

A Unifying View of Component Analysis Methods (from a Computer Vision Perspective)

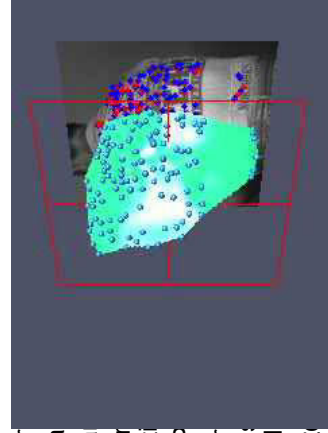
Fernando De la Torre
Carnegie Mellon
THE ROBOTICS INSTITUTE

Outline

- Introduction
- Generative models
 - Principal Component Analysis (PCA).
 - Non-negative Matrix Factorization (NMF).
 - Independent Component Analysis (ICA).
- Discriminative models
 - Linear Discriminant Analysis (LDA).
 - Canonical Correlation Analysis (CCA).
- Standard extensions of linear models.
 - Latent variable models.
 - Kernel methods.
 - Tensor factorization
- Other extensions
 - Filtered Component Analysis (active appearance models)
 - Learning Kernel expansions (support vector machines)
 - Parameterized cluster analysis (clustering)

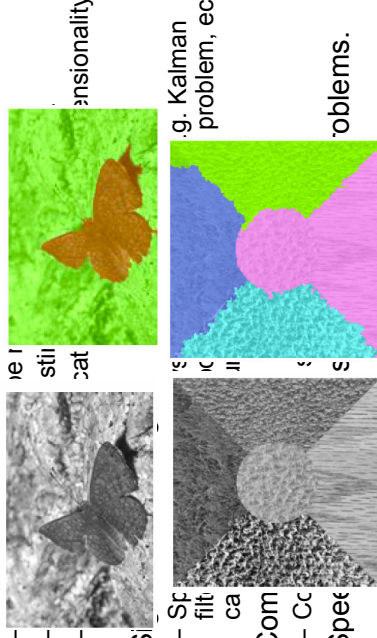
Component Analysis

- Computer Vision & Image Processing
 - **Structure from motion.**
 - Spectral graph methods for segmentation.
- Appearance
 - Fundamental
 - Compression (BRDF), synthesis,...
- Signal Processing
 - Spectral graph methods for segmentation.
- Computational
 - Compression (BRDF), synthesis,...
- Speech, bioinformatics, combinatorial problems.



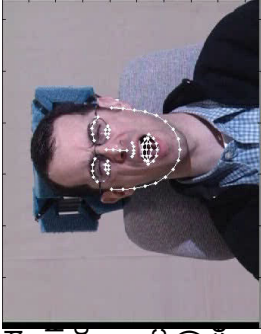
Component Analysis

- Computer Vision & Image Processing
 - Structure from motion.
 - **Spectral graph methods for segmentation.**
- Appearance
 - Fundamental
 - Compression (BRDF), synthesis,...
- Signal Processing
 - Spectral graph methods for segmentation.
- Computational
 - Compression (BRDF), synthesis,...
- Speech, bioinformatics, combinatorial problems.



Component Analysis

- Computer Vision & Image Processing
 - Structure from motion.
 - Spectral graph methods for segmentation.
 - **Appearance and shape models.**
 - Fundamental calibration.
 - Component reduction, dimensionality
- Signal processing
 - Spectral graph methods for segmentation (e.g. Kalman filter)
 - Cocktail problem, echo cancellation, ...
- Computer Graphics
 - Compression (BRDF), synthesis, ...
- Speech, bioinformatics, combinatorial problems.



Why Component Analysis?

- Learn from very high dimensional data and few samples.
 - Useful for dimensionality reduction.
- Easy to incorporate
 - Robustness to noise, missing data, outliers (Torre & Black '03)
 - Invariance to geometric transformations (Torre et al. '07, Torre & Black '02)
 - Non-linear (kernel methods) (Vapnik '95, Scholkopf '99, Cristiani & Taylor '04)
 - Probabilistic (latent variable models) (Everitt '84, Bishop et al 99, Benter '80)
 - Multi-factorial (tensors) (Rattiu '88, Vasilescu & Terzopoulos 02, Avidan & Shashua '97)
 - Exponential family PCA (Gordon '02)
- Efficient methods $O(d \times n)$
 - features d
 - samples n ($n < n^2$)

Are CA methods popular/useful/used?

- About 20% of CVPR-06 papers use CA.
- Google:
 - Results 1 - 10 of about 1,870,000 for "[principal component analysis](#)".
 - Results 1 - 10 of about 506,000 for "[independent component analysis](#)".
 - Results 1 - 10 of about 273,000 for "[linear discriminant analysis](#)".
 - Results 1 - 10 of about 46,100 for "[negative matrix factorization](#)".
 - Results 1 - 10 of about 491,000 for "[kernel methods](#)".
- Still work to do
 - Results 1 - 10 of about 65,300,000 for "[Britney Spears](#)".

