Conference on Computer Vision and Pattern Recognition*, pp.520–527 (2004).
 - 94) Yeasin, M., Bulot, B. and Sharma, R.: From Facial Expression to Level of Interest: A Spatio-Temporal Approach, *Proc. IEEE*

- Conference on Computer Vision and Pattern Recognition*, pp.922-927 (2004).
- 95) Cohen, I., Sebe, N., Garg, A., Chen, L.S. and Huang, T.S.: Facial expression recognition from video sequences: temporal and static modeling, *Computer Vision and Image Understanding*, Vol.91, No.1-2, pp.160-187 (2003).
 - 96) Nishio, S., Koyama, K. and Nakamura, T.: Temporal Differences in Eye and Mouth Movements Classifying Facial Expressions of Smiles, *Proc. IEEE Int. Conference on Automatic Face and Gesture Recognition*, pp.206-211 (1998).
 - 97) Nishiyama, M., Kawashima, H., Hirayama, T. and Matsuyama, T.: Facial Expression Representation based on Timing Structures in Faces, *IEEE International Workshop on Analysis and Modeling of Faces and Gestures (W. Zhao et al. (Eds.): AMFG 2005, LNCS 3723)*, pp.140-154 (2005).
 - 98) Gu, H. and Ji, Q.: Facial Event Classification with Task Oriented Dynamic Bayesian Network, *Proc. IEEE Conference on Computer Vision and Pattern Recognition*, pp. 870-875 (2004).
 - 99) Nefian, A.V., Liang, L., Pi, X., Xiaoxiang, L., Mao, C. and Murphy, K.: A Coupled HMM for Audio-Visual Speech Recognition, *Proc. IEEE Int. Conference on Acoustics, Speech and Signal Processing*, Vol.2, pp.2013-2016 (2002).
 - 100) Nefian, A. V., Liang, L., Pi, X., Liu, X. and Murphy, K.: Dynamic Bayesian Networks for Audio-Visual Speech Recognition, *EURASIP Journal on Applied Signal Processing*, Vol.2002, No.11, pp.1-15 (2002).
 - 101) Brand, M.: Voice Puppetry, *Proc. SIGGRAPH*, pp.21-28 (1999).
 - 102) 浅田 稔: ダイナミックシーンの理解, 電子情報通信学会 (1994).
 - 103) J.K. Aggarwal and Q. Cai: Human Motion Analysis: A Review, *Proceedings, IEEE Non-rigid and Articulated Motion Workshop* (1997).
 - 104) V.I. Pavlovic, R. Sharma, and T.S. Huang.: Visual interpretation of hand gestures for humancomputer interaction: A review, *IEEE Transaction on Pattern Analysis and Machine Intelligence*, Vol.19, No.7 (1997).
 - 105) D. Gavrilu: The Visual Analysis of Human Movement: A Survey, *Computer Vision and Image Understanding*, 73(1), pp.82-98 (1999).
 - 106) R. Polana and R. Nelson.: Low level recognition of human motion, *In Proc. of Workshop on Non-rigid Motion*, pp.77-82 (1994).
 - 107) A. Bobick and J. Davis.: Real-time recognition of activity using temporal templates, *In proc. of the Workshop on Applications of Computer Vision* (1996).
 - 108) M. Black and A. Jepson.: A probabilistic framework for matching temporal trajectories: Condensation-based recognition of gestures and expressions, *ECCV* (1998).
 - 109) 大津 展之, 栗田 多喜夫, 関田 巖: パターン認識, 朝倉書店 (1996).
 - 110) T.Kobayashi and N.Otsu: Action and Simultaneous Multiple Persons Identification Using Cubic Higher-order Local Auto-Correlation, *Proc. ICPR* (2004).
 - 111) N. Otsu : Towards Flexible and Intelligent Vision Systems - From Thresholding to CHLAC -, *IAPR Conf. on Machine Vision Applications*, Invited paper, pp.430-439 (2005).
 - 112) T.Ishihara and N.Otsu: Gesture Recognition Using Auto-Regressive Coefficients of Higher-Order Local Auto-Correlation Features, *Proc. IEEE 6th International Conference on Automatic Face and Gesture Recognition (FG 2004)* (2004).
 - 113) 村瀬 洋, V.V.Vinod: 局所色情報を用いた高速物体探索-アクティブ探索法-, 電子情報通信学会論文誌, Vol. J81-DII, No.9, pp.2035-2042 (1998).
 - 114) 柏野 邦夫, カビン スミス, 村瀬 洋: ヒストグラム特徴を用いた音響信号の高速検索-時系列アクティブ探索法-, 信学論 (D-II), Vol. J82-D-II, No.9, pp. 1365-1373 (1999).
 - 115) 柏野 邦夫, 黒住 隆行, 村瀬 洋: ヒストグラム特徴を用いた音や映像の高速 AND/OR 探索, 電子情報通信学会論文誌, Vol. J83-DII, No.12, pp.2735-2744 (2000).
 - 116) R.Bellman and R.Kalaba: Dynamic Programming and Modern Control Theory, *Academic Press* (1965).
