

Méthode GPR: Gaussian process regression et applications:

Motivation:

On rappelle le problème de régression linéaire:

$$y = f(x) + \varepsilon,$$

avec $f(x) = x^T w = \langle x, w \rangle = \sum_i w_i x_i$

où x est un vecteur, w le vecteur des poids (paramètres) du modèle linéaire. Souvent il y a un biais qu'on peut compenser en augmentant la dimension du vecteur x par élément additionnel toujours égal à 1 (Dans la notation il faut ajouter 1 aux "inputs" d'origine x $x = (1, x)$ (abus de notation)).

exple

$$x_0 = 1, x_1 = x, \dots$$

Dans ce modèle, la différence entre y et x est donnée par le bruit additif ε , qui peut être considéré comme bruit blanc fort Gaussien

(i.i.d)

$$\varepsilon \sim \mathcal{N}(0, \sigma^2)$$

Une conséquence directe de ce modèle, on peut expliciter la vraisemblance associée, sachant les paramètres, grâce à l'indépendance, la densité de probabilité de l'observation sachant les paramètres :

$$p(Y/X, w) = \prod_{i=1}^n p(y_i / x_i, w)$$

$$= \prod_{i=1}^n \frac{1}{\sqrt{2\pi} \sigma_n} \exp\left(-\frac{(y_i - x_i^T w)^2}{2\sigma_n^2}\right).$$

X la matrice formée par les colonnes $(x_i)_{i=1, \dots, n}$

$$= \frac{1}{(2\pi\sigma_n^2)^{n/2}} \exp\left(-\frac{1}{2\sigma_n^2} \|Y - X^T w\|^2\right) (*)$$

$$= \mathcal{N}(X^T w, \sigma_n^2 I_d).$$

$\|\cdot\|$ désigne la norme euclidienne.

Dans l'approche Bayésienne, on doit spécifier une loi a priori pour les paramètres qui reflètent notre intuition sur ce que peuvent être les paramètres.

Par défaut, on met une loi a priori gaussienne de matrice de covariance Σ_P sur les poids :

$$w \sim \mathcal{N}(0, \Sigma_P)$$

Rappels:

(i) la densité d'une gaussienne multivariée:

$$p(x|m, \Sigma) = \frac{1}{(2\pi)^{\frac{D}{2}} |\Sigma|^{\frac{1}{2}}} \exp \left(-\frac{1}{2} (x-m)^T \Sigma^{-1} (x-m) \right)$$

où m : vecteur moyenne (taille D)

Σ : matrice de covariance (symétrique définie positive)
de taille $D \times D$

et on écrit $x \sim N(m, \Sigma)$

(ii) Soit x, y de loi jointe gaussienne

$$\begin{pmatrix} x \\ y \end{pmatrix} \sim N \left(\begin{pmatrix} \mu_x \\ \mu_y \end{pmatrix}, \begin{bmatrix} A & C \\ C^T & B \end{bmatrix} \right) \quad \text{alors la loi marginale}$$

de x conditionnelle à y est:

$$x \sim N(\mu_x, A), \quad x/y \sim N(\mu_x + C \bar{B}^{-1} (y - \mu_y), A - C \bar{B}^{-1} C^T)$$

$$(ii) \quad p(x/y) = \frac{p(x) p(y/x)}{p(y)}$$

(rem.: applique deux fois la formule de Bayes $p(x/y)$
et $p(y/x)$).