Generative Models

$$D \approx BC$$

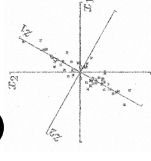
Principal Component Analysis (CA)

- **SVD** (Beltrami, 1873; Jordan 1874). **PCA** (Pearson, 1901; Hotelling, 1933)

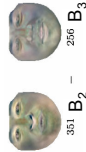
$$\mathbf{a} = \mathbf{D} = [\mathbf{d}_1 \mathbf{d}_2 \dots \mathbf{d}_n] \quad n = \text{images}$$



$$\mathbf{D}^T \mathbf{B} = \mathbf{B} \mathbf{A}$$



$$\text{Appearance, } \mathbf{B} = \mathbf{B}^0 + \mathbf{B}^1 + \mathbf{B}^2 + \mathbf{B}^3 + \dots$$

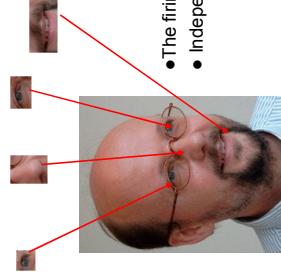


“Intercorrelations among variables are the bane of the multivariate researcher’s struggle for meaning”



Cooley and Lohnes, 1971

Part-based representation



- The firing rates of neurons are never negative.
- Independent representations.

NMF & ICA

Optimization for PCA

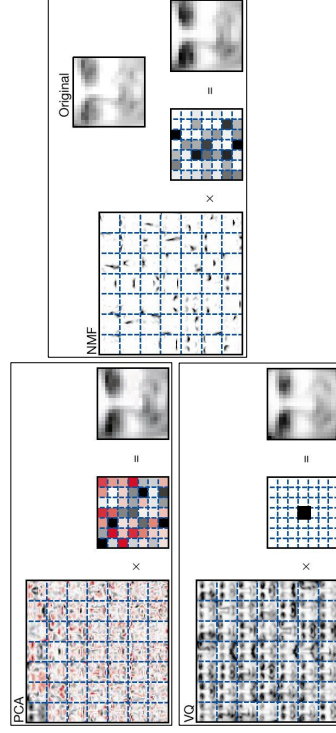
- PCA minimizes the following **CONVEX** function.

$$E_1(\mathbf{B}, \mathbf{C}) = \sum^n \|\mathbf{d}_i - \mathbf{B}\mathbf{c}_i\|_2^2 = \|\mathbf{D} - \mathbf{B}\mathbf{C}\|_F^2$$



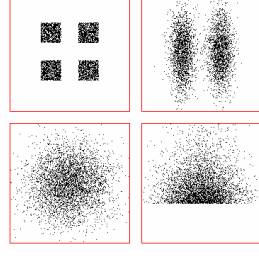
Non-negative Matrix Factorization

- Positive factorization.
 $E(\mathbf{B}, \mathbf{C}) = \|\mathbf{D} - \mathbf{BC}\|_F$ $\mathbf{B}, \mathbf{C} \geq 0$
- Leads to part-based representation.



Independent Component Analysis

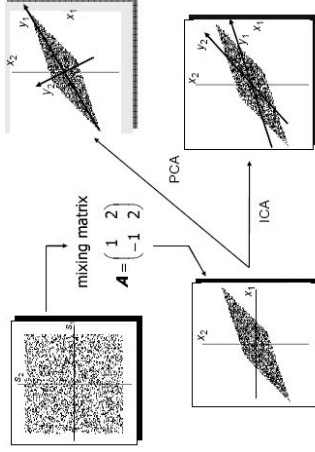
- We need more than second order statistics to represent the signal.



ICA vs PCA

(Heraut and Jutten 83; Olhansen & Field, 1996; Hyvriinen et al., 2001)

$$\mathbf{D} = \mathbf{BC} \quad \mathbf{C} \approx \mathbf{S} = \mathbf{WD} \quad \mathbf{W} \approx \mathbf{B}^{-1}$$



Many optimization criteria

- Minimize high order moments: e.g. kurtosis
 $\text{kurt}(\mathbf{W}) = E\{s^4\} - 3(E\{s^2\})^2$
- Many other information criteria.
- Also an error function: (Olhansen & Field, 1996)

$$\sum_{i=1}^n \|\mathbf{d}_i - \mathbf{B}\mathbf{c}_i\| + \sum_{i=1}^n S(\mathbf{c}_i)$$

Sparseness (e.g. $S = \|\cdot\|$)

- Other sparse PCA.

(Chennubhotla & Jepson, 2001b; Zou et al., 2005; dAspremont et al., 2004;)

Discriminative Models

- Linear Discriminant Analysis (LDA)
- Canonical Correlation Analysis (CCA)

Linear Discriminant Analysis (LDA)

(Fisher, 1938; Mardia et al., 1979; Bishop, 1995)

$$S_t = DD^T = \sum_{i=1}^n d_i d_i^T$$

$$S_b = \sum_{i=1}^C \sum_{j=1}^C (\mu_i - \mu_j)(\mu_i - \mu_j)^T$$

$$S_b B = S_t B A$$

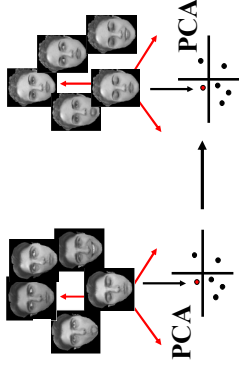
$$J(B) = \frac{|B^T S_b B|}{|B^T S_t B|}$$

$$S_w = \sum_{j=1}^C \sum_{i=1}^C (d_i - \mu_j)(d_i - \mu_j)^T$$

- Optimal linear dimensionality reduction if classes are Gaussian with equal covariance matrix.