 - 117) V.M Velichko and N.G. Zagoruyko: Automatic Recognition of 200 Words, *International Journal of Man-Machine Studies*, 2:223 (1970).
 - 118) H. Sakoe and S. Chiba: Dynamic programming optimization for spoken word recognition, *IEEE Trans. Acoustics, Speech, Signal Processing*, vol. ASSP-26, pp.43-49 (1978).
 - 119) 岡 隆一: 連続 DP を用いた連続音声認識, 音響学会音声研資料, S78-20, pp.145-152 (1978).
 - 120) K. Takahashi, S. Seki, H. Kojima, R. Oka: Recognition of Dexterous Manipulations from Time-Varying Images, *Proc. IEEE Workshop on Motion of Non-Rigid and Articulate Objects*, pp.23-28 (1994).
 - 121) 高橋 勝彦, 関 進, 小島 浩, 岡 隆一: ジェスチャー動画のスポッティング認識, 信学論 (D-II), Vol. J77-D-II, No.8, pp.1552-1561 (1994).
 - 122) 西村 拓一, 向井 理朗, 野崎 俊輔, 岡 隆一: 低解

- 像度特徴を用いた複数人物によるジェスチャの単一動画像からのスポッティング認識, 信学論 (D-II), Vol.J80-D-II, No.6, pp.1563-1570 (1997).
- 123) 佐川 浩彦, 酒匂 裕, 大平 栄二, 崎山 朝子, 安部 正博: 圧縮連続 DP 照合を用いた手話認識方式, 信学論 (D-II), Vol.J77-D-II, No.4, pp.753-763 (1994).
- 124) 中川 聖一: 拡張連続 DP 法による連続音声認識アルゴリズム, 信学論 (D), Vol.J67-D, no.10, pp.1242-1249 (1984).
- 125) 西村 拓一, 向井 理朗, 野崎 俊輔, 岡 隆一: 動作者適応のためのオンライン教示可能なジェスチャ動画像のスポッティング認識システム, 信学論 (D-II), Vol.J81-D-II, No.8, pp.1822-1830 (1998).
- 126) 西村拓一, 十河卓司, 小木しのぶ, 岡隆一, 石黒浩: 動き変化に基づく View-based Aspect Model による動作認識, 信学論 (D-II), Vol.J84-D2, No.10, pp.2212-2223 (2001).
- 127) 西村 拓一, 岡 隆一: 2 次元連続 DP による画像のスポッティング認識, 信学技報, PRMU97-55, pp.1-7 (1997).
- 128) Yuya Iwasa and Ryuichi Oka: Spotting recognition and tracking of a deformable object in a time-varying image using two-dimensional Continuous Dynamic Programming, *Proc. of the fourth Inter. Conf. on Computer and Information Technology*, pp.33-38 (2004).
- 129) Y. Itoh, J. Kiyama, and R. Oka: Speech understanding and Speech retrieval for TV program based on spotting algorithms, *Proc. of Spring meeting of ASJ*, 3-P-22(1995)[In Japanese].
- 130) 橋口 博樹, 西村 拓一, 張 建新, 滝田 順子, 岡 隆一: モデル依存傾斜制限型の連続 DP を用いた鼻歌入力による楽曲信号のスポッティング検索, 信学論 (D-II), Vol.J84-D2, No.12, pp.2479-2488 (2001).
- 131) 西村 拓一, 野崎 俊輔, 岡 隆一: Non-monotonic 連続 DP によるスポッティングに基づく移動ロボットの時系列画像を用いた大局的位置の推定, 信学論 (D-II), Vol.J81-D-II, No.8, pp.1876-1884 (1998).
- 132) 伊藤慶明, 木山次郎, 小島浩, 関進, 岡隆一: 時系列標準パターンの任意区間によるスポッティングのための Reference Interval-free 連続 DP(RIFCDP), 信学論 (D-II), J79-D-II, 9, pp.1474-1483 (1996).
- 133) 木山次郎, 伊藤慶明, 岡隆一: Incremental Reference Interval-free 連続 DP による任意話題音声の要約と話題境界検出, 信学論 (D-II), J79-D-II, 9, pp.1464-1473 (1996).
- 134) 西村 拓一, 古川 清, 向井 理朗, 岡 隆一: 時系列パターン検索のための重み減衰型 RIFCDP について, 信学論 (D-II), Vol.J81-D-II, No.3, pp.472-482 (1998).
- 135) 伊藤慶明: 時系列パターンの任意部分区間の高速マッチング手法 Shift CDP 法, 信学論 (D-II), Vol.J86-D2, No.9, pp.1267-1277 (2003).
- 136) 西村 拓一, 水野 道尚, 小木 しのぶ, 関本 信博, 岡 隆一: アクティブ探索法による時系列データ中の一致区間検出-参照区間自由時系列アクティブ探索法-, Vol.J84-D2, No.8, pp.1826-1837 (2001).
- 137) 杉山雅英: 複数時系列中の類似セグメント探索法の提案と評価, 信学技報 SP2005-196 (2006).