



# Particle approximation and the Laplace method for Bayesian filtering

Paul Bui Quang

## ► To cite this version:

Paul Bui Quang. Particle approximation and the Laplace method for Bayesian filtering. General Mathematics [math.GM]. Université de Rennes, 2013. English. NNT : 2013REN1S038 . tel-00910173

**HAL Id: tel-00910173**

**<https://theses.hal.science/tel-00910173>**

Submitted on 27 Nov 2013

**HAL** is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.



THÈSE / UNIVERSITÉ DE RENNES 1  
*sous le sceau de l'Université Européenne de Bretagne*

pour le grade de  
DOCTEUR DE L'UNIVERSITÉ DE RENNES 1

*Mention : mathématiques et applications*

Ecole doctorale MATISSE

présentée par

**Paul Bui Quang**

préparée à

ONERA centre de Palaiseau  
INRIA Rennes Bretagne-Atlantique

---

**Approximation particulière  
et méthode de Laplace  
pour le filtrage bayésien**

**Thèse soutenue à Rennes  
le 1<sup>er</sup> juillet 2013**

devant le jury composé de :

**Nicolas Chopin**

Professeur, ENSAE / rapporteur

**Branko Ristić**

Senior research scientist, DSTO / rapporteur

**James Ledoux**

Professeur, INSA de Rennes / examinateur

**Valérie Monbet**

Professeur, Université de Rennes 1 / examinateur

**Nadia Oudjane**

Ingénieur-chercheur, EDF R&D / examinateur

**François Le Gland**

Directeur de recherche, INRIA / directeur de thèse

**Christian Musso**

Maître de recherche, ONERA / co-directeur de thèse

**Myriam Vimond**

Maître de conférences, ENSAI / invitée



# Remerciements

Merci à François Le Gland et à Christian Musso pour avoir dirigé cette thèse. Leurs compétences, leur rigueur scientifique, leur imagination, ont été indispensables à sa réalisation. J'ai vivement apprécié travailler avec eux et j'ai beaucoup appris à leurs côtés – pas uniquement sur le plan scientifique.

Merci à Branko Ristić et à Nicolas Chopin pour avoir accepté d'être rapporteurs de la thèse. Leurs remarques, leurs suggestions et les échanges que nous avons eus m'ont permis d'améliorer significativement le mémoire.

Merci à James Ledoux, Valérie Monbet, Nadia Oudjane et Myriam Vimond d'avoir participé au jury de la thèse.

Merci aux permanents et aux doctorants de l'Onera auprès de qui j'ai travaillé pendant plus de trois ans dans une atmosphère agréable et chaleureuse.

Merci à mes parents, à mon frère, à mes amis, et à Léa, pour leur présence, leur attention et leur soutien.



*"The world stands on absurdities, and without them perhaps nothing at all would happen."*

Fyodor Dostoyevsky, *The Brothers Karamazov* (1880)



# Table des matières / Contents

|                                                         |           |
|---------------------------------------------------------|-----------|
| <b>Résumé étendu</b>                                    | <b>11</b> |
| <b>Echantillonnage pondéré</b>                          | <b>21</b> |
| Cas basique . . . . .                                   | 21        |
| Cas d'une loi non normalisée . . . . .                  | 23        |
| <b>Filtrage bayésien et approximation particulaire</b>  | <b>27</b> |
| Modèles de Markov caché . . . . .                       | 27        |
| Filtrage bayésien . . . . .                             | 29        |
| Filtrage particulaire . . . . .                         | 32        |
| Filtrage particulaire . . . . .                         | 32        |
| Changement de loi d'échantillonnage . . . . .           | 38        |
| <b>Introduction</b>                                     | <b>41</b> |
| <b>Notations &amp; Acronyms</b>                         | <b>45</b> |
| <b>I Preliminaries</b>                                  | <b>47</b> |
| <b>1 The Laplace method: general presentation</b>       | <b>49</b> |
| 1.1 Principle . . . . .                                 | 49        |
| 1.2 The Laplace method . . . . .                        | 50        |
| 1.2.1 Differential calculus tools . . . . .             | 50        |
| 1.2.2 Statement of the general Laplace method . . . . . | 51        |
| 1.2.3 Proofs of theorems 1.2.1 and 1.2.2 . . . . .      | 54        |

|            |                                                                                            |            |
|------------|--------------------------------------------------------------------------------------------|------------|
| <b>2</b>   | <b>Issues in importance sampling for Bayesian estimation</b>                               | <b>67</b>  |
| 2.1        | Bayesian model set-up . . . . .                                                            | 67         |
| 2.1.1      | Asymptotics . . . . .                                                                      | 67         |
| 2.1.2      | Assumptions for the Laplace method . . . . .                                               | 68         |
| 2.2        | Importance sampling for Bayesian estimation . . . . .                                      | 69         |
| 2.3        | High information . . . . .                                                                 | 72         |
| 2.4        | High dimension . . . . .                                                                   | 73         |
| <b>II</b>  | <b>The Laplace method in Bayesian statistics</b>                                           | <b>75</b>  |
| <b>3</b>   | <b>The Laplace method in static models with large observation sample size</b>              | <b>77</b>  |
| 3.1        | Model set-up, Laplace-regularity . . . . .                                                 | 77         |
| 3.2        | Approximation of moments . . . . .                                                         | 82         |
| 3.3        | Proof of theorem 3.2.5 . . . . .                                                           | 91         |
| <b>4</b>   | <b>The Laplace method in dynamic models with small observation noise</b>                   | <b>99</b>  |
| 4.1        | Model set-up . . . . .                                                                     | 99         |
| 4.2        | Algorithm . . . . .                                                                        | 105        |
| 4.2.1      | Approximation of densities . . . . .                                                       | 105        |
| 4.2.2      | Approximation of moments . . . . .                                                         | 107        |
| 4.3        | Consistency of the approximations . . . . .                                                | 107        |
| 4.4        | Propagation of the approximation error . . . . .                                           | 111        |
| <b>5</b>   | <b>The Laplace method and Kalman filtering in dynamic models with small dynamics noise</b> | <b>123</b> |
| 5.1        | Model set-up . . . . .                                                                     | 123        |
| 5.2        | The Kalman Laplace filter . . . . .                                                        | 124        |
| 5.3        | The Laplace Gaussian filter . . . . .                                                      | 128        |
| <b>III</b> | <b>Particle approximation and the Laplace method in filtering algorithms</b>               | <b>131</b> |
| <b>6</b>   | <b>Importance sampling and the Laplace method in static models</b>                         | <b>133</b> |
| 6.1        | Importance sampling based on the Laplace method . . . . .                                  | 134        |

|          |                                                                          |            |
|----------|--------------------------------------------------------------------------|------------|
| 6.2      | Example: triangulation . . . . .                                         | 136        |
| <b>7</b> | <b>Particle filtering and the Laplace method in dynamic models</b>       | <b>143</b> |
| 7.1      | Model set-up . . . . .                                                   | 144        |
| 7.2      | Standard particle filtering and regularized particle filtering . . . . . | 144        |
| 7.2.1    | Standard particle filtering . . . . .                                    | 144        |
| 7.2.2    | Regularized particle filtering . . . . .                                 | 146        |
| 7.3      | Laplace particle filtering . . . . .                                     | 146        |
| 7.3.1    | Laplace approximations of the posterior moments . . . . .                | 148        |
| 7.3.2    | Sampling particles according to the Laplace approximations . . . . .     | 148        |
| 7.3.3    | Computational issues . . . . .                                           | 149        |
| 7.3.4    | Gaussian approximation of the predictor . . . . .                        | 151        |
| <b>8</b> | <b>Simulation experiments</b>                                            | <b>157</b> |
| 8.1      | Performance evaluation in Bayesian filtering . . . . .                   | 158        |
| 8.1.1    | Simulation set-up . . . . .                                              | 158        |
| 8.1.2    | Accuracy . . . . .                                                       | 159        |
| 8.1.3    | Robustness to divergence . . . . .                                       | 160        |
| 8.2      | Bearings-only target tracking . . . . .                                  | 161        |
| 8.2.1    | Model . . . . .                                                          | 161        |
| 8.2.2    | Simulation parameters . . . . .                                          | 163        |
| 8.2.3    | Results . . . . .                                                        | 164        |
| 8.3      | Ballistic target tracking during atmospheric reentry . . . . .           | 165        |
| 8.3.1    | Model . . . . .                                                          | 167        |
| 8.3.2    | Simulation parameters . . . . .                                          | 169        |
| 8.3.3    | Results . . . . .                                                        | 169        |
| 8.4      | Neural decoding . . . . .                                                | 172        |
| 8.4.1    | Model . . . . .                                                          | 173        |
| 8.4.2    | Simulation parameters . . . . .                                          | 174        |
| 8.4.3    | Results . . . . .                                                        | 175        |
|          | <b>Conclusion</b>                                                        | <b>179</b> |
| <b>A</b> | <b>Differentiation of the log-likelihood</b>                             | <b>183</b> |
| A.1      | Bearings-only target tracking . . . . .                                  | 183        |
| A.2      | Ballistic target tracking during atmospheric reentry . . . . .           | 185        |

|                               |            |
|-------------------------------|------------|
| A.3 Neural decoding . . . . . | 185        |
| <b>Bibliography</b>           | <b>187</b> |

# Résumé étendu

Le filtrage statistique est un problème d'estimation bayésienne dans des modèles dynamiques, en utilisant des observations délivrées séquentiellement. La quantité aléatoire à estimer est appelée état, ou état caché. La dynamique de l'état obéit à un processus de Markov. Les observations sont liées à l'état à chaque instant par une fonction de vraisemblance. Le problème du filtrage existe en temps continu et en temps discret. Nous travaillons dans cette thèse exclusivement avec des modèles à temps discrets, qui sont majoritairement utilisés dans les applications que nous considérons.

Plus précisément, soit  $k \in \mathbb{N}$  l'indice temporel discret. L'état caché à l'instant  $k$  est noté  $X_k$  et le processus markovien  $\{X_k\}_{k \geq 0}$  est appelé processus d'état. L'observation délivrée à l'instant  $k$  est notée  $Y_k$  et le processus  $\{Y_k\}_{k \geq 0}$  est appelé processus d'observation. Le processus joint  $\{X_k, Y_k\}_{k \geq 0}$  est appelé modèle de Markov caché, ou modèle à espace d'état. Les modèles de Markov cachés sont utilisés pour modéliser des phénomènes dynamiques aléatoires dans de domaines variés : aérospatial et défense [Hue et al. (2002); Ristic et al. (2004); Gustafsson (2010)], traitement de la parole [Juang and Rabiner (1991)], ingénierie bio-médicale [Brockwell et al. (2004); Eddy (1996)], économétrie [Chopin and Pelgrin (2004); Rossi and Gallo (2006)], pour ne citer que quelques exemples.

Le filtrage consiste à calculer à chaque instant la loi a posteriori, c'est-à-dire la loi conditionnelle de l'état courant  $X_k$  sachant toutes les observations passées  $Y_0, \dots, Y_k$ ,

$$\mu_k(dx) = \mathbb{P}[X_k \in dx | Y_0, \dots, Y_k],$$

également appelée filtre bayésien. Les algorithmes de filtrage (ou « filtres ») peuvent être récursifs ou non, les premiers étant privilégiés car généralement plus rapides. Un filtre récursif traite les observations une par une (à l'instant  $k$ , le calcul de  $\mu_k$  ne dépend que de  $Y_k$  et de  $\mu_{k-1}$ ) plutôt que de ré-utiliser à chaque instant les observations passées. A l'inverse, les méthodes de filtrage qui utilisent plusieurs observations à chaque

instant, dites globales, demandent plus de calculs (on peut toutefois limiter le nombre d'observations traitées en ne considérant que celles délivrées à certains instants bien choisis [Musso (1993); Musso and Oudjane (2005)]).

Lorsque le modèle est linéaire et gaussien, la suite des filtres bayésiens est une suite de lois gaussiennes dont l'espérance et la matrice de covariance peuvent être calculées exactement par le filtre de Kalman [Kalman (1960)], qui est donc la méthode optimale dans ce cas. Lorsque le modèle est non-linéaire en revanche, il n'existe pas de méthode permettant d'obtenir de manière générale le filtre bayésien exact. Des algorithmes d'approximation basés sur le filtrage de Kalman, comme le filtre de Kalman étendu ou le filtre de Kalman non-parfumé [Arulampalam et al. (2002); Julier and Uhlmann (2004)], donnent une approximation gaussienne biaisée du filtre bayésien. Ils peuvent donner de bons résultats mais leur performance est dégradée lorsque le modèle s'éloigne trop du cas idéal linéaire gaussien.

Le filtrage particulaire permet de calculer récursivement, de manière approchée, le filtre bayésien dans des modèles non-linéaires et non-gaussiens. Les algorithmes de filtrage particulaire sont des méthodes de Monte Carlo séquentielles (sequential Monte Carlo, SMC) basées sur le principe de l'échantillonnage pondéré. Ils consistent à approcher, à chaque instant, la loi a posteriori  $\mu_k$  par une probabilité, notée  $\mu_k^N$ , qui s'exprime comme une somme pondérée de  $N$  masses de Dirac centrées en des points aléatoires de l'espace d'état. Ces points sont appelés « particules ».  $\mu_k^N$  converge, sous des hypothèses faibles, vers  $\mu_k$  lorsque le nombre de particules  $N$  tend vers l'infini [Crisan and Doucet (2002); Del Moral (2004)]. Ainsi, le filtrage particulaire permet d'approcher le filtre bayésien arbitrairement près lorsque la puissance de calcul disponible augmente, ce qui constitue son intérêt principal.

L'inconvénient du filtrage particulaire réside dans le fait qu'après un certain temps, on observe souvent que seules quelques particules ont un poids non nul, et que le poids de toutes les autres est numériquement évalué à zéro. Le filtrage est alors fortement dégradé, car seul un petit nombre de particules participent à l'approximation de la loi a posteriori. Ce phénomène est appelé « dégénérescence des poids ». Il est classiquement géré en ré-échantillonnant les particules selon leur poids, de manière à ce que les particules associées à un poids important soient dupliquées et que celles associées à un poids faible soient supprimées. Cette solution a été initialement proposée dans [Gordon et al. (1993)].

Le phénomène de dégénérescence des poids est particulièrement sévère lorsque que

le modèle est très informatif, par exemple lorsque l'aléa sur la dynamique markovienne est faible ou lorsque les observations sont précises [Oudjane and Musso (2000)]. Nous considérons dans cette thèse une méthode déterministe de calcul intégral, la méthode de Laplace, très utilisée en statistique bayésienne et qui, au contraire, est d'autant plus efficace que le modèle est informatif. La méthode de Laplace est classiquement utilisée pour approcher des moments a posteriori dans des modèles statiques (c'est-à-dire lorsque l'état caché n'est pas dynamique) [Tierney and Kadane (1986)]. Sous des conditions de régularité et d'identifiabilité du modèle, ces approximations sont convergentes lorsque le nombre d'observations tend vers l'infini ou, de manière équivalente, lorsque l'intensité du bruit d'observation tend vers zéro. Dans cette thèse, nous proposons d'associer la méthode de Laplace au filtrage particulaire dans le but d'améliorer la qualité du filtrage, notamment dans le cas paradoxalement difficile où le modèle est informatif.

## Méthode de Laplace

La méthode de Laplace est une méthode d'approximation d'intégrales multidimensionnelles de la forme

$$\int_{\mathbb{R}^d} b(x) e^{-\lambda h(x)} dx, \quad (1)$$

où  $\lambda$  est un paramètre réel tel que  $\lambda \gg 1$ .  $h$  est supposée admettre un minimum global en  $\hat{x}$ , être régulière dans un voisinage de  $\hat{x}$ , de sorte que  $h'(\hat{x}) = 0$  et  $\det[h''(\hat{x})] > 0$ , et vérifier la condition de coercivité suivante : pour tout  $\delta > 0$ ,

$$\inf\{h(x) - h(\hat{x}) : |x - \hat{x}| > \delta\} > 0.$$

La méthode de Laplace consiste à considérer l'intégrale (1) comme l'intégrale de  $b$  contre une mesure gaussienne de petite variance d'ordre  $1/\lambda$ . Pour ce faire, on remplace  $h(x)$  dans l'intégrande par son développement de Taylor au second ordre autour de  $\hat{x}$ ,

$$h(x) \approx h(\hat{x}) + \frac{1}{2}(x - \hat{x})^T h''(\hat{x})(x - \hat{x}).$$

L'intégrale (1) devient alors

$$\int_{\mathbb{R}^d} b(x) e^{-\lambda h(x)} dx \approx e^{-\lambda h(\hat{x})} \int_{\mathbb{R}^d} b(x) e^{-\frac{\lambda}{2}(x - \hat{x})^T h''(\hat{x})(x - \hat{x})} dx.$$

Lorsque  $\lambda$  est grand, l'intégrale de  $b$  contre la densité gaussienne

$$(2\pi)^{-d/2} \det [\lambda h''(\hat{x})]^{1/2} \exp \left( -\frac{\lambda}{2} (x - \hat{x})^T h''(\hat{x}) (x - \hat{x}) \right)$$

est proche de  $b(\hat{x})$ . En effet, la norme de la matrice de covariance  $[\lambda h''(\hat{x})]^{-1}$  est petite, ce qui implique que la densité est concentrée autour de son maximum. On a alors

$$\int_{\mathbb{R}^d} b(x) e^{-\frac{\lambda}{2} (x - \hat{x})^T h''(\hat{x}) (x - \hat{x})} dx \approx (2\pi)^{d/2} b(\hat{x}) \det [\lambda h''(\hat{x})]^{-1/2},$$

et on obtient donc l'approximation de Laplace

$$\int_E b(x) e^{-\lambda h(x)} dx \approx (2\pi)^{d/2} b(\hat{x}) e^{-\lambda h(\hat{x})} \det [\lambda h''(\hat{x})]^{-1/2}.$$

Dans les problèmes d'estimation bayésienne, les intégrales à calculer sont souvent de la forme

$$\int_{\mathbb{R}^d} b(x) e^{-\lambda h_\lambda(x)} dx \tag{2}$$

plutôt que de la forme (1). Le minimum de  $h_\lambda$  dépend alors de  $\lambda$ , on le note  $\hat{x}_\lambda$ . La méthode de Laplace est également applicable mais nécessite que l'hypothèse de coercivité sur la fonction  $h_\lambda$  soit uniforme en  $\lambda$  : pour tout  $\delta > 0$  et pour tout  $\lambda$  suffisamment grand,

$$\inf \{ h_\lambda(x) - h_\lambda(\hat{x}_\lambda) : |x - \hat{x}_\lambda| > \delta \} \geq c_\delta$$

où  $c_\delta > 0$  est indépendant de  $\lambda$ .

Dans les modèles bayésiens, le paramètre  $\lambda$  est généralement :

- la taille de l'échantillon des observations,
- l'inverse de la variance du bruit d'observation.

La méthode de Laplace est présentée dans un contexte général et applicable aux problèmes d'estimation bayésienne au chapitre 1.

## Echantillonnage pondéré et problématiques associées

L'échantillonnage pondéré consiste à approcher des lois de la forme

$$\mu \propto g\eta, \quad (3)$$

où  $\eta$  est une probabilité et  $g$  une fonction positive. Dans un contexte bayésien,  $\eta$  est la loi a priori,  $g$  la fonction de vraisemblance et  $\mu$  la loi a posteriori. Un échantillon de  $N$  « particules » indépendantes  $(\xi^1, \dots, \xi^N)$  est simulé selon  $\eta$  puis pondéré selon  $g$ . On obtient l'approximation particulaire de  $\mu$ ,

$$\mu^N = \sum_{i=1}^N w^i \delta_{\xi^i}$$

où  $w^i = \frac{g(\xi^i)}{\sum_{j=1}^N g(\xi^j)}$  et où  $\delta_{\xi^i}$  est la mesure de Dirac centrée en  $\xi^i$ .

Un indicateur de la qualité de l'échantillonnage pondéré est

$$I = \frac{\int g(x)^2 \eta(dx)}{(\int g(x) \eta(dx))^2}. \quad (4)$$

Cette quantité est liée à la divergence du  $\chi^2$  entre la loi d'intérêt  $\mu$  et la loi d'échantillonnage  $\eta$ , puisque  $\chi^2(\mu, \eta) = I - 1$ . Elle intervient dans la variance asymptotique des poids d'importance et dans la définition de la taille effective de l'échantillon (effective sample size). En appliquant la méthode de Laplace au numérateur et au dénominateur de (4), dans un cadre asymptotique et sous des hypothèses appropriées (correspondant essentiellement à l'identifiabilité du modèle statistique associé à (3)), on obtient l'approximation

$$I \approx \det \left[ \frac{-(\log g)''(\hat{x})}{4\pi} \right]^{1/2} \frac{1}{q(\hat{x})}, \quad (5)$$

où  $q$  est la densité de  $\eta$  et  $\hat{x} = \operatorname{argmax}_{x \in E} \{g(x)\}$  (maximum de vraisemblance). On quantifie l'information apportée par les observations par la matrice symétrique et positive de taille  $d \times d$   $-(\log g)''(\hat{x})$ . L'approximation (5) permet donc de voir que lorsque cette information augmente (au sens où le déterminant de  $-(\log g)''(\hat{x})$  augmente), la performance de l'échantillonnage pondéré diminue (sa convergence est plus lente). Par ailleurs, si le volume de l'ellipsoïde dans  $\mathbb{R}^d$  associée à la matrice  $-(\log g)''(\hat{x})$  augmente avec la dimension, alors la qualité de l'échantillonnage pondéré se dégrade lorsque la

dimension augmente.

Cette discussion sur l'échantillonnage pondéré dans des modèles informatifs et en grande dimension est menée au chapitre 2.

## Applications de la méthode de Laplace en statistique bayésienne

La méthode de Laplace est utilisée en statistique bayésienne pour calculer des moments a posteriori dans des modèles identifiables, où l'état à estimer  $X$  est statique. Le cadre asymptotique classique est lorsque la taille  $n$  de l'échantillon d'observations  $Y_{1:n} = \{Y_1, \dots, Y_n\}$  tend vers l'infini. Tierney et al. (1989) obtiennent des formules d'approximation pour l'espérance et la variance a posteriori en dérivant l'approximation de Laplace de la fonction génératrice des moments  $a \mapsto \mathbb{E}[e^{a^T X} | Y_{1:n}]$  (ou transformée de Laplace) et en l'évaluant en 0. Ces formules sont

$$\mathbb{E}[X | Y_{1:n}] \approx \hat{x} - \frac{1}{2} \frac{J'(\hat{x})}{J(\hat{x})^2} \quad (6)$$

et

$$\mathbb{V}[X | Y_{1:n}] \approx \frac{1}{J(\hat{x})} + \frac{J'(\hat{x})^2}{J(\hat{x})^4} - \frac{1}{2} \frac{J''(\hat{x})}{J(\hat{x})^3}. \quad (7)$$

Nous les généralisons au cas multidimensionnel au chapitre 3.

Dans le cas d'un modèle de Markov caché, où l'état obéit à un processus de Markov  $\{X_k\}_{k \geq 0}$ , la méthode de Laplace ne peut pas être paramétrée par le nombre d'observations. En effet, les moments a posteriori s'expriment comme une intégrale sur un espace dont la dimension augmente avec le temps, donc avec le nombre d'observations, et qui ne peut par conséquent pas s'écrire sous la forme (1) ou (2) avec  $\lambda = k$ . On peut alors considérer un modèle d'observation de la forme

$$Y_k = H_k(X_k) + \sqrt{\varepsilon} \sigma_k(X_k) W_k,$$

où  $\{W_k\}_{k \geq 0}$  est un bruit blanc gaussien, et on paramètre la méthode de Laplace par l'inverse de la variance du bruit d'observation, i.e.  $\lambda = 1/\varepsilon$ . Au chapitre 4, nous étudions une méthode d'approximation du filtre bayésien basée sur la méthode de Laplace, dans le cadre asymptotique  $\varepsilon \rightarrow 0$ . Malheureusement, cette étude se fait au prix de l'hypothèse que la fonction  $x \mapsto \frac{1}{2\sigma_k(x)^2} |Y_k - H_k(x)|^2$  admette un minimum global unique à

chaque instant  $k$ , qui est irréaliste car elle implique qu'à chaque instant on dispose d'un vecteur d'observation dont la dimension est au moins égale à celle de l'état, ce qui n'est généralement pas le cas.

On peut s'affranchir de cette hypothèse forte en considérant que le processus d'état a également une variance d'ordre  $\varepsilon$ , par exemple lorsque la dynamique est de la forme

$$X_k = F_k(X_{k-1}) + \sqrt{\varepsilon}V_k, \quad (8)$$

où  $\{V_k\}_{k \geq 0}$  est un bruit blanc gaussien. Au chapitre 5, nous proposons un algorithme qui associe le filtrage de Kalman étendu (pour la prédiction) et la méthode de Laplace (pour la mise à jour) lorsque le modèle d'état est de la forme (8). Cet algorithme, baptisé Kalman Laplace filter (KLF), est relativement proche du Laplace Gaussian filter proposé par Koyama et al. (2010).

## Association de la méthode de Laplace et du filtrage particulaire

L'idée de combiner la méthode de Laplace et le filtrage particulaire est justifiée par le fait que la méthode de Laplace est efficace dans le cas où le modèle est informatif, ce qui est précisément un domaine où l'échantillonnage pondéré (et donc le filtrage particulaire) est mis en difficulté.

Notre approche consiste à appliquer à la loi a priori une transformation affine, basée sur les versions multidimensionnelles des formules (6) et (7), puis de prendre cette loi transformée comme loi d'échantillonnage. Plus précisément, si  $\eta$  est la loi a priori, la loi d'échantillonnage est

$$\eta' = \eta \circ T^{-1},$$

où la transformation affine  $T$  est définie par

$$T(x) = \hat{P}^{1/2}\mathbb{V}[X]^{-1/2}(x - \mathbb{E}[X]) + \hat{m}, \quad (9)$$

où  $X \sim \eta$  et  $\hat{m}$  et  $\hat{P}$  sont les versions multidimensionnelles des approximations (6) et (7) respectivement. En pratique, il suffit d'appliquer la transformation  $T$  sur les particules échantillonnées selon  $\eta$ . Les particules sont alors déplacées de sorte que la moyenne et la matrice covariance empiriques du nuage soient respectivement égales à

$\hat{m}$  et  $\hat{P}$ . L'objectif est d'échantillonner selon une loi proche de la loi a posteriori (qui est la loi d'échantillonnage optimale), au sens où les deux premiers moments de la loi d'échantillonnage  $\eta'$  sont des approximations de ceux la loi a posteriori. Cette approche est exposée dans le cas d'un modèle statique au chapitre 6 et testée par simulation sur un exemple simple.

Dans le chapitre 7, nous adaptons l'idée d'appliquer la transformation (9) aux particules au cas d'un modèle de Markov caché (donc avec un état dynamique).  $\mathbb{E}[X]$  et  $\mathbb{V}[X]$  sont alors inconnues et doivent être remplacées par la moyenne et la matrice de covariance empiriques des particules échantillonnées. Nous proposons un nouveau filtre particulaire, le Laplace particle filter (LPF), où cette transformation est appliquée de manière adaptative, lorsque la taille effective de l'échantillon est trop faible. Le LPF est testé avec succès par simulation sur trois problèmes de filtrage non-linéaire au chapitre 8 : pistage par mesure d'angles, pistage d'un objet balistique pendant sa rentrée dans l'atmosphère, décodage neuronal.

## Contributions et perspectives

Les contributions principales de cette thèse sont les suivantes :

- le LPF, un nouveau filtre particulaire qui utilise la méthode de Laplace (chapitres 6, 7, 8) ;
- l'analyse de l'approximation du filtre bayésien uniquement avec la méthode de Laplace, quand l'intensité du bruit d'observation tend vers zéro (chapitre 4).

Les autres contributions de ce travail sont :

- l'analyse de l'échantillonnage pondéré dans des situations difficiles (grande information et grande dimension) (chapitre 2) ;
- la version multidimensionnelle des formules de Tierney et al. (1989) pour l'approximation de l'espérance et de la covariance a posteriori (chapitre 3) ;
- le KLF, un algorithme basé sur le filtrage de Kalman et la méthode de Laplace (chapitre 5) ;
- la comparaison expérimentale du KLF et du LPF dans le cadre de trois applications en ingénierie : pistage par mesures d'azimut, pistage d'objet balistique en rentrée atmosphérique, décodage neuronal (chapitre 8).

Nous évoquons pour finir quelques pistes de recherche ouvertes par cette thèse :

- l'étude asymptotique du KLF lorsque le paramètre  $\varepsilon$  tend vers zéro ;
- l'étude asymptotique du LPF lorsque  $N \rightarrow \infty$  et  $\varepsilon \rightarrow 0$  ( $\varepsilon$  représentant par exemple l'intensité du bruit d'observation et du bruit d'état) ;
- l'adaptation du LPF au cas où le filtre bayésien est très multimodal, par exemple en ajoutant à l'algorithme une étape de clustering pour identifier les modes (comme dans [Murangira et al. (2011)]) et en appliquant la transformation basée sur la méthode de Laplace localement dans chaque cluster ;
- l'étude du LPF dans des modèles de grande dimension (par exemple des modèles d'assimilation de données), le filtrage particulaire classique étant notoirement inefficace dans ce contexte (voir par exemple [Snyder et al. (2008)]) ;
- l'adaptation du LPF à des problèmes de lissage.



# Echantillonnage pondéré

Dans ce chapitre, nous présentons succinctement l'échantillonnage pondéré, qui est la méthode Monte Carlo sur laquelle est basée le filtrage particulaire. Les résultats donnés ici sont détaillés dans [Le Gland (2012)].

## Cas basique

On cherche à calculer de manière approchée un paramètre inconnu  $\theta$  défini par

$$\theta = \mathbb{E}[\phi(X)], \quad (10)$$

où  $X \sim \mu$ . On suppose la loi de  $X$  est connue au sens où l'on sait simuler selon  $\mu$ .

Les méthodes de Monte Carlo consistent à approcher la probabilité  $\mu$  par une mesure empirique  $\mu^N$  obtenue par simulation aléatoire, puis à intégrer  $\phi$  contre  $\mu^N$  pour obtenir une approximation de  $\theta$ ,

$$\hat{\theta}^N = \int \phi \mu^N \approx \int \phi \mu = \theta.$$

La méthode de Monte Carlo la plus simple consiste à approcher  $\mu$  par une moyenne de mesures de Dirac centrées sur des réalisations de  $\mu$ . Soit  $(\xi^1, \dots, \xi^N)$  un échantillon i.i.d. de loi  $\mu$ , appelé échantillon de particules, et soit  $\mu^N$  la probabilité empirique

$$\mu^N = \frac{1}{N} \sum_{i=1}^N \delta_{\xi^i}$$

de sorte que  $\mu^N \approx \mu$ . L'approximation Monte Carlo brute de  $\theta$  est alors

$$\hat{\theta}^N = \frac{1}{N} \sum_{i=1}^N \phi(\xi^i).$$

$\hat{\theta}^N$  est non biaisé et sa variance est

$$\mathbb{V} [\hat{\theta}^N] = \frac{1}{N} \int (\phi(x) - \theta)^2 \mu(dx) = \frac{\mathbb{V}[\phi(X)]}{N}.$$

L'échantillonnage pondéré (importance sampling en anglais) est une méthode de Monte Carlo qui consiste à échantillonner selon une loi  $\tilde{\mu}$  plutôt que selon  $\mu$ , dans le but de réduire la variance de l'approximation. Cette méthode est basée sur la décomposition

$$\mu(dx) = \frac{d\mu}{d\tilde{\mu}}(x) \tilde{\mu}(dx).$$

La loi d'origine  $\mu$  doit être absolument continue par rapport à la loi d'échantillonnage  $\tilde{\mu}$  pour que la densité  $\frac{d\mu}{d\tilde{\mu}}$  de  $\mu$  par rapport à  $\tilde{\mu}$  existe.  $\tilde{\mu}$  est appelée loi de proposition, loi d'importance, ou loi d'échantillonnage. Soit  $(\tilde{\xi}^1, \dots, \tilde{\xi}^N)$  un échantillon i.i.d. de loi  $\tilde{\mu}$ . La probabilité empirique  $\mu^N$  qui approche  $\mu$  est ici une moyenne pondérée de mesures de Dirac centrées sur les particules  $\tilde{\xi}^i$ ,

$$\mu^N = \frac{1}{N} \sum_{i=1}^N \frac{d\mu}{d\tilde{\mu}}(\tilde{\xi}^i) \delta_{\tilde{\xi}^i} \approx \frac{d\mu}{d\tilde{\mu}} \tilde{\mu} = \mu.$$

L'approximation de  $\theta$  par échantillonnage pondéré est alors

$$\tilde{\theta}^N = \frac{1}{N} \sum_{i=1}^N \frac{d\mu}{d\tilde{\mu}}(\tilde{\xi}^i) \phi(\tilde{\xi}^i).$$

où les  $\tilde{\xi}^i$  sont i.i.d. de loi  $\tilde{\mu}$ .  $\tilde{\theta}^N$  est non biaisé et sa variance est

$$\mathbb{V} [\tilde{\theta}^N] = \frac{1}{N} \int \left( \phi(x) \frac{d\mu}{d\tilde{\mu}}(x) - \theta \right)^2 \tilde{\mu}(dx). \quad (11)$$

L'approximation obtenue par échantillonnage pondéré a une meilleure variance que celle obtenue par échantillonnage Monte Carlo brut si

$$\begin{aligned} \mathbb{V} [\hat{\theta}^N] - \mathbb{V} [\tilde{\theta}^N] &= \frac{1}{N} \int \left( \phi(x)^2 \mu(dx) - \phi(x)^2 \left[ \frac{d\mu}{d\tilde{\mu}}(x) \right]^2 \tilde{\mu}(dx) \right) \\ &= \frac{1}{N} \mathbb{E} \left[ \phi(X)^2 \left( 1 - \frac{d\mu}{d\tilde{\mu}}(X) \right) \right] \\ &\geq 0, \end{aligned}$$

où l'espérance est prise par rapport à la loi  $\mu$ . Cette condition n'est en général pas vérifiable en pratique.

La loi d'échantillonnage optimale est

$$\tilde{\mu}^*(dx) = \frac{|\phi(x)|\mu(dx)}{\int |\phi(x)|\mu(dx)},$$

au sens où la variance l'approximation  $\tilde{\theta}^{*N}$  correspondante est minimale. Pour le justifier, notons que d'après (11) cette variance est

$$\begin{aligned} \mathbb{V} [\tilde{\theta}^{*N}] &= \frac{1}{N} \left( \int \phi(x)^2 \left[ \frac{d\mu}{d\tilde{\mu}^*}(x) \right]^2 \tilde{\mu}^*(dx) - \theta^2 \right) \\ &= \frac{1}{N} \left( \int |\phi(x)|\mu(dx) \int \frac{\phi(x)^2 \mu(dx)}{|\phi(x)|} - \theta^2 \right) \\ &= \frac{1}{N} \left( \left( \int |\phi(x)|\mu(dx) \right)^2 - \theta^2 \right). \end{aligned}$$

Il vient par conséquent

$$\begin{aligned} \mathbb{V} [\tilde{\theta}^N] - \mathbb{V} [\tilde{\theta}^{*N}] &= \frac{1}{N} \left( \int |\phi(x)|^2 \left[ \frac{d\mu}{d\tilde{\mu}}(x) \right]^2 \tilde{\mu}(dx) - \left( \int |\phi(x)|\mu(dx) \right)^2 \right) \\ &\geq \frac{1}{N} \left( \left( \int |\phi(x)| \frac{d\mu}{d\tilde{\mu}}(x) \mu(dx) \right)^2 - \left( \int |\phi(x)|\mu(dx) \right)^2 \right) \\ &= 0 \end{aligned}$$

d'après l'inégalité de Jensen. En pratique,  $\tilde{\mu}^*$  est en général inaccessible car son dénominateur est inconnu ( $\theta = \int \phi \mu$  étant inconnu, c'est a fortiori le cas de  $\int |\phi| \mu$ ).

## Cas d'une loi non normalisée

Considérons maintenant un cas plus général. On cherche à calculer le paramètre  $\theta$  défini par (10) lorsque la loi de  $X$  s'écrit sous la forme

$$\mu(dx) = \frac{g(x)\eta(dx)}{\int g(x)\eta(dx)}, \quad (12)$$

où  $\eta$  est une loi selon laquelle on sait simuler et  $g$  est une fonction réelle positive que l'on sait évaluer en tout point.

Ce type de problème se rencontre notamment en statistique bayésienne : lorsque  $\eta$  est la loi a priori et  $g$  la fonction de vraisemblance,  $\mu$  est la loi a posteriori d'après la formule de Bayes. Des estimateurs bayésiens classiques s'écrivent alors sous la forme

$$\theta = \mathbb{E}[\phi(X)] = \frac{\int \phi(x)g(x)\eta(dx)}{\int g(x)\eta(dx)},$$

par exemple l'espérance a posteriori lorsque  $\phi(x) = x$ , ou la probabilité de dépassement d'un seuil  $c$  lorsque  $\phi(x) = \mathbb{I}_{|x| \geq c}$ .

Le problème spécifique posé par la simulation selon une loi de la forme (12) est que le dénominateur  $\int g(x)\eta(dx)$  est en général inconnu. Toutefois, l'échantillonnage pondéré est adapté à la simulation selon des lois non normalisées. Soit  $(\xi^1, \dots, \xi^N)$  un échantillon i.i.d. de loi  $\eta$ . On peut approcher  $\mu$  par une probabilité empirique de la forme

$$\mu^N = \sum_{i=1}^N w^i \delta_{\xi^i}, \quad \text{où} \quad w^i = \frac{g(\xi^i)}{\sum_{i=1}^N g(\xi^i)},$$

de telle sorte que la somme des poids  $w^i$  soit égale à 1 et donc  $\int \mu^N = 1$ . L'approximation de  $\theta$  correspondante est

$$\hat{\theta}^N = \sum_{i=1}^N w^i \phi(\xi^i).$$

Cette approximation est biaisée en général. Cependant, son biais tend vers 0 lorsque la taille  $N$  de l'échantillon simulé tend vers l'infini.