Canonical Correlation Analysis

- PCA independently and general mapping



- Signals dependent signals with small energy can be lost.

Canonical Correlation Analysis (CCA)

(Mardia et al., 1979; Borga)

- Learn relations between multiple data sets? (e.g. find features in one set related to another data set)
- Given two sets $X \in \mathbb{R}^{d_x \times n}$ and $Y \in \mathbb{R}^{d_y \times n}$, CCA finds the pair of directions w_x and w_y that maximize the correlation between the projections (assume zero mean data)

$$\rho = \frac{w_x^T X^T Y w_y}{\sqrt{w_x^T X^T X w_x} \sqrt{w_y^T Y^T Y w_y}}$$

- Several ways of optimizing it:

$$A = \begin{bmatrix} 0 & X^T Y \\ X^T Y & 0 \end{bmatrix} \in \mathbb{R}^{(d_1+d_2) \times (d_1+d_2)}, \quad B = \begin{bmatrix} X^T X & 0 \\ 0 & Y^T Y \end{bmatrix} \in \mathbb{R}^{(d_1+d_2) \times (d_1+d_2)} \quad w = \begin{bmatrix} w_x \\ w_y \end{bmatrix}$$

- An stationary point of r is the solution to CCA.

$$A w = \lambda B w$$

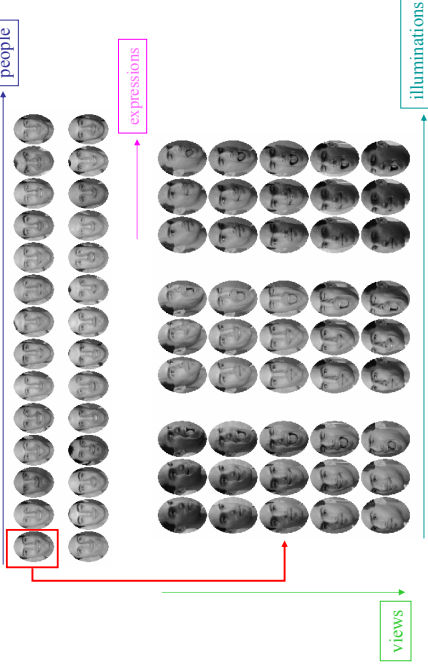
Standard extensions

- Tensor Factorization.
- Latent Variable Models
- Kernel Methods

Tensor Factorization

Tensor faces

(Vasilescu & Terzopoulos, 2002; Vasilescu & Terzopoulos, 2003)



Eigenfaces

- Facial images (identity change)
- Eigenfaces bases vectors capture the variability in facial appearance (do not decouple pose, illumination, ...)



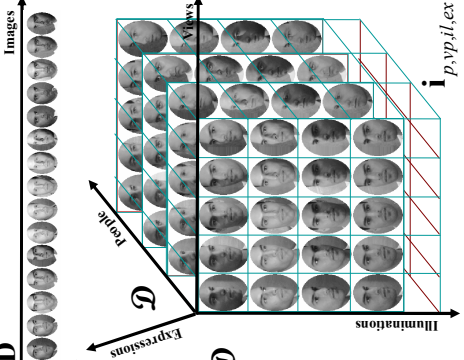
Data Organization

Linear/PCA: Data Matrix

- $R^{\text{pixels} \times \text{images}}$
- a matrix of image vectors

Multilinear: Data Tensor $\mathcal{D}$

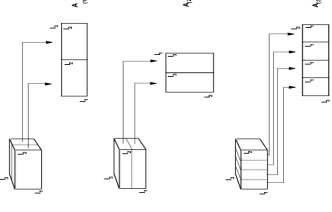
- $R^{\text{people} \times \text{views} \times \text{illums} \times \text{express} \times \text{pixels}}$
- N-dimensional matrix
- 28 people, 45 images/person
- 5 views, 3 illuminations,
- 3 expressions per person



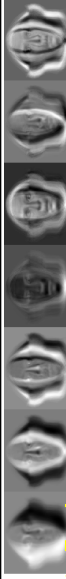
N-Mode SVD Algorithm

$$\mathcal{D} = \mathbf{Z} \mathbf{x}_1 \mathbf{U}_{\text{people}} \mathbf{x}_2 \mathbf{U}_{\text{views}} \mathbf{x}_3 \mathbf{U}_{\text{illums}} \mathbf{x}_4 \mathbf{U}_{\text{express}} \mathbf{x}_5 \mathbf{U}_{\text{pixels}}$$

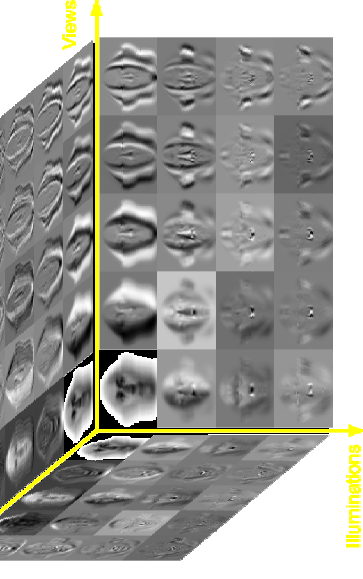
$N=3$



PCA:



TensorFaces:



Strategic Data Compression = Perceptual Quality

- TensorFaces data reduction in illumination space primarily degrades illumination effects (cast shadows, highlights)
- PCA has lower mean square error but higher perceptual error

TensorFaces		PCA	
Original	6 illum + 11 people param. 176 basis vectors	Mean Sq. Err. = 409.15 3 illum + 11 people param. 33 basis vectors	Mean Sq. Err. = 85.75 33 parameters 33 basis vectors

Latent Variable Models

Factor Analysis

- A Gaussian distribution on the coefficients and noise is added to PCA → Factor Analysis. (Mardia et al., 1979)

$$\mathbf{d} = \boldsymbol{\mu} + \mathbf{B}\mathbf{c} + \boldsymbol{\eta}$$

$$p(\mathbf{c}) = N(\mathbf{c} | \mathbf{0}, \mathbf{I}_k)$$

$$p(\mathbf{d} | \mathbf{c}, \mathbf{B}) = N(\mathbf{d} | \boldsymbol{\mu} + \mathbf{B}\mathbf{c}, \Psi)$$

$$p(\boldsymbol{\eta}) = N(\boldsymbol{\eta} | \mathbf{0}, \Psi)$$

$$\Psi = \text{diag}(\eta_1, \eta_2, \dots, \eta_d)$$

$$\text{cov}(\mathbf{d}) = E((\mathbf{d} - \boldsymbol{\mu})(\mathbf{d} - \boldsymbol{\mu})^T) = \mathbf{B}\mathbf{B}^T + \Psi$$

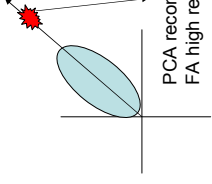
- Inference (Roweis & Ghahramani, 1999; Tipping & Bishop, 1999a)

$$p(\mathbf{c}, \mathbf{d}) \quad \text{Jointly Gaussian}$$

$$p(\mathbf{c} | \mathbf{d}) = N(\mathbf{c} | \mathbf{m}, \mathbf{V})$$

$$\mathbf{m} = \mathbf{B}^T (\mathbf{B}\mathbf{B}^T + \Psi)^{-1} (\mathbf{d} - \boldsymbol{\mu})$$

$$\mathbf{V} = (\mathbf{I} + \mathbf{B}^T \Psi^{-1} \mathbf{B})^{-1}$$

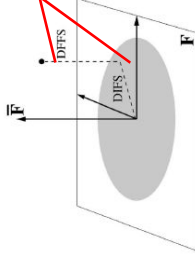


Ppca

- If $\Psi = E(\boldsymbol{\eta}\boldsymbol{\eta}^T) = \sigma_d^2 \mathbf{P}\mathbf{P}^T$.
- If $\varepsilon \rightarrow 0$ is equivalent to PCA. $\varepsilon \rightarrow 0 \quad \mathbf{B}^T (\mathbf{B}\mathbf{B}^T + \Psi)^{-1} = (\mathbf{B}^T \mathbf{B})^{-1} \mathbf{B}^T$
- Probabilistic visual learning (Moghaddam & Pentland, 1997.)

$$p(\mathbf{d}) = \int p(\mathbf{d} | \mathbf{c}) p(\mathbf{c}) d\mathbf{c} = \frac{e^{-\frac{1}{2}(\mathbf{d}-\boldsymbol{\mu})^T \Sigma^{-1}(\mathbf{d}-\boldsymbol{\mu})}}{(2\pi)^{\frac{d}{2}} |\Sigma|^{\frac{1}{2}}} = \frac{e^{-\frac{1}{2}(\mathbf{d}-\boldsymbol{\mu})^T (\mathbf{B}\mathbf{B}^T + \mathbf{I})^{-1}(\mathbf{d}-\boldsymbol{\mu})}}{(2\pi)^{\frac{d}{2}} |\Sigma|^{\frac{1}{2}}} = \frac{e^{-\frac{1}{2} \sum_{i=1}^d \frac{c_i^2}{\lambda_i}}}{(2\pi)^{\frac{d}{2}} \prod_{i=1}^d \lambda_i^{\frac{1}{2}}} \left[\frac{e^{-\frac{\varepsilon^2(\mathbf{d})}{2\sigma}}}{(2\pi\sigma)^{\frac{(d-k)}{2}}} \right]$$

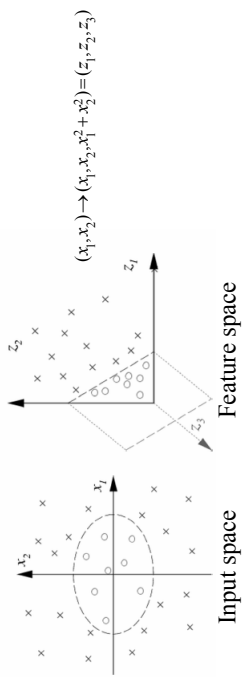
$$\mathbf{c}_i = \mathbf{B}^T \mathbf{d}_i$$