Dans le problème d'inférence Bayésienne, on utilise:

$$\text{La posteriori} = \frac{\text{la priori} \times \text{vraisemblance}}{\text{la marginale vraisemblance}}$$

$$\Rightarrow \underset{\substack{P(w/y) \\ P_x(w/y)}}{P(w/y, x)} = \frac{P(w) \cdot P(y/x, w)}{P(y/x)} = \frac{P(w) P_x(y/w)}{P_x(y)} \quad (**)$$

Notons que, le dénominateur ne dépend pas de w et on a:

$$P(y/x) = \int P(y/x, w) P(w) dw$$

ne dépend pas de w

Proposition:

Où a:

$$P(w/x, y) = \underset{\substack{\uparrow \\ \text{ne dépend pas de } w}}{c^{\text{te}}} \times \exp\left(-\frac{1}{2} (w - \bar{w})^T \left(\frac{XX^T}{\sigma_n^2} + \Sigma_p^{-1} \right) (w - \bar{w})\right)$$

$$\bar{w} = \frac{1}{\sigma_n^2} \left(\frac{XX^T}{\sigma_n^2} + \Sigma_p^{-1} \right)^{-1} X y \quad \text{!}$$

indication: Comme $P(y/x)$ ne dépend pas de w on peut le mettre de c^{te} , donc il suffit de calculer le numérateur de $(**)$ en utilisant $(*)$

Preuve: Par la formule $(**)$:

$$P(w/x, y) = c^{\text{te}} \exp\left(-\frac{1}{2} w^T \Sigma_p^{-1} w\right) \exp\left(-\frac{1}{2\sigma_n^2} \overbrace{(y - X^T w)^T (y - X^T w)}^{|y - X^T w|^2}\right)$$

Dans la preuve on fait le terme donné dans la proposition

de sorte à récupérer à nouveau une loi normale
au prix de multiplier et diviser par des constantes
qui ne dépendent pas de w , ces constantes n'impactent
pas l'optimisation en w .

C.C.: On a:

$$p(w/X, y) \sim \mathcal{N}(\bar{w}, A^{-1})$$

$$\text{où } \bar{w} = A^{-1} \frac{Xy}{\sigma_n^2} \text{ et } A^{-1} = \frac{XX^T}{\sigma_n^2} + \Sigma_p^{-1}$$

En partant d'une loi a priori $w \sim \mathcal{N}(\theta, \Sigma_p)$
on obtient une loi a posteriori $w_{|X, y} \sim \mathcal{N}(\bar{w}, A^{-1})$

I - Projection des données dans l'espace des caractéristiques

L'inconvénient du modèle précédent est qu'il a des limites puisque ça reste un modèle linéaire.

Une idée intuitive pour améliorer cela est de projeter les "inputs" dans un espace de plus haute dimension d'abord en utilisant un ensemble de bases de fonctions et puis appliquer l'idée de régression linéaire dans cet espace.

Par exemple, un scalaire x peut être projeté dans l'espace des puissances de x : $\phi(x) = (1, x, x^2, \dots)^T$, une régression linéaire avec $\phi(x)$, on obtient une régression multiple.

Tant que les projections sont des fonctions fixées au préalable (indép. de w) ceci est tiré dans la classe des modèles paramétriques.

La mise en application de cette idée nécessite de savoir comment choisir cette base de fonctions. Le formalisme des GPR permet de répondre à cette question.

Dans la suite, on suppose que la fonction de base ϕ est donnée. Plus précisément, on introduit $\phi(x)$ qui est une application d'un espace D -dimensionnel (input $x \in \mathbb{R}^D$) vers un espace N -dimensionnel des caractéristiques.

Notation: On note $\Phi(X)$ la matrice formée par les colonnes $\phi(x)$ données dans le "training set".

Maintenant, le modèle devient

$$f(x) = \phi(x)^T w \quad \text{où maintenant } w \in \mathbb{R}^N$$

→ l'analyse de ce modèle est analogue à l'étude de régression

linéaire précédente; il faut juste remplacer x par $\Phi(x)$

Par analogie, la prédiction f_* est donnée par la loi.

$$f_* / x_*, X, Y \sim \mathcal{N} \left(\frac{\Phi(x_*)^T A^{-1} \Phi^T Y}{\sigma_n^2}, \Phi(x_*)^T A^{-1} \Phi(x_*) \right)$$

on

$$\Phi := \Phi(X) \text{ et } A = \frac{\Phi \Phi^T}{\sigma_n^2} + \Sigma_p^{-1}$$

Rq: Faire des prédictions à l'aide de cette équation, nécessite d'inverser A de taille $N \times N$ qui peut être problématique quand N est grand.