**Théorème 1.**  $\hat{\theta}^N$  vérifie

$$\mathbb{E} \left[ \left| \hat{\theta}^N - \theta \right|^2 \right] = O(N^{-1}).$$

*Démonstration.* Soit  $h(x, y) = \left( \frac{x}{y} - \theta \right)^2$  pour tout  $(x, y) \in \mathbb{R} \times \mathbb{R}^*$ , de sorte que  $\mathbb{E} \left[ \left| \hat{\theta}^N - \theta \right|^2 \right] = \mathbb{E} \left[ h \left( \hat{\theta}_1^N, \hat{\theta}_2^N \right) \right]$  où  $\hat{\theta}_1^N = \frac{1}{N} \sum_{i=1}^N g(\xi^i) \phi(\xi^i)$  et  $\hat{\theta}_2^N = \frac{1}{N} \sum_{i=1}^N g(\xi^i)$ .  $h$  vérifie

$$h'(x, y) = \left( \frac{2}{y} \left( \frac{x}{y} - \theta \right) \quad -\frac{2x}{y^2} \left( \frac{x}{y} - \theta \right) \right)$$

et

$$h''(x, y) = \frac{2}{y^2} \begin{pmatrix} 1 & -2\frac{x}{y} + \theta \\ -2\frac{x}{y} + \theta & 3\frac{x^2}{y^2} - 2\theta\frac{x}{y} \end{pmatrix}.$$

Soient  $\theta_1 = \int \phi(x)g(x)\eta(dx)$  et  $\theta_2 = \int g(x)\eta(dx)$ , de sorte que  $\theta = \theta_1/\theta_2$ . En développant  $h$  à l'ordre 2 autour du point  $(\theta_1, \theta_2)$ , on obtient

$$h(\hat{\theta}_1^N, \hat{\theta}_2^N) = \frac{1}{2} \left[ (\hat{\theta}_1^N, \hat{\theta}_2^N) - (\theta_1, \theta_2) \right] h''(\alpha_1^N, \alpha_2^N) \left[ (\hat{\theta}_1^N, \hat{\theta}_2^N) - (\theta_1, \theta_2) \right]^T$$

car  $h(\theta_1, \theta_2) = 0$  et  $h'(\theta_1, \theta_2) = (0, 0)$ , avec  $(\alpha_1^N, \alpha_2^N)$  un point sur le segment formé par les points  $(\theta_1, \theta_2)$  et  $(\hat{\theta}_1^N, \hat{\theta}_2^N)$ . Comme  $(\hat{\theta}_1^N, \hat{\theta}_2^N) \xrightarrow[N \rightarrow \infty]{} (\theta_1, \theta_2)$  p.s.,  $(\alpha_1^N, \alpha_2^N) \xrightarrow[N \rightarrow \infty]{} (\theta_1, \theta_2)$  p.s. De plus,  $h''$  est continue en  $(\theta_1, \theta_2)$ , donc  $h''(\alpha_1^N, \alpha_2^N) \xrightarrow[N \rightarrow \infty]{} h''(\theta_1, \theta_2)$  p.s. Par conséquent, pour tout  $\varepsilon > 0$ ,

$$|h''(\alpha_1^N, \alpha_2^N) - h''(\theta_1, \theta_2)| \leq \varepsilon$$

p.s. pour tout  $N$  suffisamment grand. Comme  $h''(\theta_1, \theta_2) = \frac{1}{\theta_2^2} \begin{pmatrix} 1 & 1 \\ 1 & 1 \end{pmatrix}$ , on obtient

$$\begin{aligned} \mathbb{E} \left[ h(\hat{\theta}_1^N, \hat{\theta}_2^N) \right] &\leq \frac{(1 + \varepsilon)/N}{\left( \int g(x)\eta(dx) \right)^2} \left\{ \int \phi(x)^2 g(x)^2 \eta(dx) - \left( \int \phi(x)g(x)\eta(dx) \right)^2 \right. \\ &\quad + \int g(x)^2 \eta(dx) - \left( \int g(x)\eta(dx) \right)^2 \\ &\quad \left. + 2 \left( \int \phi(x)g(x)^2 \eta(dx) - \int \phi(x)g(x)\eta(dx) \int g(x)\eta(dx) \right) \right\}. \end{aligned}$$

□

$\hat{\theta}^N$  vérifie le théorème central limite suivant [Le Gland (2012)] :

**Théorème 2.**

$$\sqrt{N}(\hat{\theta}^N - \theta) \longrightarrow \mathcal{N}(0, V)$$

et

$$\sqrt{N} \left( \frac{\frac{1}{N} \sum_{i=1}^N g(\xi^i)}{\int g(x)\eta(dx)} - 1 \right) \longrightarrow \mathcal{N}(0, v)$$

lorsque  $N \rightarrow \infty$ , avec

$$V = \frac{\int (\phi(x) - \theta)^2 g(x)^2 \eta(dx)}{\left( \int g(x)\eta(dx) \right)^2} \quad \text{et} \quad v = \frac{\int g(x)^2 \eta(dx)}{\left( \int g(x)\eta(dx) \right)^2} - 1.$$

La variance asymptotique  $v$  ci-dessus est un indicateur de la performance de l'échan-

tillonnage pondéré pour la simulation selon (12). Elle est comprise entre 0 et 1 et vaut 0 lorsque la loi d'échantillonnage  $\eta$  est égale à la loi d'intérêt  $\mu$ . En effet, dans ce cas tous les poids sont égaux à  $1/N$ . De plus,  $v$  est égale à la divergence du  $\chi^2$  entre  $\mu$  et  $\eta$ .

**Définition 1.** Soient  $\pi_1$  et  $\pi_2$  deux probabilités telles que  $\pi_1$  est absolument continue par rapport à  $\pi_2$ . On appelle divergence du  $\chi^2$  entre  $\pi_1$  et  $\pi_2$  la quantité

$$\chi^2(\pi_1, \pi_2) = \int \left( \frac{d\pi_1}{d\pi_2}(x) - 1 \right)^2 \pi_2(dx).$$

La taille effective de l'échantillon est un indicateur de la qualité de l'échantillonnage pondéré défini par

$$N_{\text{eff}} = \frac{1}{\sum_{i=1}^N (w^i)^2}.$$

$N_{\text{eff}}$  vaut  $N$  lorsque toutes les particules ont un poids égal à  $1/N$  et 1 lorsqu'une des particules a un poids égal 1 et les  $N - 1$  autres ont un poids nul [Kong (1992); Arulampalam et al. (2002)].  $N_{\text{eff}}$  est liée à  $v$ , car

$$\frac{N_{\text{eff}}}{N} \longrightarrow \frac{1}{v + 1} = \frac{(\int g(x)\eta(dx))^2}{\int g(x)^2\eta(dx)}$$

lorsque  $N \rightarrow \infty$ .

# Filtrage bayésien et approximation particulière

On présente dans ce chapitre les notions et les outils qui définissent le contexte de la thèse : le filtrage bayésien et son approximation particulière. Les propriétés de convergence en loi du filtrage particulière exposées ici, ainsi que d'autres résultats de convergence, sont détaillés dans [Le Gland (2012)] ou [Del Moral (2004)].

## Modèles de Markov caché

Considérons deux processus stochastiques en temps discret  $\{X_k\}_{k \geq 0}$  et  $\{Y_k\}_{k \geq 0}$ .  $\{X_k\}_{k \geq 0}$  est inobservé et est appelé processus d'état. Les  $X_k$  sont à valeurs dans un espace  $E$ , l'espace d'état.  $\{Y_k\}_{k \geq 0}$  est appelé processus d'observation et les  $Y_k$  sont à valeurs dans un espace  $F$ .

On suppose que le processus d'état  $\{X_k\}_{k \geq 0}$  est une chaîne de Markov de noyaux de transition  $\{Q_k\}_{k \geq 1}$ , c'est-à-dire que, pour tout  $k \geq 1$ ,

$$\mathbb{P}[X_k \in dx | X_{0:k-1}] = \mathbb{P}[X_k \in dx | X_{k-1}] = Q_k(X_{k-1}, dx).$$

La loi de l'état initial  $X_0$  est notée  $Q_0$ .

On suppose que le processus des observations  $\{Y_k\}_{k \geq 0}$  vérifie l'hypothèse de *canal sans mémoire*.

### Canal sans mémoire.

- pour tout  $n \geq 0$  et pour tout  $k \in \{0, \dots, n\}$ , la loi conditionnelle de l'observation  $Y_k$  sachant les états  $\{X_0, \dots, X_n\}$  est égale à la loi conditionnelle de  $Y_k$

sachant l'état courant  $X_k$ , i.e.,

$$\mathbb{P}[Y_k \in dy_k | X_{0:n}] = \mathbb{P}[Y_k \in dy_k | X_k] = g_k(y_k, X_k);$$

- pour tout  $k \geq 0$ , les observations  $\{Y_0, \dots, Y_k\}$  sont indépendantes conditionnellement à  $\{X_0, \dots, X_k\}$ , i.e.,

$$\mathbb{P}[Y_{0:k} \in dy_{0:k} | X_{0:k}] = \mathbb{P}[Y_0 \in dy_0 | X_{0:k}] \times \dots \times \mathbb{P}[Y_k \in dy_k | X_{0:k}].$$

On suppose de plus que la loi conditionnelle de  $Y_k$  sachant  $X_k$  admet une densité  $g_k(y_k, X_k)$  par rapport à une mesure dominante sur  $F$ . Dans la suite, on omettra la dépendance de  $g_k$  en  $y_k$  en notant  $g_k(y_k, x) = g_k(x)$ . On appelle fonction de vraisemblance la fonction  $g_k(x)$ , qui mesure l'adéquation entre un état caché quelconque  $x \in E$  et l'observation  $Y_k$ .

Sous ces hypothèses, le processus joint  $\{X_k, Y_k\}_{k \geq 0}$  à valeurs dans  $E \times F$  est appelé modèle de Markov caché, ou modèle à espace d'état.

**Exemple 1** (Modèle linéaire gaussien). Soit  $E = \mathbb{R}^d$  l'espace d'état et  $F = \mathbb{R}^{d'}$  l'espace des observations. L'état initial suit la loi  $X_0 \sim \mathcal{N}(m_0, \Sigma_0)$ . Le processus d'état est défini par le modèle

$$X_k = F_k X_{k-1} + V_k$$

( $k \geq 1$ ) où  $\{V_k\}_{k \geq 1}$  est un bruit blanc gaussien, avec  $V_k \sim \mathcal{N}(0, \Sigma_k)$ , et  $F_k$  est une matrice de taille  $d \times d$ . Le processus des observations est défini par le modèle

$$Y_k = H_k X_k + W_k$$

( $k \geq 1$ ) où  $\{W_k\}_{k \geq 1}$  est un bruit blanc gaussien, avec  $W_k \sim \mathcal{N}(0, R_k)$ , et  $H_k$  est une matrice de taille  $d \times d'$ .

**Exemple 2** (Modèle non-linéaire avec bruit de mesure additif gaussien). Le processus d'état est défini par le modèle

$$X_k = F_k(X_{k-1}, V_k)$$

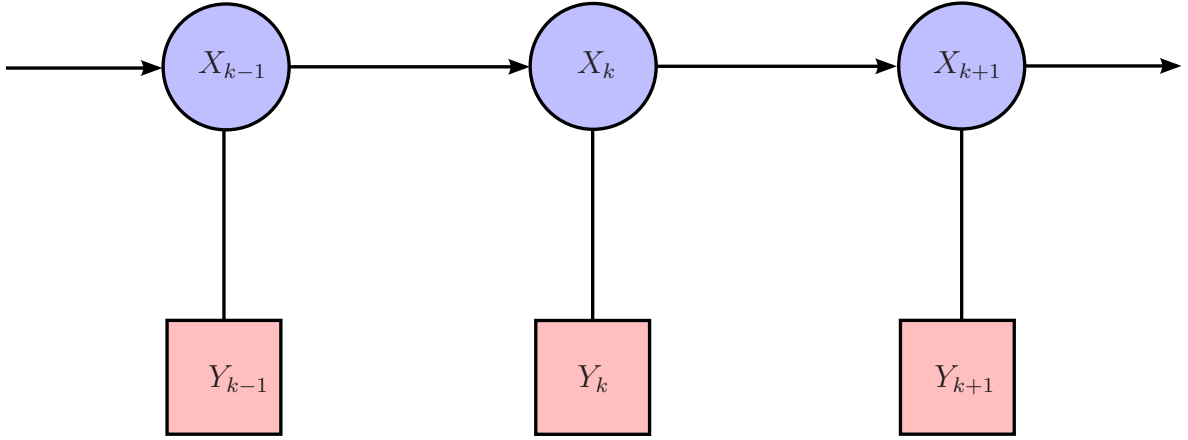


FIGURE 1 – Modèle de Markov caché.

où  $\{V_k\}_{k \geq 1}$  est une bruit blanc et  $F_k(\cdot)$  une fonction mesurable quelconque. Le processus des observations est défini par le modèle

$$Y_k = H_k(X_k) + W_k$$

où  $\{W_k\}_{k \geq 1}$  est une bruit blanc et  $H_k(\cdot)$  une fonction mesurable quelconque, appelée fonction de mesure.

**Exemple 3** (Modèle à volatilité stochastique). Le processus d'état est

$$X_k - \mu = \phi(X_{k-1} - \mu) + \sigma V_k$$

à valeurs dans  $\mathbb{R}$ , où  $(\mu, \phi, \sigma) \in \mathbb{R} \times \mathbb{R} \times \mathbb{R}_+$ , et le processus des observations est

$$Y_k = e^{X_k} W_k,$$

à valeurs dans  $\mathbb{R}$ , avec  $\{V_k\}_{k \geq 1}$  et  $\{W_k\}_{k \geq 1}$  des bruits blancs gaussiens centrés réduits.

## Filtrage bayésien

Le filtrage bayésien est un problème d'estimation statistique dans les modèles de Markov caché. Il consiste à calculer à chaque instant la loi conditionnelle de l'état caché  $X_k$

sachant les observations  $\{Y_0, \dots, Y_k\}$ . On note

$$\mu_k(dx) = \mathbb{P}[X_k \in dx | Y_{0:k}]$$

cette loi, appelée filtre bayésien. La connaissance a priori sur les états cachés est modélisée par les noyaux markoviens  $\{Q_k\}_{k \geq 1}$  et par la loi initiale  $Q_0$ . L'adéquation entre observation et état est modélisée par les fonctions de vraisemblance  $\{g_k\}_{k \geq 0}$ .  $\mu_k$  est la loi a posteriori à l'instant  $k$  à partir de laquelle on peut obtenir des estimateurs bayésiens, tels que l'espérance a posteriori

$$\mathbb{E}[X_k | Y_{0:k}] = \int_E x \mu_k(dx),$$

la variance a posteriori

$$\mathbb{V}[X_k | Y_{0:k}] = \int_E x^2 \mu_k(dx) - \left( \int_E x \mu_k(dx) \right)^2,$$

la probabilité a posteriori de dépassement d'un seuil

$$\mathbb{P}[|X_k| \geq c | Y_{0:k}] = \int_E \mathbb{I}_{\{|x| \geq c\}} \mu_k(dx),$$

etc.

Le filtre bayésien suit une relation de récurrence. En effet, en utilisant l'hypothèse de canal sans mémoire et la formule de Bayes, on obtient

$$\begin{aligned} \mu_k(dx_k) &\propto \int_{E^k} \mathbb{P}[X_{0:k} \in dx_{0:k}, Y_{0:k} = y_{0:k}] \quad (\text{où l'intégrale est prise par rapport à } dx_{0:k-1}) \\ &\propto g_k(x_k) \int_E Q_k(x_{k-1}, dx_k) \left( \int_{E^{k-1}} Q_0(dx_0) \prod_{j=1}^{k-1} Q_j(x_{j-1}, dx_j) \prod_{j=0}^{k-1} g_j(x_j) \right) \\ &\propto g_k(x_k) \int_E Q_k(x_{k-1}, dx_k) \left( \int_{E^{k-1}} \mathbb{P}[X_{0:k-1} \in dx_{0:k-1}, Y_{0:k-1} = y_{0:k-1}] \right) \\ &\quad (\text{où l'intégrale intérieure est prise par rapport à } dx_{0:k-2}) \\ &\propto g_k(x_k) \int_E Q_k(x_{k-1}, dx_k) \mu_{k-1}(dx_{k-1}). \end{aligned}$$

Notons

$$\mu_k^-(dx) = \mathbb{P}[X_k \in dx | Y_{0:k-1}]$$

la loi conditionnelle de l'état courant sachant les observations passées. Cette probabilité est appelée prédicteur. D'après la formule de Chapman–Kolmogorov, on a

$$\mu_k^-(dx') = \int_E Q_k(x, dx') \mu_{k-1}(dx), \quad (13)$$

de sorte que

$$\mu_k(dx) = \frac{g_k(x) \mu_k^-(dx)}{\int_E g_k(x) \mu_k^-(dx)}. \quad (14)$$

(13) est appelée équation de prédiction et (14) est appelée équation de mise à jour. L'équation de mise à jour correspond à la formule de Bayes, avec le prédicteur jouant le rôle de loi a priori. (13) et (14) définissent la relation de récurrence suivie par le filtre bayésien. Cette relation est résumée par le schéma suivant :

$$\mu_{k-1} \longrightarrow \mu_k^- = \mu_{k-1} Q_k \longrightarrow \mu_k = \frac{g_k \mu_k^-}{\int_E g_k \mu_k^-},$$

où on utilise la notation  $\mu_{k-1} Q_k(dx') = \int_E Q_k(x, dx') \mu_{k-1}(dx)$ .

Les algorithmes de filtrage bayésien calculent, de façon généralement approchée, le prédicteur puis la loi a posteriori à chaque instant.

Lorsque le modèle de Markov caché est linéaire et gaussien, comme dans l'exemple 1, on montre que le prédicteur et le filtre bayésien sont des lois gaussiennes. L'espérance et la matrice de covariance de ces lois sont calculées récursivement et de manière exacte par un algorithme déterministe, le filtre de Kalman [Kalman (1960)]. Il existe des versions approximatives du filtre de Kalman pour les modèles non-linéaires, le filtre de Kalman étendu et le filtre de Kalman non-parfumé, basées sur des approximations linéaires et gaussiennes du modèle. Ces méthodes peuvent donner de bons résultats, mais sont mises en défaut lorsque le modèle est trop éloigné du cas idéal linéaire gaussien.

Nous nous intéressons dans cette thèse au filtrage particulaire. L'intérêt des méthodes particulières est qu'elles permettent d'approcher arbitrairement près le filtre bayésien, au prix d'un certain nombre (potentiellement très grand) de tirages Monte Carlo. Le filtrage particulaire est relativement simple à implémenter et donne des approximations convergentes sous des hypothèses faibles [Del Moral (2004); Le Gland (2012)].

## Filtrage particulaire

Le filtrage particulaire est une méthode de Monte Carlo qui approche le filtre bayésien dans des modèles de Markov caché non-linéaires. Il consiste en une application séquentielle de l'algorithme d'échantillonnage pondéré, c'est pourquoi le filtrage particulaire fait partie des méthodes dites de Monte Carlo séquentielles. A chaque instant, la loi a posteriori est approchée par une somme pondérée de masses de Dirac.

### Echantillonnage selon le noyau markovien

Supposons que l'on dispose à l'instant  $k - 1$  d'une approximation particulaire de la loi a posteriori sous la forme

$$\mu_{k-1}^N = \sum_{i=1}^N w_{k-1}^i \delta_{\xi_{k-1}^i}.$$

D'après l'équation de prédiction (13), on obtient l'approximation particulaire  $\mu_k^{-N}$  du prédicteur à l'instant  $k$  en appliquant le noyau de transition markovien  $Q_k$  à  $\mu_{k-1}^N$ , c'est-à-dire en calculant

$$\mu_{k-1}^N Q_k. \quad (15)$$

Puis, d'après l'équation de mise à jour (14), l'approximation particulaire  $\mu_k^N$  de la loi a posteriori est donnée par

$$\frac{g_k \mu_k^{-N}}{\int_E g_k \mu_k^{-N}}. \quad (16)$$

Dans les algorithmes de filtrage particulaire, l'intégrale (15) est approchée par échantillonnage Monte Carlo. L'implémentation la plus basique est appelée algorithme d'échantillonnage pondéré séquentiel (sequential importance sampling).

### Sequential importance sampling (SIS)

A l'instant  $k$ , un nouvel échantillon de particules  $(\xi_k^1, \dots, \xi_k^N)$  est obtenu en propageant indépendamment chaque particule de l'instant précédent à travers le noyau markovien. On simule indépendamment

$$\xi_k^i \sim Q_k(\xi_{k-1}^i, dx)$$

pour  $i = 1, \dots, N$ . L'approximation particulière du prédicteur est alors

$$\mu_k^{-N} = \sum_{i=1}^N w_{k-1}^i \delta_{\xi_k^i}.$$

On obtient ensuite l'approximation de la loi a posteriori selon (16),

$$\mu_k^N = \frac{g_k \mu_k^{-N}}{\int_E g_k \mu_k^{-N}} = \frac{w_{k-1}^i g_k(\xi_k^i) \delta_{\xi_k^i}}{\sum_{i=1}^N w_{k-1}^i g_k(\xi_k^i)},$$

soit

$$\mu_k^N = \sum_{i=1}^N w_k^i \delta_{\xi_k^i}, \quad \text{où} \quad w_k^i = \frac{w_{k-1}^i g_k(\xi_k^i)}{\sum_{i=1}^N w_{k-1}^i g_k(\xi_k^i)}.$$

L'algorithme SIS est résumé dans l'algorithme 1.

---

**Algorithm 1** Sequential importance sampling.

---

- $k = 0$ .
    - Pour  $i = 1, \dots, N$ , simuler les particules  $\xi_0^i \sim Q_0(dx)$ .
    - Pour  $i = 1, \dots, N$ , calculer les poids  $w_0^i = \frac{g_0(\xi_0^i)}{\sum_{i=1}^N g_0(\xi_0^i)}$ .
  - $k \geq 1$ .
    - Pour  $i = 1, \dots, N$ , simuler les particules  $\xi_k^i \sim Q_k(\xi_{k-1}^i, dx)$ .
    - Pour  $i = 1, \dots, N$ , calculer les poids  $w_k^i = \frac{w_{k-1}^i g_k(\xi_k^i)}{\sum_{i=1}^N w_{k-1}^i g_k(\xi_k^i)}$ .
- 

Notons  $Q_{0:k}$  la loi jointe de la variable aléatoire  $X_{0:k}$  à valeurs trajectorielles, i.e.

$$Q_{0:k}(dx_{0:k}) = Q_0(dx_0)Q_1(x_0, dx_1) \cdots Q_k(x_{k-1}, dx_k).$$

Notons  $g_{0:k}$  la fonction de vraisemblance jointe, qui mesure l'adéquation entre une trajectoire d'états cachés  $x_{0:k} \in E^{k+1}$  et les observations  $Y_{0:k}$ . D'après l'hypothèse de canal sans mémoire

$$g_{0:k}(x_{0:k}) = g_0(x_0) \cdots g_k(x_k).$$

Considérons l'algorithme d'échantillonnage pondéré classique (non séquentiel) pour lequel la loi d'échantillonnage est  $Q_{0:k}$  et la fonction de pondération est  $g_{0:k}$ . On simule

indépendamment des variables aléatoires à valeurs trajectorielles

$$\xi_{0:k}^i \sim Q_{0:k}$$

pour  $i = 1, \dots, N$ , puis on associe à chaque  $\xi_{0:k}^i$  le poids

$$w_{0:k}^i = \frac{g_{0:k}(\xi_{0:k}^i)}{\sum_{i=1}^N g_{0:k}(\xi_{0:k}^i)} = \frac{\frac{g_{0:k-1}(\xi_{0:k-1}^i)}{\sum_{i=1}^N g_{0:k-1}(\xi_{0:k-1}^i)} g_k(\xi_k^i)}{\sum_{i=1}^N \frac{g_{0:k-1}(\xi_{0:k-1}^i)}{\sum_{i=1}^N g_{0:k-1}(\xi_{0:k-1}^i)} g_k(\xi_k^i)} = \frac{w_{0:k-1}^i g_k(\xi_k^i)}{\sum_{i=1}^N w_{0:k-1}^i g_k(\xi_k^i)}.$$

Par une récurrence triviale, on montre que pour tout  $i = 1, \dots, N$ , la dernière composante temporelle de la v.a.  $\xi_{0:k}^i$  est distribuée comme la particule  $\xi_k^i$  et que  $w_{0:k}^i = w_k^i$ . Par conséquent, l'algorithme SIS correspond à un algorithme d'échantillonnage pondéré d'approximation non-séquentielle dans un espace de variables aléatoires à valeurs trajectorielles.

Soit  $\theta_k = \mathbb{E}[\phi(X_k)|Y_{0:k}]$  un estimateur bayésien. Son approximation délivrée par l'algorithme SIS est

$$\theta_k^{N,\text{SIS}} = \int_E \phi \mu_k^N = \sum_{i=1}^N w_k^i \phi(\xi_k^i).$$

Cette approximation vérifie le théorème central limite suivant, équivalent au théorème 2 du chapitre précédent :

**Théorème 3.**

$$\sqrt{N} \left( \theta_k^{N,\text{SIS}} - \theta_k \right) \longrightarrow \mathcal{N}(0, V_k^{\text{SIS}})$$

et

$$\sqrt{N} \left( \frac{\frac{1}{N} \sum_{i=1}^N g_{0:k}(\xi_{0:k}^i)}{\int_{E^{k+1}} g_{0:k}(x_{0:k}) Q_{0:k}(dx_{0:k})} - 1 \right) \longrightarrow \mathcal{N}(0, v_k^{\text{SIS}})$$

lorsque  $N \rightarrow \infty$ , avec

$$V_k^{\text{SIS}} = \frac{\int_{E^{k+1}} (\phi(x_k) - \theta_k) g_{0:k}(x_{0:k})^2 Q_{0:k}(dx_{0:k})}{\left( \int_{E^{k+1}} g_{0:k}(x_{0:k}) Q_{0:k}(dx_{0:k}) \right)^2} \quad \text{et} \quad v_k^{\text{SIS}} = \frac{\int_{E^{k+1}} g_{0:k}(x_{0:k})^2 Q_{0:k}(dx_{0:k})}{\left( \int_{E^{k+1}} g_{0:k}(x_{0:k}) Q_{0:k}(dx_{0:k}) \right)^2} - 1.$$

On définit la taille effective de l'échantillon (ESS) au temps  $k$  par

$$N_{\text{eff}} = \frac{1}{\sum_{i=1}^N (w_k^i)^2}.$$

Elle quantifie le nombre de particules qui ont un poids non négligeable et qui participent à l'approximation particulaire  $\mu_k^N$  (voir chapitre précédent).

Le problème posé par l'algorithme SIS est qu'après un certain nombre d'itérations, on observe souvent que la plupart des particules ont un poids très faible et seul un très petit nombre de particules ont un poids significatif. Peu de particules participent alors de manière effective à l'approximation de la loi a posteriori, ce qui implique que  $N_{\text{eff}}$  est petit par rapport à  $N$ . Ce phénomène est appelé « dégénérescence des poids ».

La méthode la plus classique pour lutter contre la dégénérescence des poids consiste à dupliquer les particules de poids fort et supprimer celles de poids faibles. Avant l'étape de prédiction, de nouvelles particules sont sélectionnées par tirage avec remise dans l'ancienne population en fonction des poids. Chaque particule a une probabilité d'être sélectionnée égale à son poids. L'algorithme dans lequel cette étape de ré-échantillonnage des particules est effectuée à chaque instant est appelé filtre bootstrap et a été proposé dans [Gordon et al. (1993)].

### Bootstrap filtering

À l'instant  $k$ , un nouvel échantillon de particules  $(\xi_{k-1}^1, \dots, \xi_{k-1}^N)$  est sélectionné parmi  $(\xi_{k-1}^1, \dots, \xi_{k-1}^N)$  par tirage avec remise en fonction des poids  $(w_{k-1}^1, \dots, w_{k-1}^N)$ . C'est-à-dire qu'on simule selon

$$\xi_{k-1}^i \sim w_{k-1}^1 \delta_{\xi_{k-1}^1} + \dots + w_{k-1}^N \delta_{\xi_{k-1}^N},$$

indépendamment pour  $i = 1, \dots, N$ , de telle sorte la loi a posteriori à l'instant  $k - 1$  soit approchée par une approximation particulaire non pondérée,

$$\frac{1}{N} \sum_{i=1}^N \delta_{\xi_{k-1}^i} \approx \sum_{i=1}^N w_{k-1}^i \delta_{\xi_{k-1}^i} = \mu_{k-1}^N.$$

Ce nouvel échantillon est ensuite propagé à travers le noyau markovien,

$$\xi_k^i \sim Q_k(\xi_{k-1}^i, dx)$$

indépendamment pour  $i = 1, \dots, N$ . L'approximation particulaire du prédicteur est alors

$$\mu_k^{-N} = \frac{1}{N} \sum_{i=1}^N \delta_{\xi_{k-1}^i}.$$

L'approximation de la loi a posteriori est donnée par

$$\mu_k^N = \frac{g_k \mu_k^{-N}}{\int_E g_k \mu_k^{-N}} = \frac{g_k(\xi_k^i) \delta_{\xi_k^i}}{\sum_{i=1}^N g_k(\xi_k^i)},$$

soit

$$\mu_k^N = \sum_{i=1}^N w_k^i \delta_{\xi_k^i}, \quad \text{où} \quad w_k^i = \frac{g_k(\xi_k^i)}{\sum_{i=1}^N g_k(\xi_k^i)}.$$

On remarque qu'après l'étape de ré-échantillonnage, les particules sélectionnées  $(\xi_{k-1}^1, \dots, \xi_{k-1}^N)$  ont toutes le même poids,  $1/N$ .

Le filtre bootstrap est résumé dans l'algorithme 2.

---

**Algorithm 2** Bootstrap filter.

---

- $k = 0$ .
    - Pour  $i = 1, \dots, N$ , simuler les particules  $\xi_0^i \sim Q_0(dx)$ .
    - Pour  $i = 1, \dots, N$ , calculer les poids  $w_0^i = \frac{g_0(\xi_0^i)}{\sum_{i=1}^N g_0(\xi_0^i)}$ .
  - $k \geq 1$ .
    - Pour  $i = 1, \dots, N$ , simuler  $\xi_{k-1}^i \sim w_{k-1}^1 \delta_{\xi_{k-1}^1} + \dots + w_{k-1}^N \delta_{\xi_{k-1}^N}$ .
    - Pour  $i = 1, \dots, N$ , simuler les particules  $\xi_k^i \sim Q_k(\xi_{k-1}^i, dx)$ .
    - Pour  $i = 1, \dots, N$ , calculer les poids  $w_k^i = \frac{g_k(\xi_k^i)}{\sum_{i=1}^N g_k(\xi_k^i)}$ .
- 

Soit  $\theta_k^{N,\text{boot}} = \sum_{i=1}^N w_k^i \phi(\xi_k^i)$  l'approximation particulière de l'estimateur bayésien  $\theta_k$  délivrée par le filtre bootstrap.  $\theta_k^{N,\text{boot}}$  vérifie le TCL suivant [Chopin (2004); Del Moral (2004); Le Gland (2012)] :

**Théorème 4.**

$$\sqrt{N} \left( \theta_k^{N,\text{boot}} - \theta_k \right) \longrightarrow \mathcal{N}(0, V_k^{\text{boot}})$$

et

$$\sqrt{N} \left( \frac{\prod_{l=1}^k \frac{1}{N} \sum_{i=1}^N g_l(\xi_l^i)}{\prod_{l=1}^k \int_E g_l(x_l) \mu_l^-(dx_l)} - 1 \right) \longrightarrow \mathcal{N}(0, v_k^{\text{boot}})$$

lorsque  $N \rightarrow \infty$ , avec

$$v_k^{\text{boot}} = \sum_{l=0}^k \frac{\int_E (g_l(x_l) \int_{E^{k-l}} g_{l+1:k}(x_{l+1:k}) Q_{l+1}(x_l, dx_{l+1}) \cdots Q_k(x_{k-1}, dx_k))^2 \mu_l^-(dx_l)}{\left( \int_E (g_l(x_l) \int_{E^{k-l}} g_{l+1:k}(x_{l+1:k}) Q_{l+1}(x_l, dx_{l+1}) \cdots Q_k(x_{k-1}, dx_k)) \mu_l^-(dx_l) \right)^2} - 1$$

et

$$V_k^{\text{boot}} = \sum_{l=0}^k \frac{\int_E (g_l(x_l) \int_{E^{k-l}} (\phi(x_k) - \theta_k) g_{l+1:k}(x_{l+1:k}) Q_{l+1}(x_l, dx_{l+1}) \cdots Q_k(x_{k-1}, dx_k))^2 \mu_l^-(dx_l)}{\left( \int_E (g_j(x_j) \int_E g_{j+1:k}(x_{j+1:k}) Q_{j+1}(x_j, dx_{j+1}) \cdots Q_k(x_{k-1}, dx_k)) \mu_l^-(dx_l) \right)^2},$$

avec la convention  $Q_0(x_{-1}, dx_0) = Q_0(dx_0)$ .

**Remarque 1.** Soit  $D_k = \int_{E^{k+1}} g_{0:k}(x_{0:k}) Q_{0:k}(dx_{0:k})$ .  $D_k$  est le produit des constantes de normalisation des lois a posteriori  $\mu_1, \dots, \mu_k$ , c'est-à-dire le dénominateur du rapport dans le deuxième résultat de convergence ci-dessus. En effet,

$$\begin{aligned} D_k &= \int_E g_k(x_k) \int_E Q_k(x_{k-1}, dx_k) \frac{\int_{E^{k-1}} g_{0:k-1}(x_{0:k-1}) Q_{0:k-1}(dx_{0:k-1})}{\int_{E^k} g_{0:k-1}(x_{0:k-1}) Q_{0:k-1}(dx_{0:k-1})} D_{k-1} \\ &\quad (\text{où, dans le rapport d'intégrales, celle au numérateur est prise par rapport à } dx_{0:k-2}) \\ &= \int_E g_k(x_k) \int_E Q_k(x_{k-1}, dx_k) \mu_{k-1}(dx_{k-1}) D_{k-1} \\ &= \int_E g_k(x_k) \mu_k^-(dx_k) D_{k-1}, \end{aligned}$$

de sorte que  $D_k = \prod_{l=1}^k \int_E g_l(x_l) \mu_l^-(dx_l)$ , avec la convention  $\mu_0^-(dx_0) = Q_0(dx_0)$ .

Bien que le faible intérêt pratique de l'algorithme SIS soit très largement admis, il n'existe pas aujourd'hui d'argument théorique permettant de déterminer, pour un modèle donné, si il vaut mieux utiliser l'algorithme SIS ou le filtre bootstrap. Etant données une séquence de noyaux markoviens  $\{Q_k\}_{k \geq 0}$  et de fonctions de vraisemblance  $\{g_k\}_{k \geq 0}$ , il n'y a pas de moyen simple de comparer les variances asymptotiques  $v_k^{\text{SIS}}$  et  $v_k^{\text{boot}}$ , ou  $V_k^{\text{SIS}}$  et  $V_k^{\text{boot}}$ , des erreurs d'approximation des deux algorithmes.

En pratique, ré-échantillonner la population de particules à chaque instant n'est pas nécessaire et on fait généralement le compromis de ré-échantillonner à certains instants seulement. Par exemple, on ré-échantillonne lorsque la taille effective de l'échantillon

tombe sous un certain seuil fixé par l'utilisateur, c'est-à-dire si

$$N_{\text{eff}} < N_{\text{th}}$$

où  $N_{\text{th}} = \alpha N$  avec  $\alpha \in (0, 1)$ .

## Changement de loi d'échantillonnage

Comme dans l'algorithme d'échantillonnage pondéré non séquentiel (voir chapitre précédent), les particules peuvent être échantillonnées selon une loi différente de celle naturellement présente dans le modèle, c'est-à-dire ici selon un noyau markovien différent de celui du processus d'état. Pour cela, il faut que ce nouveau noyau  $M_k$  domine  $Q_k$  à tout instant  $k$ , afin que la densité  $\frac{dQ_k(x, \cdot)}{dM_k(x, \cdot)}$  existe et que les particules propagées selon  $M_k$  puissent être correctement pondérées.

Lorsque les particules sont propagées selon le noyau du processus d'état, l'approximation  $\mu_k^{-N}$  du prédicteur est une représentation particulière de la mesure

$$\mu_{k-1}^N Q_k = \sum_{i=1}^N w_{k-1}^i Q_k(\xi_{k-1}^i, dx)$$

qui s'écrit sous la forme  $\mu_k^{-N} = \sum_{i=1}^N w_{k-1}^i \delta_{\xi_k^i}$ , avec  $\xi_k^i \sim Q_k(\xi_{k-1}^i, dx)$  pour tout  $i \in \{1, \dots, N\}$ . (En cas de ré-échantillonnage, on remplace les poids  $w_{k-1}^i$  par  $1/N$  et les particules  $\xi_{k-1}^i$  par les particules ré-échantillonnées  $\xi_{k-1}^i$ .) Lorsque les particules sont propagées selon un noyau  $M_k$  différent de  $Q_k$ , l'approximation particulière de  $\mu_k^{-N}$  est basée sur la décomposition  $Q_k(x, dx') = \frac{dQ_k(x, \cdot)}{dM_k(x, \cdot)}(x') M_k(x, dx')$  qui donne

$$\mu_{k-1}^N Q_k(dx) = \sum_{i=1}^N w_{k-1}^i \frac{dQ_k(\xi_{k-1}^i, \cdot)}{dM_k(\xi_{k-1}^i, \cdot)}(x) M_k(\xi_{k-1}^i, dx).$$

L'approximation du prédicteur est alors

$$\mu_k^{-N} = \sum_{i=1}^N w_{k-1}^i \frac{dQ_k(\xi_{k-1}^i, \cdot)}{dM_k(\xi_{k-1}^i, \cdot)}(\xi_k^i) \delta_{\xi_k^i},$$

avec  $\xi_k^i \sim M_k(\xi_{k-1}^i, dx)$  pour tout  $i \in \{1, \dots, N\}$ .

La loi d'échantillonnage optimale est  $M_k(x, dx') = \mu_k(dx')$ , qui est évidemment in-

connue. La taille effective de l'échantillon est alors maximale et la variance asymptotique des poids de l'algorithme SIS est nulle : à chaque instant,  $N_{\text{eff}} = N$  et  $v_k^{\text{SIS}} = 0$ .

On conclut cette section en présentant l'algorithme de filtrage particulaire générique (algorithme 3), où les particules sont ré-échantillonnées selon le critère de la taille effective de l'échantillon et où la loi d'échantillonnage n'est pas nécessairement le noyau markovien [Arulampalam et al. (2002); Cappé et al. (2007); Doucet et al. (2000)].

---

**Algorithm 3** Filtre particulaire générique.

---

- $k = 0$ .
    - Pour  $i = 1, \dots, N$ , simuler les particules  $\xi_0^i \sim M_0(dx)$ .
    - Pour  $i = 1, \dots, N$ , calculer les poids  $w_0^i = \frac{g_0(\xi_0^i) \frac{dQ_0}{dM_0}(\xi_0^i)}{\sum_{i=1}^N g_0(\xi_0^i) \frac{dQ_0}{dM_0}(\xi_0^i)}$ .
  - $k \geq 1$ .
    - If  $N_{\text{eff}} \geq N_{\text{th}}$ .
      - \* Pour  $i = 1, \dots, N$ , simuler les particules  $\xi_k^i \sim M_k(\xi_{k-1}^i, dx)$ .
      - \* Pour  $i = 1, \dots, N$ , calculer les poids  $w_k^i = \frac{w_{k-1}^i g_k(\xi_k^i) \frac{dQ_k(\xi_{k-1}^i, \cdot)}{dM_k(\xi_{k-1}^i, \cdot)}(\xi_k^i)}{\sum_{i=1}^N w_{k-1}^i g_k(\xi_k^i) \frac{dQ_k(\xi_{k-1}^i, \cdot)}{dM_k(\xi_{k-1}^i, \cdot)}(\xi_k^i)}$ .
    - If  $N_{\text{eff}} < N_{\text{th}}$ .
      - \* Pour  $i = 1, \dots, N$ , simuler  $\xi_{k-1}^i \sim w_{k-1}^1 \delta_{\xi_{k-1}^1} + \dots + w_{k-1}^N \delta_{\xi_{k-1}^N}$ .
      - \* Pour  $i = 1, \dots, N$ , simuler les particules  $\xi_k^i \sim M_k(\xi_{k-1}^i, dx)$ .
      - \* Pour  $i = 1, \dots, N$ , calculer les poids  $w_k^i = \frac{g_k(\xi_k^i) \frac{dQ_k(\xi_{k-1}^i, \cdot)}{dM_k(\xi_{k-1}^i, \cdot)}(\xi_k^i)}{\sum_{i=1}^N g_k(\xi_k^i) \frac{dQ_k(\xi_{k-1}^i, \cdot)}{dM_k(\xi_{k-1}^i, \cdot)}(\xi_k^i)}$ .
-



# Introduction

Nonlinear statistical filtering is a Bayesian estimation problem in dynamical models, using sequentially delivered observations. The random parameter to estimate is dynamic; it is called state, or hidden state. Its dynamics obeys to a Markov process. The observations are linked to the state thanks to a likelihood function. The filtering problem exists in continuous time or discrete time. We work on this dissertation uniquely with discrete time models, which are mainly used in the applications we consider.

