

DIMENSION REDUCTION : BAYESIAN MODELING AND INFERENCE

KANIAV KAMARY

PSPM-INSTITUT CAMILLE JORDAN

INSA, LYON

June 19, 2026

Outline

- 1 Introduction
- 2 Why Bayesian Modeling in Machine Learning?
- 3 Principal Component Analysis
- 4 Probabilistic PCA
- 5 Bayesian PCA



Machine Learning?

- Machine learning is fundamentally about learning patterns from data.

How machine learning work?



Image source

At its core, machine learning is concerned with inferring latent structures from observed data.

Machine Learning?

- Input / output data : image, text, speech and tabular

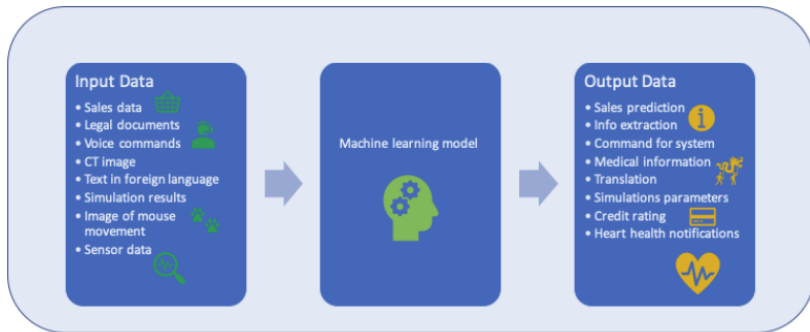


Image source

Fundamental question :

How can we model the mechanisms that generate data, and recover the hidden structures responsible for the observations?

Machine Learning?

Machine learning types by data content

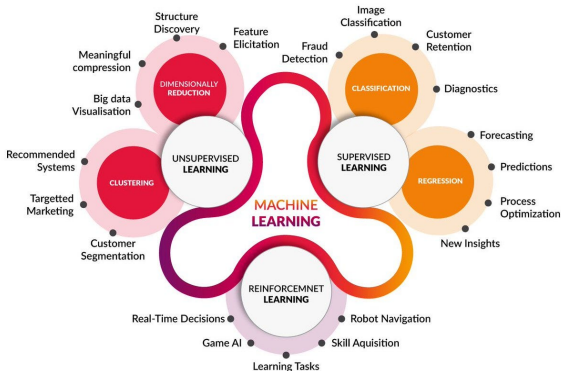


Image source

Supervised learning

- needs labeled data

Unsupervised learning

- finds groups with similarities

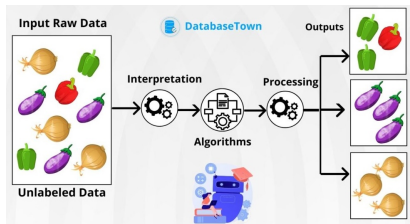
Reinforcement learning

- finds a policy to achieve a goal

Machine Learning?

Unsupervised Machine learning

Clustering



Dimensionality reduction

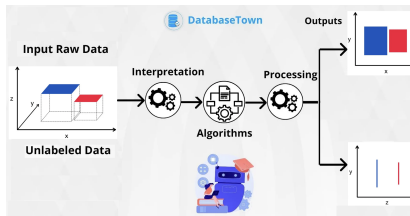


Image source

Why Bayesian Modeling in Machine Learning?

- Classical approaches : finding a single “best” estimate of unknown quantities.
- Bayesian methods : unknown quantities as random variables and quantify uncertainty about them.

Modern Bayesian machine learning builds probabilistic models of data and uses inference algorithms to learn hidden structures while quantifying uncertainty.