Other non-linear

- Mixture of Factor Analyzers
- Generative topographic mapping

Kernel methods

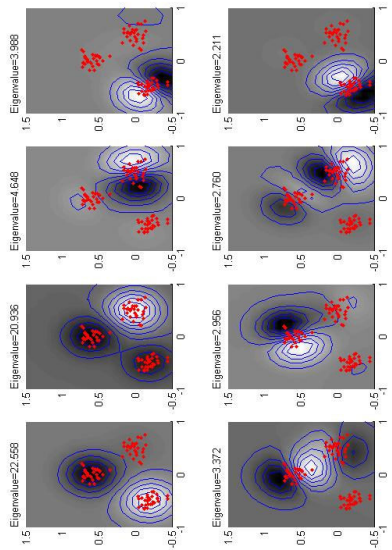


- Non-linear high dimensional mapping of the data so it becomes linearly separable in the feature space
- This is equivalent to finding a non-linear decision boundary in the input space
- Computation in the feature space can be costly because it is high dimensional
 - The feature space is typically infinite-dimensional!

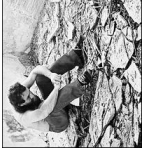
Second Part

Examples

$$\langle \mathbf{b}, \Phi(\mathbf{d}) \rangle = \sum_{i=1}^n \alpha_i K(\mathbf{d}_i, \mathbf{d})$$



Clinical Depression

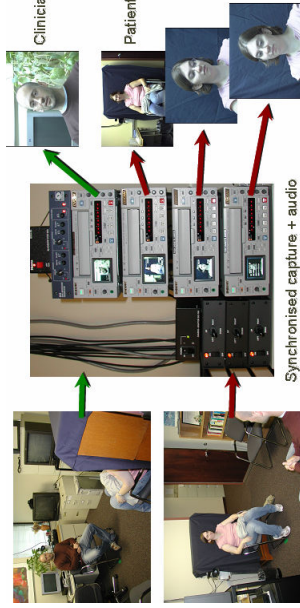


- Point Prevalence
 - 10% of the population
- High personal and social costs
- Current standards for diagnosis or therapeutic effectiveness of treatment relies on:
 - Interview with the physician.
 - Self-report measures.

1. DEPRESSED MOOD (Sadness, hopelessness, helplessness, worthlessness)

- 0= Absent
- 1= These feeling states indicated only on questioning
- 2= These feeling states spontaneously reported verbally
- 3= Communicates feeling states non-verbally—i.e., through facial expression, posture, voice, and tendency to weep
- 4= Patient reports VIRTUALLY ONLY these feeling states in his spontaneous verbal and non-verbal communication

Behavioral Analysis from Audio and Video



Problems to Solve

Tracking

Active Appearance Models

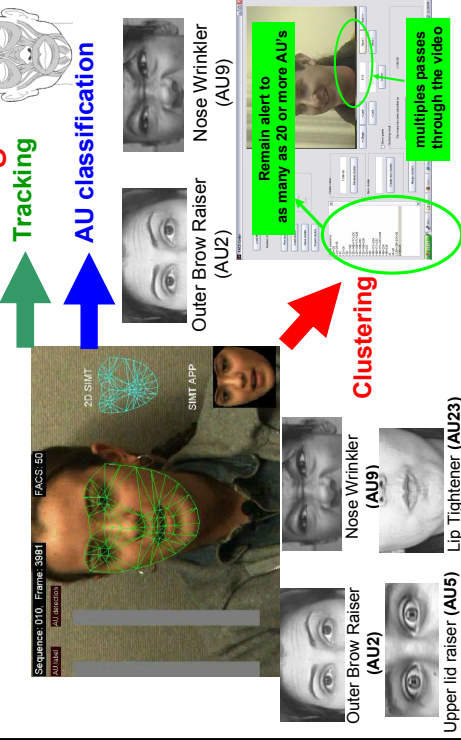
Classification
Support Vector Machine

Representation!

Spectral Graph Methods

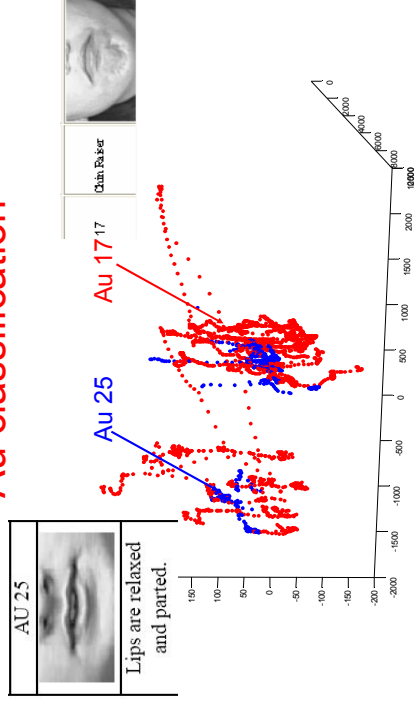
Single or multiples passes through the video