Let  $k \in \mathbb{N}$  be the discrete time index. The hidden state at time  $k$  is denoted  $X_k$  and the Markov process  $\{X_k\}_{k \geq 0}$  is called state process, or state dynamics. At time  $k$ , the current observation is denoted  $Y_k$  and the process  $\{Y_k\}_{k \geq 0}$  is called observation process. The joint process  $\{X_k, Y_k\}_{k \geq 0}$  is called hidden Markov model, or state-space model. Hidden Markov models are used in numerous fields: aerospace and defence [Hue et al. (2002); Ristic et al. (2004); Gustafsson (2010)], speech recognition [Juang and Rabiner (1991)], biomedical engineering [Eddy (1996); Brockwell et al. (2004)], econometrics [Chopin and Pelgrin (2004); Rossi and Gallo (2006)], to quote only a few examples.

Filtering consists in computing at each time step the conditional distribution of the current state  $X_k$  w.r.t. all the past observations  $Y_0, \dots, Y_k$ , i.e. the posterior distribution,

$$\mu_k(dx) = \mathbb{P}[X_k \in dx | Y_0, \dots, Y_k]$$

The sequence of conditional probability  $\mu_k$  is called the Bayesian filter. The filtering algorithms (called "filters") can be recursive or non-recursive, the former being more popular since they are generally faster. A recursive filter processes the observations one by one rather than re-using at each time step the past observations. On the other hand, the filtering algorithms that process several observations at each time step require more computations (however, one can alleviate the amount of processed observations by considering uniquely the observations delivered at well-chosen times [Musso (1993)]).

When the model is linear and Gaussian, the sequence of Bayesian filters is a sequence

of Gaussian distributions whose expectation and covariance can be computed exactly thanks to the Kalman filter [Kalman (1960)], which is therefore the optimal method in this case. When the model is nonlinear, there is no such a method providing the exact Bayesian filter. Approximation algorithms based on Kalman filtering, as the extended Kalman filter or the unscented Kalman filter [Arulampalam et al. (2002); Julier and Uhlmann (2004)], give a biased Gaussian approximation of the Bayesian filter.

Particle filtering allows to compute recursively and approximately the Bayesian filter in nonlinear non-Gaussian models. Particle filtering algorithms are sequential Monte Carlo (SMC) methods based on the importance sampling principle. At each time step, the posterior  $\mu_k$  is approximated by a weighted sum of  $N$  Dirac measures, denoted  $\mu_k^N$ , centered at random points in the state space. These points are called "particles".  $\mu_k^N$  converges, under weak assumptions, to  $\mu_k$  as  $N$  tends to infinity [Crisan and Doucet (2002); Del Moral (2004)]. Thus, particle filtering can provide more and more accurate approximations of the Bayesian filter as the computational power increases, which is its most interesting feature.

A common drawback of particle filtering is that, after some time, only a few particles have a nonzero weight and the other have a weight that is numerically evaluated at zero. Filtering is then severely impoverished, since only a few particles participate to the posterior approximation. This phenomenon is called "weight degeneracy". It is classically handled by resampling the particles according to their weight, so that particles with a large weight are duplicated and those with a small weight are deleted. This solution has initially been proposed in [Gordon et al. (1993)].

Weight degeneracy is particularly severe when the model is very informative, when the state process noise or the observation noise is small, e.g. [Oudjane and Musso (2000)]. We work in this thesis with a deterministic Bayesian method, the Laplace method, that is efficient when the model is informative, unlike particle filtering. The Laplace method is classically used to approximate posterior moments in static models (i.e., when the hidden state is not dynamic) [Tierney and Kadane (1986)]. Under regularity and identifiability assumptions, these approximations are consistent as the observation sample size tends to infinity or, equivalently, as the observation noise intensity tends to zero. In this thesis, we propose to associate the Laplace method and particle filtering in order to improve filtering, especially in the difficult case where the model is highly informative.

Several authors have associated Monte Carlo sampling and the Laplace method

for Bayesian estimation; some examples are: [Kuk (1999)], [Guihenneuc-Jouyaux and Rousseau (2005)], [Jungbacker and Koopman (2007)] et [Kleppe and Skaug (2012)]. These articles, however, do not deal with particle filtering.

In part I, we present the main notions and tools we use in the thesis. The Laplace method, which will be used throughout this dissertation, is proven in chapter 1 in the most general case. In chapter 2, we study the behaviour of importance sampling in the difficult situations where the model is highly informative or highly dimensional.

Part II deals with the application of the Laplace method to static or dynamic Bayesian estimation problems. In chapter 3, we present the estimation of posterior moments in static models by Laplace approximations, that are consistent as the number of observations tends to infinity. In chapter 4, we consider a state-space model with small observation noise and we propose a recursive filter where the integral computations involved in the Bayesian filter recursion are performed with the Laplace method. Under the unrealistic assumption that the likelihood admits an unique maximum at each time step, we study the consistency of this algorithm and the propagation of the approximation error over time. We overcome this strong assumption in chapter 5 by considering that the state dynamics noise is small as well (at the same order). We define an algorithm, the Kalman Laplace filter (KLF), associating Kalman filtering and the Laplace method in this context.

Part III deals with the association of the Laplace method and particle filtering. In chapter 6, we consider a nonlinear Bayesian estimation problem to be solved by importance sampling. We propose a transformation of the sampling distribution based on Laplace approximation formulas, which is observed to improve estimation on a simple example. From this, we propose in chapter 7 a recursive filter, the Laplace particle filter (LPF), that combines the Laplace method and particle filtering. We also detail means to alleviate the computations required by Laplace approximations. In chapter 8, the LPF and the KLF are tested by simulation on three nonlinear filtering problems: bearings-only target tracking, ballistic target tracking during atmospheric reentry, neural decoding.



# Notations & Acronyms

|                    |                                                                 |
|--------------------|-----------------------------------------------------------------|
| w.r.t.             | with respect to                                                 |
| MLE                | maximum likelihood estimator                                    |
| MAP                | maximum a posteriori                                            |
| SIR                | sequential importance resampling                                |
| RPF                | regularized particle filter                                     |
| LPF                | Laplace particle filter                                         |
| KLF                | Kalman Laplace filter                                           |
| $E$                | state space                                                     |
| $d$                | state space dimension                                           |
| $n$                | observation sample size                                         |
| $k$                | discrete time index                                             |
| $X$                | state                                                           |
| $Y$                | observation                                                     |
| $\mu(dx)$          | posterior                                                       |
| $p(x)$             | posterior density                                               |
| $\mu^-(dx)$        | predictor                                                       |
| $p^-(x)$           | predictor density                                               |
| $\eta(dx)$         | prior                                                           |
| $q(x)$             | prior density                                                   |
| $g(x)$             | likelihood function                                             |
| $\ell(x)$          | contrast function                                               |
| $\varphi_{m,Q}$    | Gaussian measure with expectation $m$ and covariance matrix $Q$ |
| $\mathcal{B}_r(a)$ | open ball with center $a$ and radius $r$                        |
| $a_{1:n}$          | $\{a_1, \dots, a_n\}$                                           |



# Part I

## Preliminaries



# Chapter 1

## The Laplace method: general presentation

In this first chapter, we provide detailed and general proofs of the Laplace method, which is extensively used throughout the thesis. These proofs are adapted from [Kass et al. (1990)], where the Laplace method is established in the unidimensional case.

We first expose the basic principle of the method in section 1.1. Then, we provide proofs for the multidimensional Laplace method in section 1.2.

### 1.1 Principle

The Laplace method is an approximation method to compute multidimensional integrals in the form of

$$\int_{\mathbb{R}^d} b(x) e^{-\lambda h(x)} dx, \quad (1.1)$$

where  $\lambda$  is a large real positive parameter. Suppose that  $h$  admits a unique global minimum at  $\hat{x}$  and that  $h$  is regular in a neighbourhood of  $\hat{x}$ . Then,  $h'(\hat{x}) = 0$  and  $\det[h''(\hat{x})] > 0$ .

The Laplace method considers integral (1.1) as the integral of  $b$  w.r.t. a Gaussian measure with small variance of order  $1/\lambda$ . To do so, one replaces  $h(x)$  in the integrand by its second-order Taylor expansion at  $\hat{x}$ ,

$$h(x) \approx h(\hat{x}) + \frac{1}{2}(x - \hat{x})^T h''(\hat{x})(x - \hat{x}),$$

thus yielding

$$\int_{\mathbb{R}^d} b(x) e^{-\lambda h(x)} dx \approx e^{-\lambda h(\hat{x})} \int_{\mathbb{R}^d} b(x) e^{-\frac{\lambda}{2}(x-\hat{x})^T h''(\hat{x})(x-\hat{x})} dx.$$

When  $\lambda$  is large, the integral of  $b$  w.r.t. the Gaussian density

$$(2\pi)^{-d/2} \det [\lambda h''(\hat{x})]^{1/2} \exp \left( -\frac{\lambda}{2} (x - \hat{x})^T h''(\hat{x})(x - \hat{x}) \right)$$

is close to  $b(\hat{x})$ , since the norm of the covariance matrix  $[\lambda h''(\hat{x})]^{-1}$  is small, so that the density is concentrated around its maximum. Then,

$$\int_{\mathbb{R}^d} b(x) e^{-\frac{\lambda}{2}(x-\hat{x})^T h''(\hat{x})(x-\hat{x})} dx \approx (2\pi)^{d/2} b(\hat{x}) \det [\lambda h''(\hat{x})]^{-1/2},$$

so that one obtains the approximation

$$\int_{\mathbb{R}^d} b(x) e^{-\lambda h(x)} dx \approx (2\pi)^{d/2} b(\hat{x}) e^{-\lambda h(\hat{x})} \det [\lambda h''(\hat{x})]^{-1/2},$$

called the Laplace approximation.

## 1.2 The Laplace method

### 1.2.1 Differential calculus tools

Let  $E$  be an open subset of  $\mathbb{R}^d$ . If  $E'$  is a vector space, we denote by  $\mathcal{L}(E, E')$  the set of linear mappings defined over  $E$  and taking values in  $E'$ . Let  $f : E \rightarrow \mathbb{R}$  be a  $k$  times continuously differentiable function, with  $k > 1$ .

Suppose that the derivative of order  $k-1$  of  $f$  at  $x_0 \in E$  is a  $(k-1)$ -linear form  $f^{(k-1)}(x_0) : E \rightarrow \mathbb{R}$ . Consider the mapping

$$\begin{aligned} f^{(k-1)} : E &\rightarrow \mathcal{L}_{k-1}(E, \mathbb{R}) \\ x &\mapsto f^{(k-1)}(x). \end{aligned}$$

Then, the (first-order) derivative of  $f^{(k-1)}$  at  $x_0$  is a linear mapping  $f^{(k)}(x_0) : E \rightarrow \mathcal{L}_{k-1}(E, \mathbb{R})$ , i.e.  $f^{(k)}(x_0) \in \mathcal{L}(E, \mathcal{L}_{k-1}(E, \mathbb{R}))$ . Hence,  $f^{(k)}(x_0)$  can be identified with an element of  $\mathcal{L}_k(E, \mathbb{R})$ , i.e. a  $k$ -linear form.

Thus, the derivative of order  $k$  of  $f$  at  $x_0$  is a  $k$ -linear form defined by

$$\begin{aligned} f^{(k)}(x_0) : \quad E^k &\rightarrow \mathbb{R} \\ (e_1, \dots, e_k) &\mapsto f^{(k)}(x_0) \cdot (e_1, \dots, e_k). \end{aligned}$$

(The application of  $f^{(k)}(x_0)$  to  $(e_1, \dots, e_k) \in E^k$  is denoted by the dot product.) For all fixed  $(e_2, \dots, e_k) \in E^{k-1}$ , the linear form

$$\begin{aligned} E &\rightarrow \mathbb{R} \\ e_1 &\mapsto f^{(k)}(x_0) \cdot (e_1, \dots, e_k) \end{aligned}$$

matches with the derivative of  $x \mapsto f^{(k-1)}(x) \cdot (e_2, \dots, e_k)$  at  $x_0$ . For  $k = 1$ ,  $f'(x)$  is naturally the linear form defined by

$$\begin{aligned} f'(x) : E &\rightarrow \mathbb{R} \\ e &\mapsto f'(x) \cdot e. \end{aligned}$$

See the chapter 8 of [Dieudonné (1960)] for more details.

To avoid to introduce additional notations,  $f'(x)$  (resp.  $f''(x)$ ) indifferently denotes in this chapter a linear (resp. bilinear) form as defined above, as well as a  $1 \times d$  Jacobian matrix (resp.  $d \times d$  Hessian matrix); the context will make explicit what object is considered. Besides, we use the notation  $e^{(k)} = (e, \dots, e) \in E^k$  for all  $e \in E$ .

Lastly, we define the norm  $\|\cdot\|$  of  $f^{(k)}(x)$  by

$$\|f^{(k)}(x)\| = \sup_{|u|=1} \{|f^{(k)}(x) \cdot u^{(k)}|\}.$$

### 1.2.2 Statement of the general Laplace method

We now state two general theorems to be applied in Bayesian estimation problems in different contexts in the following chapters. Theorem 1.2.2 will be used in static models, when the observation sample size is large (chapter 3), whereas theorem 1.2.1 will be used in dynamic state-space models, when the observation noise intensity is small (chapter 4). In particular, we provide in theorem 1.2.1 an explicit upper bound for the approximation error, which will be useful in a filtering context to study the propagation of the approximation error over time. Theorem 1.2.2 is proven by Kass, Tierney and Kadane in [Kass et al. (1990)] when  $d = 1$ . The proofs in section 1.2.3

below are adapted from their article and valid when  $d \geq 2$ .

Consider the integral

$$\int_E b(x) e^{-\lambda h_\lambda(x)} dx$$

where  $E$  is a open subset of  $\mathbb{R}^d$  and  $\lambda > 0$  a real-valued parameter. For all  $a \in \mathbb{R}^d$  and  $r > 0$ , the open ball centered at  $a$  and with radius  $r$  is denoted  $\mathcal{B}_r(a)$ .

**Theorem 1.2.1.** *Suppose that there exist  $\lambda_0 > 0$ ,  $\Delta > 0$ ,  $m > 0$ ,  $p > 0$ ,  $K > 0$ ,  $M > 0$  and  $M' > 0$  such that, for all  $\lambda \geq \lambda_0$ :*

$$(i) \int_E |b(x)| e^{-\lambda_0 h_\lambda(x)} dx \leq K \lambda^p;$$

(ii) *there exists  $\hat{x}_\lambda$  such that, for all  $\delta \in (0, \Delta)$ , there exists  $c_\delta > 0$  independent of  $\lambda$  such that*

$$\inf_{x \notin \mathcal{B}_\delta(\hat{x}_\lambda)} \{h_\lambda(x) - h_\lambda(\hat{x}_\lambda)\} \geq c_\delta;$$

(iii)  *$h_\lambda$  is four times continuously differentiable over  $\mathcal{B}_\Delta(\hat{x}_\lambda)$  and, for all  $x \in \mathcal{B}_\Delta(\hat{x}_\lambda)$  and  $j \in \{0, \dots, 4\}$ ,  $\|h_\lambda^{(j)}(x)\| \leq M$ ;*

$$(iv) \det[h_\lambda''(\hat{x}_\lambda)] \geq m;$$

(v)  *$b$  is twice continuously differentiable over  $\mathcal{B}_\Delta(\hat{x}_\lambda)$  and, for all  $x \in \mathcal{B}_\Delta(\hat{x}_\lambda)$  and  $j \in \{0, 1, 2\}$ ,  $\|b^{(j)}(x)\| \leq M'$ .*

*Then, there exist positive constants  $C_0$ ,  $C$ ,  $C'$  and  $C''$  such that*

$$\begin{aligned} & \left| \frac{\int_E b(x) e^{-\lambda h_\lambda(x)} dx}{(2\pi)^{d/2} \det[\lambda h_\lambda''(\hat{x}_\lambda)]^{-1/2} e^{-\lambda h_\lambda(\hat{x}_\lambda)}} - b(\hat{x}_\lambda) \right| \\ & \leq C_0 \lambda^{d/2} e^{-\lambda c_\delta} \int_E |b(x)| e^{-\lambda_0 h_\lambda(x)} dx + \lambda^{-1} \left( C |b(\hat{x}_\lambda)| + C' \|b'(\hat{x}_\lambda)\| + C'' \sup_{x \in \mathcal{B}_\delta(\hat{x}_\lambda)} \|b''(x)\| \right) \end{aligned}$$

*for all  $\lambda \geq \lambda_0$  and for all  $\delta \in (0, \Delta)$  verifying  $\frac{\delta}{3} + \frac{\delta^2}{12} < \frac{m}{M^d}$ . In particular,*

$$\int_E b(x) e^{-\lambda h_\lambda(x)} dx = (2\pi)^{d/2} \det[\lambda h_\lambda''(\hat{x}_\lambda)]^{-1/2} e^{-\lambda h_\lambda(\hat{x}_\lambda)} \{b(\hat{x}_\lambda) + O(\lambda^{-1})\}$$

*when  $\lambda \rightarrow \infty$ .*

The purpose of the next result is to obtain a second-order expansion, under additional regularity assumptions.

**Theorem 1.2.2.** *Suppose that there exist  $\lambda_0 > 0$ ,  $\Delta > 0$ ,  $m > 0$ ,  $p > 0$ ,  $K > 0$ ,  $M > 0$  and  $M' > 0$  such that, for all  $\lambda \geq \lambda_0$ :*

$$(i) \int_E |b(x)| e^{-\lambda_0 h_\lambda(x)} dx \leq K \lambda^p;$$

(ii) *there exists  $\hat{x}_\lambda$  such that, for all  $\delta \in (0, \Delta)$ , there exists  $c_\delta > 0$  independent of  $\lambda$  such that*

$$\inf_{x \notin \mathcal{B}_\delta(\hat{x}_\lambda)} \{h_\lambda(x) - h_\lambda(\hat{x}_\lambda)\} \geq c_\delta;$$

(iii)  *$h_\lambda$  is six times continuously differentiable over  $\mathcal{B}_\Delta(\hat{x}_\lambda)$  and, for all  $x \in \mathcal{B}_\Delta(\hat{x}_\lambda)$  and  $j \in \{0, \dots, 6\}$ ,  $\|h_\lambda^{(j)}(x)\| \leq M$ ;*

$$(iv) \det[h_\lambda''(\hat{x}_\lambda)] \geq m;$$

(v)  *$b$  is four times continuously differentiable over  $\mathcal{B}_\Delta(\hat{x}_\lambda)$  and, for all  $x \in \mathcal{B}_\Delta(\hat{x}_\lambda)$  and  $j \in \{0, \dots, 4\}$ ,  $\|b^{(j)}(x)\| \leq M'$ .*

Then,

$$\begin{aligned} & \int_E b(x) e^{-\lambda h_\lambda(x)} dx \\ &= (2\pi)^{d/2} \det[\lambda h_\lambda''(\hat{x}_\lambda)]^{-1/2} e^{-\lambda h_\lambda(\hat{x}_\lambda)} \\ & \quad \times \left\{ b(\hat{x}_\lambda) + \lambda^{-1} \left( \frac{1}{2} \mathbb{E}[b''(\hat{x}_\lambda) \cdot X^{(2)}] - \frac{1}{6} \mathbb{E}[(b'(\hat{x}_\lambda) \cdot X)(h_\lambda'''(\hat{x}_\lambda) \cdot X^{(3)})] \right. \right. \\ & \quad \left. \left. + \frac{1}{72} b(\hat{x}_\lambda) \mathbb{E}[(h_\lambda'''(\hat{x}_\lambda) \cdot X^{(3)})^2] - \frac{1}{24} b(\hat{x}_\lambda) \mathbb{E}[h_\lambda^{(4)}(\hat{x}_\lambda) \cdot X^{(4)}] \right) + O(\lambda^{-2}) \right\}, \end{aligned}$$

where  $X \sim \mathcal{N}(0, h_\lambda''(\hat{x}_\lambda)^{-1})$ .

For theorems 1.2.2 and 1.2.2 above, the following remarks can be made.

**Remark 1.2.3.** *The coercivity condition (ii) implies that  $\hat{x}_\lambda$  is the unique global minimum of  $h_\lambda$  when  $\lambda \geq \lambda_0$ .*

**Remark 1.2.4.** *Condition (i) is verified as soon as  $\int_E |b(x)| dx < \infty$ , since*

$$\int_E |b(x)| e^{-\lambda_0 h_\lambda(x)} dx = e^{-\lambda_0 h_\lambda(\hat{x}_\lambda)} \int_E |b(x)| e^{-\lambda_0 (h_\lambda(x) - h_\lambda(\hat{x}_\lambda))} dx \leq e^{\lambda_0 M} \int_E |b(x)| dx$$

when  $\lambda \geq \lambda_0$ .

**Remark 1.2.5.** Let  $\sigma[h''_\lambda(\hat{x}_\lambda)]$  be the spectrum of the  $d \times d$  symmetric positive matrix  $h''_\lambda(\hat{x}_\lambda)$ . We have that  $\inf_{|u|=1} \{u^T h''_\lambda(\hat{x}_\lambda) u\} = \min \sigma[h''_\lambda(\hat{x}_\lambda)]$ , so that  $u^T h''_\lambda(\hat{x}_\lambda) u \geq \min \sigma[h''_\lambda(\hat{x}_\lambda)] |u|^2$ . Besides, for all  $\nu \in \sigma[h''_\lambda(\hat{x}_\lambda)]$ ,  $\nu \leq \sup_{|u|=1} \{u^T h''_\lambda(\hat{x}_\lambda) u\} = \|h''_\lambda(\hat{x}_\lambda)\| \leq M$ . Thus,  $m \leq \det[h''_\lambda(\hat{x}_\lambda)] \leq \min \sigma[h''_\lambda(\hat{x}_\lambda)] M^{d-1}$ , which implies  $\min \sigma[h''_\lambda(\hat{x}_\lambda)] \geq \frac{m}{M^{d-1}}$ . Hence,

$$u^T h''_\lambda(\hat{x}_\lambda) u \geq \frac{m}{M^{d-1}} |u|^2.$$

This inequality will be useful in the proofs in the next section.

**Example 1.2.6** (Euler's gamma function). Consider the real gamma function defined by

$$\Gamma(\lambda) = \int_0^\infty x^{\lambda-1} e^{-x} dx,$$

for all  $\lambda > 0$ . Let us apply theorem 1.2.1 to compute the integral  $\Gamma(\lambda)$ , letting  $b(x) = 1$  and  $h_\lambda(x) = -\frac{\lambda-1}{\lambda} \log x + \frac{x}{\lambda}$  for all  $x \in \mathbb{R}_+^*$ . We obtain Stirling's approximation formula

$$\Gamma(\lambda) = \sqrt{\frac{2\pi}{\lambda}} \left(\frac{\lambda}{e}\right)^\lambda \{1 + O(\lambda^{-1})\},$$

which is very precise (see figure 1.1).

### 1.2.3 Proofs of theorems 1.2.1 and 1.2.2

*Proof of theorem 1.2.1.* Let  $\lambda \geq \lambda_0$  and  $\delta \in (0, \Delta)$ . The integral of interest can be decomposed as

$$\mathcal{I}_\lambda = \int_{E - \mathcal{B}_\delta(\hat{x}_\lambda)} b(x) e^{-\lambda h_\lambda(x)} dx + \int_{\mathcal{B}_\delta(\hat{x}_\lambda)} b(x) e^{-\lambda h_\lambda(x)} dx. \quad (1.2)$$

Consider the first term of (1.2). For all  $x \in E - \mathcal{B}_\delta(\hat{x}_\lambda)$ ,

$$\begin{aligned} -\lambda h_\lambda(x) &= -\lambda h_\lambda(x) + \lambda_0 h_\lambda(x) - \lambda_0 h_\lambda(x) \\ &= -(\lambda - \lambda_0)(h_\lambda(x) - h_\lambda(\hat{x}_\lambda)) - (\lambda - \lambda_0)h_\lambda(\hat{x}_\lambda) - \lambda_0 h_\lambda(x), \end{aligned}$$

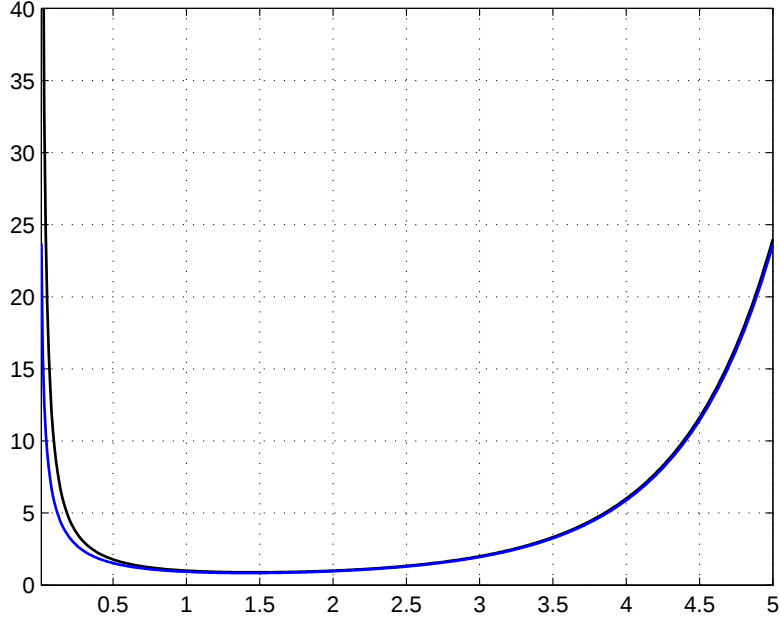


Figure 1.1: Euler's gamma function (black line) and Stirling's approximation obtained thanks to the Laplace method (blue line).

so that

$$\begin{aligned}
 \left| \int_{E-\mathcal{B}_\delta(\hat{x}_\lambda)} b(x) e^{-\lambda h_\lambda(x)} dx \right| &\leq \exp \left( -(\lambda - \lambda_0) \inf_{x \notin \mathcal{B}_\delta(\hat{x}_\lambda)} \{h_\lambda(x) - h_\lambda(\hat{x}_\lambda)\} \right) e^{-(\lambda - \lambda_0) h_\lambda(\hat{x}_\lambda)} \\
 &\quad \times \int_{E-\mathcal{B}_\delta(\hat{x}_\lambda)} |b(x)| e^{-\lambda_0 h_\lambda(x)} dx \\
 &\leq e^{-(\lambda - \lambda_0) c_\delta} e^{-\lambda h_\lambda(\hat{x}_\lambda)} e^{\lambda_0 h_\lambda(\hat{x}_\lambda)} \int_{E-\mathcal{B}_\delta(\hat{x}_\lambda)} |b(x)| e^{-\lambda_0 h_\lambda(x)} dx.
 \end{aligned}$$

Therefore, we obtain

$$\left| \frac{\int_{E-\mathcal{B}_\delta(\hat{x}_\lambda)} b(x) e^{-\lambda h_\lambda(x)} dx}{(2\pi)^{d/2} \det[\lambda h''_\lambda(\hat{x}_\lambda)]^{-1/2} e^{-\lambda h_\lambda(\hat{x}_\lambda)}} \right| \tag{1.3}$$

$$\begin{aligned}
 &\leq (2\pi)^{-d/2} \det[\lambda h''_\lambda(\hat{x}_\lambda)]^{1/2} e^{-(\lambda - \lambda_0) c_\delta} e^{\lambda_0 M} \int_E |b(x)| e^{-\lambda_0 h_\lambda(x)} dx \\
 &\leq (2\pi)^{-d/2} \sqrt{M} K \lambda^{d/2+p} e^{-(\lambda - \lambda_0) c_\delta} e^{\lambda_0 M} \\
 &= o(\lambda^{-1}). \tag{1.4}
 \end{aligned}$$

Consider now the second term of (1.2). Let  $x \in \mathcal{B}_\delta(\hat{x}_\lambda)$ . The first-order Taylor expansion of  $b$  and the third-order Taylor expansion of  $\lambda h_\lambda$  at  $\hat{x}_\lambda$ , with integral form of the remainder, are respectively [Dieudonné (1960)]

$$b(x) = b(\hat{x}_\lambda) + b'(\hat{x}_\lambda) \cdot (x - \hat{x}_\lambda) + \int_0^1 b''(\hat{x}_\lambda + \theta(x - \hat{x}_\lambda)) \cdot (x - \hat{x}_\lambda)^{(2)}(1 - \theta) d\theta$$

and

$$\begin{aligned} \lambda h_\lambda(x) &= \lambda h_\lambda(\hat{x}_\lambda) + \frac{\lambda}{2} h_\lambda''(\hat{x}_\lambda) \cdot (x - \hat{x}_\lambda)^{(2)} + \frac{\lambda}{6} h_\lambda'''(\hat{x}_\lambda) \cdot (x - \hat{x}_\lambda)^{(3)} \\ &\quad + \frac{\lambda}{6} \int_0^1 h_\lambda^{(4)}(\hat{x}_\lambda + \theta(x - \hat{x}_\lambda)) \cdot (x - \hat{x}_\lambda)^{(4)}(1 - \theta)^3 d\theta, \end{aligned}$$

because  $h'_\lambda(\hat{x}_\lambda) = 0$ . Let  $u = \lambda^{1/2}(x - \hat{x}_\lambda)$ , so that  $x \in \mathcal{B}_\delta(\hat{x}_\lambda)$  implies  $u \in \mathcal{B}_{\lambda^{1/2}\delta}(0)$ . Then,

$$\begin{aligned} &\int_{\mathcal{B}_\delta(\hat{x}_\lambda)} b(x) e^{-\lambda h_\lambda(x)} dx \\ &= \lambda^{-1/2} e^{-\lambda h_\lambda(\hat{x}_\lambda)} \int_{\mathcal{B}_{\lambda^{1/2}\delta}(0)} \exp \left( -\frac{1}{2} h_\lambda''(\hat{x}_\lambda) \cdot u^{(2)} + R_\lambda(u) \right) \\ &\quad \times \left( b(\hat{x}_\lambda) + \lambda^{-1/2} b'(\hat{x}_\lambda) \cdot u + \lambda^{-1} \int_0^1 b''(\hat{x}_\lambda + \theta(x - \hat{x}_\lambda)) \cdot u^{(2)}(1 - \theta) d\theta \right) du, \end{aligned}$$

where  $R_\lambda(u) = -\frac{\lambda^{-1/2}}{6} h_\lambda'''(\hat{x}_\lambda) \cdot u^{(3)} - \frac{\lambda^{-1}}{6} \int_0^1 h_\lambda^{(4)}(\hat{x}_\lambda + \theta(x - \hat{x}_\lambda)) \cdot u^{(4)}(1 - \theta)^3 d\theta$ . The Taylor expansion of  $\exp$  at 0 to the order 2 yields

$$\exp R_\lambda(u) = 1 + R_\lambda(u) + \frac{1}{2} R_\lambda(u)^2 + \frac{1}{6} (\exp S_\lambda(u)) R_\lambda(u)^3$$

where  $S_\lambda(u)$  is such that  $|S_\lambda(u)| \leq |R_\lambda(u)|$ . Let  $v = \frac{u}{|u|}$ , so that  $|v| = 1$ . Using the fact

that  $|u| \leq \lambda^{1/2}\delta$ , we have that

$$\begin{aligned}
& |S_\lambda(u)| \\
& \leq \left| - \left( \frac{\lambda^{-1/2}}{6} h_\lambda'''(\hat{x}_\lambda) \cdot v^{(3)} |u| + \frac{\lambda^{-1}}{6} |u|^2 \int_0^1 h_\lambda^{(4)}(\hat{x}_\lambda + \theta(x - \hat{x}_\lambda)) \cdot v^{(4)} (1 - \theta)^3 d\theta \right) |u|^2 \right| \\
& \leq \left( \frac{\delta \|h_\lambda'''(\hat{x}_\lambda)\|}{6} + \frac{\delta^2}{24} \sup_{x \in \mathcal{B}_\delta(\hat{x}_\lambda)} \left\{ \|h_\lambda^{(4)}(x)\| \right\} \right) |u|^2 \\
& \leq \left( \frac{\delta M}{6} + \frac{\delta^2 M}{24} \right) |u|^2,
\end{aligned}$$

and that

$$\begin{aligned}
|R_\lambda(u)|^3 & \leq \left( \frac{\lambda^{-1/2}}{6} \|h_\lambda'''(\hat{x}_\lambda)\| |u|^3 + \frac{\lambda^{-1}}{24} \sup_{x \in \mathcal{B}_\delta(\hat{x}_\lambda)} \left\{ \|h_\lambda^{(4)}(x)\| \right\} |u|^4 \right)^3 \\
& \leq \lambda^{-3/2} |u|^9 \left( \frac{M}{6} + \frac{\lambda^{-1/2}}{24} M |u| \right)^3 \\
& \leq \lambda^{-3/2} |u|^9 \left( \frac{M}{6} + \frac{\delta M}{24} \right)^3.
\end{aligned}$$

Hence, using the inequality  $h_\lambda''(\hat{x}_\lambda) \cdot u^{(2)} \geq \frac{m}{M^{d-1}} |u|^2$  (see remark 1.2.5), we have

$$\begin{aligned}
& \left| \int_{\mathcal{B}_{\lambda^{1/2}\delta}(0)} \left( b(\hat{x}_\lambda) + \lambda^{-1/2} b'(\hat{x}_\lambda) \cdot v |u| + \lambda^{-1} |u|^2 \int_0^1 b''(\hat{x}_\lambda + \theta(x - \hat{x}_\lambda)) \cdot v^{(2)} (1 - \theta) d\theta \right) \right. \\
& \times \exp \left( -\frac{1}{2} h_\lambda''(\hat{x}_\lambda) \cdot u^{(2)} \right) (\exp S_\lambda(u)) R_\lambda(u)^3 du \Big| \\
& \leq \int_{\mathcal{B}_{\lambda^{1/2}\delta}(0)} \left( |b(\hat{x}_\lambda)| + \lambda^{-1/2} \|b'(\hat{x}_\lambda)\| |u| + \frac{\lambda^{-1}}{2} \sup_{x \in \mathcal{B}_\delta(\hat{x}_\lambda)} \left\{ \|b''(x)\| \right\} |u|^2 \right) \\
& \times \exp \left( -\frac{1}{2} \left( \frac{m}{M^{d-1}} - \frac{\delta M}{3} - \frac{\delta^2 M}{12} \right) |u|^2 \right) \left( \lambda^{-3/2} |u|^9 \left( \frac{M}{6} + \frac{\delta M}{24} \right)^3 \right) \\
& \leq c |b(\hat{x}_\lambda)| \lambda^{-3/2} + c' \|b'(\hat{x}_\lambda)\| \lambda^{-2} + c'' \sup_{x \in \mathcal{B}_\delta(\hat{x}_\lambda)} \left\{ \|b''(x)\| \right\} \lambda^{-5/2} \tag{1.5}
\end{aligned}$$

where  $c$ ,  $c'$  and  $c''$  do not depend on  $\lambda$ . We can always choose  $\delta \in (0, \Delta)$  such that

$\frac{\delta}{3} + \frac{\delta^2}{12} < \frac{m}{M^d}$  so that the integral above is finite for all  $\lambda$ . Therefore,

$$\begin{aligned} & \int_{\mathcal{B}_\delta(\hat{x}_\lambda)} b(x) e^{-\lambda h_\lambda(x)} dx \\ &= \lambda^{-1/2} e^{-\lambda h_\lambda(\hat{x}_\lambda)} \left\{ \int_{\mathcal{B}_{\lambda^{1/2}\delta}(0)} \exp\left(-\frac{1}{2} h_\lambda''(\hat{x}_\lambda) \cdot u^{(2)}\right) \left(1 + R_\lambda(u) + \frac{1}{2} R_\lambda(u)^2\right) \right. \\ & \quad \times \left( b(\hat{x}_\lambda) + \lambda^{-1/2} b'(\hat{x}_\lambda) \cdot u + \lambda^{-1} \int_0^1 b''(\hat{x}_\lambda + \theta(x - \hat{x}_\lambda)) \cdot u^{(2)}(1 - \theta) d\theta \right) du + O(\lambda^{-3/2}) \left. \right\}, \end{aligned}$$

where

$$\begin{aligned} & 1 + R_\lambda(u) + \frac{1}{2} R_\lambda(u)^2 \\ &= 1 - \frac{\lambda^{-1/2}}{6} h_\lambda'''(\hat{x}_\lambda) \cdot u^{(3)} - \frac{\lambda^{-1}}{6} \int_0^1 h_\lambda^{(4)}(\hat{x}_\lambda + \theta(x - \hat{x}_\lambda)) \cdot u^{(4)}(1 - \theta)^3 d\theta \\ & \quad + \frac{1}{2} \left( -\frac{\lambda^{-1/2}}{6} h_\lambda'''(\hat{x}_\lambda) \cdot u^{(3)} - \frac{\lambda^{-1}}{6} \int_0^1 h_\lambda^{(4)}(\hat{x}_\lambda + \theta(x - \hat{x}_\lambda)) \cdot u^{(4)}(1 - \theta)^3 d\theta \right)^2 \\ &= 1 - \frac{\lambda^{-1/2}}{6} h_\lambda'''(\hat{x}_\lambda) \cdot u^{(3)} \\ & \quad + \lambda^{-1} \left( \frac{1}{72} (h_\lambda'''(\hat{x}_\lambda) \cdot u^{(3)})^2 - \frac{1}{6} |u|^4 \int_0^1 h_\lambda^{(4)}(\hat{x}_\lambda + \theta(x - \hat{x}_\lambda)) \cdot v^{(4)}(1 - \theta)^3 d\theta \right) \\ & \quad + \frac{\lambda^{-3/2}}{36} \left( h_\lambda'''(\hat{x}_\lambda) \cdot u^{(3)} \int_0^1 h_\lambda^{(4)}(\hat{x}_\lambda + \theta(x - \hat{x}_\lambda)) \cdot v^{(4)}(1 - \theta)^3 d\theta |u|^4 \right) \\ & \quad + \frac{\lambda^{-2}}{72} \left( \int_0^1 h_\lambda^{(4)}(\hat{x}_\lambda + \theta(x - \hat{x}_\lambda)) \cdot v^{(4)}(1 - \theta)^3 d\theta \right)^2 |u|^8. \end{aligned}$$