Principal Component Analysis

Classical PCA :

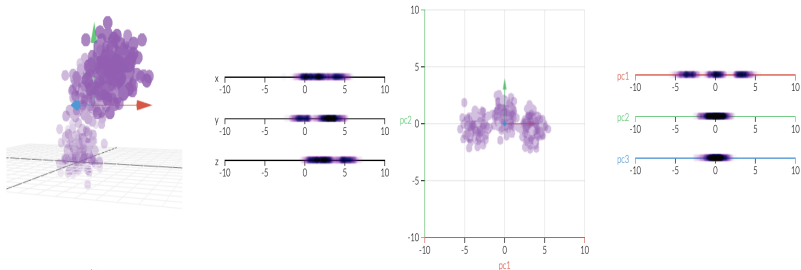
- a widely covered machine learning method frequently used for analyzing unsupervised datasets,
- linear projection of a data set,
- finding directions explaining maximum variance,
- for which the sum-of-squares reconstruction cost is minimized.

[Bishop, C. (1998)]

Used for several purposes :

- Data visualization : visualize multidimensional data in two or three dimensions,
- Uncorrelated variables : generates new uncorrelated variables,
- Dimension reduction : perform data dimension reduction in machine learning.

Principal Component Analysis



Data source

- horizontal axis PC1 has the most variation : 73% of the total variance,
- vertical axis PC2, the second-most : 25% of the total variance,
- third axis PC3, the least : 2% of the total variance.

Principal Component Analysis

For a data set of observed d -dimensional vectors $\{t_n\}$ where $n \in \{1, \dots, N\}$,

Classical PCA can be broken down into five steps :

- I. Standardize the range of continuous initial variables :

$$t'_n = \frac{t_n - \bar{t}}{\sigma(t)},$$

- II. Compute the covariance matrix to identify correlations

$$S = \frac{1}{N} \sum_{n=1}^N (t_n - \bar{t})(t_n - \bar{t})^T,$$

- III. Compute the eigenvectors u_i and eigenvalues λ_i of the covariance matrix to identify the principal components

$$\forall i = 1, \dots, d: \quad Su_i = \lambda_i u_i,$$

- IV. Create a feature vector to decide which principal components to keep ($q < d$),
- V. Recast the data along the principal components axes : reduced-dimensionality representation of the data set is defined by $x_n = U_q^T (t_n - \bar{t})$ where $U_q = (u_1, \dots, u_q)$

Principal Component Analysis

Classical PCA :

- But does not define a probability distribution (significant limitation),
- Where does uncertainty enter?
- How do we choose the number of components?
- How can PCA handle missing data?
- Can PCA be embedded into a probabilistic machine learning framework?

PCA can be reformulated as the maximum likelihood solution of a specific **latent variable model**.

[Tipping, M. & Bishop, C. (1999)]

Latent variable models

Suppose we observe $\mathcal{D} = \{\mathbf{t}\}$ with $\mathbf{t} \in \mathbb{R}^d$. Many datasets contain hidden factors :

- Intelligence behind exam scores
- Customer preferences behind purchases
- Biological processes behind gene expression
- Hidden directions of variation in high-dimensional data

General form of a latent variable model :

$$\mathbf{t} = f(\mathbf{x}) + \epsilon$$

where

- $\mathbf{x}$: latent variable
- ϵ : measurement noise

Probabilistic PCA

For $q < d$, $\{\mathbf{x}_i\}$'s are q -dimensional latent variables :

$$\mathbf{t}_i = \mathbf{W}\mathbf{x}_i + \boldsymbol{\mu} + \boldsymbol{\epsilon}_i,$$

$\mathbf{W} : d \times q$ matrix, $\boldsymbol{\mu} : d - \text{dimensional vector}$, $\boldsymbol{\epsilon}_i \sim \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I}_d)$.

[Tipping, M. & Bishop, C. (1999)]

Thus,

$$\mathbf{t}_i | \mathbf{x}_i, \mathbf{W}, \boldsymbol{\mu}, \sigma^2 \sim \mathcal{N}(\mathbf{W}\mathbf{x}_i + \boldsymbol{\mu}, \sigma^2 \mathbf{I}_d).$$

By the formulation above, PCA becomes a probabilistic generative model.