Afin de résoudre le problème, on peut réécrire l'équation précédente comme:

$$f_* / x_*, X, Y \sim \mathcal{N} \left(\Phi_*^T \Sigma_p \Phi (K + \sigma_n^2 I)^{-1} Y, \right. \\ \left. \Phi_*^T \Sigma_p \Phi_* - \Phi_*^T \Sigma_p \Phi (K + \sigma_n^2 I)^{-1} \Phi^T \Sigma_p \Phi_* \right)$$

on

$$\Phi_* = \Phi(x_*) \text{ et } K = \Phi^T \Sigma_p \Phi$$

Exercice 17.9 Les termes d'espérance dans les deux
formules précédentes sont égaux.
Montrer l'égalité des variances en utilisant le lemme
ci-dessous :

Lemme

$$(Z + UV^T)^{-1} = Z^{-1} - Z^{-1}U(W^{-1} + V^T Z^{-1}U)^{-1}V^T Z^{-1}$$

$$Z \in \mathbb{R}^{n \times n}, \quad W \in \mathbb{R}^{m \times m}, \quad U, V \in \mathbb{R}^{n \times m}$$

Il faut appliquer ce lemme avec

$$\left[\begin{array}{l} Z^{-1} = \Sigma_p, \quad W^{-1} = \sigma_n^2 I \quad ; \quad V = U = \Phi. \end{array} \right.$$

(i.) d'avantage de cette deuxième représentation
de la prédiction f_x est qu'on inverse des
matrices de taille $n \times n$ ce qui est mieux quand
 $n < N$.

(ii.) Dans cette nouvelle expression de la prédiction f_x
on voit tjrs apparaître :

$$\Phi^T \Sigma_p \Phi, \quad \Phi_*^T \Sigma_p \Phi \quad \text{ou} \quad \Phi_*^T \Sigma_p \Phi_*, \quad \text{en d'autres} \\ \text{termes on voit tjrs apparaître la forme } \Phi(x)^T \Sigma_p \Phi(x')$$

où $x, x' \in \{x, x_0\}$.

En partant de cette déf. on peut introduire

$$k(x, x') = \Phi(x)^T \Sigma_p \Phi(x').$$

Dans la suite, cette fonction k est appelée **fonction de covariance** ou tout simplement **noyau**.

→ Comme Σ_p est symétrique définie positive, on peut trouver $\Sigma_p^{1/2}$ t.q. $(\Sigma_p^{1/2})^2 = \Sigma_p$.

On peut alors introduire $\psi(x) = \Sigma_p^{1/2} \Phi(x)$, on peut écrire.

$$k(x, x') = \psi(x)^T \psi(x') \text{ comme produit scalaire.}$$

II - Point de vue espace fonctionnel.

Déf: Un processus gaussien est une famille de v.a. t.q. toute sous famille de dimension finie forme un vecteur gaussien.

Un processus gaussien est complètement caractérisé par sa fonction moyenne et sa fonction covariance.

On définit, une fonction moyenne $m(x)$ et une fonction de covariance $k(x, x')$ d'un processus réel

$$f(x) \text{ comme : } m(x) = \mathbb{E}[f(x)]$$

$$k(x, x') = \mathbb{E}[(f(x) - m(x))(f(x') - m(x')))]$$

et on écrit le processus gaussien comme

$$f(x) \sim GP(m(x), k(x, x'))$$

Dans la pratique $m(x)=0$; mais on peut proposer d'autres moyennes.

⚠ Généralement, les processus gaussiens sont définis comme v.a dépendant du temps t ($(x_t)_{t \geq 0}$ est un processus gaussien).

Dans l'utilisation GPR ce n'est pas le cas, les v.a sont indexés par un ensemble $\mathcal{X}$: ensemble de "inputs" possibles qui est généralement $\mathbb{R}^D$.