We can decompose the remaining integral of interest as

$$\begin{aligned} & \int_{\mathcal{B}_{\lambda^{1/2}\delta}(0)} \left( b(\hat{x}_\lambda) + \lambda^{-1/2} b'(\hat{x}_\lambda) \cdot u + \lambda^{-1} \int_0^1 b''(\hat{x}_\lambda + \theta(x - \hat{x}_\lambda)) \cdot u^{(2)}(1 - \theta) d\theta \right) \\ & \quad \times \exp\left(-\frac{1}{2} h_\lambda''(\hat{x}_\lambda) \cdot u^{(2)}\right) \left(1 + R_\lambda(u) + \frac{1}{2} R_\lambda(u)^2\right) du \\ &= I_{\lambda,0} + \lambda^{-1/2} I_{\lambda,1} + \lambda^{-1} I_{\lambda,2} + \lambda^{-3/2} I_{\lambda,3} + \lambda^{-2} I_{\lambda,4} + \lambda^{-5/2} I_{\lambda,4} + \lambda^{-3} I_{\lambda,6}, \quad (1.6) \end{aligned}$$

where the  $I_{\lambda,j}$ 's are uniquely defined. Then, (1.5) and (1.6) yield

$$\begin{aligned}
& \left| \frac{\int_{\mathcal{B}_\delta(\hat{x}_\lambda)} b(x) e^{-\lambda h_\lambda(x)} dx}{(2\pi)^{d/2} \det[h''_\lambda(\hat{x}_\lambda)]^{-1/2} e^{-\lambda h_\lambda(\hat{x}_\lambda)}} - b(\hat{x}_\lambda) \right| \\
& \leq \left| (2\pi)^{d/2} \det[h''_\lambda(\hat{x}_\lambda)]^{-1/2} I_{\lambda,0} - b(\hat{x}_\lambda) \right| + (2\pi)^{d/2} \det[h''_\lambda(\hat{x}_\lambda)]^{-1/2} \left\{ \lambda^{-1/2} |I_{\lambda,1}| + \lambda^{-1} |I_{\lambda,2}| \right. \\
& \quad \left. + \lambda^{-3/2} |I_{\lambda,3}| + \lambda^{-2} |I_{\lambda,4}| + \lambda^{-3/2} |I_{\lambda,5}| + \lambda^{-3} |I_{\lambda,6}| \right\} \\
& \quad + \sqrt{\frac{M}{2\pi}} \left( c |b(\hat{x}_\lambda)| \lambda^{-3/2} + c' \|b'(\hat{x}_\lambda)\| \lambda^{-2} + c'' \sup_{x \in \mathcal{B}_\delta(\hat{x}_\lambda)} \{\|b''(x)\|\} \lambda^{-5/2} \right).
\end{aligned}$$

First of all, we have that

$$I_{\lambda,0} = b(\hat{x}_\lambda) \int_{\mathcal{B}_\delta(\hat{x}_\lambda)} \exp \left( -\frac{1}{2} h''_\lambda(\hat{x}_\lambda) \cdot u^{(2)} \right) du,$$

so that

$$\begin{aligned}
& \left| (2\pi)^{-d/2} \det[h''_\lambda(\hat{x}_\lambda)]^{1/2} I_{\lambda,0} - b(\hat{x}_\lambda) \right| \\
& = \left| b(\hat{x}_\lambda) (2\pi)^{-d/2} \det[h''_\lambda(\hat{x}_\lambda)]^{1/2} \int_{\mathcal{B}_\delta(\hat{x}_\lambda)} \exp \left( -\frac{1}{2} h''_\lambda(\hat{x}_\lambda) \cdot u^{(2)} \right) du - b(\hat{x}_\lambda) \right| \\
& = |b(\hat{x}_\lambda)| \mathbb{P} [|X| \geq \lambda^{1/2} \delta],
\end{aligned}$$

where  $X \sim \mathcal{N}(0, h''_\lambda(\hat{x}_\lambda)^{-1})$ . Bienaymé–Tchebychev's inequality then yields

$$\left| (2\pi)^{-d/2} \det[h''_\lambda(\hat{x}_\lambda)]^{1/2} I_{\lambda,0} - b(\hat{x}_\lambda) \right| \leq \frac{|b(\hat{x}_\lambda)| \det[h''_\lambda(\hat{x}_\lambda)]^{-1/2}}{\lambda \delta^2}.$$

Besides,

$$I_{\lambda,1} = \int_{\mathcal{B}_\delta(\hat{x}_\lambda)} \left\{ -\frac{1}{6} b(\hat{x}_\lambda) h'''_\lambda(\hat{x}_\lambda) \cdot u^{(3)} + b'(\hat{x}_\lambda) \cdot u \right\} \exp \left( -\frac{1}{2} h''_\lambda(\hat{x}_\lambda) \cdot u^{(2)} \right) du = 0,$$

because it is the integral of an odd function over a ball centered at the origin. Moreover,

$$\begin{aligned}
I_{\lambda,2} &= \int_{\mathcal{B}_\delta(\hat{x}_\lambda)} \left\{ \frac{1}{72} b(\hat{x}_\lambda) (h_\lambda'''(\hat{x}_\lambda) \cdot u^{(3)})^2 - \frac{1}{6} b(\hat{x}_\lambda) |u|^4 \int_0^1 h_\lambda^{(4)}(\hat{x}_\lambda + \theta(x - \hat{x}_\lambda)) \cdot v^{(4)}(1 - \theta)^3 d\theta \right. \\
&\quad \left. - \frac{1}{6} (b'(\hat{x}_\lambda) \cdot u) (h_\lambda'''(\hat{x}_\lambda) \cdot u^{(3)}) + \left( \int_0^1 b''(\hat{x}_\lambda + \theta(x - \hat{x}_\lambda)) \cdot v^{(2)}(1 - \theta) d\theta \right) u^2 \right\} \\
&\quad \times \exp \left( -\frac{1}{2} h_\lambda''(\hat{x}_\lambda) \cdot u^{(2)} \right) du,
\end{aligned}$$

so that

$$\begin{aligned}
&|(2\pi)^{-d/2} \det[h_\lambda''(\hat{x}_\lambda)]^{1/2} I_{\lambda,2}| \\
&\leq (2\pi)^{-d/2} \det[h_\lambda''(\hat{x}_\lambda)]^{1/2} \left\{ \frac{1}{72} |b(\hat{x}_\lambda)| (h_\lambda'''(\hat{x}_\lambda) \cdot v^{(3)})^2 \int_{\mathcal{B}_\delta(\hat{x}_\lambda)} |u|^6 \exp \left( -\frac{1}{2} h_\lambda''(\hat{x}_\lambda) \cdot u^{(2)} \right) du \right. \\
&\quad + \frac{1}{6} |b(\hat{x}_\lambda)| \left| \int_0^1 h_\lambda^{(4)}(\hat{x}_\lambda + \theta(x - \hat{x}_\lambda)) \cdot v^{(4)}(1 - \theta)^3 d\theta \right| \int_{\mathcal{B}_\delta(\hat{x}_\lambda)} |u|^4 \exp \left( -\frac{1}{2} h_\lambda''(\hat{x}_\lambda) \cdot u^{(2)} \right) du \\
&\quad + \frac{1}{6} (b'(\hat{x}_\lambda) \cdot v) (h_\lambda'''(\hat{x}_\lambda) \cdot v^{(3)}) \int_{\mathcal{B}_\delta(\hat{x}_\lambda)} |u|^4 \exp \left( -\frac{1}{2} h_\lambda''(\hat{x}_\lambda) \cdot u^{(2)} \right) du \\
&\quad \left. + \left| \int_0^1 b''(\hat{x}_\lambda + \theta(x - \hat{x}_\lambda)) \cdot v^{(2)}(1 - \theta) d\theta \right| \int_{\mathcal{B}_\delta(\hat{x}_\lambda)} |u|^2 \exp \left( -\frac{1}{2} h_\lambda''(\hat{x}_\lambda) \cdot u^{(2)} \right) du \right\} \\
&\leq \frac{1}{2} \sup_{x \in \mathcal{B}_\delta(\hat{x}_\lambda)} \{ \|b''(x)\| \} m_2 + \frac{1}{6} \left( \|b'(\hat{x}_\lambda)\| \|h_\lambda'''(\hat{x}_\lambda)\| + \frac{1}{4} |b(\hat{x}_\lambda)| \sup_{x \in \mathcal{B}_\delta(\hat{x}_\lambda)} \{ \|h_\lambda^{(4)}(x)\| \} \right) m_4 \\
&\quad + \frac{1}{72} |b(\hat{x}_\lambda)| \|h_\lambda'''(\hat{x}_\lambda)\|^2 m_6
\end{aligned}$$

where  $m_j = \mathbb{E}[|X|^j]$ , with  $X \sim \mathcal{N}(0, h_\lambda''(\hat{x}_\lambda)^{-1})$ , for  $j \in \{2, 4, 6\}$ . Lastly, after some similar derivations, it can be shown that there exist positive constant numbers  $c_j, c'_j, c''_j$ , for  $j \in \{3, 4, 5, 6\}$ , such that

$$|I_{\lambda,j}| \leq c_j |b(\hat{x}_\lambda)| + c'_j \|b'(\hat{x}_\lambda)\| + c''_j \sup_{x \in \mathcal{B}_\delta(\hat{x}_\lambda)} \{ \|b''(x)\| \}.$$

Thus, we have

$$\begin{aligned}
& \left| \frac{\int_{\mathcal{B}_\delta(\hat{x}_\lambda)} b(x) e^{-\lambda h_\lambda(x)} dx}{(2\pi)^{d/2} \det[h_\lambda''(\hat{x}_\lambda)]^{-1/2} e^{-\lambda h_\lambda(\hat{x}_\lambda)}} - b(\hat{x}_\lambda) \right| \\
& \leq \lambda^{-1} \left\{ \frac{|b(\hat{x}_\lambda)| \det[h_\lambda''(\hat{x}_\lambda)]^{-1/2}}{\delta^2} + \frac{1}{2} \sup_{x \in \mathcal{B}_\delta(\hat{x}_\lambda)} \{ \|b''(x)\| \} m_2 \right. \\
& \quad + \frac{1}{6} \left( \|b'(\hat{x}_\lambda)\| \|h_\lambda'''(\hat{x}_\lambda)\| + \frac{1}{4} |b(\hat{x}_\lambda)| \sup_{x \in \mathcal{B}_\delta(\hat{x}_\lambda)} \{ \|h_\lambda^{(4)}(x)\| \} \right) m_4 \\
& \quad \left. + \frac{1}{72} |b(\hat{x}_\lambda)| \|h_\lambda'''(\hat{x}_\lambda)\|^2 m_6 + \lambda^{-1/2} \left( c |b(\hat{x}_\lambda)| + c' \|b'(\hat{x}_\lambda)\| + c'' \sup_{x \in \mathcal{B}_\delta(\hat{x}_\lambda)} \{ \|b''(x)\| \} \right) \right\} \\
& = O(\lambda^{-1}) \tag{1.7}
\end{aligned}$$

where  $c, c', c''$  are positive constants. Finally, regrouping (1.4) and (1.7) we obtain

$$\begin{aligned}
& \left| \frac{\mathcal{I}_\lambda}{(2\pi)^{d/2} \det[h_\lambda''(\hat{x}_\lambda)]^{-1/2} e^{-\lambda h_\lambda(\hat{x}_\lambda)}} - b(\hat{x}_\lambda) \right| \\
& \leq \left| \frac{\int_{E-\mathcal{B}_\delta(\hat{x}_\lambda)} b(x) e^{-\lambda h_\lambda(x)} dx}{(2\pi)^{d/2} \det[h_\lambda''(\hat{x}_\lambda)]^{-1/2} e^{-\lambda h_\lambda(\hat{x}_\lambda)}} \right| + \left| \frac{\int_{\mathcal{B}_\delta(\hat{x}_\lambda)} b(x) e^{-\lambda h_\lambda(x)} dx}{(2\pi)^{d/2} \det[h_\lambda''(\hat{x}_\lambda)]^{-1/2} e^{-\lambda h_\lambda(\hat{x}_\lambda)}} - b(\hat{x}_\lambda) \right| \\
& = O(\lambda^{-1}).
\end{aligned}$$

□

*Proof of theorem 1.2.2.* Proceeding as in the proof of theorem 1.2.1, we have that

$$\left| \frac{\int_{E-\mathcal{B}_\delta(\hat{x}_\lambda)} b(x) e^{-\lambda h_\lambda(x)} dx}{(2\pi)^{d/2} \det[\lambda h_\lambda''(\hat{x}_\lambda)]^{-1/2} e^{-\lambda h_\lambda(\hat{x}_\lambda)}} \right| = O(\lambda^{-k})$$

for all  $k \in \mathbb{N}$ .

Let us consider the other part of the integral. Let  $x \in \mathcal{B}_\delta(\hat{x}_\lambda)$ ,  $u = \lambda^{1/2}(x - \hat{x}_\lambda)$  (so that  $u \in \mathcal{B}_{\lambda^{1/2}\delta}(0)$ ) and  $v = \frac{u}{|u|}$  (so that  $|v| = 1$ ). By expanding  $b$  at  $\hat{x}_\lambda$  to the order 3 and  $h_\lambda$  at  $\hat{x}_\lambda$  to the order 5, we obtain

$$b(x) = b(\hat{x}_\lambda) + \lambda^{-1/2} b'(\hat{x}_\lambda) \cdot u + \frac{\lambda^{-1}}{2} b''(\hat{x}_\lambda) \cdot u^{(2)} + \frac{\lambda^{-3/2}}{2} |u|^3 \int_0^1 b'''(\hat{x}_\lambda + \theta(x - \hat{x}_\lambda)) \cdot v^{(3)} (1 - \theta)^2 d\theta \tag{1.8}$$

and

$$\begin{aligned}
\lambda h_\lambda(x) &= \lambda h_\lambda(\hat{x}_\lambda) + \frac{1}{2} h_\lambda''(\hat{x}_\lambda) \cdot u^{(2)} + \frac{\lambda^{-1/2}}{6} h_\lambda'''(\hat{x}_\lambda) \cdot u^{(3)} + \frac{\lambda^{-1}}{24} h_\lambda^{(4)}(\hat{x}_\lambda) \cdot u^{(4)} \\
&\quad + \frac{\lambda^{-3/2}}{120} h_\lambda^{(5)}(\hat{x}_\lambda) \cdot u^{(5)} + \frac{\lambda^{-2}}{120} |u|^6 \int_0^1 h_\lambda^{(6)}(\hat{x}_\lambda + \theta(x - \hat{x}_\lambda)) \cdot v^{(6)}(1 - \theta)^5 d\theta \\
&= \lambda h_\lambda(\hat{x}_\lambda) + \frac{1}{2} h_\lambda''(\hat{x}_\lambda) \cdot u^{(2)} - R_\lambda(u),
\end{aligned}$$

where

$$\begin{aligned}
R_\lambda(u) &= -\frac{\lambda^{-1/2}}{6} h_\lambda'''(\hat{x}_\lambda) \cdot u^{(3)} - \frac{\lambda^{-1}}{24} h_\lambda^{(4)}(\hat{x}_\lambda) \cdot u^{(4)} - \frac{\lambda^{-3/2}}{120} h_\lambda^{(5)}(\hat{x}_\lambda) \cdot u^{(5)} \\
&\quad - \frac{\lambda^{-2}}{120} |u|^6 \int_0^1 h_\lambda^{(6)}(\hat{x}_\lambda - \theta(x - \hat{x}_\lambda)) \cdot v^{(6)}(1 - \theta)^5 d\theta \\
&= -\frac{\lambda^{-1/2}}{6} h_\lambda'''(\hat{x}_\lambda) \cdot v^{(3)} |u|^3 - \frac{\lambda^{-1}}{24} h_\lambda^{(4)}(\hat{x}_\lambda) \cdot v^{(4)} |u|^4 - \frac{\lambda^{-3/2}}{120} h_\lambda^{(5)}(\hat{x}_\lambda) \cdot v^{(5)} |u|^5 \\
&\quad - \frac{\lambda^{-2}}{120} |u|^6 \int_0^1 h_\lambda^{(6)}(\hat{x}_\lambda - \theta(x - \hat{x}_\lambda)) \cdot v^{(6)}(1 - \theta)^5 d\theta.
\end{aligned}$$

The Taylor expansion of  $\exp$  at 0 to the order 4 yields

$$\exp R_\lambda(u) = 1 + R_\lambda(u) + \frac{1}{2} R_\lambda(u)^2 + \frac{1}{6} R_\lambda(u)^3 + \frac{1}{24} (\exp S_\lambda(u)) R_\lambda(u)^4$$

where  $|S_\lambda(u)| \leq |R_\lambda(u)|$ . We have that

$$|S_\lambda(u)| \leq \left( \frac{\delta M}{6} + \frac{\delta^2 M}{24} + \frac{\delta^3 M}{120} + \frac{\delta^4 M}{720} \right) |u|^2$$

and that

$$R_\lambda(u)^4 \leq \lambda^{-2} |u|^{12} \left( \frac{M}{6} + \frac{\delta M}{24} + \frac{\delta^2 M}{120} + \frac{\delta^3 M}{720} \right)^4.$$

Let us consider the integral of the expansion of  $b$  (1.8) times  $(\exp S_\lambda(u)) R_\lambda(u)^4$  w.r.t. the measure  $\lambda^{-1/2} e^{-\lambda h_\lambda(\hat{x}_\lambda)} \exp\left(-\frac{1}{2} h_\lambda''(\hat{x}_\lambda) \cdot u^{(2)}\right)$  over the domain  $\{u \in \mathcal{B}_{\lambda^{1/2}\delta}(0)\}$ . It

can be bounded by

$$\begin{aligned} & \lambda^{-2} \left( \frac{M}{6} + \frac{\delta M}{24} + \frac{\delta^2 M}{120} + \frac{\delta^3 M}{720} \right)^4 \\ & \times \int_{\mathcal{B}_{\lambda^{1/2}\delta}(0)} \left( |b(\hat{x}_\lambda)| + \lambda^{-1/2} \|b'(\hat{x}_\lambda)\| |u| + \frac{\lambda^{-1}}{2} \|b''(\hat{x}_\lambda)\| |u|^2 + \frac{\lambda^{-3/2}}{2} \sup_{x \in \mathcal{B}_\delta(\hat{x}_\lambda)} \{\|b'''(x)\|\} |u|^3 \right) \\ & \times |u|^{12} \exp \left( -\frac{1}{2} \left( m - \frac{\delta M}{3} - \frac{\delta^2 M}{12} - \frac{\delta^3 M}{60} - \frac{\delta^4 M}{360} \right) |u|^2 \right) du, \end{aligned}$$

which is of order  $O(\lambda^{-2})$  provided  $\delta$  is small enough such that  $\frac{\delta}{3} + \frac{\delta^2}{12} + \frac{\delta^3}{60} + \frac{\delta^4}{360} < \frac{m}{M^d}$  (using the inequality in remark 1.2.5). Moreover, we have that

$$\begin{aligned} & 1 + R_\lambda(u) + \frac{1}{2} R_\lambda(u)^2 + \frac{1}{6} R_\lambda(u)^3 \\ & = 1 - \frac{\lambda^{-1/2}}{6} h_\lambda'''(\hat{x}_\lambda) \cdot u^{(3)} + \lambda^{-1} \left( \frac{1}{72} (h_\lambda'''(\hat{x}_\lambda) \cdot u^{(3)})^2 - \frac{1}{24} h_\lambda^{(4)}(\hat{x}_\lambda) \cdot u^{(4)} \right) \\ & \quad + \lambda^{-3/2} \left( \frac{1}{144} (h_\lambda'''(\hat{x}_\lambda) \cdot u^{(3)}) (h_\lambda^{(4)}(\hat{x}_\lambda) \cdot u^{(4)}) - \frac{1}{120} h_\lambda^{(5)}(\hat{x}_\lambda) \cdot u^{(5)} - \frac{1}{1296} (h_\lambda'''(\hat{x}_\lambda) \cdot u^{(3)})^3 \right) \\ & \quad + \sum_{j=4}^{12} \lambda^{-j/2} T_{\lambda,j}(u), \end{aligned} \tag{1.9}$$

where the  $T_{\lambda,j}(u)$ 's are polynomials in  $|u|$  whose coefficients can be bounded by constants independent of  $\lambda$ . Hence,  $\int_{\mathcal{B}_\delta(\hat{x}_\lambda)} b(x) e^{-\lambda h_\lambda(x)} dx$  is the integral of (1.8) times (1.9) w.r.t. the measure  $\lambda^{-1/2} e^{-\lambda h_\lambda(\hat{x}_\lambda)} \exp \left( -\frac{1}{2} h_\lambda''(\hat{x}_\lambda) \cdot u^{(2)} \right)$  over the domain  $\{u \in \mathcal{B}_{\lambda^{1/2}\delta}(0)\}$ , plus the remainder

$$\lambda^{-1/2} e^{-\lambda h_\lambda(\hat{x}_\lambda)} \int_{\mathcal{B}_{\lambda^{1/2}\delta}(0)} (\exp S_\lambda(u)) R_\lambda(u)^4 \exp \left( -\frac{1}{2} h_\lambda''(\hat{x}_\lambda) \cdot u^{(2)} \right) du$$

which is of order  $O(\lambda^{-2})$ . By developing the product of (1.8) times (1.9) and regrouping the terms by powers of  $\lambda^{-1/2}$ , it can be written  $\sum_{j=0}^{15} \lambda^{-j/2} P_{\lambda,j}(u)$  where  $P_{\lambda,0}(u) \equiv 1$ ,

$$P_{\lambda,1}(u) = b'(\hat{x}_\lambda) \cdot u + b(\hat{x}_\lambda) h_\lambda'''(\hat{x}_\lambda) \cdot u^{(3)},$$

$$\begin{aligned}
P_{\lambda,2}(u) &= \frac{1}{72}b(\hat{x}_\lambda)(h_\lambda'''(\hat{x}_\lambda) \cdot u^{(3)})^2 - \frac{1}{24}b(\hat{x}_\lambda)h_\lambda^{(4)}(\hat{x}_\lambda) \cdot u^{(4)} - \frac{1}{6}(b'(\hat{x}_\lambda) \cdot u)(h_\lambda'''(\hat{x}_\lambda) \cdot u^{(3)}) \\
&\quad + \frac{1}{2}b''(\hat{x}_\lambda) \cdot u^{(2)}
\end{aligned}$$

and

$$\begin{aligned}
P_{\lambda,3}(u) &= \frac{1}{144}(h_\lambda'''(\hat{x}_\lambda) \cdot u^{(3)})(h_\lambda^{(4)}(\hat{x}_\lambda) \cdot u^{(4)}) - \frac{1}{120}h_\lambda^{(5)}(\hat{x}_\lambda) \cdot u^{(5)} - \frac{1}{1296}(h_\lambda'''(\hat{x}_\lambda) \cdot u^{(3)})^3 \\
&\quad + (b'(\hat{x}_\lambda) \cdot u) \left( \frac{1}{72}(h_\lambda'''(\hat{x}_\lambda) \cdot u^{(3)})^2 - \frac{1}{24}h_\lambda^{(4)}(\hat{x}_\lambda) \cdot u^{(4)} \right) + \frac{1}{12}(b''(\hat{x}_\lambda) \cdot u^{(2)})(h_\lambda'''(\hat{x}_\lambda) \cdot u^{(3)}) \\
&\quad + \frac{1}{2}|u|^3 \int_0^1 b'''(\hat{x}_\lambda + \theta(x - \hat{x}_\lambda)) \cdot v^{(3)}(1 - \theta)^2 d\theta.
\end{aligned}$$

We have that

$$\int_{\mathcal{B}_{\lambda^{1/2}\delta}(0)} P_{\lambda,1}(u) \exp\left(-\frac{1}{2}h_\lambda''(\hat{x}_\lambda) \cdot u^{(2)}\right) du = 0$$

and

$$\int_{\mathcal{B}_{\lambda^{1/2}\delta}(0)} P_{\lambda,3}(u) \exp\left(-\frac{1}{2}h_\lambda''(\hat{x}_\lambda) \cdot u^{(2)}\right) du = 0$$

because the integrand is an odd function and the integral is taken over a ball centered at the origin. Regarding  $P_{\lambda,2}(u)$ , we have the decomposition

$$\begin{aligned}
&\int_{\mathcal{B}_{\lambda^{1/2}\delta}(0)} P_{\lambda,2}(u) \exp\left(-\frac{1}{2}h_\lambda''(\hat{x}_\lambda) \cdot u^{(2)}\right) du \\
&= \int_{\mathbb{R}^d} P_{\lambda,2}(u) \exp\left(-\frac{1}{2}h_\lambda''(\hat{x}_\lambda) \cdot u^{(2)}\right) du - \int_{\mathcal{B}_{\lambda^{1/2}\delta}(0)^c} P_{\lambda,2}(u) \exp\left(-\frac{1}{2}h_\lambda''(\hat{x}_\lambda) \cdot u^{(2)}\right) du
\end{aligned}$$

where the second integral can be bounded by a linear combination of integrals of the form  $\int_{\mathcal{B}_{\lambda^{1/2}\delta}(0)^c} |u|^{2p} \exp\left(-\frac{1}{2}h_\lambda''(\hat{x}_\lambda) \cdot u^{(2)}\right) du$  with  $p \in \mathbb{N}$ . Let  $\sigma_\lambda^{\min} > 0$  be the smallest

eigenvalue of  $h''_\lambda(\hat{x}_\lambda)$ , so that  $h''_\lambda(\hat{x}_\lambda) \cdot u^{(2)} \geq \sigma_\lambda^{\min} |u|^2$ . Let  $z = |u|^2$ .

$$\begin{aligned}
& \int_{\mathcal{B}_{\lambda^{1/2}\delta}(0)^c} |u|^{2p} \exp\left(-\frac{1}{2}h''_\lambda(\hat{x}_\lambda) \cdot u^{(2)}\right) du \\
& \leq \int_{\mathcal{B}_{\lambda^{1/2}\delta}(0)^c} |u|^{2p} \exp\left(-\frac{1}{2}\sigma_\lambda^{\min}|u|^2\right) du \\
& = \int_{\Omega} \int_{r=\lambda^{1/2}\delta}^{r=\infty} r^{2p+1} \exp\left(-\frac{1}{2}\sigma_\lambda^{\min}r^2\right) dr d\Omega \\
& = A \int_{\lambda\delta^2}^{\infty} s^p \exp\left(-\frac{1}{2}\sigma_\lambda^{\min}s\right) ds
\end{aligned}$$

where  $\Omega$  is the solid angle in  $\mathbb{R}^d$ ,  $A$  is a constant independent of  $\lambda$ , and  $s = |r|^2$ . A straightforward recursion shows that the last integral above is of order  $O(\lambda^{-k})$  for any  $k > 0$ . Thus,

$$\begin{aligned}
& \lambda^{-1/2} \int_{\mathcal{B}_{\lambda^{1/2}\delta}(0)} P_{\lambda,2}(u) \exp\left(-\frac{1}{2}h''_\lambda(\hat{x}_\lambda) \cdot u^{(2)}\right) du \\
& = \lambda^{-1/2} \int_{\mathbb{R}^d} \left( \frac{1}{72}b(\hat{x}_\lambda)(h'''_\lambda(\hat{x}_\lambda) \cdot u^{(3)})^2 - \frac{1}{24}b(\hat{x}_\lambda)h^{(4)}_\lambda(\hat{x}_\lambda) \cdot u^{(4)} \right. \\
& \quad \left. - \frac{1}{6}(b'(\hat{x}_\lambda) \cdot u)(h'''_\lambda(\hat{x}_\lambda) \cdot u^{(3)}) + \frac{1}{2}b''(\hat{x}_\lambda) \cdot u^{(2)} \right) \exp\left(-\frac{1}{2}h''_\lambda(\hat{x}_\lambda) \cdot u^{(2)}\right) du + O(\lambda^{-2}) \\
& = (2\pi)^{d/2} \det[\lambda h''_\lambda(\hat{x}_\lambda)]^{-1/2} \left( \frac{1}{2}\mathbb{E}[b''(\hat{x}_\lambda) \cdot X^{(2)}] - \frac{1}{6}\mathbb{E}[(b'(\hat{x}_\lambda) \cdot X)(h'''_\lambda(\hat{x}_\lambda) \cdot X^{(3)})] \right. \\
& \quad \left. + \frac{1}{72}b(\hat{x}_\lambda)\mathbb{E}[(h'''_\lambda(\hat{x}_\lambda) \cdot X^{(3)})^2] - \frac{1}{24}b(\hat{x}_\lambda)\mathbb{E}[h^{(4)}_\lambda(\hat{x}_\lambda) \cdot X^{(4)}] \right) + O(\lambda^{-2}),
\end{aligned}$$

where  $X \sim \mathcal{N}(0, h''_\lambda(\hat{x}_\lambda)^{-1})$ . This integral is of order  $O(1)$  as  $\lambda \rightarrow \infty$ . Finally, we have that

$$\sum_{j=4}^{12} \lambda^{-j/2} \int_{\mathcal{B}_{\lambda^{1/2}\delta}(0)} P_{\lambda,j}(u) \exp\left(-\frac{1}{2}h''_\lambda(\hat{x}_\lambda) \cdot u^{(2)}\right) du$$

is of order  $O(\lambda^{-2})$ , which concludes the proof.  $\square$



# Chapter 2

## Issues in importance sampling for Bayesian estimation

We propose in this chapter an analysis of the degradation of importance sampling when used in highly informative or high dimensional nonlinear Bayesian models. Our analysis is based on using the Laplace method to compute a quantity describing the quality of importance sampling. Therefore, we must define the expansion parameter of the Laplace method and build an asymptotic set-up that allows to apply it.

Section 2.1 is devoted to the definition of the model. In section 2.2, we present the importance sampling algorithm and its application to Bayesian estimation. We discuss the issues of high information and high dimensionality in sections 2.3 and 2.4 respectively.

### 2.1 Bayesian model set-up

Let  $X$  be the hidden state of a system taking values in an open subset  $E$  of  $\mathbb{R}^d$ .  $X$  is distributed according to a prior distribution  $\eta$  with density  $q$  with respect to the Lebesgue measure. (Throughout the thesis, densities will be defined w.r.t. the Lebesgue measure.)

#### 2.1.1 Asymptotics

The hidden state is observed through a nonlinear observation function corrupted by an additive Gaussian white noise. We consider in this chapter two types of observation

models, associated with two types of asymptotic regime.

(a) Large observation sample size asymptotics:

$$Y_i = H_i(X) + W_i, \quad i \in \{1, \dots, n\} \quad (n \text{ large}),$$

where  $(Y_1, \dots, Y_n)$  is a sample of observations and  $(W_1, \dots, W_n)$  is a Gaussian white noise. Hence, the likelihood function is

$$g_n(x) \propto \exp \left( -\frac{1}{2} \sum_{i=1}^n |Y_i - H_i(x)|^2 \right)$$

up to a multiplicative constant.

(b) Small observation noise asymptotics:

$$Y = H(X) + \sqrt{\varepsilon}W, \quad \varepsilon > 0 \quad (\varepsilon \text{ small}),$$

where  $Y$  is an observation vector and  $W$  is a Gaussian white noise. Hence, the likelihood function is

$$g^\varepsilon(x) \propto \exp \left( -\frac{1}{2\varepsilon} |Y - H(x)|^2 \right)$$

up to a multiplicative constant.

### 2.1.2 Assumptions for the Laplace method

These two asymptotic regimes are compatible with the Laplace method, when the expansion parameter  $\lambda$  is set to  $\lambda = n$  in case (a) and  $\lambda = 1/\varepsilon$  in case (b).

We thus define the likelihood function parametrized by  $\lambda$  by  $g_\lambda(x) = g_n(x)$  in case (a) and  $g_\lambda(x) = g^\varepsilon(x)$  in case (b). Let  $\ell_\lambda(x) = -\frac{1}{\lambda} \log g_\lambda(x)$  be the contrast function and

$$\hat{x}_\lambda = \operatorname{argmax}_{x \in E} \{g_\lambda(x)\}.$$

be the maximum likelihood estimator (MLE).

To apply the Laplace method, we suppose throughout the chapter that the following assumptions are verified.

There exist  $\Delta > 0$ ,  $m > 0$ ,  $M > 0$ ,  $M' > 0$  and  $\lambda_0 > 0$  such that, for all  $\lambda > \lambda_0$ :

- the MLE  $\hat{x}_\lambda$  exists and is such that, for all  $\delta \in (0, \Delta)$ , there exists  $c_\delta > 0$  independent of  $\lambda$  such that  $\inf_{x \notin \mathcal{B}_\delta(\hat{x}_\lambda)} \{\ell_\lambda(x) - \ell_\lambda(\hat{x}_\lambda)\} \geq c_\delta$ ;
- $\ell_\lambda$  is four-times continuously differentiable over  $\mathcal{B}_\Delta(\hat{x}_\lambda)$  and, for all  $x \in \mathcal{B}_\Delta(\hat{x}_\lambda)$  and  $j \in \{0, \dots, 4\}$ ,  $\|\ell_\lambda^{(j)}(x)\| \leq M$ ;
- $\det[\ell_\lambda''(\hat{x}_\lambda)] \geq m$ ;
- $q$  is twice continuously differentiable over  $\mathcal{B}_\Delta(\hat{x}_\lambda)$  and, for all  $x \in \mathcal{B}_\Delta(\hat{x}_\lambda)$  and  $j \in \{0, 1, 2\}$ ,  $\|q^{(j)}(x)\| \leq M'$ .

These conditions are called Laplace-regularity assumptions in [Kass et al. (1990)] (see chapter 3). They essentially require that:

- the MLE exists and the posterior becomes more concentrated around it as  $\lambda \rightarrow \infty$  (this will be shown in chapter 3) (identifiability condition);
- the constraint function  $\ell_\lambda$  remains smooth around the MAP as  $\lambda \rightarrow \infty$  (local regularity condition).

Let  $\mu_\lambda$  be the posterior distribution and  $p_\lambda$  its density. The Bayes formula yields

$$p_\lambda(x) = \frac{g_\lambda(x)q(x)}{\int_E g_\lambda(x)q(x)dx}. \quad (2.1)$$

## 2.2 Importance sampling for Bayesian estimation

Let  $\phi$  be a function integrable w.r.t. the posterior and let the conditional moment

$$\theta = \mathbb{E}[\phi(X)|Y]$$

be a Bayesian estimator.  $\theta$  describes the state of the system given the available observations and we aim at computing it.

Given the nonlinear nature of the model, importance sampling is a suitable technique to approximate  $\theta$ . The simplest way to apply it is to take the prior  $\eta$  as a sampling

distribution and the likelihood  $g_\lambda$  as a weighting function. The corresponding algorithm is algorithm 4 [Rubinstein (1981)].

---

**Algorithm 4** Importance sampling with prior sampling.

---

- For  $i = 1, \dots, N$ , sample  $\xi^i \sim \eta$ .
  - For  $i = 1, \dots, N$ , compute  $w^i = \frac{g_\lambda(\xi^i)}{\sum_{i=1}^N g_\lambda(\xi^i)}$ .
  - Compute  $\theta^N = \sum_{i=1}^N w^i \phi(\xi^i)$ .
- 

The importance sampling algorithm delivers an approximation  $\theta^N$  of the Bayesian estimator  $\theta$  in the form of a so-called "particle" approximation

$$\theta^N = \sum_{i=1}^N w^i \phi(\xi^i).$$

This approximation is consistent (in the mean squared error sense, e.g.) at the rate  $O(N^{-1})$  [Crisan and Doucet (2002)].  $\theta^N$  is equal to the integral of  $\phi$  w.r.t. to the empirical measure

$$\mu^N = \sum_{i=1}^N w^i \delta_{\xi^i},$$

which is an approximation of the true posterior  $\mu_\lambda$ . Thus, importance sampling consists in approximating the posterior distribution by a weighted sum of Dirac masses centered at the particles.

A feature of importance sampling is that, in certain difficult situations, a very few particles get a positive weight whereas most of the particles get a weight that is numerically evaluated to 0. In this context, a very few particles effectively participate to the approximation of  $\theta$  and most of the computational power is wasted. This phenomenon is called weight degeneracy and it is an important issue in importance sampling and particle filtering [Doucet et al. (2000); Arulampalam et al. (2002)].

The quality of importance sampling for a given model, i.e. the severity of weight degeneracy, can be quantified by the  $\chi^2$ -divergence between the targeted distribution  $\mu_\lambda$  and the sampling distribution  $\eta$ . The  $\chi^2$ -divergence between two probability measures

$\nu$  and  $\nu'$ , such that  $\nu$  is absolutely continuous w.r.t.  $\nu'$ , is defined as

$$\chi^2(\nu, \nu') = \int_E \left( \frac{d\nu}{d\nu'}(x) - 1 \right)^2 \nu'(dx) = \int_E \frac{d\nu}{d\nu'}(x)^2 \nu'(dx) - 1$$

(see for example [Keziou (2003)]). As a divergence between probability measures, the  $\chi^2$ -divergence is nonnegative. Besides,  $\chi^2(\nu, \nu') = 0$  implies  $\nu = \nu'$ . From (2.1), the  $\chi^2$ -divergence  $\chi^2(\mu_\lambda, \eta)$  between the sampling distribution (the prior) and the distribution of interest (the posterior) is well defined because the latter is absolutely continuous w.r.t. the former, and the derivative  $\frac{d\mu_\lambda}{d\eta}$  is equal to  $\frac{g_\lambda}{\int_E g_\lambda \eta}$ .  $\chi^2(\mu_\lambda, \eta) = 0$  means that the prior exactly matches the posterior, which is never true in practice.

The central limit theorem on the importance weights involves the divergence  $\chi^2(\mu_\lambda, \eta)$ . Indeed, we have that

$$\sqrt{N} \left( \frac{\frac{1}{N} \sum_{i=1}^N g_\lambda(\xi^i)}{\int_E g_\lambda(x) \eta(dx)} - 1 \right) \longrightarrow \mathcal{N}(0, \chi^2(\mu_\lambda, \eta))$$

in distribution as  $N \rightarrow \infty$ , thanks to a straightforward application of the classical central limit theorem [Le Gland (2012)]. This means that the more  $\chi^2(\mu_\lambda, \eta)$  is large, the more the variability of the weights is large, asymptotically in the number of particles.

Moreover,  $\chi^2(\mu_\lambda, \eta)$  is involved in the definition of the effective sample size (ESS). The ESS is a criterion introduced by [Kong (1992)] which quantifies the number of particles that are significantly weighted. It is defined here by

$$N_{\text{eff}} = \frac{1}{\sum_{i=1}^N (w^i)^2},$$

and it verifies

$$\frac{N_{\text{eff}}}{N} \longrightarrow \frac{\left( \int_E g_\lambda(x) \eta(dx) \right)^2}{\int_E g_\lambda(x)^2 \eta(dx)} = \frac{1}{\chi^2(\mu_\lambda, \eta) + 1}$$

as  $N \rightarrow \infty$ .  $N_{\text{eff}} = N$  and  $\chi^2(\mu_\lambda, \eta) = 0$  when  $\eta = \mu_\lambda$ , i.e. when the prior matches the posterior.

Thus, we consider  $\chi^2(\mu_\lambda, \eta)$  as a measure of weight degeneracy. In the sequel, we analyze its behaviour in situations that are known to be difficult in importance sampling.

## 2.3 High information

A typical situation where importance sampling fails, i.e. weight degeneracy is severe, is when the model is very informative. This phenomenon might seem rather paradoxical: one could expect that a precise model improves Monte Carlo approximation. However, what actually happens in this case is that the likelihood function is peaked around its maximum and takes significantly nonzero values only in a small area of the state space. Yet the particles are sampled from the prior distribution, without taking the observations into account. Consequently, the probability that a particle lies in an area where the likelihood takes significant values can be very small. Most of the particles are then numerically weighted at zero and weight degeneracy occurs.

We propose here a more formal illustration of this phenomenon. On the one hand, as  $\lambda \rightarrow \infty$ , the model becomes more informative (in the Fisher sense, e.g.). On the other hand, the  $\chi^2$ -divergence between  $\mu_\lambda$  and  $\eta$  becomes larger as  $\lambda \rightarrow \infty$  (hence the asymptotic variance of the weights goes to infinity and the ESS goes to 0).

The observed information matrix of the state  $x \in E$  depends on the observation sample (unlike the Fisher information matrix) and is defined as

$$J_\lambda(x) = -(\log g_\lambda)''(x) - (\log q)''(x) = \lambda \ell_\lambda''(x) - (\log q)''(x)$$

(see for example [Efron and Hinkley (1978)]). From the assumptions in section 2.1.2, we have that  $\det[\ell_\lambda''(\hat{x}_\lambda)] \geq m$  and  $|q(\hat{x}_\lambda)| \leq M$  when  $\lambda > \lambda_0$ . Hence, the norm of  $J_\lambda(\hat{x}_\lambda)$ , the observed information evaluated at the MAP  $\hat{x}_\lambda$ , increases at speed  $\lambda$  as  $\lambda \rightarrow \infty$ .