Automatic FACS coding



Classification

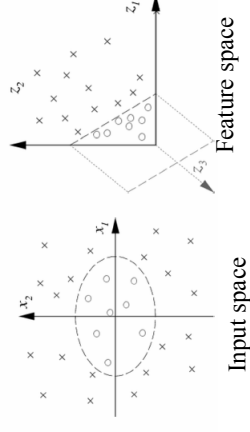
Au classification



Kernels

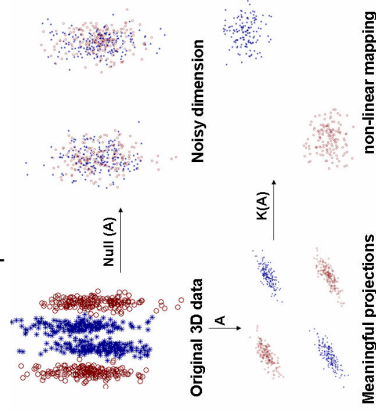
- Support Vector Machines
 - Large margin (QP- no local minima).
 - Kernels.**

$$(x_1, x_2) \rightarrow (x_1, x_2, x_1^2 + x_2^2) = (z_1, z_2, z_3)$$



Problem

- Learning a non-linear representation for classification

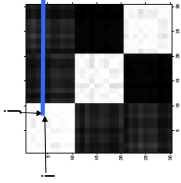


Previous work

- Neural networks
 - Hard to train, local minima, difficult interpretation of hidden units.
- Metric learning (Xing et al. '02, Goldberger et al. '05, He et al. '03, Weinberger '06, Lebanon '03, Shental '02)
- Kernels (Cristiani et al '01, Lanchkriet et al. '04, Chapelle et al. '02, Kwok et al. '03)
- Learning kernel expansions
- Kernel with dimensionality reduction
- Computational cost $O(dn^2)$ $\rightarrow$ samples features

Learning Kernel Expansions

$$\mathbf{K}(\mathbf{B}_1, \dots, \mathbf{B}_p, \lambda_1, \dots, \lambda_p, \boldsymbol{\alpha}) = \sum_{t=0}^p \alpha_t \hat{\mathbf{M}}_t \quad \alpha_t > 0$$

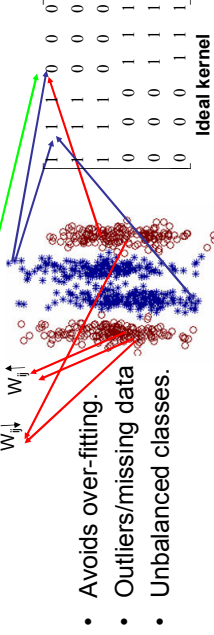


$$k_{ij} = \sum_{t=0}^p \alpha_t (\hat{m}_{ij})^t \quad \alpha_t > 0$$

$$\hat{m}_{ij}^t = \frac{\mathbf{d}_i^T (\mathbf{B}_t \mathbf{A}_t + \lambda_t \mathbf{I}) \mathbf{d}_j}{\sqrt{\mathbf{d}_i^T (\mathbf{B}_t \mathbf{B}_t^T \mathbf{A}_t + \lambda_t \mathbf{I}) \mathbf{d}_i} \sqrt{\mathbf{d}_j^T (\mathbf{B}_t \mathbf{B}_t^T \mathbf{A}_t + \lambda_t \mathbf{I}) \mathbf{d}_j}}$$

Learning from an Ideal Kernel

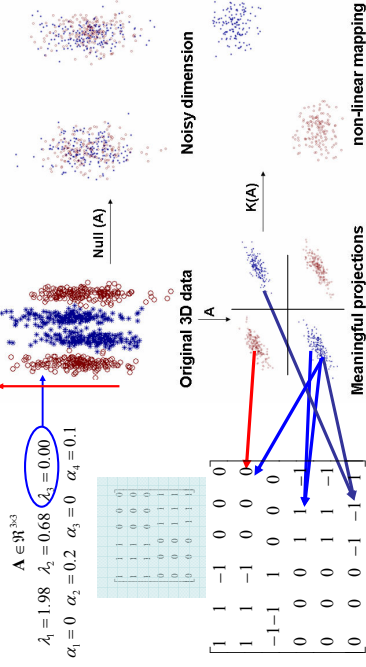
$$E = \left\| \mathbf{W} \circ (\mathbf{F} - \mathbf{K}(\mathbf{B}_1, \dots, \mathbf{B}_p, \lambda_1, \dots, \lambda_p, \boldsymbol{\alpha})) \right\|_F$$



- Avoids over-fitting.
- Outliers/missing data
- Unbalanced classes.

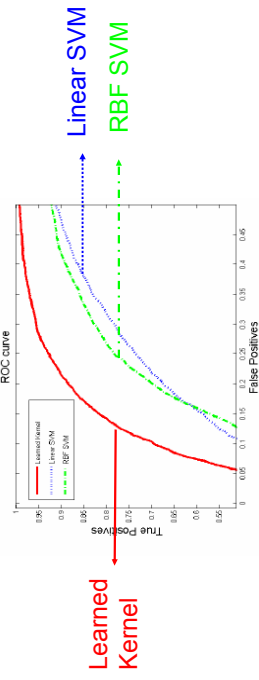
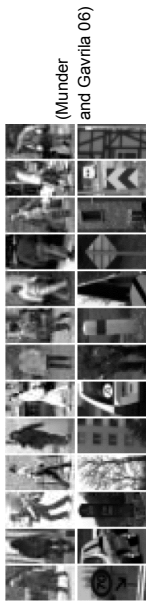
- First order (t=1) $E = \left\| \mathbf{F} - \mathbf{D}^T \mathbf{B} \mathbf{B}^T \mathbf{D} \right\|_F \rightarrow \text{LDA}$