→ On se permet l'abus de notation

$f_i = f(x_i)$ est la v.a correspondant au couple (x_i, y_i)

$$(y_i = f(x_i) + \epsilon)$$

⚠ Dans les méthodes GP, on suppose l'hypothèse de consistance vérifiée:

$$\text{Si } (U, V) \sim N(\mu, \Sigma) \Rightarrow U \sim N(\mu_1, \Sigma_{11}) \quad \tilde{U} \quad \Sigma_{11} \text{ est}$$

la sous-matrice adéquate de Σ

Exemple: GP par le modèle Bayésien linéaire.

Dans ce cas, $f(x) = \phi(x)^T w$ et $w \sim N(0, \Sigma_p)$ li à priori

$$\mathbb{E}[f(x)] = \phi(x)^T \mathbb{E}[w] = 0$$

$$\mathbb{E}[f(x)f(x')] = \phi(x)^T \mathbb{E}[ww^T] \phi(x') = \phi(x)^T \Sigma_p \phi(x')$$

i.e $f(x)$ et $f(x')$ sont conjointement gaussien de moyenne 0 et de covariance:

$$\Phi(X)^T \Sigma_p \Phi(X)$$

Ainsi, les valeurs des fonctions $f(x_1), \dots, f(x_n)$ correspondent à une loi gaussienne multivariée.

→ Dans la suite, on se focalise sur une fonction de covariance spécifique donnée par :

$$\text{Cov}(f(x_p), f(x_q)) = k(x_p, x_q) = \exp(-\frac{1}{2} |x_p - x_q|^2)$$

Notons que la covariance entre les outputs $f(x_p)$ et $f(x_q)$ sont écrites comme une fonction des inputs x_p et x_q

§ Prediction avec des observations sans bruit.

On considère, le cas simple où les observations ne contiennent pas de bruit i.e on observe directement

$$\{(x_i, \overset{\text{"}y_i\text{"}}{f_i}), \text{ i.e } i = 1, \dots, n\} \quad \text{i.e.} \quad \boxed{y = f(x)}$$

donc la loi jointe des "training outputs" f et des "testing outputs" f_*

$$\begin{bmatrix} f \\ f_* \end{bmatrix} \sim \mathcal{N}\left(0, \begin{bmatrix} k(x, x) & k(x, x_*) \\ k(x_*, x) & k(x_*, x_*) \end{bmatrix}\right)$$

S'il y a n : "training points" et n_* : "testing points"

alors $k(x, x_*)$ n'est autre que la matrice $n \times n_*$ des covariances calculées entre toutes les paires "training" et "testing"

et donc on peut en déduire que

$$f_{x_* | x_*, x, f} \sim \mathcal{N} \left(K(x_*, x) K(x, x)^{-1} \overset{y}{f}, \right. \\ \left. K(x_*, x_*) - K(x_*, x) K(x, x)^{-1} K(x, x_*) \right)$$

$y \equiv f$ observable

Prediction avec des observations bruitées:

Pour une modélisation plus réaliste, c'est plus naturel de supposer qu'on a pas accès directement aux fonctions valeurs mais des versions bruitées:

$$y = f(x) + \epsilon.$$

On suppose que le bruit $\epsilon \sim \mathcal{N}(0, \sigma_n^2)$. i.i.d

Dans ce cas,

$$\text{Cov}(y_p, y_q) = K(x_p, x_q) + \sigma_n^2 \delta_{pq}$$

$$\delta_{pq} = 1_{p=q}. \quad \text{Donc}$$

$$\text{Cov}(y) = K(x, x) + \sigma_n^2 I.$$

Pour la même procédure qu'avant, on a la loi jointe

$$\begin{bmatrix} y \\ f_* \end{bmatrix} \sim \mathcal{N} \left(0, \begin{bmatrix} K(x, x) + \sigma_n^2 I & K(x, x_*) \\ K(x_*, x) & K(x_*, x_*) \end{bmatrix} \right)$$