Moreover, the  $\chi^2$ -divergence between  $\mu_\lambda$  and  $\eta$  is

$$\chi^2(\mu_\lambda, \eta) = I_\lambda - 1$$

where

$$I_\lambda = \frac{\int_E g_\lambda(x)^2 q(x) dx}{\left(\int_E g_\lambda(x) q(x) dx\right)^2}.$$

Applying the Laplace method (theorem 1.2.1) to the numerator and the denominator

of  $I_\lambda$  yields

$$\begin{aligned} \frac{\int_E g_\lambda(x)^2 q(x) dx}{\left(\int_E g_\lambda(x) q(x) dx\right)^2} &= \frac{(2\pi)^{d/2} \det[-2(\log g_\lambda)''(\hat{x}_\lambda)]^{-1/2}}{(2\pi)^d \det[-(\log g_\lambda)''(\hat{x}_\lambda)]^{-1}} \cdot \frac{q(\hat{x}_\lambda) + O(\lambda^{-1})}{\{q(\hat{x}_\lambda) + O(\lambda^{-1})\}^2} \\ &= (4\pi)^{-d/2} \det[\lambda \ell''_\lambda(\hat{x}_\lambda)]^{1/2} \cdot \left\{ \frac{1}{q(\hat{x}_\lambda)} + O(\lambda^{-1}) \right\} \\ &= \left( \frac{\lambda}{4\pi} \right)^{d/2} \det[\ell''_\lambda(\hat{x}_\lambda)]^{1/2} \cdot \left\{ \frac{1}{q(\hat{x}_\lambda)} + O(\lambda^{-1}) \right\}. \end{aligned}$$

Hence,  $\chi^2(\mu_\lambda, \eta)$  increases at speed  $\lambda^{d/2}$ , i.e.  $N_{\text{eff}}/N$  decreases at speed  $\lambda^{-d/2}$ , when  $\lambda \rightarrow \infty$ . (Recall that  $\lambda$  is equivalently interpreted as the observation sample size  $n$  or as the inverse of the observation noise variance  $1/\varepsilon$ .)

## 2.4 High dimension

The problem of the poor behaviour of importance sampling, hence particle filtering, when the state dimension is large has been addressed by many authors. Some examples are: [Daum et al. (2003)], [Bengtsson et al. (2003)], [Chorin and Krause (2004)], [Berliner and Wickle (2007)], [Bengtsson et al. (2008)], [Bickel et al. (2008)], [Snyder et al. (2008)], [van Leeuwen (2009)], [Bui Quang et al. (2010)], [Beskos et al. (2011)]. In particular, it raises high interest in the field of data assimilation in geosciences.

The analysis we lead in this chapter is based on the Laplace method applied on a quantity that describes the quality of importance sampling. However, the remainder  $O(\lambda^{-1})$  in the derivations above depends on the state space dimension  $d$ , that is it should be written  $O_d(\lambda^{-1})$ . Unfortunately, in [Shun and McCullagh (1995)], the authors argue that the Laplace method cannot be trusted when the dimension of the integral increases at the same rate as the expansion parameter  $\lambda$ . This means that  $O_d(\lambda^{-1})$  might increase in absolute value when  $d \rightarrow \infty$ . Despite that, Shun and McCullagh also argue that the Laplace method is reliable when  $d$  and  $\lambda$  increase at rates such that  $d = o(\lambda^{1/3})$ .

In this section, we study the asymptotic behaviour of  $I_\lambda$  as  $d \rightarrow \infty$  based on its approximation by the Laplace method. Thus, we suppose that  $\lambda$  goes to  $\infty$  sufficiently fast so that this approximation remains valid. We also suppose that  $q(\hat{x}_\lambda)$ , the prior density evaluated at the MAP, is bounded uniformly in  $d$ .

We have that

$$I_\lambda \sim \det \left[ \frac{-(\log g_\lambda)''(\hat{x}_\lambda)}{4\pi} \right]^{1/2} \frac{1}{q(\hat{x}_\lambda)}$$

when  $\lambda \rightarrow \infty$  (note that here  $\sim$  denotes asymptotic equivalence rather than equality in distribution). Hence, the behaviour of the approximation of  $I_\lambda$  as  $d \rightarrow \infty$  is determined by the volume of the ellipsoid associated with the symmetric positive matrix  $-(\log g_\lambda)''(\hat{x}_\lambda)$ . This volume somehow quantifies the amount of information brought by the data. Thus, if the amount of information coming from the observation, quantified by  $-(\log g_\lambda)''(\hat{x}_\lambda)$ , increases when  $d \rightarrow \infty$ , then  $I_\lambda$  increases as well. In particular, letting  $\sigma[\ell''_\lambda(\hat{x}_\lambda)]$  be the spectrum of the symmetric positive  $d \times d$  matrix  $\ell''_\lambda(\hat{x}_\lambda)$ , we have that

$$\begin{aligned} \det \left[ \frac{-(\log g_\lambda)''(\hat{x}_\lambda)}{4\pi} \right]^{1/2} \frac{1}{q(\hat{x}_\lambda)} &\geq \left( \frac{\lambda \min \{\sigma[\ell''_\lambda(\hat{x}_\lambda)]\}}{4\pi} \right)^{d/2} \frac{1}{q(\hat{x}_\lambda)} \\ &\geq \left( \frac{d \min \{\sigma[\ell''_\lambda(\hat{x}_\lambda)]\}}{4\pi} \right)^{d/2} \frac{1}{q(\hat{x}_\lambda)}, \end{aligned}$$

so that the approximation of  $I_\lambda$  increases when  $d \rightarrow \infty$  as soon as

$$\frac{d^\alpha}{\min \{\sigma[\ell''_\lambda(\hat{x}_\lambda)]\}} = O(1)$$

where  $\alpha > -1$ . This is a sufficient condition which implies that  $\chi^2(\mu_\lambda, \eta)$  increases at speed  $\left(\frac{d^{1+\alpha}}{4\pi}\right)^{d/2}$ , i.e.  $N_{\text{eff}}/N$  decreases at speed  $\left(\frac{d^{1+\alpha}}{4\pi}\right)^{-d/2}$ , as  $d \rightarrow \infty$ . This condition is fulfilled when  $\min \{\sigma[\ell''_\lambda(\hat{x}_\lambda)]\}$  is bounded from below uniformly in  $d$ , e.g.

## Conclusion

We have proposed in this chapter an asymptotic framework to study the behaviour of importance sampling for Bayesian inference in difficult situations. From our analysis, we can conclude that importance sampling approximation is often poor in high dimensional or high information models. Our derivations are based on the Laplace method. They are valid under assumptions that can essentially be interpreted as identifiability conditions, and when the observation sample size is large or the observation noise intensity is small.

Taking the observation into account in the sampling distribution, rather than sampling from the prior, allows to improve importance sampling in these situations (especially in the high information case). We propose in this thesis a way to do so in particle filtering algorithms (see chapters 6 and 7).

## Part II

# The Laplace method in Bayesian statistics



# Chapter 3

## The Laplace method in static models with large observation sample size

In Bayesian statistics, the Laplace method is used to compute expectations of nonlinear functions w.r.t. the posterior distribution [Tierney and Kadane (1986); Tierney et al. (1989)] or marginal posterior densities [Tierney and Kadane (1986)]. The asymptotic regime in which it yields consistent approximations is classically the observation sample size asymptotics [Kass et al. (1988)]. That is, the number of observations  $n$  plays the role of the parameter  $\lambda$  in theorems 1.2.1 and 1.2.2 in chapter 1.

In section 3.1, we define the Laplace–regularity assumption and give some properties of the model implied by this assumption (which are somehow identifiability properties). In section 3.2, we provide the approximation formulas established by Tierney, Kass and Kadane under the Laplace–regularity assumption [Tierney et al. (1989); Kass et al. (1990, 1991)]. In particular, we generalize to the multidimensional case their approximation formulas for the posterior expectation and variance.

In this chapter, we use the differential calculus principles presented in section 1.2.1 of chapter 1.

### 3.1 Model set-up, Laplace–regularity

Let  $Y_{1:n} := \{Y_1, \dots, Y_n\}$  be a sequence of observations and  $X$  be the (static) hidden state of interest taking values in an open subset  $E$  of  $\mathbb{R}^d$ . The likelihood function  $g_n$  is the function that associates to  $x$  the density of the conditional probability of  $Y_{1:n}$  given

$X = x$ ,

$$x \in E \longmapsto g_n(x, Y_{1:n}) \in \mathbb{R}^+.$$

Let  $q$  be the prior density of  $X$ . Let  $\mu_n$  be the posterior, i.e. the conditional probability of  $X$  given  $Y_{1:n}$ , and  $p_n$  its density,

$$p_n(x, Y_{1:n}) = \frac{g_n(x, Y_{1:n})q(x)}{\int_E g_n(x, Y_{1:n})q(x)dx}$$

according to Bayes' rule. To alleviate notation, we do not denote the dependency of  $g_n$  and  $p_n$  on the observations sequence in the sequel. The dependency of any quantity on  $Y_{1:n}$  will be denoted by the subscript  $n$ .

The contrast function is defined as

$$\ell_n(x) = -\frac{1}{n} \log g_n(x).$$

Kass, Tierney and Kadane [Kass et al. (1990)] introduce the Laplace–regularity assumptions, to characterize the Bayesian models on which the Laplace method can be applied. A sequence of contrast functions  $\{\ell_n\}_{n \geq 1}$  is said to be  $(n_0, \Delta, m, M)$ –Laplace–regular, or simply Laplace–regular, if it verifies the following assumption.

**Assumption  $L_n$ : Laplace–regularity.** There exist  $n_0 > 0$ ,  $\Delta > 0$ ,  $m > 0$  and  $M > 0$  such that, for all  $n \geq n_0$ :

- there exists  $\tilde{x}_n$  such that, for all  $\delta \in (0, \Delta)$ , there exists  $c_\delta > 0$  independent of  $n$  such that  $\inf_{x \notin \mathcal{B}_\delta(\tilde{x}_n)} \{\ell_n(x) - \ell_n(\tilde{x}_n)\} \geq c_\delta$ ;
- $\ell_n$  is six–times continuously differentiable over  $\mathcal{B}_\Delta(\tilde{x}_n)$  and for all  $x \in \mathcal{B}_\Delta(\tilde{x}_n)$ , for all  $j \in \{0, \dots, 6\}$ ,  $|\ell_n^{(j)}(x)| \leq M$ ;
- $\det [\ell_n''(\tilde{x}_n)] \geq m$ .

This corresponds to the conditions (ii), (iii) and (iv) of theorem 1.2.2.

Under the Laplace–regularity assumptions, the maximum likelihood estimator (MLE)

$$\tilde{x}_n = \operatorname{argmin}_{x \in E} \{\ell_n(x)\} = \operatorname{argmax}_{x \in E} \{g_n(x)\}$$

exists and is unique for all  $n \geq n_0$ . The maximum a posteriori estimator (MAP) is defined by

$$\hat{x}_n = \operatorname{argmax}_{x \in E} \{p_n(x)\} = \operatorname{argmax}_{x \in E} \{g_n(x)q(x)\}.$$

**Proposition 3.1.1.** *Suppose that  $\{\ell_n\}_{n \geq 0}$  is  $(n_0, \Delta, m, M)$ -Laplace-regular. Suppose that  $\log q$  is continuously differentiable over  $\mathcal{B}_\Delta(\tilde{x}_n)$  for all  $n \geq n_0$ . Suppose that  $|(\log q)(x)| \leq M$  and  $\|(\log q)'(x)\| \leq M$  for all  $x \in \mathcal{B}_\Delta(\tilde{x}_n)$ . Then, the MAP exists for all  $n \geq n_0$  and*

$$|\hat{x}_n - \tilde{x}_n| = O(n^{-1})$$

as  $n \rightarrow \infty$ .

*Proof.* Let  $h_n(x) = \ell_n(x) - \frac{1}{n} \log q(x)$  and let  $x \in \mathcal{B}_\Delta(\tilde{x}_n)$ . Then,

$$h_n(x) \geq \ell_n(\tilde{x}_n) - \frac{1}{n}M$$

so that  $h_n$  admits a minimum in the ball  $\mathcal{B}_\Delta(\tilde{x}_n)$  at, say,  $\hat{x}_n$ . Besides,

$$\begin{aligned} 0 &\leq \ell_n(\hat{x}_n) - \ell_n(\tilde{x}_n) \\ &\leq \ell_n(\hat{x}_n) - h_n(\hat{x}_n) + h_n(\tilde{x}_n) - \ell_n(\tilde{x}_n) \\ &= \frac{1}{n}(\log q(\hat{x}_n) - \log q(\tilde{x}_n)) \\ &\leq \frac{2M}{n}. \end{aligned}$$

Let  $\delta \in (0, \Delta)$ . Because of Laplace-regularity, there exists  $c_\delta$  such that  $|x - \tilde{x}_n| \geq \delta$  implies  $\ell_n(x) - \ell_n(\tilde{x}_n) \geq c_\delta$ . Thanks to the above inequality, we can always choose  $n$  large enough so that  $\ell_n(\hat{x}_n) - \ell_n(\tilde{x}_n) \leq c_\delta$ , which implies  $|\hat{x}_n - \tilde{x}_n| \leq \delta$ . This holds for any  $\delta > 0$ , hence we obtain  $|\hat{x}_n - \tilde{x}_n| \xrightarrow[n \rightarrow \infty]{} 0$ .

From the definition of  $\hat{x}_n$  and  $\tilde{x}_n$ , for all  $u \in E$  we have that

$$\begin{aligned}
0 &= h'_n(\hat{x}_n) \cdot u \\
&= \ell'_n(\hat{x}_n) \cdot u - \frac{1}{n}(\log q)'(\hat{x}_n) \cdot u \\
&= \left( \ell'_n(\tilde{x}_n) + \ell''_n(\tilde{x}_n) \cdot (\hat{x}_n - \tilde{x}_n) + \int_0^1 \ell'''_n(\tilde{x}_n + \theta(\hat{x}_n - \tilde{x}_n)) \cdot (\hat{x}_n - \tilde{x}_n)^{(2)}(1 - \theta) d\theta \right) \cdot u \\
&\quad - \frac{1}{n}(\log q)'(\hat{x}_n) \cdot u \\
&= \ell''_n(\tilde{x}_n) \cdot (\hat{x}_n - \tilde{x}_n, u) + \int_0^1 \ell'''_n(\tilde{x}_n + \theta(\hat{x}_n - \tilde{x}_n)) \cdot (\hat{x}_n - \tilde{x}_n, \hat{x}_n - \tilde{x}_n, u)(1 - \theta) d\theta \\
&\quad - \frac{1}{n}(\log q)'(\hat{x}_n) \cdot u.
\end{aligned}$$

In particular, for  $u = \hat{x}_n - \tilde{x}_n$ ,

$$0 \leq \ell''_n(\tilde{x}_n) \cdot (\hat{x}_n - \tilde{x}_n)^{(2)} \leq \frac{1}{2} \sup_{|x - \tilde{x}_n| \leq |\hat{x}_n - \tilde{x}_n|} \{\|\ell'''_n(x)\|\} |\hat{x}_n - \tilde{x}_n|^3 + \frac{M}{n} |\hat{x}_n - \tilde{x}_n|.$$

Besides,  $\ell''_n(\tilde{x}_n) \cdot (\hat{x}_n - \tilde{x}_n)^{(2)} \geq \frac{m}{M^{d-1}} |\hat{x}_n - \tilde{x}_n|^2$  according to remark 1.2.5 in chapter 1, and  $|\hat{x}_n - \tilde{x}_n| \leq \Delta$  when  $n$  is sufficiently large so that  $\sup_{|x - \tilde{x}_n| \leq |\hat{x}_n - \tilde{x}_n|} \{\|\ell'''_n(x)\|\} \leq \sup_{x \in \mathcal{B}_\Delta(\tilde{x}_n)} \{\|\ell'''_n(x)\|\} \leq M$ . Then,

$$\frac{m}{M^{d-1}} |\hat{x}_n - \tilde{x}_n| \leq \frac{M}{2} |\hat{x}_n - \tilde{x}_n|^2 + \frac{M}{n},$$

so that

$$\frac{m}{M^{d-1}} |\hat{x}_n - \tilde{x}_n| \left( 1 - \frac{M^d}{2m} |\hat{x}_n - \tilde{x}_n| \right) \leq \frac{M}{n}$$

for all sufficiently large  $n$ . Hence,  $|\hat{x}_n - \tilde{x}_n| = O(n^{-1})$ . □

**Proposition 3.1.2.** *Suppose that  $\{\ell_n\}_{n \geq 0}$  is Laplace-regular, that  $\int_E g_n(x) q(x) dx < \infty$  for all sufficiently large  $n$ , and that  $\liminf_{n \rightarrow \infty} \int_{\mathcal{B}_r(\tilde{x}_n)} q(x) dx > 0$  for all  $r > 0$ . Then, for all bounded continuous function  $\phi$ ,*

$$\left| \int_E \phi(x) \mu_n(dx) - \int_E \phi(x) \delta_{\tilde{x}_n}(dx) \right| \longrightarrow 0$$

as  $n \rightarrow \infty$ .

*Proof.* Let  $\delta > 0$ . We have that

$$\mu_n(\mathcal{B}_\delta(\tilde{x}_n)^c) = \frac{\int_{\mathcal{B}_\delta(\tilde{x}_n)^c} g_n(x)q(x)dx}{\int_E g_n(x)q(x)dx}. \quad (3.1)$$

We can bound from above the numerator of (3.1) as

$$\begin{aligned} \int_{\mathcal{B}_\delta(\tilde{x}_n)^c} g_n(x)q(x)dx &= g_n(\tilde{x}_n) \int_{\mathcal{B}_\delta(\tilde{x}_n)^c} \exp(-n(\ell_n(x) - \ell_n(\tilde{x}_n))) q(x)dx \\ &\leq g_n(\tilde{x}_n) e^{-nc_\delta} \int_{\mathcal{B}_\delta(\tilde{x}_n)^c} q(x)dx, \end{aligned}$$

for some  $c_\delta > 0$  whose existence is insured by the Laplace-regularity of  $\{\ell_n\}_{n \geq 0}$ . Let now  $A_{\delta,n}$  be the subset of  $E$  defined by

$$A_{\delta,n} = \left\{ x \in E : \ell_n(x) - \ell_n(\tilde{x}_n) \leq \frac{c_\delta}{2} \right\}.$$

The denominator of (3.1) can be bounded from below as

$$\begin{aligned} \int_E g_n(x)q(x)dx &\geq \int_{A_{\delta,n}} g_n(x)q(x)dx \\ &= g_n(\tilde{x}_n) \int_{A_{\delta,n}} \exp(-n(\ell_n(x) - \ell_n(\tilde{x}_n))) q(x)dx \\ &\geq g_n(\tilde{x}_n) e^{-\frac{nc_\delta}{2}} \int_{A_{\delta,n}} q(x)dx. \end{aligned}$$

Consequently, we have

$$\mu_n(\mathcal{B}_\delta(\tilde{x}_n)^c) \leq e^{-\frac{nc_\delta}{2}} \frac{\int_{\mathcal{B}_\delta(\tilde{x}_n)^c} q(x)dx}{\int_{A_{\delta,n}} q(x)dx}.$$

Besides, for all  $x \in E$ ,

$$0 \leq \ell_n(x) - \ell_n(\tilde{x}_n) = \int_0^1 \ell_n''(\tilde{x}_n + \theta(x - \tilde{x}_n)) \cdot (x - \tilde{x}_n)^{(2)}(1 - \theta)d\theta \leq \frac{M}{2}|x - \tilde{x}_n|^2,$$

so that  $|x - \tilde{x}_n| \leq \sqrt{\frac{c_\delta}{M}}$  implies  $\ell_n(x) - \ell_n(\tilde{x}_n) \leq \frac{c_\delta}{2}$ . Hence  $\mathcal{B}_{\sqrt{\frac{c_\delta}{M}}}(\tilde{x}_n) \subset A_{\delta,n}$  and thus  $\int_{A_{\delta,n}} q(x)dx \geq \int_{\mathcal{B}_{\sqrt{\frac{c_\delta}{M}}}(\tilde{x}_n)} q(x)dx$ . Since  $\liminf_{n \rightarrow \infty} \int_{\mathcal{B}_{\sqrt{\frac{c_\delta}{M}}}(\tilde{x}_n)} q(x)dx > 0$ , we have that  $\mu_n(\mathcal{B}_\delta(\tilde{x}_n)^c) \xrightarrow{n \rightarrow \infty} 0$  for all  $\delta > 0$ , which yields the result.  $\square$

## 3.2 Approximation of moments

In order to derive approximation formulas for posterior moments, we make the following assumption so that theorem 1.2.2 is applicable.

**Assumption  $L'_n$ .** There exist  $n_0 > 0$ ,  $\Delta > 0$ ,  $m > 0$ ,  $p > 0$ ,  $K > 0$  and  $M > 0$  such that, for all  $n \geq n_0$ :

- $\{\ell_n\}_{n \geq 0}$  is  $(n_0, \Delta, m, M)$ -Laplace-regular;
- for all  $n \geq n_0$ ,  $\int_E (g_n(x)q(x))^{n_0/n} dx \leq Kn^p$ ;
- for all  $n \geq n_0$ ,  $\log q \in C^6(\mathcal{B}_\Delta(\tilde{x}_n))$  and  $(\log q)^{(j)}$  is bounded over  $\mathcal{B}_\Delta(\tilde{x}_n)$  uniformly in  $n$  for all  $j \in \{0, \dots, 6\}$ ;
- $\sup_{x \in E} \{q(x)\} < \infty$ .

The second item in assumption  $L'_n$  above corresponds to condition (i) in theorem 1.2.2. It is verified in particular when  $q$  is a Gaussian density. Indeed, in this case we have that

$$\begin{aligned}
 \int_E (g_n(x)q(x))^{n_0/n} dx &= \int_E e^{-n_0 \ell_n(x)} q(x)^{n_0/n} dx \\
 &= e^{-n_0 \ell_n(\tilde{x}_n)} \int_E e^{-n_0 (\ell_n(x) - \ell_n(\tilde{x}_n))} q(x)^{n_0/n} dx \\
 &\leq e^{n_0 M} \int_E q(x)^{n_0/n} dx \\
 &\leq e^{n_0 M} \left( \frac{n}{n_0} \right)^{d/2}
 \end{aligned}$$

when  $n \geq n_0$ . Weaker conditions on  $q$  can also be derived from lemma 7 in [Holmström and Klemelä (1992)].

Firstly, we establish in the lemma below that a sequence of Laplace-regular functions remains Laplace-regular when it is perturbed by a smooth function of order  $1/n$ .

**Lemma 3.2.1.** *Let  $\{f_n : E \rightarrow \mathbb{R}\}_{n \geq 1}$  be a Laplace-regular collection of functions admitting minima at  $\{x_n\}_{n \geq 1}$ . Let  $k : E \mapsto \mathbb{R}$  such that  $k \in C^6(\mathcal{B}_\Delta(x_n))$ . Suppose that*

$\inf_{x \in E} \{k(x)\} > -\infty$  and  $\|k^{(j)}(x)\| \leq M$  for all  $x \in \mathcal{B}_\Delta(x_n)$  and  $j \in \{0, \dots, 6\}$ . Let  $h_n(x) = f_n(x) + \frac{1}{n}k(x)$ . Then,  $\{h_n\}_{n \geq 0}$  is Laplace-regular.

*Proof.* Let  $x_n^* = \operatorname{argmin}_{x \in E} \{h_n(x)\}$ . Proceeding as in the proof of proposition 3.1.1, we have that  $x_n^*$  exists as soon as  $n$  is large enough and  $|x_n^* - x_n| = O(n^{-1})$ . In particular, this implies that  $k$  and  $h_n$  and their derivatives are bounded uniformly in  $n$  in a neighbourhood of  $x_n^*$ .

The Laplace-regularity of  $\{f_n\}_{n \geq 0}$  implies that

$$f_n(x_n^*) = f_n(x_n) + \int_0^1 f'_n(x_n + \theta(x_n^* - x_n)) \cdot (x_n^* - x_n)(1 - \theta) d\theta \leq f_n(x_n) + M|x_n^* - x_n|,$$

so that  $|f_n(x_n^*) - f_n(x_n)| \xrightarrow{n \rightarrow \infty} 0$ .

Let  $x \notin \mathcal{B}_\delta(x_n^*)$ . Then,  $|x - x_n| \geq |x - x_n^*| - |x_n^* - x_n| \geq \delta - |x_n^* - x_n| \geq \frac{\delta}{2} = \delta'$  when  $n$  is large enough. Therefore, there exists  $c_{\delta'}$  such that  $f_n(x) - f_n(x_n) \geq c_{\delta'}$ . Thus,

$$\begin{aligned} h_n(x) - h_n(x_n^*) &\geq (h_n(x) - f_n(x)) + (f_n(x) - f_n(x_n)) + (f_n(x_n) - f_n(x_n^*)) \\ &\quad - (h_n(x_n^*) - f_n(x_n^*)) \\ &\geq \frac{1}{n} \inf_{x \in E} \{k(x)\} + c_{\delta'} + (f_n(x_n) - f_n(x_n^*)) - \frac{1}{n}k(x_n^*) \\ &\geq \frac{c_{\delta'}}{2} \end{aligned}$$

when  $n$  is large enough.

Besides, for all  $u \in \mathbb{R}^d$ ,

$$\begin{aligned} h_n''(x_n^*) \cdot u^{(2)} &= f_n''(x_n^*) \cdot u^{(2)} + \frac{1}{n}k''(x_n^*) \cdot u^{(2)} \\ &= f_n''(x_n) \cdot u^{(2)} + \int_0^1 f_n'''(x_n + \theta(x_n^* - x_n)) \cdot (x_n^* - x_n, u, u)(1 - \theta) d\theta \\ &\quad + \frac{1}{n}k''(x_n^*) \cdot u^{(2)}. \end{aligned}$$

The Laplace-regularity of  $f_n$  and the assumptions on  $k$  imply that  $h_n''(x_n^*) \cdot u^{(2)} - f_n''(x_n^*) \cdot u^{(2)} = O(n^{-1})$ , so that  $\det[h_n''(x_n^*)]$  is bounded away from 0 uniformly in  $n$ . Hence, we have verified that  $\{h_n\}_{n \geq 1}$  is Laplace-regular.  $\square$

We now derive approximation formulas for the conditional expectation of a positive function given the observations. The theorem below is established in [Tierney and Kadane (1986)].

**Theorem 3.2.2** (Fully exponential Laplace approximation). *Suppose that assumption  $\mathbf{L}'_n$  is verified. Let  $\phi \in L^1(\mu_n)$  be a positive function such that the function  $x \mapsto \phi(x)q(x)$  is bounded over  $E$ . Suppose that  $\log \phi \in C^6(\mathcal{B}_\Delta(\tilde{x}_n))$  and  $(\log \phi)^{(j)}(x)$  is bounded over  $\mathcal{B}_\Delta(\tilde{x}_n)$  uniformly in  $n$  for all  $j \in \{0, \dots, 6\}$ . Suppose that  $\int_E (\phi(x)g_n(x)q(x))^{n_0/n} dx \leq K^\phi n^p$  for some  $p > 0$  and  $K^\phi > 0$  when  $n \geq n_0$ . Then,*

$$\mathbb{E}[\phi(X)|Y_{1:n}] = \frac{\det[J_n(\hat{x}_n)]^{1/2} e^{-nh_n^\phi(\hat{x}_n^\phi)}}{\det[J_n^\phi(\hat{x}_n^\phi)]^{1/2} e^{-nh_n(\hat{x}_n)}} \{1 + O(n^{-2})\}$$

where

$$\begin{aligned} h_n(x) &= -\frac{1}{n}(\log g_n(x) + \log q(x)), \quad h_n^\phi(x) = -\frac{1}{n}(\log \phi(x) + \log g_n(x) + \log q(x)), \\ \hat{x}_n &= \operatorname{argmin}_{x \in E} \{h_n(x)\}, \quad \hat{x}_n^\phi = \operatorname{argmin}_{x \in E} \{h_n^\phi(x)\}, \\ J_n(x) &= nh_n''(x) \quad \text{and} \quad J_n^\phi(x) = n(h_n^\phi)''(x). \end{aligned}$$

The approximation in theorem 3.2.2 is called *fully exponential Laplace approximation* of  $\mathbb{E}[\phi(X)|Y_{1:n}]$ . It is obtained by applying the Laplace method (theorem 1.2.2) to the numerator and the denominator of the ratio of integrals defining  $\mathbb{E}[\phi(X)|Y_{1:n}]$ . The proof below is a detailed version of the one available in the appendix of [Tierney and Kadane (1986)]. A proof can also be found in [Schervish (1995)].

*Proof.* Thanks to lemma 3.2.1, we have that  $h_n$  and  $h_n^\phi$  are Laplace-regular. All the conditions of theorem 1.2.2 are fulfilled, so that we get

$$\begin{aligned} \mathbb{E}[\phi(X)|Y_{1:n}] &= \frac{\int_E \phi(x)g_n(x)q(x)dx}{\int_E g_n(x)q(x)dx} \\ &= \frac{\det[J_n(\hat{x}_n)]^{1/2} e^{-nh_n^\phi(\hat{x}_n^\phi)}}{\det[J_n^\phi(\hat{x}_n^\phi)]^{1/2} e^{-nh_n(\hat{x}_n)}} \times \frac{1 + n^{-1}A_n^\phi + O(n^{-2})}{1 + n^{-1}A_n + O(n^{-2})} \\ &= \frac{\det[J_n(\hat{x}_n)]^{1/2} e^{-nh_n^\phi(\hat{x}_n^\phi)}}{\det[J_n^\phi(\hat{x}_n^\phi)]^{1/2} e^{-nh_n(\hat{x}_n)}} \times \{1 + n^{-1}(A_n^\phi - A_n) + O(n^{-2})\} \end{aligned}$$

where

$$A_n = -\frac{1}{24}\mathbb{E}[h_n^{(4)}(\hat{x}_n) \cdot X^{(4)}] + \frac{1}{72}\mathbb{E}[(h_n'''(\hat{x}_n) \cdot X^{(3)})^2]$$

and

$$A_n^\phi = -\frac{1}{24}\mathbb{E}[(h_n^\phi)^{(4)}(\hat{x}_n) \cdot (X^\phi)^{(4)}] + \frac{1}{72}\mathbb{E}[(h_n^\phi)'''(\hat{x}_n) \cdot (X^\phi)^{(3)}]^2,$$

with  $X \sim \mathcal{N}(0, h_n''(\hat{x}_n)^{-1})$  and  $X^\phi \sim \mathcal{N}(0, (h_n^\phi)''(\hat{x}_n^\phi)^{-1})$  respectively.

Let  $j \in \{1, \dots, 4\}$  and let  $u \in E$ . Then,

$$\begin{aligned} & (h_n^\phi)^{(j)}(\hat{x}_n^\phi) \cdot u^{(j)} - h_n^{(j)}(\hat{x}_n) \cdot u^{(j)} \\ &= -\frac{1}{n}((\log \phi)^{(j)}(\hat{x}_n^\phi) + (\log q)^{(j)}(\hat{x}_n^\phi) - (\log q)^{(j)}(\hat{x}_n)) \cdot u^{(j)} + \ell_n^{(j)}(\hat{x}_n) \cdot u^{(j)} - \ell_n^{(j)}(\hat{x}_n^\phi) \cdot u^{(j)}. \end{aligned}$$

With arguments similar than in the proof of proposition 3.1.1, we have that  $|\hat{x}_n^\phi - \hat{x}_n| = O(n^{-1})$ . In particular, when  $n$  is large enough,  $\hat{x}_n^\phi$  and  $\hat{x}_n$  lie in  $\mathcal{B}_\Delta(\tilde{x}_n)$ . Then,

$$\ell_n^{(j)}(\hat{x}_n^\phi) \cdot u^{(j)} = (\ell_n^{(j)}(\hat{x}_n) + \ell_n^{(j+1)}(\hat{x}_n) \cdot (\hat{x}_n^\phi - \hat{x}_n)) \cdot u^{(j)} + O(|\hat{x}_n^\phi - \hat{x}_n|^2)$$

because Laplace-regularity implies that  $\|\ell_n^{(j+2)}(x)\|$  is bounded uniformly in  $n$  over  $\mathcal{B}_\Delta(\tilde{x}_n)$ . Hence,

$$(h_n^\phi)^{(j)}(\hat{x}_n^\phi) \cdot u^{(j)} - h_n^{(j)}(\hat{x}_n) \cdot u^{(j)} = \ell_n^{(j+1)}(\hat{x}_n) \cdot (\hat{x}_n^\phi - \hat{x}_n, u, \dots, u) + O(n^{-1}).$$

Since  $\|\ell_n^{(j+1)}(\hat{x}_n)\| = O(1)$ , then  $(h_n^\phi)^{(j)}(\hat{x}_n^\phi) \cdot u^{(j)} - h_n^{(j)}(\hat{x}_n) \cdot u^{(j)} = O(n^{-1})$ . Consequently, there exist  $e_n^0, e_n^1, e_n^2$  and  $e_n^3$  of order  $O(1)$ , such that

$$\det[(h_n^\phi)''(\hat{x}_n^\phi)]^{1/2} = \det[h_n''(\hat{x}_n)]^{1/2} + e_n^0 n^{-1},$$

$$(h_n^\phi)^{(4)}(\hat{x}_n) \cdot u^{(4)} = h_n^{(4)}(\hat{x}_n) \cdot u^{(4)} + e_n^1 n^{-1}$$

and

$$\begin{aligned} \exp\left(-\frac{1}{2}(h_n^\phi)''(\hat{x}_n^\phi) \cdot u^{(2)}\right) &= \exp\left(-\frac{1}{2}h_n''(\hat{x}_n) \cdot u^{(2)} + e_n^2 n^{-1}\right) \\ &= \exp\left(-\frac{1}{2}h_n''(\hat{x}_n) \cdot u^{(2)}\right) (1 + e_n^3 n^{-1}). \end{aligned}$$

Therefore,

$$\begin{aligned}
& \mathbb{E} [(h_n^\phi)^{(4)}(\hat{x}_n) \cdot (X^\phi)^{(4)}] \\
&= \frac{\det[(h_n^\phi)''(\hat{x}_n)]^{1/2}}{(2\pi)^{d/2}} \int_E (h_n^\phi)^{(4)}(\hat{x}_n) \cdot u^{(4)} \exp\left(-\frac{1}{2}(h_n^\phi)''(\hat{x}_n) \cdot u^{(2)}\right) du \\
&= \frac{\det[h_n''(\hat{x}_n)]^{1/2} + e_n^0}{(2\pi)^{d/2}} \int_E (h_n^{(4)}(\hat{x}_n) \cdot u^{(4)} + e_n^1 n^{-1}) \exp\left(-\frac{1}{2}h_n''(\hat{x}_n) \cdot u^{(2)}\right) (1 + e_n^3 n^{-1}) du \\
&= \frac{\det[h_n''(\hat{x}_n)]^{1/2}}{(2\pi)^{d/2}} \int_E h_n^{(4)}(\hat{x}_n) \cdot u^{(4)} \exp\left(-\frac{1}{2}h_n''(\hat{x}_n) \cdot u^{(2)}\right) du + O(n^{-1}).
\end{aligned}$$

Similarly, we have that  $\mathbb{E} [((h_n^\phi)'''(\hat{x}_n) \cdot (X^\phi)^{(3)})^2] - \mathbb{E} [(h_n'''(\hat{x}_n) \cdot X^{(3)})^2] = O(n^{-1})$ , so that we obtain  $A_n^\phi - A_n = O(n^{-1})$ .  $\square$

Let  $M_n$  be the moment generating function (mgf) associating with the posterior  $\mu_n$ . It is defined on  $\mathbb{R}^d$  by

$$a \longmapsto M_n(a) = \mathbb{E}[e^{a^T X} | Y_{1:n}]$$

and it verifies, when it exists,  $\mathbb{E}[X | Y_{1:n}] = M_n'(0)^T$  and  $\mathbb{V}[X | Y_{1:n}] = M_n''(0) - M_n'(0)^T M_n'(0)$ .

Tierney, Kass and Kadane propose in [Tierney et al. (1989)] to approximate the posterior expectation and variance by differentiating the fully exponential Laplace approximation of the mgf and evaluate it at 0. Let  $\hat{M}_n$  be the fully exponential Laplace approximation of the mgf, i.e.

$$a \longmapsto \hat{M}_n(a) = \frac{e^{a^T \hat{x}_n^a} g_n(\hat{x}_n^a) q(\hat{x}_n^a)}{g_n(\hat{x}_n) q(\hat{x}_n)} \left( \frac{\det[J_n(\hat{x}_n)]}{\det[J_n(\hat{x}_n^a)]} \right)^{1/2}$$

where  $\hat{x}_n^a = \operatorname{argmax}_{x \in E} \{e^{a^T x} g_n(x) q(x)\}$  for all  $a \in \mathbb{R}^d$ . Thus, the Laplace approximations of the posterior expectation and covariance matrix are respectively

$$\hat{m}_n = \hat{M}_n'(0)^T$$

and

$$\hat{P}_n = \hat{M}_n''(0) - \hat{M}_n'(0)^T \hat{M}_n'(0).$$

Note that  $\hat{M}_n'(0)$  is a  $1 \times d$  Jacobian matrix and  $\hat{M}_n''(0)$  is a  $d \times d$  Hessian matrix, so that  $\hat{m}_n$  is a  $d$ -dimensional column vector and  $\hat{P}_n$  is a  $d \times d$  matrix. The proposition below insures that these approximations exist under assumption  $\mathbf{L}_n'$ . Recall that the

observed information matrix at  $x \in E$ , when it exists, is defined by

$$J_n(x) = -(\log g_n)''(x) - (\log q)(x) = -(\log p_n)''(x).$$

**Proposition 3.2.3.** *Suppose that assumption  $\mathbf{L}'_n$  is verified. Then,  $\hat{M}_n$  is twice continuously differentiable in a neighbourhood of 0 with*

$$\hat{M}'_n(0) \cdot u = \hat{x}_n^T u - \frac{1}{2} \text{tr}[J_n(\hat{x}_n)^{-1} \cdot (J'_n(\hat{x}_n) \cdot (J_n(\hat{x}_n)^{-1} u))]$$

for all  $u \in \mathbb{R}^d$ .