Experiments



Pedestrian Detection

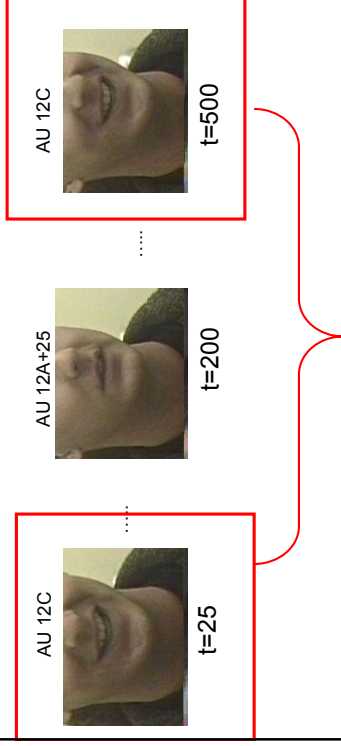
- Training (20,000) Cross validation (10,000) Testing (20,000)



Clustering

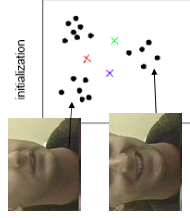
Clustering AUs

- How to cluster invariantly to geometric transformations to capture subtle facial behavior?



The Clustering Problem

- Partition the data set in c-disjoint "clusters" of data points.



- Number of possible partitions

$$S(n, c) = \frac{1}{c} \sum_{i=1}^c (-1)^i \binom{c}{i} i^n \approx 10^{12} \quad \begin{cases} n=21 \\ c=4 \end{cases}$$
- NP-hard and approximate algorithms (k-means, hierarchical clustering, mog, ...)

K-means

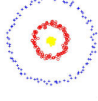
(MacQueen 67)

- K-means criterion:

$$E(\mathbf{B}, \mathbf{G}^T) = \|\mathbf{d}_1 \dots \mathbf{d}_n\| - \|\mathbf{b}_1 \dots \mathbf{b}_k\| \mathbf{G}^T \|\mathbf{F}\|$$

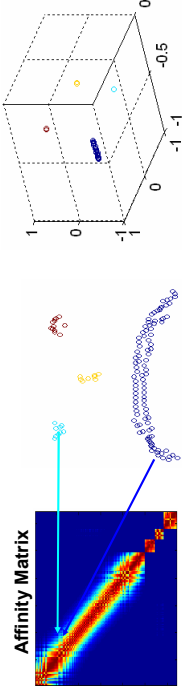
$$\mathbf{G}^T = \begin{bmatrix} 1 & \dots & 0 \\ 0 & \dots & 1 \\ 0 & \dots & 0 \end{bmatrix}$$

- Problems:
 - Local minima and sensitive to initial conditions
 - Optimal for spherical clusters



A Unifying View of Clustering

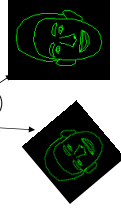
- K-means (Gu et al. 01*) $E(\mathbf{G}, \mathbf{B}) = \|\mathbf{D} - \mathbf{B}\mathbf{G}^T\|_F$
- Energy-based Spectral Clustering (Dhillon et al., '04; Torre et al '06).
 $E(\mathbf{G}, \mathbf{B}) = \|\mathbf{G} - \mathbf{B}\mathbf{G}^T\|_F \rightarrow$ Normalized Cuts (Shi & Malik '00)
Ratio-cuts (Hagen & Kahng '02)
 $\mathbf{\Gamma} = [\varphi(\mathbf{d}_1) \ \varphi(\mathbf{d}_2) \ \dots \ \varphi(\mathbf{d}_n)] \quad \mathbf{\Gamma} = \varphi(\mathbf{D})$



Parameterized Cluster Analysis

- $E_1(\mathbf{G}, \mathbf{B}) = \|\mathbf{\Gamma} - \mathbf{B}\mathbf{G}^T\|_F \quad \mathbf{\Gamma} = [\varphi(\mathbf{d}_1) \ \varphi(\mathbf{d}_2) \ \dots \ \varphi(\mathbf{d}_n)]$
- Parameterize the shape.