If $\mathbf{x}_i \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_q)$, then the marginal distribution of the observed variable is :

$$\mathbf{t}_i | \mathbf{W}, \boldsymbol{\mu}, \sigma^2 \sim \mathcal{N}(\boldsymbol{\mu}, \mathbf{M}), \quad \mathbf{M} = \mathbf{W}^T \mathbf{W} + \sigma^2 \mathbf{I}_q,$$

and the conditional distribution of $\mathbf{x}_i$:

$$\mathbf{x}_i | \mathbf{t}_i, \mathbf{W}, \boldsymbol{\mu}, \sigma^2 \sim \mathcal{N}(\mathbf{M}^{-1} \mathbf{W}^T (\mathbf{t}_i - \boldsymbol{\mu}), \sigma^2 \mathbf{M}^{-1}).$$

Probabilistic PCA

Maximum likelihood estimation :

The log likelihood under the observed data set $\mathcal{D} = \{\mathbf{t}_i\}$:

$$\mathcal{L}(\boldsymbol{\mu}, \mathbf{W}, \sigma^2; \mathcal{D}) = -\frac{N}{2} \left[d \ln(2\pi) + \ln |\mathbf{M}| + \text{Tr}[\mathbf{M}^{-1} \mathbf{S}] \right],$$

where $\mathbf{S}$ is the sample covariance matrix and

$$\boldsymbol{\mu}_{\text{ML}} = \bar{\mathbf{t}}, \quad \mathbf{W}_{\text{ML}} = \mathbf{U}_q \sqrt{\boldsymbol{\Lambda}_q - \sigma_{\text{ML}}^2 \mathbf{I}_q}, \quad \sigma_{\text{ML}}^2 = \frac{1}{d-q} \sum_{i=q+1}^d \lambda_i.$$

where the columns of $\mathbf{U}_q$ are eigenvectors of $\mathbf{S}^2$, with corresponding eigenvalues in the diagonal matrix $\boldsymbol{\Lambda}_q$.

The maximum of the likelihood is achieved when the q largest eigenvalues are chosen.

Question : But how to choose q ? **Solution :** Bayesian PCA.

Bayesian PCA

Bayesian method to determine automatically an appropriate effective dimensionality :

Under the probabilistic reformulation of the PCA,

Step.1 introduce a prior distribution over the parameters of the model : μ, W, σ^2 ,

Step.2 compute and infer the posterior distribution $p(\mu, W, \sigma^2 | \mathcal{D})$,

by focusing on the posterior distribution of W to obtain the effective dimensionality of the latent space, e.i., number of retained principal components.

Instead of estimating a single projection matrix, Bayesian PCA learns a distribution over possible projection matrices.

Bayesian PCA

Bishop, (1998)'s strategy to determine an appropriate effective dimensionality :

- start with many components $q = d - 1$,
- hierarchical prior $p(W|\alpha)$ over the matrix W ,

$$\forall j \in 1, \dots, q : \mathbf{w}_j | \alpha_j \sim \mathcal{N}(0, \alpha_j^{-1} \mathbf{I}_d); \quad \alpha_j \sim \mathcal{G}(a_{\alpha_j}, b_{\alpha_j})$$

- infer the posterior distribution of the parameters.

Interpretation :

- each hyper-parameter α_j controls a column of the matrix W ,
- large $\alpha_j \rightarrow \infty$, implies $\mathbf{w}_j \rightarrow 0$ (since $\text{Var}(\mathbf{w}_j) = 1/\alpha_j$),
- irrelevant components automatically shrink toward zero.

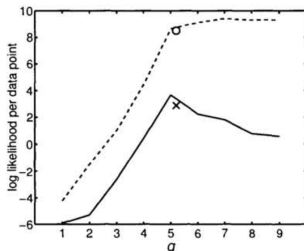
The effective dimensionality of the model is the number of vectors $\mathbf{w}_j$ whose values remain non-zero.