*Proof.* Let  $\varphi_n : a \mapsto \operatorname{argmax}_{x \in E} \{e^{a^T x} p_n(x)\}$  defined over  $\mathbb{R}^d$ , so that

$$\hat{M}_n(a) = \frac{e^{a^T \varphi_n(a)} p_n(\varphi_n(a))}{p_n(\hat{x}_n)} \left( \frac{\det[J_n(\hat{x}_n)]}{\det[J_n(\varphi_n(a))]} \right)^{1/2}.$$

Let  $f_n : (a, x) \mapsto a^T + (\log p_n)'(x)$  defined over  $\mathbb{R}^d \times E$ , so that  $f_n(a, \varphi_n(a)) = 0$  and  $\frac{\partial^2 f_n}{\partial x^2}(a, x = \varphi_n(a)) = (\log p_n)''(\hat{x}_n) > 0$ . According to the implicit function theorem,  $\varphi_n$  has the same order of regularity as  $f_n$  in a neighbourhood of 0. In particular, it implies that  $\hat{M}_n$  is twice continuously differentiable in a neighbourhood of 0. Besides, the first-order derivative of  $\varphi_n$  is  $\varphi'_n(a) = -\left[\frac{\partial f_n}{\partial x}(a, \varphi_n(a))\right]^{-1} \frac{\partial f_n}{\partial a}(a, \varphi_n(a)) = -(\log p_n)''(\hat{x}_n^a)^{-1} = J_n(\varphi_n(a))^{-1}$  (here,  $\frac{\partial f_n}{\partial x}(a, \varphi_n(a))$  is a  $d \times d$  Jacobian matrix). Thus,

$$(\det \circ J_n \circ \varphi_n)'(a) \cdot u = \det[J_n(\varphi_n(a))] \text{tr}[J_n(\varphi_n(a))^{-1} \cdot (J'_n(\varphi_n(a)) \cdot (J_n(\varphi_n(a))^{-1} \cdot u))],$$

where, for all  $v \in \mathbb{R}^d$ ,  $J'_n(\varphi_n(a)) \cdot v = \left[ \left( \frac{\partial^2 \log p_n}{\partial x_i \partial x_j} \right)'(a) \cdot v \right]_{i,j}$  is a  $d \times d$  matrix. Hence,

$$\begin{aligned} & \hat{M}'_n(a) \cdot u \\ &= \frac{p_n(\hat{x}_n)}{\det[J_n(\hat{x}_n)]^{1/2}} \left( - \frac{\det[J_n(\varphi_n(a))]^{1/2} \text{tr}[J_n(\varphi_n(a))^{-1} \cdot (J'_n(\varphi_n(a)) \cdot (J_n(\varphi_n(a))^{-1} \cdot u))]}{2e^{a^T \varphi_n(a)} p_n(\varphi_n(a))} \right. \\ & \quad \left. + \frac{\det[J_n(\varphi_n(a))]^{1/2} ((\varphi_n(a))^T + a^T \varphi'_n(a)) e^{a^T \varphi_n(a)} p_n(\varphi_n(a)) + e^{a^T \varphi_n(a)} p'_n(\varphi_n(a)) \varphi'_n(a) \cdot u}{e^{2a^T \varphi_n(a)} p_n(\varphi_n(a))^2} \right). \end{aligned}$$

We obtain the result by evaluating the above expression at  $a = 0$ .  $\square$

Let

$$e_n(a) = n^2 \left( \frac{M_n(a)}{\hat{M}_n(a)} - 1 \right).$$

From proposition 3.2.3,  $e_n$  is twice continuously differentiable in a neighbourhood of the origin. The next theorem establishes the consistency of approximations  $\hat{m}_n$  and  $\hat{P}_n$  and the speed of convergence as  $n \rightarrow \infty$ , under the ad hoc assumption that  $e'_n$  and  $e''_n$  are bounded uniformly in  $n$  in a neighbourhood of 0.

**Theorem 3.2.4.** *Suppose that assumption  $\mathbf{L}'_n$  is verified. Suppose that the MLE  $\tilde{x}_n$  do not diverge, i.e.  $|\tilde{x}_n|$  is bounded uniformly in  $n$ . Suppose that  $e'_n(a)$  and  $e''_n(a)$  are of order  $O(1)$  as  $n \rightarrow \infty$  for all  $a$  in a neighbourhood of the origin. Suppose that  $\int_E \left( e^{a^T x} g_n(x) q(x) \right)^{n_0/n} dx \leq K^a n^p$  for some  $p > 0$  and  $K^a > 0$  when  $n \geq n_0$  and for all  $a \in \mathbb{R}^d$ . Then,*

$$\mathbb{E}[X|Y_{1:n}] = \hat{m}_n + O(n^{-2})$$

and

$$\mathbb{V}[X|Y_{1:n}] = \hat{P}_n + O(n^{-2}).$$

Note that  $\int_E \left( e^{a^T x} g_n(x) q(x) \right)^{n_0/n} dx < \infty$  insures that the mgf exists, by letting  $n = n_0$ .

*Proof.* We have that

$$M_n(a) = \hat{M}_n(a)(1 + e_n(a)n^{-2}).$$

The fact that  $|\tilde{x}_n|$  is bounded uniformly in  $n$  insures that the function  $x \mapsto a^T x$  is bounded uniformly in  $n$  in a neighbourhood of  $\tilde{x}_n$ , which is required to apply theorem 3.2.2. From theorem 3.2.2,  $e_n(a)$  is of order  $O(1)$  as  $n \rightarrow \infty$  for all  $a \in \mathbb{R}^d$ . Taking the logarithm and differentiating on both sides of (3.2) yields

$$\frac{M'_n(a)}{M_n(a)} = \frac{\hat{M}'_n(a)}{\hat{M}_n(a)} + \frac{e'_n(a)n^{-2}}{1 + e_n(a)n^{-2}}.$$

Evaluating at  $a = 0$  then gives

$$\mathbb{E}[X|Y_{1:n}] = \hat{m}_n + O(n^{-2})$$

because  $M_n(0) = \hat{M}_n(0) = 1$ . Differentiating once again yields

$$\begin{aligned} \frac{M_n''(a)M_n(a) - M_n'(a)^2}{M_n(a)^2} &= \frac{\hat{M}_n''(a)\hat{M}_n(a) - \hat{M}_n'(a)^2}{\hat{M}_n(a)^2} \\ &\quad + \frac{e_n''(a)(1 + e_n(a)n^{-2}) - e_n'(a)e_n'(a)n^{-2}}{(1 + e_n(a)n^{-2})^2}n^{-2}, \end{aligned}$$

so that we obtain

$$\mathbb{V}[X|Y_{1:n}] = \hat{P}_n + O(n^{-2})$$

after evaluation at  $a = 0$ . □

Theorems 3.2.2 and 3.2.4 above state that the convergence of the Laplace approximations to the true posterior moments is of order  $O(n^{-2})$ . On the other hand, the convergence of the MAP to the MLE is only of order  $O(n^{-1})$  (proposition 3.1.1). Hence, when the model is sufficiently regular, the Laplace method provides approximations converging much faster than the more classical MAP estimator. This illustrates the power of the Laplace method as an estimation technique in Bayesian statistics.

The theorem below provides ready to use multidimensional formulas for the Laplace approximations  $\hat{m}_n$  and  $\hat{P}_n$ .

**Theorem 3.2.5** (Multidimensional Laplace approximations for posterior moments). *Suppose that  $\log p_n$  admits a maximum at  $\hat{x}_n$  and is four-times continuously differentiable in a neighbourhood of  $\hat{x}_n$ . Using the matrix calculus rules defined in [Fackler (2005)] and [Magnus (2010)], we have that*

$$\hat{m}_n = \hat{x}_n - \frac{1}{2}J_n(\hat{x}_n)^{-1}J_n'(\hat{x}_n)^T \text{vec}[J_n(\hat{x}_n)^{-1}] \quad (3.2)$$

and

$$\begin{aligned} \hat{P}_n &= J_n(\hat{x}_n)^{-1} + \frac{1}{2}J_n(\hat{x}_n)^{-1}J_n'(\hat{x}_n)^T(J_n(\hat{x}_n)^{-1} \otimes J_n(\hat{x}_n)^{-1})J_n'(\hat{x}_n)J_n(\hat{x}_n)^{-1} \\ &\quad + \frac{1}{2}(I_d \otimes \text{vec}(J_n(\hat{x}_n)^{-1})^T J_n'(\hat{x}_n))(J_n'(\hat{x}_n)^{-1} \otimes J_n(\hat{x}_n)^{-1})J_n'(\hat{x}_n)J_n(\hat{x}_n)^{-1} \\ &\quad - \frac{1}{2}J_n(\hat{x}_n)^{-1}(I_d \otimes \text{vec}(J_n(\hat{x}_n)^{-1})^T)J_n''(\hat{x}_n)J_n(\hat{x}_n)^{-1}. \end{aligned} \quad (3.3)$$

In theorem 3.2.5 above,  $\otimes$  denotes the Kronecker product and  $\text{vec}$  denotes a matrix operator that vectorizes a matrix by stacking its columns. Thus, when  $M$  is a  $d \times d$  matrix,  $\text{vec } M$  is a  $d^2$ -dimensional column vector. Under the matrix calculus rules we

use, that can be found in [Fackler (2005); Magnus (2010)],  $J'_n(x)$  is a Jacobian matrix defined by

$$J'_n(x) = \frac{d \operatorname{vec} J_n}{dx} = \left[ \frac{d(\operatorname{vec} J_n)_i}{dx_j} \right]_{i \in \{1, \dots, d^2\}, j \in \{1, \dots, d\}},$$

hence it is a  $d^2 \times d$  matrix. Consequently,  $J''_n(x)$  is a Jacobian matrix as well, defined by

$$J''_n(x) = \frac{d \operatorname{vec} J'_n}{dx},$$

and its dimensions are  $d^3 \times d$ .

The proof of theorem 3.2.5 is provided in section 3.3.

**Remark 3.2.6.** When the state dimension is  $d = 1$ , the formulas in theorem 3.2.5 become

$$\hat{m}_n = \hat{x}_n - \frac{1}{2} \frac{J'_n(\hat{x}_n)}{J_n(\hat{x}_n)^2} \quad (3.4)$$

and

$$\hat{P}_n = \frac{1}{J_n(\hat{x}_n)} + \frac{J'_n(\hat{x}_n)^2}{J_n(\hat{x}_n)^4} - \frac{1}{2} \frac{J''_n(\hat{x}_n)}{J_n(\hat{x}_n)^3}. \quad (3.5)$$

**Remark 3.2.7.** Other approximations have been proposed in the literature for the first two posterior moments (mean and covariance matrix), see for instance Theorem 5.1b in [Ghosh (1994)] or [Johnson (1970)]. These other approximations have been obtained with a different approach, based on the second-order, not fully exponential, Laplace approximation restated in our theorem 1.2.2. In particular, in these other approximations, coefficients are evaluated at the MLE whereas in our theorem 3.2.5 coefficients are evaluated at the MAP. This may be seen as a possible advantage of our proposed approximations, since in a few applications, including the target tracking problems considered in chapter 8, the MLE is not uniquely defined whereas the MAP is uniquely defined

**Example 3.2.8** (Gamma distribution). In Bayesian statistics, the gamma distribution is often used as a conjugate prior (for the inference on rate parameters, e.g). When the likelihood model is Poisson, exponential, log-normal, Pareto, gamma, or inverse gamma, then a gamma prior model yields a gamma posterior [Fink (1997)].

The density of the gamma distribution  $\Gamma(\alpha, \beta)$  is defined on  $\mathbb{R}_+$  by

$$f(x) = \frac{\beta^\alpha}{\Gamma(\alpha)} x^{\alpha-1} e^{-\beta x},$$

with  $\alpha > 0$  and  $\beta > 0$ . When  $\alpha > 1$ , its mode exists and is  $\hat{x} = \frac{\alpha-1}{\beta}$ . The observed information at  $x \in \mathbb{R}_+^*$  is  $J(x) = -(\log f)''(x) = \frac{\alpha-1}{x^2}$ . Then,  $J(\hat{x}) = \frac{\beta^2}{\alpha-1}$ ,  $J'(\hat{x}) = -\frac{2\beta^3}{(\alpha-1)^2}$  and  $J''(\hat{x}) = \frac{6\beta^4}{(\alpha-1)^3}$ . The Laplace approximations (3.4) and (3.5) respectively give  $\alpha/\beta$  and  $\alpha/\beta^2$ , which are the exact values of the expectation and the variance of the gamma distribution [Bui Quang et al. (2012)].

Hence, the Laplace approximations are exact when the posterior is gamma, for  $\alpha > 1$  and  $\beta > 0$ , even for  $n = 1$ . Notice that a gamma density can present various shapes and be seriously asymmetric, as show in figure 3.1.

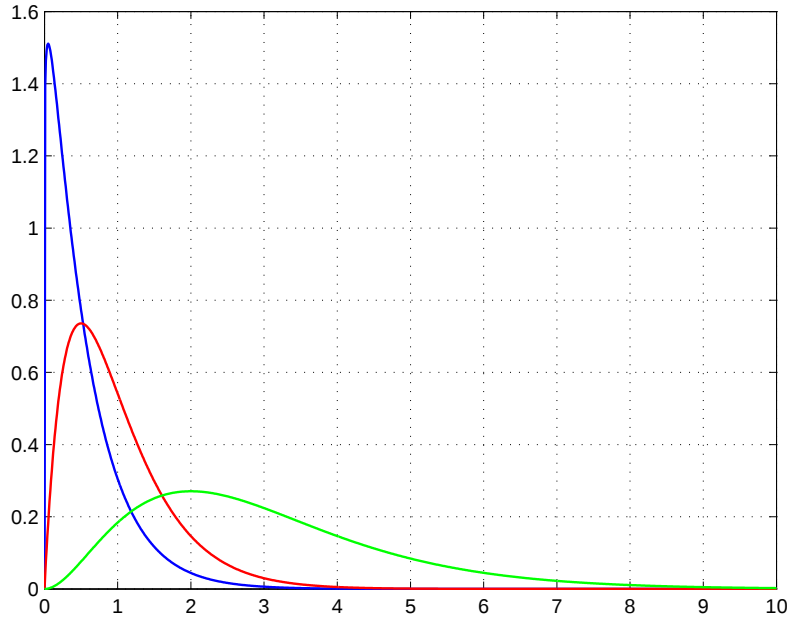


Figure 3.1: Density of the gamma distribution  $\Gamma(\alpha, \beta)$ . Blue:  $\alpha = 1.1$ ,  $\beta = 0.5$ ; red:  $\alpha = 2$ ,  $\beta = 0.5$ ; green:  $\alpha = 3$ ,  $\beta = 1$ .

### 3.3 Proof of theorem 3.2.5

*Proof of theorem 3.2.5.* We have that

$$(\hat{M}_n)'(a) = \frac{\partial \hat{M}_n}{\partial \hat{x}_n^a} \frac{d\hat{x}_n^a}{da} + \frac{\partial \hat{M}_n}{\partial a}.$$

Since  $\hat{x}_n^a = \hat{x}_n$  when  $a = 0$ ,

$$\hat{M}'_n(0)^T = \frac{\partial \hat{M}_n}{\partial \hat{x}_n^a}(\hat{x}_n^a = \hat{x}_n) \frac{d\hat{x}_n^a}{da}(a=0) + \frac{\partial \hat{M}_n}{\partial a}(a=0)$$

First of all,

$$\frac{\partial \hat{M}_n}{\partial a} = \hat{M}_n(a)(\hat{x}_n^a)^T,$$

so that

$$\frac{\partial \hat{M}_n}{\partial a}(a=0) = \hat{x}_n^T.$$

Besides,

$$\frac{\partial \hat{M}_n}{\partial \hat{x}_n^a} = \frac{\det[J(\hat{x}_n)]^{1/2}}{p_n(\hat{x}_n)} \frac{\partial}{\partial \hat{x}_n^a} \left( \frac{e^{a^T \hat{x}_n} p_n(\hat{x}_n^a)}{\det[J(\hat{x}_n^a)]^{1/2}} \right),$$

with

$$\begin{aligned} & \frac{\partial}{\partial \hat{x}_n^a} \left( \frac{e^{a^T \hat{x}_n} p_n(\hat{x}_n^a)}{\det[J(\hat{x}_n^a)]^{1/2}} \right) \\ &= \frac{a^T e^{a^T \hat{x}_n} p_n(\hat{x}_n^a) + e^{a^T \hat{x}_n} p'_n(\hat{x}_n^a)}{\det[J(\hat{x}_n^a)]^{1/2}} + e^{a^T \hat{x}_n} p_n(\hat{x}_n^a) \frac{d}{d\hat{x}_n^a} \left( \frac{1}{\det[J(\hat{x}_n^a)]^{1/2}} \right) \\ &= \frac{a^T e^{a^T \hat{x}_n} p_n(\hat{x}_n^a) + e^{a^T \hat{x}_n} p'_n(\hat{x}_n^a)}{\det[J(\hat{x}_n^a)]^{1/2}} - \frac{1}{2} \frac{e^{a^T \hat{x}_n} p_n(\hat{x}_n^a)}{\det[J(\hat{x}_n^a)]^{1/2}} \text{vec}(J(\hat{x}_n^a)^{-1})^T J'(\hat{x}_n^a) \end{aligned}$$

so that,

$$\frac{\partial \hat{M}_n}{\partial \hat{x}_n^a}(\hat{x}_n^a = \hat{x}_n) = -\frac{1}{2} \text{vec}(J(\hat{x}_n)^{-1})^T J'(\hat{x}_n).$$

Moreover, the definition of  $\hat{x}_n^a$  implies that

$$\varphi(a) := a + (\log p_n)'(\hat{x}_n^a) = 0$$

for all  $a \in \mathbb{R}^d$ , so that

$$\begin{aligned} \varphi'(a) &= \frac{\partial \varphi}{\partial x}(\hat{x}_n^a(a)) \frac{d\hat{x}_n^a}{da} + \frac{\partial \varphi}{\partial a} \\ &= -J(\hat{x}_n^a) \frac{d\hat{x}_n^a}{da} + I_d = 0, \end{aligned}$$

which gives

$$\frac{d\hat{x}_n^a}{da} = J(\hat{x}_n^a)^{-1}.$$

Finally,

$$\hat{m}_n = \hat{x}_n - \frac{1}{2} J(\hat{x}_n)^{-1} J'(\hat{x}_n)^T \text{vec}(J(\hat{x}_n)^{-1}).$$

Let us derive now the second derivative of  $\hat{M}_n$ . We have that

$$\begin{aligned} (\hat{M}_n)''(a) &= \frac{d}{da} \left( \frac{\partial \hat{M}_n}{\partial \hat{x}_n^a} \frac{d\hat{x}_n^a}{da} \right) + \frac{d}{da} \left( \frac{\partial \hat{M}_n}{\partial a} \right) \\ &= \left[ \frac{d\hat{x}_n^a}{da} \right]^T \frac{d}{da} \left( \frac{\partial \hat{M}_n}{\partial \hat{x}_n^a} \right) + \left( I_d \otimes \frac{\partial \hat{M}_n}{\partial \hat{x}_n^a} \right) \frac{d^2 \hat{x}_n^a}{da^2} + \frac{d}{da} \left( \frac{\partial \hat{M}_n}{\partial a} \right), \end{aligned}$$

where

$$\frac{d}{da} \left( \frac{\partial \hat{M}_n}{\partial \hat{x}_n^a} \right) = \frac{\partial^2 \hat{M}_n}{\partial (\hat{x}_n^a)^2} \frac{d\hat{x}_n^a}{da} + \frac{\partial^2 \hat{M}_n}{\partial a \partial \hat{x}_n^a}$$

and

$$\frac{d}{da} \left( \frac{\partial \hat{M}_n}{\partial a} \right) = \frac{\partial^2 \hat{M}_n}{\partial \hat{x}_n^a \partial a} \frac{d\hat{x}_n^a}{da} + \frac{\partial^2 \hat{M}_n}{\partial a^2},$$

so that

$$\begin{aligned} (\hat{M}_n)''(a) &= \left[ \frac{d\hat{x}_n^a}{da} \right]^T \frac{\partial^2 \hat{M}_n}{\partial (\hat{x}_n^a)^2} \frac{d\hat{x}_n^a}{da} + \left( I_d \otimes \frac{\partial \hat{M}_n}{\partial \hat{x}_n^a} \right) \frac{d^2 \hat{x}_n^a}{da^2} \\ &\quad + \left[ \frac{d\hat{x}_n^a}{da} \right]^T \frac{\partial^2 \hat{M}_n}{\partial a \partial \hat{x}_n^a} + \frac{\partial^2 \hat{M}_n}{\partial \hat{x}_n^a \partial a} \frac{d\hat{x}_n^a}{da} + \frac{\partial^2 \hat{M}_n}{\partial a^2}. \end{aligned} \quad (3.6)$$

Let us calculate the terms of (3.6) and evaluate them at  $a = 0$ .

**The first term**  $\left[ \frac{d\hat{x}_n^a}{da} \right]^T \frac{\partial^2 \hat{M}_n}{\partial (\hat{x}_n^a)^2} \frac{d\hat{x}_n^a}{da}$

$$\frac{\partial^2 \hat{M}_n}{\partial (\hat{x}_n^a)^2} = \frac{\det[J_n(\hat{x}_n)]^{1/2}}{p_n(\hat{x}_n)} \frac{\partial^2}{\partial (\hat{x}_n^a)^2} \left( \frac{e^{a^T \hat{x}_n^a} p(\hat{x}_n^a)}{\det[J_n(\hat{x}_n^a)]^{1/2}} \right)$$

with

$$\begin{aligned} \frac{\partial^2}{\partial(\hat{x}_n^a)^2} \left( \frac{e^{a^T \hat{x}_n^a} p_n(\hat{x}_n^a)}{\det[J_n(\hat{x}_n^a)]^{1/2}} \right) &= \frac{aa^T e^{a^T \hat{x}_n^a} p_n(\hat{x}_n^a) + 2ae^{a^T \hat{x}_n^a} p'_n(\hat{x}_n^a) + e^{a^T \hat{x}_n^a} p''_n(\hat{x}_n^a)}{\det[J_n(\hat{x}_n^a)]^{1/2}} \\ &\quad - \frac{a^T e^{a^T \hat{x}_n^a} p_n(\hat{x}_n^a) + e^{a^T \hat{x}_n^a} p'_n(\hat{x}_n^a)}{\det[J_n(\hat{x}_n^a)]^{1/2}} \text{vec}(J_n(\hat{x}_n^a)^{-1})^T J'_n(\hat{x}_n^a) \\ &\quad - \frac{1}{2} e^{a^T \hat{x}_n^a} p_n(\hat{x}_n^a) \frac{d}{d\hat{x}_n^a} \left( \frac{\text{vec}(J_n(\hat{x}_n^a)^{-1})^T J'_n(\hat{x}_n^a)}{\det[J_n(\hat{x}_n^a)]^{1/2}} \right) \end{aligned}$$

where

$$\begin{aligned} &\frac{d}{d\hat{x}_n^a} \left( \frac{\text{vec}(J_n(\hat{x}_n^a)^{-1})^T J'_n(\hat{x}_n^a)}{\det[J_n(\hat{x}_n^a)]^{1/2}} \right) \\ &= -\frac{1}{2} \frac{1}{\det[J_n(\hat{x}_n^a)]^{1/2}} J'_n(\hat{x}_n^a)^T \text{vec}(J_n(\hat{x}_n^a)^{-1}) \text{vec}(J_n(\hat{x}_n^a)^{-1})^T J'_n(\hat{x}_n^a) \\ &\quad + \left( I_d \otimes \frac{1}{\det[J_n(\hat{x}_n^a)]^{1/2}} \right) \frac{d}{d\hat{x}_n^a} \left( \text{vec}(J_n(\hat{x}_n^a)^{-1})^T J'_n(\hat{x}_n^a) \right), \\ &\frac{d}{d\hat{x}_n^a} \left( \text{vec}(J_n(\hat{x}_n^a)^{-1})^T J'_n(\hat{x}_n^a) \right) \\ &= J'_n(\hat{x}_n^a)^T (J_n^{-1})'(\hat{x}_n^a) + (I_d \otimes \text{vec}(J_n(\hat{x}_n^a)^{-1})^T) J''_n(\hat{x}_n^a), \end{aligned}$$

and

$$(J_n^{-1})'(\hat{x}_n^a) = -(J_n(\hat{x}_n^a)^{-1} \otimes J_n(\hat{x}_n^a)^{-1}) J'_n(\hat{x}_n^a),$$

so that

$$\begin{aligned} &\frac{d}{d\hat{x}_n^a} \left( \frac{\text{vec}(J_n(\hat{x}_n^a)^{-1})^T J'_n(\hat{x}_n^a)}{\det[J_n(\hat{x}_n^a)]^{1/2}} \right) \\ &= -\frac{1}{2} \frac{1}{\det[J_n(\hat{x}_n^a)]^{1/2}} J'_n(\hat{x}_n^a)^T \text{vec}(J_n(\hat{x}_n^a)^{-1}) \text{vec}(J_n(\hat{x}_n^a)^{-1})^T J'_n(\hat{x}_n^a) \\ &\quad - \left( I_d \otimes \frac{1}{\det[J_n(\hat{x}_n^a)]^{1/2}} \right) J'_n(\hat{x}_n^a)^T (J_n(\hat{x}_n^a)^{-1} \otimes J_n(\hat{x}_n^a)^{-1}) J'_n(\hat{x}_n^a) \\ &\quad + \left( I_d \otimes \frac{1}{\det[J_n(\hat{x}_n^a)]^{1/2}} \right) (I_d \otimes \text{vec}(J_n(\hat{x}_n^a)^{-1})^T) J''_n(\hat{x}_n^a). \end{aligned}$$

Noticing that

$$J_n(\hat{x}_n) = -(\log p_n)''(\hat{x}_n) = -\frac{p''_n(\hat{x}_n)}{p_n(\hat{x}_n)},$$

we get after evaluation at  $a = 0$

$$\begin{aligned}
& \frac{d\hat{x}_n^a}{da}(a=0) \frac{\partial^2 \hat{M}_n}{\partial(\hat{x}_n^a)^2}(\hat{x}_n^a = \hat{x}_n) \frac{d\hat{x}_n}{da}(a=0) \\
&= -J_n(\hat{x}_n)^{-1} + \frac{1}{4} J_n(\hat{x}_n)^{-1} J'_n(\hat{x}_n)^T \text{vec}(J_n(\hat{x}_n)^{-1}) \text{vec}(J_n(\hat{x}_n)^{-1})^T J'_n(\hat{x}_n) J_n(\hat{x}_n)^{-1} \\
&\quad + \frac{1}{2} J_n(\hat{x}_n)^{-1} J'_n(\hat{x}_n)^T (J_n(\hat{x}_n)^{-1} \otimes J_n(\hat{x}_n)^{-1}) J'_n(\hat{x}_n) J_n(\hat{x}_n)^{-1} \\
&\quad - \frac{1}{2} J_n(\hat{x}_n)^{-1} (I_d \otimes \text{vec}(J_n(\hat{x}_n)^{-1})^T) J''_n(\hat{x}_n) J_n(\hat{x}_n)^{-1}.
\end{aligned}$$

**The second term**  $\left( I_d \otimes \frac{\partial \hat{M}_n}{\partial \hat{x}_n^a} \right) \frac{d^2 \hat{x}_n^a}{da^2}$

$$\begin{aligned}
\frac{d^2 \hat{x}_n^a}{da^2} &= \frac{d}{da} (J_n(\hat{x}_n^a(a))^{-1}) = \frac{dJ_n^{-1}}{dJ_n}(J_n(\hat{x}_n^a)) J'_n(\hat{x}_n^a) \frac{d\hat{x}_n^a}{da} \\
&= -(J_n(\hat{x}_n^a)^{-1} \otimes J_n(\hat{x}_n^a)^{-1}) J'_n(\hat{x}_n^a) J_n^{-1}(\hat{x}_n^a),
\end{aligned}$$

so that

$$\begin{aligned}
& \left( I_d \otimes \frac{\partial \hat{M}_n}{\partial \hat{x}_n^a}(\hat{x}_n^a = \hat{x}_n) \right) \frac{d^2 \hat{x}_n^a}{da^2}(a=0) \\
&= \frac{1}{2} (I_d \otimes \text{vec}(J_n(\hat{x}_n)^{-1})^T J'_n(\hat{x}_n)) (J_n(\hat{x}_n)^{-1} \otimes J_n(\hat{x}_n)^{-1}) J'_n(\hat{x}_n) J_n^{-1}(\hat{x}_n).
\end{aligned}$$

**The third term**  $\left[ \frac{d\hat{x}_n^a}{da} \right]^T \frac{\partial^2 \hat{M}_n}{\partial a \partial \hat{x}_n^a} + \frac{\partial^2 \hat{M}_n}{\partial \hat{x}_n^a \partial a} \frac{d\hat{x}_n^a}{da}$

$$\begin{aligned}
\frac{\partial^2 \hat{M}_n}{\partial \hat{x}_n^a \partial a} &= \hat{x}_n^a \frac{\partial \hat{M}_n}{\partial \hat{x}_n^a} + \hat{M}_n(\hat{x}_n^a) I_d \\
&= \frac{\det[J_n(\hat{x}_n)]^{1/2}}{p_n(\hat{x}_n)} \left( \frac{(\hat{x}_n^a a^T + I_d) e^{a^T \hat{x}_n^a} p_n(\hat{x}_n^a) + \hat{x}_n^a e^{a^T \hat{x}_n^a} p'_n(\hat{x}_n^a)}{\det[J_n(\hat{x}_n^a)]^{1/2}} \right. \\
&\quad \left. - \frac{1}{2} \frac{\hat{x}_n^a e^{a^T \hat{x}_n^a} p_n(\hat{x}_n^a)}{\det[J_n(\hat{x}_n^a)]^{1/2}} \text{vec}(J_n(\hat{x}_n^a)^{-1})^T J'_n(\hat{x}_n^a) \right)
\end{aligned}$$

so that

$$\frac{\partial^2 \hat{M}_n}{\partial \hat{x}_n^a \partial a}(a=0, \hat{x}_n^a = \hat{x}_n) = I_d - \frac{1}{2} \hat{x}_n \text{vec}(J_n(\hat{x}_n)^{-1})^T J'_n(\hat{x}_n).$$

Thus,

$$\begin{aligned} & \left[ \frac{d\hat{x}_n^a}{da}(a=0) \right]^T \frac{\partial^2 \hat{M}_n}{\partial a \partial \hat{x}_n^a}(a=0, \hat{x}_n^a = \hat{x}_n) + \frac{\partial^2 \hat{M}_n}{\partial \hat{x}_n^a \partial a}(a=0, \hat{x}_n^a = \hat{x}_n) \frac{d\hat{x}_n^a}{da}(a=0) \\ &= 2J_n(\hat{x}_n)^{-1} - \frac{1}{2} \left( J_n(\hat{x}_n)^{-1} J_n'(\hat{x}_n)^T \text{vec}(J_n(\hat{x}_n)^{-1}) \hat{x}_n^T + \hat{x}_n \text{vec}(J_n(\hat{x}_n)^{-1})^T J_n'(\hat{x}_n) J_n(\hat{x}_n)^{-1} \right). \end{aligned}$$

**The fourth term**  $\frac{\partial^2 \hat{M}_n}{\partial a^2}$

$$\frac{\partial^2 \hat{M}_n}{\partial a^2} = \hat{x}_n^a (\hat{x}_n^a)^T \hat{M}_n(a)$$

which evaluated at  $a = 0$  gives

$$\frac{\partial^2 \hat{M}_n}{\partial a^2}(a=0) = \hat{x}_n \hat{x}_n^T.$$

Finally,

$$\begin{aligned} \hat{P}_n &= (\hat{M}_n)''(0) - \hat{m}_n \hat{m}_n^T \\ &= J_n(\hat{x}_n)^{-1} + \frac{1}{2} J_n(\hat{x}_n)^{-1} J_n'(\hat{x}_n)^T (J_n(\hat{x}_n)^{-1} \otimes J_n(\hat{x}_n)^{-1}) J_n'(\hat{x}_n) J_n(\hat{x}_n)^{-1} \\ &\quad + \frac{1}{2} (I_d \otimes \text{vec}(J_n(\hat{x}_n)^{-1})^T J_n'(\hat{x}_n)) (J_n(\hat{x}_n)^{-1} \otimes J_n(\hat{x}_n)^{-1}) J_n'(\hat{x}_n) J_n(\hat{x}_n)^{-1} \\ &\quad - \frac{1}{2} J_n(\hat{x}_n)^{-1} (I_d \otimes \text{vec}(J_n(\hat{x}_n)^{-1})^T) J_n''(\hat{x}_n) J_n(\hat{x}_n)^{-1}. \end{aligned}$$

□

## Conclusion

We have presented in this chapter the application of the Laplace method in Bayesian statistics, to approximate posterior moments when the hidden state is static. Under regularity and identifiability assumptions, the approximations are consistent as the observation sample size  $n$  goes to infinity. These assumptions essentially require that the posterior tends to a Dirac measure centered at the MLE as  $n \rightarrow \infty$ .

In particular, we have derived multidimensional approximation formulas for the posterior expectation and covariance matrix (theorem 3.2.5). These formulas have been established by Tierney, Kass and Kadane in [Tierney et al. (1989)] in the unidimensional case. Their multidimensional version is useful in practice when working with array programming languages, such as MATLAB or R. Indeed, formulas using matrix

operators are more convenient than component-wise formulas in these languages. The latter requires the use of for-loops whereas the former can be more concisely coded (see the conclusion of chapter 7).

In the next chapters, we apply the Laplace the method in a dynamic framework, i.e. in the case where the state obeys to a Markov process.



# Chapter 4

## The Laplace method in dynamic models with small observation noise

We propose in this chapter a recursive filtering algorithm for state-space nonlinear models based on the Laplace method. The integrals involved in the prediction and update steps are computed thanks to the Laplace method. This filter provides at each time step an approximation of the posterior density. As for asymptotics, we consider that the observation noise intensity goes to zero. Thanks to a strong assumption on the likelihood model, we show the consistency of the approximated density and we study the stability of the approximation error over time. This assumption, which is generally not fulfilled in practice, is that likelihood function associated to each observation  $Y_k$  admits a unique global maximum when the observation noise is small enough.

We present the model we consider and derive some properties in section 4.1. The algorithm we study is defined in section 4.2. It allows to approximate the density of Bayesian filter as well as the posterior expectation and covariance matrix at each time step. We study the consistency of the approximations in section 4.3 and the propagation of the approximation error over time in section 4.4.

### 4.1 Model set-up

Let  $\{X_k, Y_k\}_{k \geq 0}$  be a state-space model. The state space is an open subset  $E$  of  $\mathbb{R}^d$ , the state process  $\{X_k\}_{k \geq 0}$  is a Markov process whose transition kernel admits a density, denoted  $q_k$ , so that

$$\mathbb{P}[X_k \in dx' | X_{k-1} = x] = q_k(x, x')dx'$$

for  $k \geq 1$ . The initial distribution of the Markov process is  $q_0(x)dx$ . The observations take values in  $\mathbb{R}^d$  and the observation model is

$$Y_k = H_k(X_k) + \sqrt{\varepsilon}\sigma_k(X_k)W_k$$

for  $k \geq 0$ , with  $\{W_k\}_{k \geq 0}$  a Gaussian standardized white noise,  $\varepsilon$  a small positive parameter and  $\sigma_k : \mathbb{R}^d \rightarrow \mathbb{R}_+^*$  a positive function. The likelihood function is

$$g_k^\varepsilon(x) = |2\pi\varepsilon\sigma_k(x)^2|^{-d/2} \exp\left(-\frac{1}{2\varepsilon\sigma_k(x)^2}|Y_k - H_k(x)|^2\right),$$

which is proportional to the conditional density of  $Y_k$  w.r.t.  $X_k$ . We define the contrast function as

$$\ell_k^\varepsilon(x) = \frac{1}{2\sigma_k(x)^2}|Y_k - H_k(x)|^2 + d\varepsilon \log \sigma_k(x) = -\varepsilon \log g_k^\varepsilon(x) + \text{constant}.$$

Note that when  $\sigma_k(x)$  does not depend on  $x$ , the contrast function  $\ell_k^\varepsilon(x)$  does not depend on  $\varepsilon$ .

Let  $\mu_k^{\varepsilon-}$  be the predictor, i.e. the conditional distribution of  $X_k$  given  $Y_{0:k-1}$ . Let  $\mu_k^\varepsilon$  be the posterior, i.e. the conditional distribution of  $X_k$  given  $Y_{0:k}$ . The densities of the predictor and the posterior are respectively denoted  $p_k^{\varepsilon-}$  and  $p_k^\varepsilon$ . The sequence of predictors and posteriors obeys to the following recursive relation:

$$p_k^{\varepsilon-}(x') = \int q_k(x, x') p_{k-1}^\varepsilon(x) dx \quad (4.1)$$

and

$$p_k^\varepsilon(x') = \frac{g_k^\varepsilon(x') p_k^{\varepsilon-}(x')}{\int g_k^\varepsilon(x') p_k^{\varepsilon-}(x') dx'}, \quad (4.2)$$

with the initial condition  $p_0^\varepsilon(x) = \frac{g_0^\varepsilon(x) q_0(x)}{\int g_0^\varepsilon(x) q_0(x) dx}$ . We use the convention  $q_0(x, x') = q_0(x')$  in the sequel of the chapter. Let  $m_k^\varepsilon = \mathbb{E}[X_k | Y_{0:k}]$  and  $P_k^\varepsilon = \mathbb{V}[X_k | Y_{0:k}]$  respectively denote the posterior expectation and covariance at time  $k$ .

We define below the coercivity condition **C** that will be used throughout the chapter.

**Condition C.** Let  $f : E \rightarrow \mathbb{R}$ .  $f$  is said to satisfy condition **C** when there exists  $x^*$  such that:

- $f$  is continuous at  $x^*$ ;
- for all  $\delta > 0$ , there exists  $c_\delta > 0$  such that for all  $x \in E$ ,  $|x - x^*| \geq \delta$  implies  $f(x) - f(x^*) \geq c_\delta$ .

Condition **C** implies that  $f$  admits an unique global minimum at  $x^*$ . The lemma below investigates the situation when the function  $f$  is corrupted by a perturbation of order  $\varepsilon$ .

**Lemma 4.1.1.** *Let  $f : E \rightarrow \mathbb{R}$  be a function admitting a minimum at  $x^*$  and satisfying condition **C**. Let  $g : E \rightarrow \mathbb{R}$  be a bounded function and  $\{g^\varepsilon : E \rightarrow \mathbb{R}\}_{\varepsilon > 0}$  be a collection of functions continuous at  $x^*$ , such that  $\inf_{x \in E} \{g^\varepsilon(x)\} \geq m > -\infty$  for all  $\varepsilon > 0$  and such that  $g^\varepsilon(x) \rightarrow g(x)$  as  $\varepsilon \rightarrow 0$  for all  $x \in E$ . Let  $f^\varepsilon(x) = f(x) + \varepsilon g^\varepsilon(x)$ . Then, there exists a collection  $\{x^{\varepsilon*}\}_{\varepsilon > 0}$  such that  $x^{\varepsilon*} \rightarrow x^*$  as  $\varepsilon \rightarrow 0$  and  $f^\varepsilon$  is continuous at  $x^{\varepsilon*}$ .*

*If additionally,  $g^\varepsilon(x^{\varepsilon*}) \rightarrow g(x^*)$  as  $\varepsilon \rightarrow 0$ , then for all sufficiently small  $\varepsilon$  and for all  $\delta > 0$ , there exists  $c_\delta > 0$  independent of  $\varepsilon$  such that:*

$$\text{for all } x \in E, |x - x^{\varepsilon*}| \geq \delta \text{ implies } f^\varepsilon(x) - f^\varepsilon(x^{\varepsilon*}) \geq c_\delta;$$

*i.e.,  $f^\varepsilon$  satisfies condition **C** uniformly in  $\varepsilon$ .*

In particular, lemma 4.1.1 implies that  $x^{\varepsilon*} = \operatorname{argmin}_{x \in E} \{f^\varepsilon(x)\}$  when  $\varepsilon$  is sufficiently small.