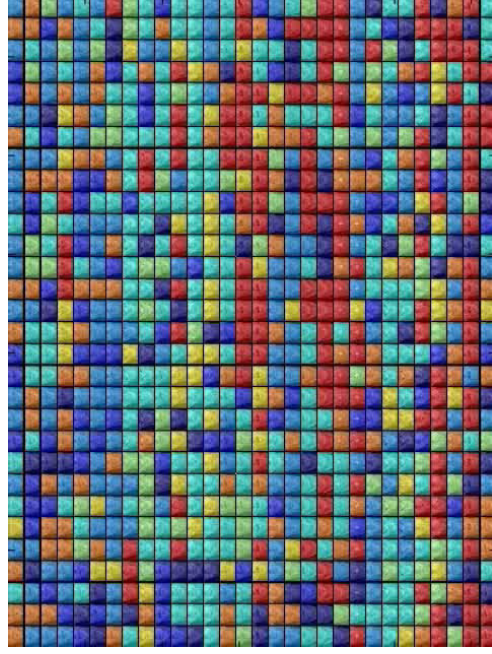
$$\mathbf{\Gamma} = [\varphi(\mathbf{A}\mathbf{d}_1) \ \varphi(\mathbf{A}_2\mathbf{d}_2) \ \dots \ \varphi(\mathbf{A}_n\mathbf{d}_n)]$$



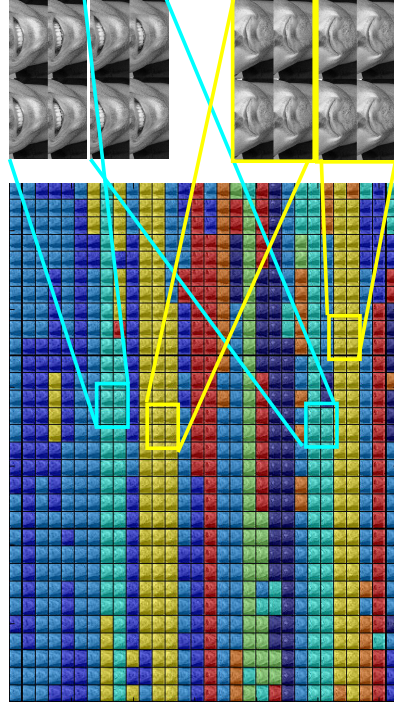
parameterized kernel

$$E_2(\mathbf{A}, \mathbf{G}) = \text{tr}(\mathbf{\Gamma}(\mathbf{A})^T \mathbf{\Gamma}(\mathbf{A})(\mathbf{I} - \mathbf{G}(\mathbf{G}^T \mathbf{G})^{-1} \mathbf{G}^T))$$

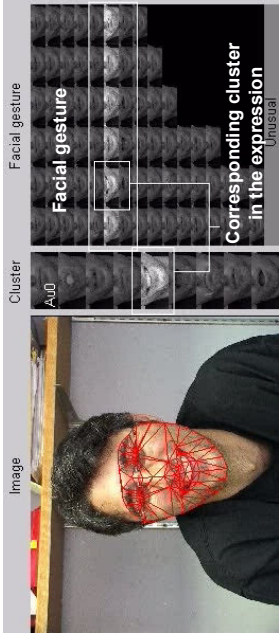
$$\mathbf{G} \geq \mathbf{0} \quad \mathbf{G}\mathbf{1}_c = \mathbf{1}_n$$



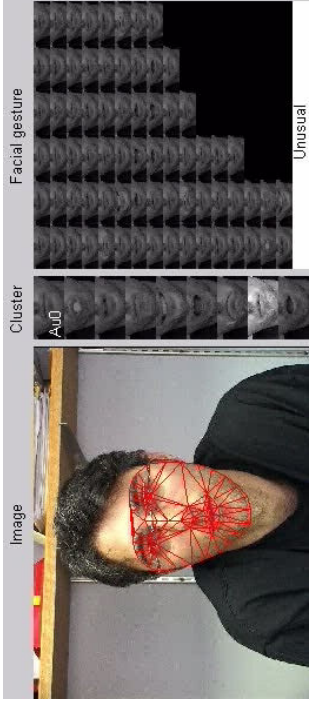
Example



Experiments



Experiments



An Energy-based Framework for Component Analysis

Modeling

- Filtered CA [CVPR-Torre et al. 07]
- Robust PCA [JCV-Torre & Black 03]
- Convolutional CA [Torre et al. in prep.]
- Parameterized CA [JVC-Torre & Black 03]
- Unsupervised bilinear [ICCV-Gross et al. 07]

$$\begin{bmatrix} \mathbf{D} \otimes \mathbf{f}_1 \\ \vdots \\ \mathbf{D} \otimes \mathbf{f}_F \end{bmatrix}_F$$

$$\|(\varphi(\mathbf{D}(\mathbf{x}, \mathbf{a})) - \mathbf{B}\mathbf{G})\mathbf{W}\|_F$$

Clustering

- Parameterized CA [CVPR-Torre et al. 07]
- Discriminative CA [ICML-Torre & Kanade 06]

Classification

- Learning kernel [CVPR-Torre & Vinyals 07]
- Multimodal Oriented DA [ICML-Torre et al. 05]
- Representational OCA [CVPR-Torre et al. 05]

$$E(\mathbf{B}, \mathbf{a}) = \left\| \sum_{i=1}^n \alpha_i \left[\frac{\mathbf{D}' \mathbf{B} \mathbf{B}' \mathbf{D}}{\text{normalization}} \right] - \mathbf{F} \right\|_F$$

$$E(\mathbf{B}, \mathbf{C}) = \|\mathbf{D} - \mathbf{B}\mathbf{C}\|_F$$

Visualization

- DCCA [CVPR-Torre & Black 02]
- Dynamic MDS [ICML-Cabero et al. 07]



Thanks

1st Workshop on **Component Analysis**
<http://www.cs.cmu.edu/~ftorre/ca/>