[Neal & MacKay, (1995)], [Bishop, (1998)]

Bayesian PCA : Simulation Study

Example

- 10 data sets in 10 dimensions,
- generated from a Gaussian distribution,
- having standard deviations in 5 directions given by $\{1.0, 0.8, 0.6, 0.4, 0.2\}$ and standard deviation 0.04 in the remaining 5 directions,



Training set (dashed curve), Test set (solid curve)

- q_{eff} is averaged over 10 independent experiments,

[Bishop, (1998)]

Bayesian PCA

Mixture prior distribution for $\mathbf{w}_j$:

$$\mathbf{w}_j | p, \alpha \sim (1 - p) \delta_0(\mathbf{w}_j) + p(1 - \delta_0(\mathbf{w}_j)) \mathcal{N} \left(0, \frac{1}{\alpha} \mathbf{I}_d \right), \quad j = 1, \dots, q,$$

$$p \sim \text{Beta}(c_0, c_1), \quad \alpha \sim \mathcal{G}(a_\alpha, b_\alpha),$$

$$\boldsymbol{\mu} | \beta \sim \mathcal{N} \left(0, \frac{1}{\beta} \mathbf{I}_d \right), \quad \beta \sim \mathcal{G}(a_\beta, b_\beta), \quad \tau = \frac{1}{\sigma^2} \sim \mathcal{G}(a_\tau, b_\tau).$$

■ Advantages :

- posterior inference on the parameters by applying a simple Gibbs sampler
- data dimension reduction is performed automatically during the MCMC procedure

■ Disadvantages :

- high number of hyper-parameters to be identified
- how to select the hyper-parameters a_α, b_α ?

Bayesian PCA

Noninformative prior modeling :

Mixture prior distribution for $\mathbf{w}_j$:

$$\mathbf{w}_j | p, \alpha \sim (1 - p)\delta_0(\mathbf{w}_j) + p(1 - \delta_0(\mathbf{w}_j))\mathcal{N}\left(0, \frac{1}{\alpha}\mathbf{I}_d\right), \quad j = 1, \dots, q,$$

Noninformative prior for p and exponential distribution for α :

$$p \sim \mathcal{Beta}(c_0, c_1),$$

Jeffreys prior for $(\boldsymbol{\mu}, \sigma^2)$:

$$\pi(\boldsymbol{\mu}, \tau) \propto \frac{1}{\sqrt{\tau}}, \quad \tau = \frac{1}{\sigma^2}$$

And for α :

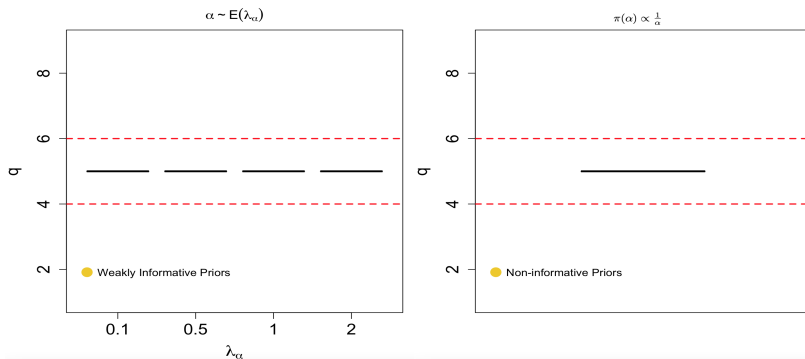
$$\alpha \sim \mathcal{Exp}(\lambda_\alpha), \quad \text{or} \quad \pi(\alpha) \propto 1/\alpha.$$

Resulting conditional posterior distributions with closed-form PDF, let us to use a Gibbs sampler to obtain the posterior estimates of the parameters.

Bayesian PCA : Illustrations

50 independent data simulations, each of size $n = 20$ points with $d = 10$ from $\mathcal{N}_d(\mathbf{0}, \sigma^2 \mathbf{I}_d)$:

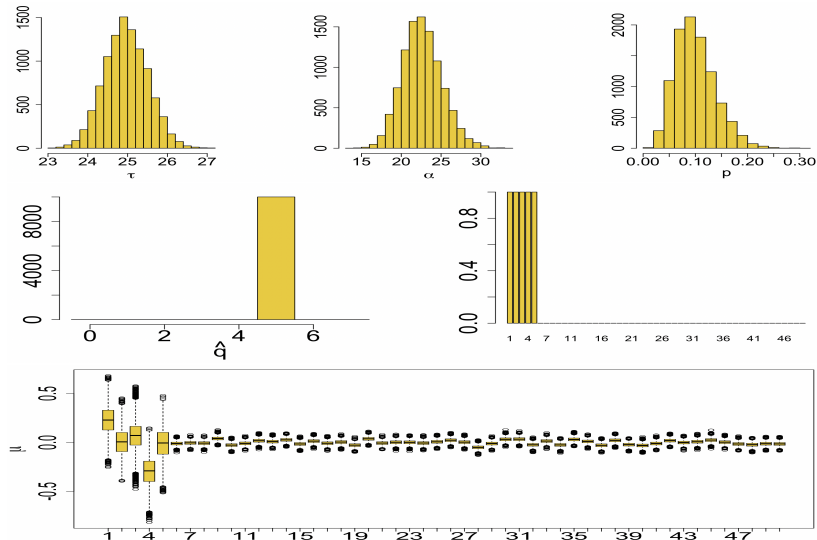
- σ^2 : in 5 directions given by $(1.0, 0.8, 0.6, 0.4, 0.2)^2$ and 0.04^2 in the remaining 5 directions



Note that for each λ_{α} , 50 values of $\hat{q}$ are the sample average of the MCMC draws (with 9000 discarded draws as burn-in).

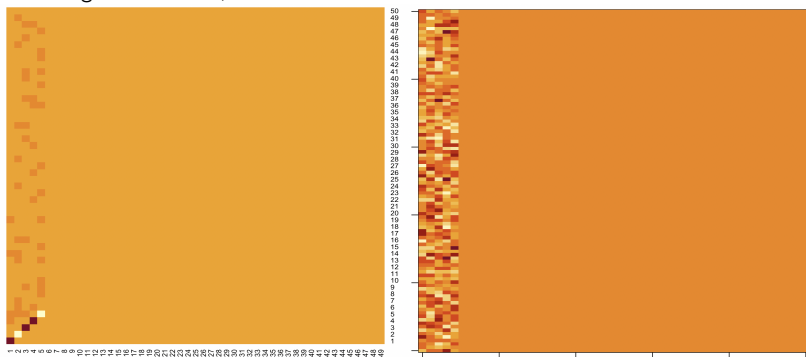
Bayesian PCA : Illustrations

$n = 100$ with $d = 50$ from $\mathcal{N}_d(\mathbf{0}, \sigma^2 \mathbf{I}_d)$ with $\sigma^2 = 2$: in 5 directions and 0.04 in the remaining 45 directions;



Bayesian PCA : Illustrations

$n = 100$ with $d = 50$ from $\mathcal{N}_d(\mathbf{0}, \sigma^2 \mathbf{I}_d)$ with $\sigma^2 = 2$: in 5 directions and 0.04 in the remaining 45 directions;

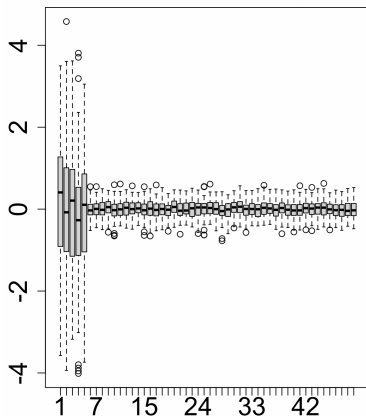


Posterior mean of :
(Left) $\hat{W}$; (Right) $\hat{X}$.

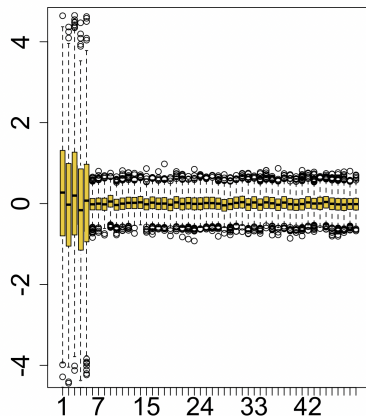
Bayesian PCA : Illustrations

$n = 100$ with $d = 50$ from $\mathcal{N}_d(\mathbf{0}, \sigma^2 \mathbf{I}_d)$ with $\sigma^2 = 2$: in 5 directions and 0.04 in the remaining 45 directions;

Uncertainty quantification?