*Proof.* Let  $\delta > 0$ . Because of condition **C**, there exists  $c_\delta > 0$  such that  $|x - x^*| \geq \delta$  implies  $f(x) - f(x^*) \geq c_\delta$ . Let  $x^{\varepsilon*} \in E$  such that  $f^\varepsilon(x^{\varepsilon*}) \leq f^\varepsilon(x^*)$ . Then,

$$\begin{aligned} 0 &\leq f(x^{\varepsilon*}) - f(x^*) \\ &\leq f(x^{\varepsilon*}) + f^\varepsilon(x^*) - f^\varepsilon(x^{\varepsilon*}) - f(x^*) \\ &= \varepsilon(g^\varepsilon(x^*) - g^\varepsilon(x^{\varepsilon*})) \\ &\leq \varepsilon(g^\varepsilon(x^*) - g(x^*)) + \varepsilon(g(x^*) - \inf_{x \in E} \{g^\varepsilon(x)\}). \end{aligned}$$

Since  $g^\varepsilon$  converges pointwise to  $g$  and  $\inf_{x \in E} \{g^\varepsilon(x)\} \geq m$ , the right-hand side above converges to zero as  $\varepsilon \rightarrow 0$ . Thus, we can always choose  $\varepsilon$  small enough so that

$f(x^{\varepsilon*}) - f(x^*) \leq c_\delta$ , which implies  $|x^{\varepsilon*} - x^*| \leq \delta$ . This holds for any  $\delta > 0$ , hence we obtain  $|x^{\varepsilon*} - x^*| \xrightarrow{\varepsilon \rightarrow 0} 0$ . Since  $f^\varepsilon$  is continuous at  $x^*$ ,  $f^\varepsilon$  is also continuous at  $x^{\varepsilon*}$ .

Suppose now that  $g^\varepsilon(x^{\varepsilon*}) \xrightarrow{\varepsilon \rightarrow 0} g(x^*)$  and let  $x \notin \mathcal{B}_\delta(x^{\varepsilon*})$ . Then,  $|x - x^*| \geq |x - x^{\varepsilon*}| - |x^{\varepsilon*} - x^*| \geq \delta - |x^{\varepsilon*} - x^*| \geq \frac{\delta}{2} = \delta'$  when  $\varepsilon$  is small enough, since  $|x^{\varepsilon*} - x^*| \xrightarrow{\varepsilon \rightarrow 0} 0$ . Therefore, there exists  $c_{\delta'}$  such that  $f(x) - f(x^*) \geq c_{\delta'}$ . Thus,

$$\begin{aligned} f^\varepsilon(x) - f^\varepsilon(x^{\varepsilon*}) &= f^\varepsilon(x) - f(x) + f(x) - f(x^*) + f(x^*) - f^\varepsilon(x^{\varepsilon*}) \\ &\geq \varepsilon g^\varepsilon(x) + c_{\delta'} + f(x^*) - f(x^{\varepsilon*}) - \varepsilon g^\varepsilon(x^{\varepsilon*}) \\ &\geq \frac{c_{\delta'}}{2} \end{aligned}$$

where the last inequality holds when  $\varepsilon$  is small enough because  $g^\varepsilon(x) \geq m$  and  $\varepsilon g^\varepsilon(x^{\varepsilon*}) \xrightarrow{\varepsilon \rightarrow 0} 0$ . Hence,  $f^\varepsilon$  satisfies conditions **C** uniformly in  $\varepsilon$ .  $\square$

We now make the following assumption on the model.

**Assumption  $\mathbf{L}^\varepsilon$ .** For all  $k \geq 0$ :

- $\sup_{(x, x') \in E^2} \{q_k(x, x')\} < \infty$ ;
- $\sigma_k^+ := \sup_{x \in E} \{\sigma_k(x)\} < \infty$  and  $\sigma_k^- := \inf_{x \in E} \{\sigma_k(x)\} > 0$ ;
- the function  $x \mapsto \frac{1}{2\sigma_k(x)^2} |Y_k - H_k(x)|^2$  admits a minimum at  $x_k^*$  and satisfies condition **C**;
- $q_0(\cdot) > 0$  in a neighbourhood of  $x_0^*$  and  $q_k(x_{k-1}^*, \cdot) > 0$  for all  $k \geq 1$ ;
- for all  $x' \in E$ , the function  $x \mapsto q_k(x, x')$  is continuous;
- for all Borel subset  $B \subset E$ , the function  $x \mapsto \int_B q_k(x, x') dx'$  is continuous.

The restrictive assumption on the likelihood function which has been announced is the third bullet of assumption **L** $^\varepsilon$ .

The maximum likelihood estimator (MLE) is

$$x_k^{\varepsilon*} = \operatorname{argmax}_{x \in E} \{g_k^\varepsilon(x)\} = \operatorname{argmin}_{x \in E} \{\ell_k^\varepsilon(x)\}.$$

Since  $\ell_k^\varepsilon(x) = \frac{1}{2\sigma_k(x)^2} |Y_k - H_k(x)|^2 + d\varepsilon \log \sigma_k(x)$  and  $\inf_{x \in E} \{\log \sigma_k(x)\} = \log \sigma_k^- > -\infty$ ,

assumption  $\mathbf{L}^\varepsilon$  and lemma 4.1.1 imply that  $x_k^{\varepsilon*}$  exists and verifies  $|x_k^{\varepsilon*} - x_k^*| \rightarrow 0$  as  $\varepsilon \rightarrow 0$ .

**Proposition 4.1.2.** *If assumption  $\mathbf{L}^\varepsilon$  is verified, then*

- for all  $k \geq 1$  and for all  $x \in E$ ,  $p_k^{\varepsilon-}(x) \rightarrow q_k(x_{k-1}^*, x)$  as  $\varepsilon \rightarrow 0$ ,
- for all  $k \geq 0$  and for all bounded continuous function  $\phi$ ,  $\int_E \phi(x) \mu_k^\varepsilon(dx) \rightarrow \phi(x_k^*)$  as  $\varepsilon \rightarrow 0$ , i.e.  $\mu_k^\varepsilon \Rightarrow \delta_{x_k^*}$  as  $\varepsilon \rightarrow 0$ .

*Proof.* Let  $k \geq 1$  and suppose that  $\mu_{k-1}^\varepsilon \Rightarrow \delta_{x_{k-1}^*}$ . Then, for all  $x' \in E$ ,

$$p_k^{\varepsilon-}(x') = \int_E q_k(x, x') \mu_{k-1}^\varepsilon(dx) \xrightarrow{\varepsilon \rightarrow 0} \int_E q_k(x, x') \delta_{x_{k-1}^*}(dx) = q_k(x_{k-1}^*, x').$$

Let  $\delta > 0$ . We have that

$$\mu_k^\varepsilon(\mathcal{B}_\delta(x_k^*)^c) = \frac{\int_{\mathcal{B}_\delta(x_k^*)^c} g_k^\varepsilon(x) p_k^{\varepsilon-}(x) dx}{\int_E g_k^\varepsilon(x) p_k^{\varepsilon-}(x) dx}.$$

We can bound from above the numerator of this ratio as

$$\begin{aligned} & \int_{\mathcal{B}_\delta(x_k^*)^c} g_k^\varepsilon(x) p_k^{\varepsilon-}(x) dx \\ &= g_k^\varepsilon(x_k^*) \int_{\mathcal{B}_\delta(x_k^*)^c} \exp\left(-\frac{1}{\varepsilon} \left( \frac{1}{2\sigma_k(x)} |Y_k - H_k(x)|^2 - \frac{1}{2\sigma_k(x_k^*)} |Y_k - H_k(x_k^*)|^2 \right)\right) \\ & \quad \times \exp\left(d \log \frac{\sigma_k(x_k^*)}{\sigma_k(x)}\right) p_k^{\varepsilon-}(x) dx \\ &\leq g_k^\varepsilon(x_k^*) \exp\left(-\frac{c_\delta}{\varepsilon} + d \log \frac{\sigma_k(x_k^*)}{\sigma_k(x)}\right) \int_{\mathcal{B}_\delta(x_k^*)^c} p_k^{\varepsilon-}(x) dx, \end{aligned}$$

for some  $c_\delta > 0$ , where the inequality holds because of assumption  $\mathbf{L}^\varepsilon$ . Let now  $A_\delta$  be the subset of  $E$  defined as

$$A_\delta = \left\{ x \in E : \frac{1}{2\sigma_k(x)^2} |Y_k - H_k(x)|^2 - \frac{1}{2\sigma_k(x_k^*)^2} |Y_k - H_k(x_k^*)|^2 \leq \frac{c_\delta}{2} \right\}.$$

$A_\delta$  contains a nonempty open set, hence it has positive Lebesgue measure. The denom-

inator can be bounded from below as

$$\begin{aligned}
& \int_E g_k^\varepsilon(x') p_k^{\varepsilon-}(x') dx' \\
& \geq \int_{A_\delta} g_k^\varepsilon(x') p_k^{\varepsilon-}(x') dx' \\
& = g_k^\varepsilon(x_k^*) \int_{A_\delta} \exp\left(-\frac{1}{\varepsilon} \left( \frac{1}{2\sigma_k(x)} |Y_k - H_k(x)|^2 - \frac{1}{2\sigma_k(x_k^*)} |Y_k - H_k(x_k^*)|^2 \right)\right) \\
& \quad \times \exp\left(d \log \frac{\sigma_k(x_k^*)}{\sigma_k(x)}\right) p_k^{\varepsilon-}(x) dx \\
& \geq g_k^\varepsilon(x_k^*) \exp\left(-\frac{c_\delta}{2\varepsilon} + d \log \frac{\sigma_k(x_k^*)}{\sigma_k^+}\right) \int_{A_\delta} p_k^{\varepsilon-}(x) dx.
\end{aligned}$$

Therefore, we have that

$$\mu_k^\varepsilon(\mathcal{B}_\delta(x_k^*)^c) \leq \exp\left(-\frac{c_\delta}{2\varepsilon} + d \log \frac{\sigma_k^+}{\sigma_k^-}\right) \frac{\int_{\mathcal{B}_\delta(x_k^*)^c} p_k^{\varepsilon-}(x) dx}{\int_{A_\delta} p_k^{\varepsilon-}(x) dx}.$$

Besides, for all  $B \subset E$ , we have that

$$\int_B p_k^{\varepsilon-}(x) dx = \int_B \int_E q_k(x, x') \mu_{k-1}^\varepsilon(dx) dx' = \int_E v(x) \mu_{k-1}^\varepsilon(dx),$$

where  $v : x \mapsto \int_B q_k(x, x') dx'$  is continuous and bounded. Hence,

$$\int_E v(x) \mu_{k-1}^\varepsilon(x) dx \xrightarrow{\varepsilon \rightarrow 0} \int_E v(x) \delta_{x_{k-1}^*}(dx) = v(x_{k-1}^*) = \int_B q_k(x_{k-1}^*, x') dx'.$$

Taking  $B = \mathcal{B}_\delta(x_k^*)^c$  in the numerator and  $B = A_\delta$  in the denominator yields

$$\frac{\int_{\mathcal{B}_\delta(x_k^*)^c} p_k^{\varepsilon-}(x) dx}{\int_{A_\delta} p_k^{\varepsilon-}(x) dx} \xrightarrow{\varepsilon \rightarrow 0} \frac{\int_{\mathcal{B}_\delta(x_k^*)^c} q_k(x_{k-1}^*, x) dx}{\int_{A_\delta} q_k(x_{k-1}^*, x) dx},$$

where the limit is finite because  $\int_{A_\delta} q_k(x_k^*, x) dx > 0$ . Thus, for all  $\delta > 0$ ,

$$\mu_k^\varepsilon(\mathcal{B}_\delta(x_k^*)^c) \xrightarrow{\varepsilon \rightarrow 0} 0.$$

This implies that  $\mu_k^\varepsilon \xrightarrow{\varepsilon \rightarrow 0} \delta_{x_k^*}$ .

Using the same arguments, we have that

$$\mu_0^\varepsilon(\mathcal{B}_\delta(x_0^*)^c) \leq \exp\left(-\frac{c_\delta}{2\varepsilon} + d \log \frac{\sigma_0^+}{\sigma_0^-}\right) \frac{\int_{\mathcal{B}_\delta(x_0^*)^c} q_0(x) dx}{\int_{A_\delta} q_0(x) dx}.$$

Hence, the recursion is initialized.  $\square$

The maximum a posteriori (MAP) is

$$x_k^\varepsilon = \operatorname{argmax}_{x \in E} \{p_k^\varepsilon(x)\} = \operatorname{argmin}_{x \in E} \{\ell_k^\varepsilon(x) - \varepsilon \log p_k^{\varepsilon-}(x)\}.$$

It follows from (4.1) that  $\log p_k^{\varepsilon-}$  is bounded from above by  $\log \sup_{(x,x') \in E^2} \{q_k(x, x')\}$ , which is finite under assumption  $\mathbf{L}^\varepsilon$  (first bullet). Besides, proposition 4.1.2 and assumption  $\mathbf{L}^\varepsilon$  (fourth bullet) insure that  $\log p_k^{\varepsilon-}(x) \xrightarrow{\varepsilon \rightarrow 0} \log q_k(x_{k-1}^*, x)$  for all  $x \in E$ . Then, lemma 4.1.1 implies that  $x_k^{\varepsilon*}$  exists when  $\varepsilon$  is small enough and verifies  $|x_k^\varepsilon - x_k^*| \xrightarrow{\varepsilon \rightarrow 0} 0$ , so that  $|x_k^\varepsilon - x_k^*| \xrightarrow{\varepsilon \rightarrow 0} 0$ , i.e. the difference between the MLE and the MAP vanishes as  $\varepsilon \rightarrow 0$ .

The observed information matrix of the state  $x \in E$  is defined, when it exists, as

$$J_k^\varepsilon(x) = -(\log g_k^\varepsilon)''(x) - (\log p_k^{\varepsilon-})''(x).$$

## 4.2 Algorithm

In this section, we heuristically derive a recursive filtering algorithm where the prediction and the update step are performed thanks to the Laplace method.

### 4.2.1 Approximation of densities

Suppose we have an approximation  $\hat{p}_{k-1}^\varepsilon$  of the posterior density at time  $k-1$ . Then, according to (4.1), an approximation of the predictor density at time  $k$  must verify

$$\hat{p}_k^{\varepsilon-}(x') \approx \int_E q_k(x, x') \hat{p}_{k-1}^\varepsilon(x) dx. \quad (4.3)$$

The above integral being analytically unknown in general, let us apply the Laplace method on it. It yields

$$\int q_k(x, x') \hat{p}_{k-1}^\varepsilon(x) dx \approx (2\pi)^{d/2} \det[-(\log \hat{p}_{k-1}^\varepsilon)''(\hat{x}_{k-1}^\varepsilon)]^{-1/2} q_k(\hat{x}_{k-1}^\varepsilon, x') \hat{p}_{k-1}^\varepsilon(\hat{x}_{k-1}^\varepsilon)$$

where  $\hat{x}_{k-1}^\varepsilon = \operatorname{argmax}_{x \in E} \{\hat{p}_{k-1}^\varepsilon\} \approx x_{k-1}^\varepsilon$ . The approximation of the predictor density should integrate to 1. Hence, we get after normalization

$$\hat{p}_k^{\varepsilon-}(x') = q_k(\hat{x}_{k-1}^\varepsilon, x'), \quad (4.4)$$

which is an actual density. We will see below that this normalization comes naturally because  $(2\pi)^{d/2} \det[-(\log \hat{p}_{k-1}^\varepsilon)''(\hat{x}_{k-1}^\varepsilon)]^{-1/2} \hat{p}_{k-1}^\varepsilon(\hat{x}_{k-1}^\varepsilon) = 1$ .

According to Bayes' formula (4.2), an approximation of the posterior density must then verify

$$\hat{p}_k^\varepsilon(x') \approx \frac{g_k(x') \hat{p}_k^{\varepsilon-}(x')}{\int_E g_k(x') \hat{p}_k^{\varepsilon-}(x') dx'}. \quad (4.5)$$

The normalizing constant is generally unknown. Applying once again the Laplace method on the denominator of (4.5) and using the approximation of the predictor density (4.4) yields

$$\hat{p}_k^\varepsilon(x') = (2\pi)^{-d/2} \det \left[ -(\log g_k)''(\hat{x}_k^\varepsilon) - \frac{\partial^2 \log q_k}{\partial x'^2}(\hat{x}_{k-1}^\varepsilon, x' = \hat{x}_k^\varepsilon) \right]^{1/2} \frac{g_k(x') q_k(\hat{x}_{k-1}^\varepsilon, x')}{g_k(\hat{x}_k^\varepsilon) q_k(\hat{x}_{k-1}^\varepsilon, \hat{x}_k^\varepsilon)},$$

where  $\hat{x}_k^\varepsilon = \operatorname{argmax}_{x \in E} \{g_k(x') q_k(\hat{x}_{k-1}^\varepsilon, x')\} = \operatorname{argmax}_{x \in E} \{\hat{p}_k^\varepsilon(x')\} \approx x_k^\varepsilon$ . This approximation of the posterior density is then recursively defined and relies on the computation of the approximated MAP  $\hat{x}_k^\varepsilon$  at each time step.

These derivations provide an approximation of the information matrix in the form of

$$\hat{J}_k^\varepsilon = -(\log g_k)'' - (\log \hat{p}_k^{\varepsilon-})'' = -(\log g_k)'' - (\log q_k)''(\hat{x}_{k-1}^\varepsilon, \cdot). \quad (4.6)$$

Hence  $\hat{p}_k^\varepsilon$  can be written as

$$\hat{p}_k^\varepsilon(x') = (2\pi)^{-d/2} \det[\hat{J}_k^\varepsilon(\hat{x}_k^\varepsilon)]^{1/2} \frac{g_k^\varepsilon(x') q_k(\hat{x}_{k-1}^\varepsilon, x')}{g_k^\varepsilon(\hat{x}_k^\varepsilon) q_k(\hat{x}_{k-1}^\varepsilon, \hat{x}_k^\varepsilon)}. \quad (4.7)$$

Note that  $(2\pi)^{d/2} \det[\hat{J}_k^\varepsilon(\hat{x}_k^\varepsilon)]^{-1/2} \hat{p}_k^\varepsilon(\hat{x}_k^\varepsilon) = 1$ , so that, by recursion, the normalization that allows to obtain (4.4) is valid.

### 4.2.2 Approximation of moments

Approximations (4.4) and (4.7) are pointwise approximations of the posterior and the predictor densities that can be recursively computed. However, one is often interested in moments rather than in pointwise expressions of densities.

The posterior expectation and covariance matrix can be expressed as

$$\int x' \hat{p}_k^\varepsilon(x') dx' \quad \text{and} \quad \int x' x'^T \hat{p}_k^\varepsilon(x') dx'.$$

When these integrals are difficult to compute, we can use the fully exponential Laplace approximations for the posterior expectation and variance, which are respectively (when  $d = 1$ ):

$$\mathbb{E}[X_k | Y_{0:k}] \approx x_k^\varepsilon - \frac{1}{2} \frac{(J_k^\varepsilon)'(x_k^\varepsilon)}{J_k^\varepsilon(x_k^\varepsilon)^2}$$

and

$$\mathbb{V}[X_k | Y_{0:k}] \approx J_k(x_k^\varepsilon)^{-1} + \frac{(J_k^\varepsilon)'(x_k^\varepsilon)^2}{J_k^\varepsilon(x_k^\varepsilon)^4} - \frac{1}{2} \frac{(J_k^\varepsilon)''(x_k^\varepsilon)}{J_k^\varepsilon(x_k^\varepsilon)^3},$$

replacing  $J_k^\varepsilon$  by  $\hat{J}_k^\varepsilon$  and evaluating at  $\hat{x}_k^\varepsilon$  in place of  $x_k^\varepsilon$ . The multidimensional versions of these approximations are provided in theorem 3.2.5 in chapter 3, mutatis mutandis.

Besides, the predicted expectation and covariance matrix can be expressed as

$$\int x' \hat{p}_k^{\varepsilon-}(x') dx' \quad \text{and} \quad \int x' x'^T \hat{p}_k^{\varepsilon-}(x') dx'.$$

They cannot be computed with the Laplace method, since  $\hat{p}_k^{\varepsilon-}(x') = q_k(\hat{x}_{k-1}^\varepsilon, x')$  does not fulfill the conditions needed to apply it. However, it is often the case that the density  $q_k$  of the Markov kernel is known sufficiently well so that its moments are known (when the state noise is additive and Gaussian, e.g.).

The algorithm we have described is summarized in algorithm 5. Note that the computation of the approximated posterior density  $\hat{p}_k^\varepsilon$  is not necessary in the recursion. Only the approximated MAP is recursively computed as  $\hat{x}_k^\varepsilon = \underset{x \in E}{\operatorname{argmax}} \{g_k^\varepsilon(x) q_k(\hat{x}_{k-1}^\varepsilon, x)\}$ .

## 4.3 Consistency of the approximations

We show in this section the consistency of the approximated posterior density  $\hat{p}_k^\varepsilon$  as the observation noise intensity  $\varepsilon$  goes to zero under assumption  $\mathbf{L}^{\varepsilon'}$  below. These results are

**Algorithm 5**

- 
- $k = 0$ .
    - Compute  $\hat{x}_0^\varepsilon = \operatorname{argmax}_{x \in E} \{g_0^\varepsilon(x)q_0^\varepsilon(x)\}$ .
    - Compute  $\hat{J}_0^\varepsilon(\hat{x}_0^\varepsilon) = -(\log g_0^\varepsilon)''(\hat{x}_0^\varepsilon) - (\log q_0)''(\hat{x}_0^\varepsilon)$ .
    - Compute  $\hat{p}_0^\varepsilon(x) = (2\pi)^{-d/2} \det \left[ \hat{J}_0^\varepsilon(\hat{x}_0^\varepsilon) \right]^{1/2} \frac{g_0^\varepsilon(x)q_0(x)}{g_0^\varepsilon(\hat{x}_0^\varepsilon)q_0(\hat{x}_0^\varepsilon)}$ .
    - Compute  $\hat{m}_0^\varepsilon$  and  $\hat{P}_0^\varepsilon$ .
  - $k \geq 1$ .
    - Compute  $\hat{p}_k^{\varepsilon-}(x) = q_k(\hat{x}_{k-1}^\varepsilon, x)$ .
    - Compute  $\hat{x}_k^\varepsilon = \operatorname{argmax}_{x \in E} \{g_k^\varepsilon(x)\hat{p}_k^{\varepsilon-}(x)\}$ .
    - Compute  $\hat{J}_k^\varepsilon(\hat{x}_k^\varepsilon) = -(\log g_k^\varepsilon)''(\hat{x}_k^\varepsilon) - (\log \hat{p}_k^{\varepsilon-})''(\hat{x}_k^\varepsilon)$ .
    - Compute  $\hat{p}_k^\varepsilon(x) = (2\pi)^{-d/2} \det \left[ \hat{J}_k^\varepsilon(\hat{x}_k^\varepsilon) \right]^{1/2} \frac{g_k^\varepsilon(x)\hat{p}_k^{\varepsilon-}(x)}{g_k^\varepsilon(\hat{x}_k^\varepsilon)\hat{p}_k^{\varepsilon-}(\hat{x}_k^\varepsilon)}$ .
    - Compute  $\hat{m}_k^\varepsilon$  and  $\hat{P}_k^\varepsilon$ .
- 

essentially an application of the Laplace method, with  $1/\varepsilon$  as an asymptotic parameter.

**Assumption  $\mathbf{L}^{\varepsilon'}$ .** There exists  $\varepsilon_0 > 0$  and there exist sequences of positive numbers  $\{\Delta_k\}_{k \geq 0}$ ,  $\{m_k\}_{k \geq 0}$ ,  $\{p_k\}_{k \geq 0}$ ,  $\{K_k\}_{k \geq 0}$ ,  $\{M_k\}_{k \geq 0}$  and  $\{M'_k\}_{k \geq 0}$ , such that the following conditions hold for all  $\varepsilon \leq \varepsilon_0$  and for all  $k \geq 0$ :

- for all  $x \in \mathcal{B}_{\Delta_{k-1}}(x_{k-1}^{\varepsilon*})$ ,  $\int_E (g_k^\varepsilon(x')q_k(x, x'))^{\varepsilon/\varepsilon_0} dx \leq K_k \varepsilon^{-p_k}$ ;
- $\ell_k^\varepsilon \in C^4(\mathcal{B}_{\Delta_k}(x_k^{\varepsilon*}))$  and  $\log q_k(x, \cdot) \in C^4(\mathcal{B}_{\Delta_k}(x_k^{\varepsilon*}))$  for all  $x \in \mathcal{B}_{\Delta_{k-1}}(x_{k-1}^{\varepsilon*})$ ;
- $\det [(\ell_k^\varepsilon)''(x_k^{\varepsilon*})] \geq m_k$ ;
- for all  $j \in \{0, \dots, 4\}$  and for all  $x' \in \mathcal{B}_{\Delta_k}(x_k^{\varepsilon*})$ ,  $\|(\ell_k^\varepsilon)_k^{(j)}(x')\| \leq M_k$  and  $\left\| \frac{\partial^j \log q_k}{\partial x'^j}(x, x') \right\| \leq M_k$  for all  $x \in \mathcal{B}_{\Delta_{k-1}}(x_{k-1}^{\varepsilon*})$ ;
- for all  $x' \in E$ ,  $\log q_k(\cdot, x') \in C^2(\mathcal{B}_{\Delta_{k-1}}(x_{k-1}^{\varepsilon*}))$ ;

- for all  $x' \in E$ , for all  $j \in \{0, 1, 2\}$  and for all  $x \in \mathcal{B}_{\Delta_{k-1}}(x_{k-1}^{\varepsilon*})$ ,  

$$\left\| \frac{\partial^j \log q_k}{\partial x^j}(x, x') \right\| \leq M'_k.$$

The norm  $\|\cdot\|$  is defined in section 4.2.1. Assumption  $\mathbf{L}^{\varepsilon'}$  is needed to insure that the collection of contrast functions  $\{\ell_k^\varepsilon\}_{\varepsilon>0}$  is Laplace-regular in the sense defined in chapter 3.

As in chapter 3, the first bullet of assumption  $\mathbf{L}^{\varepsilon'}$  is verified, for example, when the Markov kernel is Gaussian, i.e. when the state model is in the form of  $X_k = F_k(X_{k-1}) + V_k$  where  $\{V_k\}_{k \geq 0}$  is a Gaussian white noise.

For all  $k \geq 0$  and for all  $x \in E$ , let

$$h_k^\varepsilon(x) = -\varepsilon (\log g_k^\varepsilon(x) + \log \hat{p}_k^{\varepsilon-}(x)) = \ell_k^\varepsilon(x) - \varepsilon \log q_k(\hat{x}_{k-1}^\varepsilon, x) + \text{constant}$$

so that  $\hat{x}_k^\varepsilon = \operatorname{argmin}_{x' \in E} \{h_k^\varepsilon(x')\}$  and  $\hat{J}_k^\varepsilon(x) = \frac{1}{\varepsilon} (h_k^\varepsilon)''(x)$ .

**Proposition 4.3.1.** *If assumption  $\mathbf{L}^\varepsilon$  is verified, then  $|\hat{x}_k^\varepsilon - x_k^\varepsilon| \rightarrow 0$  as  $\varepsilon \rightarrow 0$  for all  $k \geq 0$ .*

*Proof.* Let  $k \geq 1$  and suppose that  $|\hat{x}_{k-1}^\varepsilon - x_{k-1}^\varepsilon| \xrightarrow{\varepsilon \rightarrow 0} 0$ . Since  $|x_{k-1}^\varepsilon - x_{k-1}^*| \xrightarrow{\varepsilon \rightarrow 0} 0$ , then  $|\hat{x}_{k-1}^\varepsilon - x_{k-1}^*| \xrightarrow{\varepsilon \rightarrow 0} 0$ , so that  $\log q_k(\hat{x}_{k-1}^\varepsilon, \cdot)$  converges pointwise to  $\log q_k(x_{k-1}^*, \cdot)$ . According to assumption  $\mathbf{L}^\varepsilon$ ,  $\sup_{(x, x') \in E^2} \{\log q_k(x, x')\} < \infty$ . Hence, applying lemma 4.1.1 to

$$h_k^\varepsilon(x) = \frac{1}{2\sigma_k(x)^2} |Y_k - H_k(x)|^2 - \varepsilon (\log q_k(\hat{x}_{k-1}^\varepsilon, x) - d \log \sigma_k(x)) + \text{constant}$$

yields  $|\hat{x}_k^\varepsilon - x_k^*| \xrightarrow{\varepsilon \rightarrow 0} 0$ . Since  $|x_k^\varepsilon - x_k^*| \xrightarrow{\varepsilon \rightarrow 0} 0$ , then  $|\hat{x}_k^\varepsilon - x_k^\varepsilon| \xrightarrow{\varepsilon \rightarrow 0} 0$ . The recursion is initialized with the same arguments.  $\square$

**Proposition 4.3.2.** *If assumptions  $\mathbf{L}^\varepsilon$  and  $\mathbf{L}^{\varepsilon'}$  are verified, then*

$$\int_E q_k(x, x') \hat{p}_{k-1}^\varepsilon(x) dx = \hat{p}_k^{\varepsilon-}(x') + \varepsilon r_k^{\varepsilon-}(x')$$

for all  $k \geq 1$ , where

$$\begin{aligned} |r_k^\varepsilon(x')| &\leq \alpha_k^0 \varepsilon^{-d/2} e^{-\frac{c_k}{\varepsilon}} \int_E q_k(x, x') e^{-\frac{1}{\varepsilon_0} h_{k-1}^\varepsilon(x)} dx + \alpha_k q_k(\hat{x}_{k-1}^\varepsilon, x') + \alpha'_k \left\| \frac{\partial q_k}{\partial x}(x = \hat{x}_{k-1}^\varepsilon, x') \right\| \\ &\quad + \alpha''_k \left\| \frac{\partial^2 q_k}{\partial x^2}(x = \hat{x}_{k-1}^\varepsilon, x') \right\| \\ &< \infty \end{aligned}$$

for all sufficiently small  $\varepsilon$ , with  $\alpha_k^0$ ,  $\alpha_k$ ,  $\alpha'_k$  and  $\alpha''_k$  positive constants.

*Proof.* Let  $k \geq 0$  and consider the integral

$$\int_E q_{k+1}(x, x') \hat{p}_k^\varepsilon(x) dx \propto \int_E q_{k+1}(x, x') e^{-\frac{1}{\varepsilon} h_k^\varepsilon(x)} dx.$$

One has just to check that the integrand verifies the conditions (i)–(v) needed to apply the Laplace method (theorem 1.2.1), using the assumptions on the model.

Let  $\Delta'_k > 0$  such that  $\Delta'_k \leq \Delta_k$ . Since  $|\hat{x}_k^\varepsilon - x_k^\varepsilon| \xrightarrow{\varepsilon \rightarrow 0} 0$  (proposition 4.3.1) and  $|x_k^\varepsilon - x_k^{\varepsilon*}| \xrightarrow{\varepsilon \rightarrow 0} 0$ , we have that  $|\hat{x}_k^\varepsilon - x_k^{\varepsilon*}| \xrightarrow{\varepsilon \rightarrow 0} 0$  so that  $\mathcal{B}_{\Delta'_k}(\hat{x}_k^\varepsilon) \subset \mathcal{B}_{\Delta_k}(x_k^{\varepsilon*})$  whenever  $\varepsilon$  is small enough. This is also true at time  $k-1$ , i.e.  $\mathcal{B}_{\Delta'_{k-1}}(\hat{x}_{k-1}^\varepsilon) \subset \mathcal{B}_{\Delta_{k-1}}(x_{k-1}^{\varepsilon*})$ , with  $\Delta'_{k-1} \leq \Delta_{k-1}$ .

Using assumption  $\mathbf{L}^{\varepsilon'}$ , we have that

$$\begin{aligned} \int_E q_{k+1}(x, x') e^{-\frac{1}{\varepsilon} h_k^\varepsilon(x)} dx &= \int_E (g_k^\varepsilon(x) q_k(\hat{x}_{k-1}^\varepsilon, x))^{\varepsilon/\varepsilon_0} q_{k+1}(x, x') dx \\ &\leq \sup_{(x, x') \in E^2} \{q_{k+1}(x, x')\} K_k \varepsilon^{-p_k} \end{aligned}$$

because  $\hat{x}_{k-1}^\varepsilon \in \mathcal{B}_{\Delta'_{k-1}}(\hat{x}_{k-1}^\varepsilon) \subset \mathcal{B}_{\Delta_{k-1}}(x_{k-1}^{\varepsilon*})$  when  $\varepsilon$  is small enough. (Condition (i).)

$\log q_k(\hat{x}_{k-1}^\varepsilon, \cdot)$  converges pointwise to  $\log q_k(x_{k-1}^*, \cdot)$ ,  $\log q_k(\cdot, \cdot)$  is bounded from above w.r.t. its two arguments (assumption  $\mathbf{L}^\varepsilon$ ), and  $\log q_k(\hat{x}_{k-1}^\varepsilon, \hat{x}_k^\varepsilon) \xrightarrow{\varepsilon \rightarrow 0} \log q_k(x_{k-1}^*, x_k^*)$  because  $\log q_k(\cdot, \cdot)$  is continuous w.r.t. its two arguments in a neighbourhood of  $(x_{k-1}^*, x_k^*)$  (assumption  $\mathbf{L}^{\varepsilon'}$ ). Hence, lemma 4.1.1 implies that  $h_k^\varepsilon$  satisfies condition  $\mathbf{C}$  uniformly in  $\varepsilon$  for all sufficiently small  $\varepsilon$ . (Conditions (ii).)

Then, since

$$\ell_k^\varepsilon(\cdot) - \varepsilon \log q_k(x, \cdot) \in C^4(\mathcal{B}_{\Delta_k}(x_k^{\varepsilon*})) \subset C^4(\mathcal{B}_{\Delta'_k}(\hat{x}_k^\varepsilon))$$

for all  $x \in \mathcal{B}_{\Delta_{k-1}}(x_{k-1}^{\varepsilon*})$ , this also holds for all  $x \in \mathcal{B}_{\Delta'_{k-1}}(\hat{x}_{k-1}^\varepsilon)$  and therefore  $h_k^\varepsilon \in$

$C^4(\mathcal{B}_{\Delta'_k}(\hat{x}_k^\varepsilon))$ . With the same argument, we have that, for  $j \in \{0, \dots, 4\}$ ,  $(h_k^\varepsilon)^{(j)}$  is bounded over  $\mathcal{B}_{\Delta'_k}(\hat{x}_k^\varepsilon)$  uniformly in  $\varepsilon$ . (Condition (iii).)

Besides, since  $C^2(\mathcal{B}_{\Delta'_{k-1}}(\hat{x}_{k-1}^\varepsilon)) \subset C^2(\mathcal{B}_{\Delta_{k-1}}(x_{k-1}^{\varepsilon*}))$ , we have that, for all  $x' \in E$ ,  $\log q_k(\cdot, x') \in C^2(\mathcal{B}_{\Delta'_{k-1}}(\hat{x}_{k-1}^\varepsilon))$  and  $\frac{\partial^j \log q_k}{\partial x^j}(\cdot, x')$  is bounded over  $\mathcal{B}_{\Delta'_{k-1}}(\hat{x}_{k-1}^\varepsilon)$  uniformly in  $\varepsilon$  for all  $j \in \{0, 1, 2\}$ . (Conditions (v).)

Lastly, there exists  $m'_k > 0$  such that

$$(h_k^\varepsilon)''(\hat{x}_k^\varepsilon) = (\ell_k^\varepsilon)''(\hat{x}_k^\varepsilon) - \varepsilon \frac{\partial^2 \log q_k}{\partial x^2}(\hat{x}_{k-1}^\varepsilon, x = \hat{x}_k^\varepsilon) \geq m'_k I_d$$

because  $|\hat{x}_k^\varepsilon - x_k^{\varepsilon*}| \xrightarrow{\varepsilon \rightarrow 0} 0$ ,  $\det[(\ell_k^\varepsilon)''(x_k^{\varepsilon*})] > m_k$ , and  $\frac{\partial^2 \log q_k}{\partial x^2}(\cdot, x')$  is bounded over  $\mathcal{B}_{\Delta'_{k-1}}(\hat{x}_{k-1}^\varepsilon)$  uniformly in  $\varepsilon$ . (Condition (iv).)  $\square$

**Proposition 4.3.3.** *If assumptions  $\mathbf{L}^\varepsilon$  and  $\mathbf{L}^{\varepsilon'}$  are verified, then*

$$\int_E g_k^\varepsilon(x') \hat{p}_k^{\varepsilon-}(x') dx' = (2\pi)^{d/2} \det[\hat{J}_k(\hat{x}_k^\varepsilon)]^{-1/2} g_k^\varepsilon(\hat{x}_k^\varepsilon) \hat{p}_k^{\varepsilon-}(\hat{x}_k^\varepsilon) (1 + \varepsilon r_k^\varepsilon)$$

for all  $k \geq 1$ , where

$$|r_k^\varepsilon| \leq \beta_k$$

for all sufficiently small  $\varepsilon$ , with  $\beta_k$  a positive constant.

*Proof.* Here, the integral of interest is

$$\int_E g_k^\varepsilon(x') \hat{p}_k^{\varepsilon-}(x') dx' = \int_E e^{-\frac{1}{\varepsilon} h_k^\varepsilon(x')} dx'.$$

It has been shown in the proof of proposition 4.3.2 that the integrand verifies the conditions needed to apply the Laplace method.  $\square$

## 4.4 Propagation of the approximation error

In this section, we study the propagation of the approximation error over time, under the additional assumption  $\mathbf{L}^{\varepsilon''}$  below concerning the Markov kernel density  $q_k$ .

**Assumption  $\mathbf{L}^{\varepsilon''}$ .**

- For all  $k \geq 0$ ,  $\sup_{(x, x') \in E^2} \left\{ \frac{\partial q_k}{\partial x'}(x, x') \right\} < \infty$  and  $\sup_{(x, x') \in E^2} \left\{ \frac{\partial^2 q_k}{\partial x'^2}(x, x') \right\} < \infty$ ;

- for all  $k \geq 0$  and for all  $x' \in E$ , the functions  $x \mapsto \frac{\partial q_k}{\partial x'}(x, x')$  and  $x \mapsto \frac{\partial^2 q_k}{\partial x'^2}(x, x')$  are continuous;
- for all  $k \geq 0$  and for all  $j \in \{0, 1, 2\}$ , the function  $x' \mapsto \frac{\partial^j q_k}{\partial x'^j}(x, x')$  is continuous at  $x' = x_k^*$  uniformly (w.r.t.  $x$ ) on any compact; i.e., for all  $\nu > 0$  and all compact subset  $K \subset E$ , there exists  $\delta_{\nu, K}$  such that  $|x' - x_k^*| \leq \delta_{\nu, K}$  implies  $\sup_{x \in K} \left\{ \left\| \frac{\partial^j q_k}{\partial x'^j}(x, x') - \frac{\partial^j q_k}{\partial x'^j}(x, x_k^*) \right\| \right\} \leq \nu$ .

Let  $\hat{\mu}_k^\varepsilon$  denote the approximation of the posterior  $\mu_k^\varepsilon$ , i.e. the probability measure having density  $\hat{p}_k^\varepsilon$ . Let  $\|\cdot\|_{\text{TV}}$  be the total variation norm of probability measures.

**Theorem 4.4.1.** *Suppose that assumptions  $\mathbf{L}^\varepsilon$ ,  $\mathbf{L}^{\varepsilon'}$  and  $\mathbf{L}^{\varepsilon''}$  are verified. Then*

$$\|\hat{\mu}_k^\varepsilon - \mu_k^\varepsilon\|_{\text{TV}} \leq \varepsilon A_k$$

for all  $k \geq 0$ , where  $A_k$  is a positive constant. For all  $k > 1$ ,  $A_k$  verifies

$$A_k \leq 2 \frac{\sup_{(x, x') \in E^2} \{q_k(x, x')\}}{q_k(x_{k-1}^*, x_k^*)} A_{k-1} + B_k + s_k^\varepsilon$$

where  $B_k > 0$  is independent of  $\varepsilon$  and  $s_k^\varepsilon \rightarrow 0$  as  $\varepsilon \rightarrow 0$ .

Theorem 4.4.1 states in particular that, for any bounded function  $\phi$ ,

$$\left| \int_E \phi(x') \hat{p}_k^\varepsilon(x') dx' - \int_E \phi(x') p_k^\varepsilon(x') dx' \right| = O(\varepsilon).$$

Before proving it, we establish three useful results below.