Dnc

$$f_* | x, y, x_* \sim \mathcal{N}(\bar{f}_*, \text{Cov}(f_*)) \text{ avec}$$

$$\bar{f}_* = \mathbb{E}[f_* | x, y, x_*] = K(x_*, x) [K(x, x) + \sigma_n^2 I]^{-1} y$$

$$\text{Cov}(f_*) = K(x_*, x_*) - K(x_*, x) [K(x, x) + \sigma_n^2 I]^{-1} K(x, x_*)$$

Dans la suite, pour simplifier ces écritures on introduit

$$K = K(x, x) ; \quad K_* = K(x, x_*)$$

et s'il y a un seul point test x_* , on écrit $K(x_*) = K_*$

On obtient,

$$\bar{f}_* = K_*^T (K + \sigma_n^2 I)^{-1} y \quad (1)$$

$$\text{Var}(f_*) = K(x_*, x_*) - K_*^T (K + \sigma_n^2 I)^{-1} K_* \quad (2)$$

d'équation précédente de $\bar{f}_*$ équation (1) s'écrit :

$$\bar{f}_* = \sum_{i=1}^n \alpha_i k(x_i, x_*)$$

où

$$\alpha = (K + \sigma_n^2 I)^{-1} y.$$

$\bar{f}_*$ s'exprime dans une base de fonction représentée par le noyau k .

§ Mise en place de méthodes GPR :

On rappelle que $f: \mathbb{R}^p \rightarrow \mathbb{R}$ suit un modèle GPR
noté $f \sim \text{GP}(\mu, k)$

Si on a des "input points" $x_1, x_2, \dots, x_n \in \mathbb{R}^p$

alors

$$(f(x_1), f(x_2), \dots, f(x_n)) \sim \mathcal{N}(\mu, K_{xx})$$

avec $(K_{xx})_{ij} = k(x_i, x_j)$

→ les noyaux k peuvent être n'importe quelle fonction
symétrique positive i.e. vérifiant :

$$\sum_{i,j=1}^n k(x_i, x_j) \xi_i \xi_j \geq 0 \quad \forall x_k \in \mathbb{R}^p ; \xi_k \in \mathbb{R}.$$

$$k: \mathbb{R}^p \times \mathbb{R}^p \rightarrow \mathbb{R}$$

[p : nbre de caractéristiques]
[n : nbre d'observations.]

Classe des noyaux RBF : "Radial Basis Function"

Les noyaux du type RBF dépendent seulement de
 $\|x - x'\|$, comme par ex. le noyau gaussien.
ou exponentiel :

$$k(x, x') = \exp\left(-\frac{1}{2\ell^2} \|x - x'\|^2\right)$$

ℓ : "length scale" mesure la distance à avoir entre les données pour être décorrélées.

→ il existe d'autres noyaux : noyau Matern (MA)

$$k(x, x') = \frac{\sigma^2 2^{1-\nu}}{\Gamma(\nu)} \left(\sqrt{2\nu} \frac{\|x-x'\|}{\ell} \right)^\nu K_\nu \left(\sqrt{2\nu} \frac{\|x-x'\|}{\ell} \right)$$

où Γ : c'est la fonction gamma.

$$(\dots) \Gamma(z) = \int_0^\infty t^{z-1} e^{-t} dt.$$

(...) K_ν : la fonction modifiée de Bessel de second type

$$K_\nu(z) = \frac{\Gamma(\nu + \frac{1}{2}) (2z)^\nu}{\sqrt{\pi}} \int_0^\infty \frac{\cos t}{(t^2 + z^2)^{\nu + \frac{1}{2}}} dt$$

Étant donné, les données de test: x_*

$f_*/x, x_*, y \sim$ loi gaussienne

d'espérance $\mathbb{E}[f_*/x, x_*, y] = K_{x_*, x} [K_{x, x} + \sigma_n^2 I]^{-1} y$

et de variance.