(Left) Data distribution



(Right) Posterior predictive distribution

Bayesian PCA : Illustrations

$n = 100$ with $d = 50$ from $\mathcal{N}_d(\mathbf{0}, \sigma^2 \mathbf{I}_d)$ with $\sigma^2 = 2$: in 5 directions and 0.04 in the remaining 45 directions;

Uncertainty quantification?

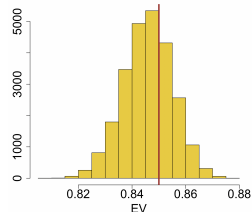
- Posterior distribution of q :

d	1	2	3	4	5	6	7	> 7
$\mathbb{P}(q = d \dots)$	$2e-5$	$3.6e-4$	$6.4e-4$	$3.32e-3$	0.995	$2.2e-4$	$6e-5$	0

- Bayesian explained variance :

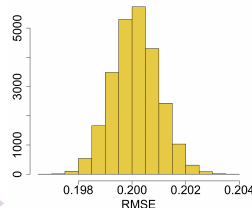
$$EV = \frac{1}{L} \sum_{l=1}^L \frac{\text{tr}(\mathbf{W}^{(l)} \mathbf{W}^{(l)T})}{\text{tr}(\mathbf{W}^{(l)} \mathbf{W}^{(l)T}) + d\sigma^2}; \quad \hat{EV} = 0.85$$

L : number of MCMC iterations.



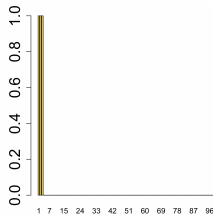
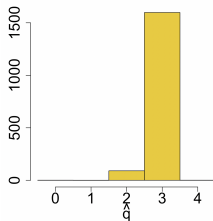
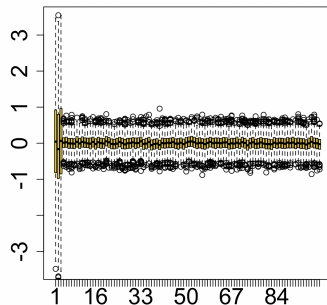
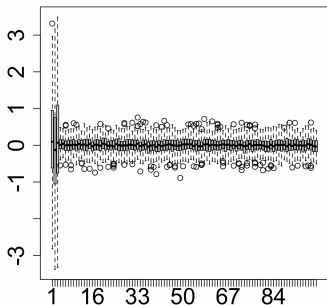
- Posterior distribution of the RMSE :

$$\hat{\mathbf{t}}_i = \hat{\mu}^{(l)} + \mathbf{W}^{(l)} \mathbf{x}_i; \quad \text{RMSE}^{(l)} = \sqrt{\frac{1}{nd} \sum_{i=1}^n \|\mathbf{t}_i - \hat{\mathbf{t}}_i\|^2}.$$



Bayesian PCA : Illustrations

$n = 120$ with $d = 100$ from $\mathcal{N}_d(\mathbf{0}, \sigma^2 \mathbf{I}_d)$ with $\sigma^2 = 2$: in 3 directions and 0.04 in the remaining 97 directions;



Conclusion

- Classical PCA :
 - choose the latent space dimensionality q manually,
- Bayesian PCA :
 - offers a significant advantage in allowing the model discover hidden structure in the data automatically,
 - quantify uncertainty,
 - using noninformative prior modeling leads to reduce the number of model hyper-parameter,
 - easy implementation by applying a Gibbs sampler algorithm,
 - particularly advantageous for small data sets in high dimensions.

References

- Bishop, C. (1998). *Bayesian pca*. Advances in neural information processing systems, 11.
- Tipping, M., & Bishop, C. (1999). *Probabilistic principal component analysis*. Journal of the Royal Statistical Society, 11(3), 611-622.
- Oh, H. S., & Kim, D. G. (2010). *Bayesian principal component analysis with mixture priors*. Journal of the Korean Statistical Society, 39(3), 387-396.

Thank you