$$\text{Var}[f_*/x, x_*, y] = K_{x_*, x_*} - K_{x_*, x} [K_{x, x} + \sigma_n^2 I]^{-1} K_{x, x_*}$$

§ Optimization de l'hyperparamètre du noyau:

On suppose que le noyau k dépend d'un paramètre λ
 $\lambda = (\ell, \sigma_n)$ dans le cas du noyau exponentiel.

pour trouver λ^* : on a : output training

$$\lambda^* = \underset{\lambda}{\operatorname{argmax}} \operatorname{Log} p(\underset{\substack{\uparrow \\ \text{input training}}}{Y/X}, \underset{\substack{\uparrow \\ \text{hyperparamètre}}}{\lambda})$$

Comme on est de un modèle gaussien on a :

$$\log p(Y/X, \lambda) = -\frac{1}{2} \left[Y^T (K_{X,X} + \sigma_n^2 I)^{-1} Y + \log \det (K_{X,X} + \sigma_n^2 I) \right] - \frac{n}{2} \log 2\pi$$

§ Calcul de sensibilité:

On rappelle que la prédiction GPR

$$f_* / X, Y, X_* \sim \mathcal{N}(\mathbb{E}[f_* / X, Y, X_*], \operatorname{Var}[f_* / X, Y, X_*])$$

avec

$$\mathbb{E}[f_* / X, Y, X_*] = \mathbb{E}[X_*] + K_{X_*, X} [K_{X, X} + \sigma_n^2 I]^{-1} Y$$

$$\operatorname{Var}[f_* / X, Y, X_*] = K_{X_*, X_*} - K_{X_*, X} [K_{X, X} + \sigma_n^2 I]^{-1} K_{X, X_*}$$

On rappelle que

$(K_{X,X})_{ij} = k(x_i, x_j)$. Dans le cas de noyau RBF

on a $k(x_i, x_j) = \sigma_f \exp\left(-\frac{1}{2\ell^2} \|x_i - x_j\|^2\right)$

→ Une fois le GPR entraîné, sa prédiction est donnée par $\mathbb{E}[f_* | X, Y, X_*]$, on peut en déduire directement $\frac{\partial}{\partial X_*}$.

En particulier, dans le cas d'un noyau exponentiel gaussien:

$$\frac{\partial}{\partial X_*} \mathbb{E}[f_* | X, Y, X_*] = \frac{\partial}{\partial X_*} \mathbb{E}[x_*^*] + \underbrace{\frac{\partial}{\partial X_*} K_{X_*, X}}_{\text{}} [K_{X, X} + \sigma_v^2 I]^{-1} Y$$

$$(K_{X_*, X})_{ij} = \sigma_f \exp\left(-\frac{1}{2\ell^2} \langle x_i^* - x_j, x_i^* - x_j \rangle\right)$$

$$= -\frac{\sigma_f}{2\ell^2} \frac{\partial}{\partial X_*} \langle x_i^* - x_j, x_i^* - x_j \rangle$$

$$\times \underbrace{\exp\left(-\frac{1}{2\ell^2} \|x_i^* - x_j\|^2\right)}$$

$$= -\frac{1}{2\ell^2} 2(x_i^* - x_j) (K_{X_*, X})_{ij}$$

$$= \frac{1}{\ell^2} (x_j - x_i^*) (K_{X_*, X})_{ij}$$

Donc en prenant $\mathbb{E}[x_*] = 0$

$$\frac{\partial}{\partial x_*} \mathbb{E}[f_* | X, y, x_*] = \frac{1}{\sigma^2} (x - x_*) K_{x_*, x} [K_{x, x} + \sigma_n^2 I]^{-1} y$$

Donc avec cette formule, il n'y a pas besoin d'entraîner à nouveau le GPR pour calculer la dérivée de la fonction cible initiale qu'on a déjà apprise par la méthode GPR. L'estimation de la dérivée est presque gratuite.