**Lemma 4.4.2.** *Suppose that assumptions  $\mathbf{L}^\varepsilon$  and  $\mathbf{L}^{\varepsilon''}$  are verified. Then,*

$$p_k^{\varepsilon-}(x_k^\varepsilon) \rightarrow q_k(x_{k-1}^*, x_k^*),$$

$$(p_k^{\varepsilon-})'(x_k^\varepsilon) \rightarrow \frac{\partial q_k}{\partial x'}(x_{k-1}^*, x_k^*),$$

and

$$(p_k^{\varepsilon-})''(x_k^\varepsilon) \rightarrow \frac{\partial^2 q_k}{\partial x'^2}(x_{k-1}^*, x_k^*)$$

as  $\varepsilon \rightarrow 0$  for all  $k \geq 0$ .

*Proof.* Let  $k \geq 1$  (the case  $k = 0$  is straightforward because  $p_0^{\varepsilon-} = q_0$ ) and let  $j \in \{0, 1, 2\}$ . It follows from assumption  $\mathbf{L}^{\varepsilon''}$  (first bullet) that the function  $p_k^{\varepsilon-}$  is twice differentiable, with  $j$ th derivative given by

$$(p_k^{\varepsilon-})^{(j)}(x') = \int_E \frac{\partial^j q_k}{\partial x'^j}(x, x') \mu_{k-1}^{\varepsilon}(dx).$$

Thus,

$$\begin{aligned} (p_k^{\varepsilon-})^{(j)}(x_k^{\varepsilon}) - \frac{\partial^j q_k}{\partial x'^j}(x_{k-1}^*, x_k^*) &= \int_E \left( \frac{\partial^j q_k}{\partial x'^j}(x, x_k^{\varepsilon}) - \frac{\partial^j q_k}{\partial x'^j}(x, x_k^*) \right) \mu_{k-1}^{\varepsilon}(dx) \\ &\quad + \left( \int_E \frac{\partial^j q_k}{\partial x'^j}(x, x_k^*) \mu_{k-1}^{\varepsilon}(dx) - \frac{\partial^j q_k}{\partial x'^j}(x_{k-1}^*, x_k^*) \right) \end{aligned}$$

where the second term tends to 0 as  $\varepsilon \rightarrow 0$  because  $\mu_{k-1}^{\varepsilon} \xrightarrow[\varepsilon \rightarrow 0]{\delta} \delta_{k-1}^*$  (proposition 4.1.2) and the function  $x \mapsto \int_E \frac{\partial^j q_k}{\partial x'^j}(x, x_k^*)$  is continuous and bounded (assumption  $\mathbf{L}^{\varepsilon}$  for  $j = 0$ , or assumption  $\mathbf{L}^{\varepsilon''}$  for  $j = 1$  or  $2$ ).

Let  $\eta > 0$ .  $\mu_{k-1}^{\varepsilon} \xrightarrow[\varepsilon \rightarrow 0]{\delta} \delta_{k-1}^*$  implies that there exists a compact  $K_{\eta}$  independent of  $\varepsilon$  such that  $\mu_{k-1}^{\varepsilon}(K_{\eta}^c) \leq \eta$  when  $\varepsilon$  is sufficiently small (tightness property).

The first term of the sum above can be bounded as

$$\begin{aligned} &\left\| \int_E \left( \frac{\partial^j q_k}{\partial x'^j}(x, x_k^{\varepsilon}) - \frac{\partial^j q_k}{\partial x'^j}(x, x_k^*) \right) \mu_{k-1}^{\varepsilon}(dx) \right\| \\ &\leq \int_{K_{\eta}} \left\| \frac{\partial^j q_k}{\partial x'^j}(x, x_k^{\varepsilon}) - \frac{\partial^j q_k}{\partial x'^j}(x, x_k^*) \right\| \mu_{k-1}^{\varepsilon}(dx) + \int_{K_{\eta}^c} \left( \frac{\partial^j q_k}{\partial x'^j}(x, x_k^{\varepsilon}) + \frac{\partial^j q_k}{\partial x'^j}(x, x_k^*) \right) \mu_{k-1}^{\varepsilon}(dx) \\ &\leq \sup_{x \in K_{\eta}} \left\{ \left\| \frac{\partial^j q_k}{\partial x'^j}(x, x_k^{\varepsilon}) - \frac{\partial^j q_k}{\partial x'^j}(x, x_k^*) \right\| \right\} + 2 \sup_{(x, x') \in E^2} \left\{ \left\| \frac{\partial^j q_k}{\partial x'^j}(x, x') \right\| \right\} \eta. \end{aligned}$$

The first term of the upper bound tends to 0 as  $\varepsilon \rightarrow 0$ . The second term can be made arbitrarily small since  $\eta > 0$  is arbitrary, which allows to conclude.  $\square$

Proposition 4.4.3 below, which is similar to proposition 3.1.1 in chapter 3, establishes that the difference between the MLE and the MAP vanishes at rate  $\varepsilon$  when  $\varepsilon \rightarrow 0$ .

**Proposition 4.4.3.** *Suppose that assumptions  $\mathbf{L}^{\varepsilon}$ ,  $\mathbf{L}^{\varepsilon'}$  and  $\mathbf{L}^{\varepsilon''}$  are verified. Then*

$$|x_k^{\varepsilon} - x_k^{\varepsilon*}| = O(\varepsilon)$$

as  $\varepsilon \rightarrow 0$ .

*Proof.* Recall  $x_k^\varepsilon = \operatorname{argmax}_{x \in E} \{p_k^\varepsilon(x)\}$ . Thus, for all  $u \in E$ , we have that

$$\begin{aligned}
0 &= -\varepsilon(\log p_k^\varepsilon)'(x_k^\varepsilon) \cdot u \\
&= (\ell_k^\varepsilon)'(x_k^\varepsilon) \cdot u - \varepsilon(\log p_k^{\varepsilon-})'(x_k^\varepsilon) \cdot u \\
&= (\ell_k^\varepsilon)'(x_k^{\varepsilon*}) \cdot u + (\ell_k^\varepsilon)''(x_k^{\varepsilon*}) \cdot (x_k^\varepsilon - x_k^{\varepsilon*}, u) \\
&\quad + \left( \int_0^1 (\ell_k^\varepsilon)'''(x_k^{\varepsilon*} + \theta(x_k^\varepsilon - x_k^{\varepsilon*})) \cdot (x_k^\varepsilon - x_k^{\varepsilon*})^{(2)} (1 - \theta) d\theta \right) \cdot u - \varepsilon(\log p_k^{\varepsilon-})'(x_k^\varepsilon) \cdot u \\
&= (\ell_k^\varepsilon)''(x_k^{\varepsilon*}) \cdot (x_k^\varepsilon - x_k^{\varepsilon*}, u) + \int_0^1 (\ell_k^\varepsilon)'''(x_k^{\varepsilon*} + \theta(x_k^\varepsilon - x_k^{\varepsilon*})) \cdot ((x_k^\varepsilon - x_k^{\varepsilon*})^{(2)}, u) (1 - \theta) d\theta \\
&\quad - \varepsilon(\log p_k^{\varepsilon-})'(x_k^\varepsilon) \cdot u.
\end{aligned}$$

In particular, for  $u = x_k^\varepsilon - x_k^{\varepsilon*}$ ,

$$\begin{aligned}
0 &\leq (\ell_k^\varepsilon)''(x_k^{\varepsilon*}) \cdot (x_k^\varepsilon - x_k^{\varepsilon*})^{(2)} \\
&\leq \frac{1}{2} \sup_{|x - x_k^{\varepsilon*}| \leq |x_k^\varepsilon - x_k^{\varepsilon*}|} \{ \|(\ell_k^\varepsilon)'''(x)\| \} |x_k^\varepsilon - x_k^{\varepsilon*}|^3 + \varepsilon \|(\log p_k^{\varepsilon-})'(x_k^\varepsilon)\| |x_k^\varepsilon - x_k^{\varepsilon*}|.
\end{aligned}$$

Besides,  $(\ell_k^\varepsilon)''(x_k^{\varepsilon*}) \cdot (x_k^\varepsilon - x_k^{\varepsilon*})^{(2)} \geq \frac{m_k}{M_k^{d-1}} |x_k^\varepsilon - x_k^{\varepsilon*}|^2$  (see remark 1.2.5 in chapter 1) and  $\sup_{|x - x_k^{\varepsilon*}| \leq |x_k^\varepsilon - x_k^{\varepsilon*}|} \{ \|(\ell_k^\varepsilon)'''(x)\| \} \leq \sup_{x \in \mathcal{B}_{\Delta_k}(x_k^{\varepsilon*})} \{ \|(\ell_k^\varepsilon)'''(x)\| \} \leq M_k$ . Then,

$$\frac{m_k}{M_k^{d-1}} |x_k^\varepsilon - x_k^{\varepsilon*}| \leq \frac{M_k}{2} |x_k^\varepsilon - x_k^{\varepsilon*}|^2 + \varepsilon \|(\log p_k^{\varepsilon-})'(x_k^\varepsilon)\|,$$

so that

$$\frac{m_k}{M_k^{d-1}} |x_k^\varepsilon - x_k^{\varepsilon*}| \left( 1 - \frac{M_k^d}{2m_k} |x_k^\varepsilon - x_k^{\varepsilon*}| \right) \leq \varepsilon \|(\log p_k^{\varepsilon-})'(x_k^\varepsilon)\|$$

for all sufficiently small  $\varepsilon$ . Lastly,

$$(\log p_k^{\varepsilon-})'(x_k^\varepsilon) = \frac{\frac{\partial q_k}{\partial x'}(x_{k-1}^\varepsilon, x_k^\varepsilon)}{q_k(x_{k-1}^\varepsilon, x_k^\varepsilon)} \xrightarrow{\varepsilon \rightarrow 0} \frac{\frac{\partial q_k}{\partial x'}(x_k^*, x_{k-1}^*)}{q_k(x_{k-1}^*, x_k^*)}$$

according to lemma 4.4.2. Hence,  $|x_k^\varepsilon - x_k^{\varepsilon*}| = O(\varepsilon)$ . □

**Lemma 4.4.4.** *Suppose that assumptions  $\mathbf{L}^\varepsilon$ ,  $\mathbf{L}^{\varepsilon'}$  and  $\mathbf{L}^{\varepsilon''}$  are verified. Then, for any compact  $K \subset E$ ,*

$$\sup_{x \in K} \left\{ \left| \frac{\int_E g_k^\varepsilon(x') q_k(x, x') dx'}{\int_E g_k^\varepsilon(x') p_k^{\varepsilon-}(x') dx'} - \frac{q_k(x, x_k^*)}{q_k(x_{k-1}^*, x_k^*)} \right| \right\} \rightarrow 0$$

as  $\varepsilon \rightarrow 0$  for all  $k \geq 0$ . Besides, for all  $k \geq 0$ ,

$$\lim_{\varepsilon \rightarrow 0} \sup_{x \in E} \left\{ \frac{\int_E g_k^\varepsilon(x') q_k(x, x') dx'}{\int_E g_k^\varepsilon(x') p_k^{\varepsilon-}(x') dx'} \right\} \leq \frac{\sup_{(x, x') \in E^2} \{q_k(x, x')\}}{q_k(x_{k-1}^*, x_k^*)}.$$

*Proof.* The Laplace method (theorem 1.2.1) yields

$$\begin{aligned} \frac{\int_E g_k^\varepsilon(x') q_k(x, x') dx'}{\int_E g_k^\varepsilon(x') p_k^{\varepsilon-}(x') dx'} &= \left( \frac{\det[(\ell_k^\varepsilon)''(x_k^\varepsilon) - \varepsilon(\log p_k^{\varepsilon-})''(x_k^\varepsilon)]}{\det[(\ell_k^\varepsilon)''(x_k^{\varepsilon*})]} \right)^{1/2} \frac{g_k^\varepsilon(x_k^{\varepsilon*})}{g_k^\varepsilon(x_k^\varepsilon) p_k^{\varepsilon-}(x_k^\varepsilon)} \\ &\quad \times \frac{q_k(x, x_k^{\varepsilon*}) + \varepsilon R_k^\varepsilon(x)}{1 + \varepsilon r_k^\varepsilon}, \end{aligned}$$

where  $r_k^\varepsilon = O(1)$  and  $R_k^\varepsilon(x) = O(1)$  for all  $x \in E$  as  $\varepsilon \rightarrow 0$ .

First of all,  $\frac{g_k^\varepsilon(x_k^{\varepsilon*})}{g_k^\varepsilon(x_k^\varepsilon)} = \exp\left(\frac{1}{\varepsilon}(\ell_k^\varepsilon(x_k^\varepsilon) - \ell_k^\varepsilon(x_k^{\varepsilon*}))\right)$  and

$$\ell_k^\varepsilon(x_k^\varepsilon) = \ell_k^\varepsilon(x_k^{\varepsilon*}) + \frac{1}{2} \int_0^1 (\ell_k^\varepsilon)''(x_k^{\varepsilon*} + \theta(x_k^\varepsilon - x_k^{\varepsilon*})) \cdot (x_k^\varepsilon - x_k^{\varepsilon*})^2 (1 - \theta) d\theta$$

so that

$$0 \leq \ell_k^\varepsilon(x_k^\varepsilon) - \ell_k^\varepsilon(x_k^{\varepsilon*}) \leq \frac{1}{4} M_k |x_k^\varepsilon - x_k^{\varepsilon*}|^2$$

when  $\varepsilon$  is sufficiently small. According to proposition 4.4.3,  $|x_k^\varepsilon - x_k^{\varepsilon*}|^2 = O(\varepsilon^2)$ , hence  $\frac{1}{\varepsilon}(\ell_k^\varepsilon(x_k^\varepsilon) - \ell_k^\varepsilon(x_k^{\varepsilon*})) = O(\varepsilon)$ , which implies that  $\frac{g_k^\varepsilon(x_k^{\varepsilon*})}{g_k^\varepsilon(x_k^\varepsilon)} \xrightarrow{\varepsilon \rightarrow 0} 1$ . Thanks to assumption  $\mathbf{L}^{\varepsilon''}$  (third bullet) and to lemma 4.4.2,

$$\begin{aligned} &\sup_{x \in K} \left\{ \left| \frac{g_k^\varepsilon(x_k^{\varepsilon*}) q_k(x, x_k^{\varepsilon*})}{g_k^\varepsilon(x_k^\varepsilon) p_k^{\varepsilon-}(x_k^\varepsilon) (1 + \varepsilon r_k^\varepsilon)} - \frac{q_k(x, x_k^*)}{q_k(x_{k-1}^*, x_k^*)} \right| \right\} \\ &\leq \sup_{x \in K} \left\{ \left| \frac{g_k^\varepsilon(x_k^{\varepsilon*}) q_k(x, x_k^{\varepsilon*})}{g_k^\varepsilon(x_k^\varepsilon) p_k^{\varepsilon-}(x_k^\varepsilon) (1 + \varepsilon r_k^\varepsilon)} - \frac{q_k(x, x_k^{\varepsilon*})}{p_k^{\varepsilon-}(x_k^\varepsilon) (1 + \varepsilon r_k^\varepsilon)} \right| \right\} \\ &\quad + \sup_{x \in K} \left\{ \left| \frac{q_k(x, x_k^{\varepsilon*})}{p_k^{\varepsilon-}(x_k^\varepsilon) (1 + \varepsilon r_k^\varepsilon)} - \frac{q_k(x, x_k^*)}{q_k(x_{k-1}^*, x_k^*)} \right| \right\} \\ &\leq \frac{\sup_{x \in K} \{q_k(x, x_k^{\varepsilon*})\}}{p_k^{\varepsilon-}(x_k^\varepsilon) (1 + \varepsilon r_k^\varepsilon)} \left| 1 - \frac{g_k^\varepsilon(x_k^{\varepsilon*})}{g_k^\varepsilon(x_k^\varepsilon)} \right| + \sup_{x \in K} \left\{ \left| \frac{q_k(x, x_k^{\varepsilon*})}{p_k^{\varepsilon-}(x_k^\varepsilon) (1 + \varepsilon r_k^\varepsilon)} - \frac{q_k(x, x_k^*)}{q_k(x_{k-1}^*, x_k^*)} \right| \right\} \\ &\xrightarrow{\varepsilon \rightarrow 0} 0. \end{aligned}$$

for any compact  $K \subset E$ . (Notice that  $p_k^{\varepsilon-}(x) > 0$  for all  $x \in E$  when  $\varepsilon$  is small enough because  $q_k(x_{k-1}^*, x) > 0$  for all  $x \in E$ .)

Then,  $R_k^\varepsilon(x)$  can be bounded as

$$\begin{aligned} |R_k^\varepsilon(x)| \leq & \alpha_0 \int_E q_k(x, x') e^{-\frac{1}{\varepsilon_0} h_k^\varepsilon(x')} dx' + \alpha |q_k(x, x_k^{\varepsilon*})| + \alpha' \left\| \frac{\partial q_k}{\partial x'}(x, x' = x_k^{\varepsilon*}) \right\| \\ & + \alpha'' \sup_{x' \in E} \left\| \frac{\partial^2 q_k}{\partial x'^2}(x, x') \right\| \end{aligned}$$

for all  $x \in E$  according to theorem 1.2.1. Jensen's inequality yields

$$\int_E q_k(x, x') e^{-\frac{1}{\varepsilon_0} h_k^\varepsilon(x')} dx' = \int_E g_k^\varepsilon(x')^{\varepsilon/\varepsilon_0} q_k(x, x') dx' \leq \left( \int_E g_k^\varepsilon(x') q_k(x, x') dx' \right)^{\varepsilon/\varepsilon_0},$$

and

$$\begin{aligned} \int_E g_k^\varepsilon(x') q_k(x, x') dx' &= [2\pi\varepsilon\sigma_k(x)^2]^{-d/2} \int_E \exp\left(-\frac{1}{2\varepsilon\sigma_k(x')^2} |Y_k - H_k(x')|^2\right) q_k(x, x') dx' \\ &\leq [2\pi\varepsilon(\sigma_k^-)^2]^{-d/2} \end{aligned}$$

so that

$$\left( \int_E g_k^\varepsilon(x') q_k(x, x') dx' \right)^{\varepsilon/\varepsilon_0} \leq \exp\left(-\frac{d\varepsilon}{2\varepsilon_0} \log(2\pi\varepsilon(\sigma_k^-)^2)\right) \xrightarrow{\varepsilon \rightarrow 0} 0.$$

Thus,  $R_k^\varepsilon(x)$  is bounded uniformly in  $x \in E$ .

Besides, it follows from lemma 4.4.2 that

$$\begin{aligned} (\log p_k^{\varepsilon-})''(x_k^\varepsilon) &= \frac{(p_k^{\varepsilon-})''(x_k^\varepsilon)}{p_k^{\varepsilon-}(x_k^\varepsilon)} - \frac{(p_k^{\varepsilon-})'(x_k^\varepsilon)^T (p_k^{\varepsilon-})'(x_k^\varepsilon)}{p_k^{\varepsilon-}(x_k^\varepsilon)^2} \\ &\xrightarrow{\varepsilon \rightarrow 0} \frac{1}{q_k(x_{k-1}^*, x_k^*)} \frac{\partial^2 q_k}{\partial x'^2}(x_{k-1}^*, x_k^*) + \frac{1}{q_k(x_{k-1}^*, x_k^*)^2} \frac{\partial q_k}{\partial x'}(x_{k-1}^*, x_k^*)^T \frac{\partial q_k}{\partial x'}(x_{k-1}^*, x_k^*) \end{aligned}$$

so that  $\varepsilon \|(\log p_k^{\varepsilon-})''(x_k^\varepsilon)\| \xrightarrow{\varepsilon \rightarrow 0} 0$ , which implies that

$$\left| \left( \frac{\det[(\ell_k^\varepsilon)''(x_k^\varepsilon) - \varepsilon(\log p_k^{\varepsilon-})''(x_k^\varepsilon)]}{\det[(\ell_k^\varepsilon)''(x_k^{\varepsilon*})]} \right)^{1/2} - \left( \frac{\det[(\ell_k^\varepsilon)''(x_k^\varepsilon)]}{\det[(\ell_k^\varepsilon)''(x_k^{\varepsilon*})]} \right)^{1/2} \right| \xrightarrow{\varepsilon \rightarrow 0} 0.$$

Let  $u \in \mathbb{R}^d$  such that  $|u| = 1$ . Then,

$$\begin{aligned} (\ell_k^\varepsilon)''(x_k^\varepsilon) \cdot u^{(2)} &= (\ell_k^\varepsilon)''(x_k^{\varepsilon*}) \cdot u^{(2)} + \int_0^1 (\ell_k^\varepsilon)'''(x_k^{\varepsilon*} + \theta(x_k^\varepsilon - x_k^{\varepsilon*})) \cdot (x_k^\varepsilon - x_k^{\varepsilon*}, u, u)(1 - \theta) d\theta \\ &\leq (\ell_k^\varepsilon)''(x_k^{\varepsilon*}) \cdot u^{(2)} + M_k |x_k^\varepsilon - x_k^{\varepsilon*}|, \end{aligned}$$

which implies that  $\left( \frac{\det[(\ell_k^\varepsilon)''(x_k^\varepsilon)]}{\det[(\ell_k^\varepsilon)''(x_k^{\varepsilon*})]} \right)^{1/2} \xrightarrow{\varepsilon \rightarrow 0} 1$ , since  $|x_k^\varepsilon - x_k^{\varepsilon*}| \xrightarrow{\varepsilon \rightarrow 0} 0$ . Therefore,

$$\left( \frac{\det[(\ell_k^\varepsilon)''(x_k^\varepsilon) - \varepsilon(\log p_k^{\varepsilon-})''(x_k^\varepsilon)]}{\det[(\ell_k^\varepsilon)''(x_k^{\varepsilon*})]} \right)^{1/2} \xrightarrow{\varepsilon \rightarrow 0} 1,$$

which concludes the proof of the first part of the lemma.

Lastly,

$$\begin{aligned} \sup_{x \in E} \left\{ \frac{\int_E g_k^\varepsilon(x') q_k(x, x') dx'}{\int_E g_k^\varepsilon(x') p_k^{\varepsilon-}(x') dx'} \right\} &\leq \left( \frac{\det[(\ell_k^\varepsilon)''(x_k^\varepsilon) - \varepsilon(\log p_k^{\varepsilon-})''(x_k^\varepsilon)]}{\det[(\ell_k^\varepsilon)''(x_k^{\varepsilon*})]} \right)^{1/2} \frac{g_k^\varepsilon(x_k^{\varepsilon*})}{g_k^\varepsilon(x_k^\varepsilon) p_k^{\varepsilon-}(x_k^\varepsilon)} \\ &\quad \times \frac{\sup_{(x, x') \in E^2} \{q_k(x, x')\} + \varepsilon \sup_{x \in E} \{|R_k^\varepsilon(x)|\}}{1 + \varepsilon r_k^\varepsilon} \\ &\xrightarrow{\varepsilon \rightarrow 0} \frac{\sup_{(x, x') \in E^2} \{q_k(x, x')\}}{q_k(x_{k-1}^*, x_k^*)}, \end{aligned}$$

which proves the second part of the lemma.  $\square$

*Proof of theorem 4.4.1.* Let  $k \geq 1$ . Suppose that at time  $k - 1$ ,  $\hat{p}_{k-1}^\varepsilon$  can be written as

$$\hat{p}_{k-1}^\varepsilon(x) = p_{k-1}^\varepsilon(x) + \varepsilon a_{k-1}^\varepsilon(x),$$

where  $\int_E |a_{k-1}^\varepsilon(x)| dx \leq A_{k-1} < \infty$  for all sufficiently small  $\varepsilon$ .

According to proposition 4.3.2,

$$\int_E q_k(x, x') \hat{p}_{k-1}^\varepsilon(x) dx = q_k(\hat{x}_{k-1}^\varepsilon, x') + \varepsilon r_k^{\varepsilon-}(x')$$

where

$$\begin{aligned} |r_k^{\varepsilon-}(x')| &\leq \alpha_k^0 \varepsilon^{-d/2} e^{-\frac{c_k}{\varepsilon}} \int_E q_k(x, x') e^{-\frac{1}{\varepsilon_0} h_{k-1}^\varepsilon(x)} dx + \alpha_k q_k(\hat{x}_{k-1}^\varepsilon, x') \\ &\quad + \alpha'_k \left\| \frac{\partial q_k}{\partial x}(x = \hat{x}_{k-1}^\varepsilon, x') \right\| + \alpha''_k \left\| \frac{\partial^2 q_k}{\partial x^2}(x = \hat{x}_{k-1}^\varepsilon, x') \right\| \\ &< \infty \end{aligned}$$

for all sufficiently small  $\varepsilon$ . Therefore,

$$\int_E q_k(x, x') p_{k-1}^\varepsilon(x) dx + \varepsilon \int_E a_{k-1}^\varepsilon(x) q_k(x, x') dx = q_k(\hat{x}_{k-1}^\varepsilon, x') + \varepsilon r_k^{\varepsilon-}(x'),$$

so that

$$\hat{p}_k^{\varepsilon-}(x') = p_k^{\varepsilon-}(x') + \varepsilon a_k^{\varepsilon-}(x') \quad (4.8)$$

where  $a_k^{\varepsilon-}(x') = \int_E a_{k-1}^\varepsilon(x) q_k(x, x') dx - r_k^{\varepsilon-}(x')$ . For all  $x' \in E$ , we have that

$$\begin{aligned} |a_k^{\varepsilon-}(x')| &\leq \left| \int_E a_{k-1}^\varepsilon(x) q_k(x, x') dx \right| + |r_k^{\varepsilon-}(x')| \\ &\leq \int_E |a_{k-1}^\varepsilon(x)| q_k(x, x') dx + \alpha_k^0 \varepsilon^{-d/2} e^{-\frac{c_k}{\varepsilon}} \int_E q_k(x, x') e^{-\frac{1}{\varepsilon_0} h_{k-1}^\varepsilon(x)} dx \\ &\quad + \alpha_k |q_k(\hat{x}_{k-1}^\varepsilon, x')| + \alpha'_k \left\| \frac{\partial q_k}{\partial x}(x = \hat{x}_{k-1}^\varepsilon, x') \right\| + \alpha''_k \left\| \frac{\partial^2 q_k}{\partial x^2}(x = \hat{x}_{k-1}^\varepsilon, x') \right\|. \end{aligned}$$

Then, according to proposition 4.3.3,

$$\int_E g_k^\varepsilon(x') \hat{p}_k^{\varepsilon-}(x') dx' = (2\pi)^{d/2} \det[\hat{J}_k(\hat{x}_k^\varepsilon)]^{-1/2} g_k^\varepsilon(\hat{x}_k^\varepsilon) \hat{p}_k^{\varepsilon-}(\hat{x}_k^\varepsilon) (1 + \varepsilon r_k^\varepsilon)$$

where  $|r_k^\varepsilon| \leq \beta_k$  for all sufficiently small  $\varepsilon$ . Let  $l_k^\varepsilon = \int_E g_k^\varepsilon(x') p_k^{\varepsilon-}(x') dx'$ . Then, from (4.8),

$$\int_E g_k^\varepsilon(x') (p_k^{\varepsilon-}(x') + \varepsilon a_k^{\varepsilon-}(x')) dx' = (2\pi)^{d/2} \det[\hat{J}_k(\hat{x}_k^\varepsilon)]^{-1/2} g_k^\varepsilon(\hat{x}_k^\varepsilon) \hat{p}_k^{\varepsilon-}(\hat{x}_k^\varepsilon) (1 + \varepsilon r_k^\varepsilon),$$

so that

$$(2\pi)^{d/2} \det[\hat{J}_k(\hat{x}_k^\varepsilon)]^{-1/2} g_k^\varepsilon(\hat{x}_k^\varepsilon) \hat{p}_k^{\varepsilon-}(\hat{x}_k^\varepsilon) = \frac{l_k^\varepsilon + \varepsilon b_k^\varepsilon}{1 + \varepsilon r_k^\varepsilon}, \quad (4.9)$$

where  $b_k^\varepsilon = \int_E g_k^\varepsilon(x') a_k^{\varepsilon-}(x') dx'$ .

Then,

$$\hat{p}_k^\varepsilon(x') = (2\pi)^{d/2} \det[\hat{J}_k(\hat{x}_k^\varepsilon)]^{-1/2} \frac{g_k^\varepsilon(x') \hat{p}_k^{\varepsilon-}(x')}{g_k^\varepsilon(\hat{x}_k^\varepsilon) \hat{p}_k^{\varepsilon-}(\hat{x}_k^\varepsilon)} \quad (4.10)$$

yields

$$\hat{p}_k^\varepsilon(x') = \frac{g_k^\varepsilon(x') (p_k^{\varepsilon-}(x') + \varepsilon a_k^{\varepsilon-}(x'))}{l_k^\varepsilon + \varepsilon b_k^\varepsilon} (1 + \varepsilon r_k^\varepsilon).$$

We get

$$\hat{p}_k^\varepsilon(x') = p_k^\varepsilon(x') + \varepsilon a_k^\varepsilon(x')$$

where

$$\begin{aligned}
\varepsilon a_k^\varepsilon(x') &= \hat{p}_k^\varepsilon(x') - \frac{g_k^\varepsilon(x') p_k^{\varepsilon-}(x')}{\int g_k^\varepsilon(x') p_k^{\varepsilon-}(x') dx'} \\
&= \frac{(g_k^\varepsilon(x') p_k^{\varepsilon-}(x') + \varepsilon a_k^{\varepsilon-}(x') g_k^\varepsilon(x')) (1 + \varepsilon r_k^\varepsilon)}{l_k^\varepsilon + \varepsilon b_k^\varepsilon} - \frac{g_k^\varepsilon(x') p_k^{\varepsilon-}(x') (l_k^\varepsilon + \varepsilon b_k^\varepsilon)}{l_k^\varepsilon (l_k^\varepsilon + \varepsilon b_k^\varepsilon)} \\
&= \frac{g_k^\varepsilon(x') p_k^{\varepsilon-}(x') + \varepsilon r_k^\varepsilon g_k^\varepsilon(x') p_k^{\varepsilon-}(x') + \varepsilon a_k^{\varepsilon-}(x') g_k^\varepsilon(x') + \varepsilon^2 r_k^\varepsilon a_k^{\varepsilon-}(x') g_k^\varepsilon(x')}{l_k^\varepsilon + \varepsilon b_k^\varepsilon} \\
&\quad - \frac{g_k^\varepsilon(x') p_k^{\varepsilon-}(x') + \varepsilon b_k^\varepsilon p_k^\varepsilon(x')}{l_k^\varepsilon + \varepsilon b_k^\varepsilon} \\
&= \varepsilon \frac{a_k^{\varepsilon-}(x') g_k^\varepsilon(x') + r_k^\varepsilon g_k^\varepsilon(x') p_k^{\varepsilon-}(x') - b_k^\varepsilon p_k^\varepsilon(x') + \varepsilon r_k^\varepsilon a_k^{\varepsilon-}(x') g_k^\varepsilon(x')}{l_k^\varepsilon + \varepsilon b_k^\varepsilon}.
\end{aligned}$$

The triangular inequality yields

$$\begin{aligned}
\int_E |a_k^\varepsilon(x')| dx' &\leq \frac{\int_E |a_k^{\varepsilon-}(x')| g_k^\varepsilon(x') dx' + |b_k^\varepsilon| + |r_k^\varepsilon| (l_k^\varepsilon + \varepsilon \int_E |a_k^{\varepsilon-}(x')| g_k^\varepsilon(x') dx')}{|l_k^\varepsilon + \varepsilon b_k^\varepsilon|} \\
&\leq \left( 2 \frac{\int_E |a_k^{\varepsilon-}(x')| g_k^\varepsilon(x') dx'}{l_k^\varepsilon} + \beta_k \left( 1 + \varepsilon \frac{\int_E |a_k^{\varepsilon-}(x')| g_k^\varepsilon(x') dx'}{l_k^\varepsilon} \right) \right) \\
&\quad \times \left( 1 - \varepsilon \frac{\int_E |a_k^{\varepsilon-}(x')| g_k^\varepsilon(x') dx'}{l_k^\varepsilon} \right)^{-1}, \tag{4.11}
\end{aligned}$$

where the second inequality holds because  $|b_k^\varepsilon| \leq \int_E |a_k^{\varepsilon-}(x')| g_k^\varepsilon(x') dx'$ . We have

$$\begin{aligned}
&\int_E |a_k^{\varepsilon-}(x')| g_k^\varepsilon(x') dx' \\
&\leq \int_E \int_E |a_{k-1}^\varepsilon(x)| g_k^\varepsilon(x') q_k(x, x') dx dx' + \int_E |r_k^{\varepsilon-}(x')| g_k^\varepsilon(x') dx' \\
&\leq \int_E \int_E |a_{k-1}^\varepsilon(x)| g_k^\varepsilon(x') q_k(x, x') dx dx' + \alpha_k^0 \varepsilon^{-d/2} e^{-\frac{c_k}{\varepsilon}} \int_E \int_E g_k^\varepsilon(x') q_k(x, x') e^{-\frac{1}{\varepsilon_0} h_{k-1}^\varepsilon(x)} dx dx' \\
&\quad + \alpha_k \int_E q_k(\hat{x}_{k-1}^\varepsilon, x') g_k^\varepsilon(x') dx' + \alpha'_k \int_E \left\| \frac{\partial q_k}{\partial x}(x = \hat{x}_{k-1}^\varepsilon, x') \right\| g_k^\varepsilon(x') dx' \\
&\quad + \alpha''_k \int_E \left\| \frac{\partial^2 q_k}{\partial x^2}(x = \hat{x}_{k-1}^\varepsilon, x') \right\| g_k^\varepsilon(x') dx'.
\end{aligned}$$

Notice that

$$\left\| \frac{\partial q_k}{\partial x}(x = \hat{x}_{k-1}^\varepsilon, x') \right\| = \left\| \frac{\partial \log q_k}{\partial x}(x = \hat{x}_{k-1}^\varepsilon, x') \right\| q_k(\hat{x}_{k-1}^\varepsilon, x') \leq M'_k q_k(\hat{x}_{k-1}^\varepsilon, x'),$$

$$\begin{aligned}
& \left\| \frac{\partial^2 q_k}{\partial x^2}(x = \hat{x}_{k-1}^\varepsilon, x') \right\| \\
&= \left( \left\| \frac{\partial^2 \log q_k}{\partial x^2}(x = \hat{x}_{k-1}^\varepsilon, x') \right\| + \left\| \left( \frac{\partial \log q_k}{\partial x}(x = \hat{x}_{k-1}^\varepsilon, x') \right) \right\|^2 \right) q_k(\hat{x}_{k-1}^\varepsilon, x') \\
&\leq (M'_k + M_k'^2) q_k(\hat{x}_{k-1}^\varepsilon, x'),
\end{aligned}$$

and

$$\begin{aligned}
\int_E \int_E g_k^\varepsilon(x') q_k(x, x') e^{-\frac{1}{\varepsilon_0} h_{k-1}^\varepsilon(x)} dx dx' &\leq \sup_{x \in E} \left\{ \int_E g_k^\varepsilon(x') q_k(x, x') dx' \right\} \int_E e^{-\frac{1}{\varepsilon_0} h_{k-1}^\varepsilon(x)} dx \\
&\leq \sup_{x \in E} \left\{ \int_E g_k^\varepsilon(x') q_k(x, x') dx' \right\} K_{k-1} \varepsilon^{-p_{k-1}}
\end{aligned}$$

when  $\varepsilon$  is sufficiently small, where the last inequality follows from assumption  $\mathbf{L}^{\varepsilon'}$  (first bullet).

Thus, for all sufficiently small  $\varepsilon > 0$ ,

$$\begin{aligned}
& \int_E |a_k^{\varepsilon-}(x')| g_k^\varepsilon(x') dx' \\
&\leq \int_E |a_{k-1}^\varepsilon(x)| \left( \int_E g_k^\varepsilon(x') q_k(x, x') dx' \right) dx \\
&\quad + \alpha_k^0 \sup_{x \in E} \left\{ \int_E g_k^\varepsilon(x') q_k(x, x') dx' \right\} \varepsilon^{-(d/2+p_{k-1})} K_{k-1} e^{-\frac{c_k}{\varepsilon}} \\
&\quad + (\alpha_k + \alpha'_k M'_k + \alpha''_k (M'_k + M_k'^2)) \int_E g_k^\varepsilon(x') q_k(\hat{x}_{k-1}^\varepsilon, x') dx',
\end{aligned}$$

so that

$$\begin{aligned}
& \frac{\int_E |a_k^{\varepsilon-}(x')| g_k^\varepsilon(x') dx'}{l_k^\varepsilon} \\
&\leq \sup_{x \in E} \left\{ \frac{\int_E g_k^\varepsilon(x') q_k(x, x') dx'}{\int_E g_k^\varepsilon(x') p_k^{\varepsilon-}(x') dx'} \right\} \int_E |a_{k-1}^\varepsilon(x)| dx \\
&\quad + \alpha_k^0 \sup_{x \in E} \left\{ \frac{\int_E g_k^\varepsilon(x') q_k(x, x') dx'}{\int_E g_k^\varepsilon(x') p_k^{\varepsilon-}(x') dx'} \right\} K_{k-1} \varepsilon^{-(d/2+p_{k-1})} e^{-\frac{c_k}{\varepsilon}} \\
&\quad + (\alpha_k + \alpha'_k M'_k + \alpha''_k (M'_k + M_k'^2)) \frac{\int_E g_k^\varepsilon(x') q_k(\hat{x}_{k-1}^\varepsilon, x') dx'}{\int_E g_k^\varepsilon(x') p_k^{\varepsilon-}(x') dx'}.
\end{aligned}$$

Since  $\hat{x}_{k-1}^\varepsilon \xrightarrow{\varepsilon \rightarrow 0} x_{k-1}^*$ ,  $\hat{x}_{k-1}^\varepsilon$  belongs to a compact subset  $K \subset E$  when  $\varepsilon$  is small enough.

Then,

$$\begin{aligned}
& \left| \frac{\int_E g_k^\varepsilon(x') q_k(\hat{x}_{k-1}^\varepsilon, x') dx'}{\int_E g_k^\varepsilon(x') p_k^{\varepsilon-}(x') dx'} - 1 \right| \\
& \leq \left| \frac{\int_E g_k^\varepsilon(x') q_k(\hat{x}_{k-1}^\varepsilon, x') dx'}{\int_E g_k^\varepsilon(x') p_k^{\varepsilon-}(x') dx'} - \frac{q_k(\hat{x}_{k-1}^\varepsilon, x_k^*)}{q_k(x_{k-1}^*, x_k^*)} \right| + \left| \frac{q_k(\hat{x}_{k-1}^\varepsilon, x_k^*)}{q_k(x_{k-1}^*, x_k^*)} - 1 \right| \\
& \leq \sup_{x \in K} \left\{ \left| \frac{\int_E g_k^\varepsilon(x') q_k(x, x') dx'}{\int_E g_k^\varepsilon(x') p_k^{\varepsilon-}(x') dx'} - \frac{q_k(x, x_k^*)}{q_k(x_{k-1}^*, x_k^*)} \right| \right\} + \frac{|q_k(\hat{x}_{k-1}^\varepsilon, x_k^*) - q_k(x_{k-1}^*, x_k^*)|}{q_k(x_{k-1}^*, x_k^*)},
\end{aligned}$$

so that

$$\frac{\int_E g_k^\varepsilon(x') q_k(\hat{x}_{k-1}^\varepsilon, x') dx'}{\int_E g_k^\varepsilon(x') p_k^{\varepsilon-}(x') dx'} \xrightarrow{\varepsilon \rightarrow 0} 1$$

thanks to lemma 4.4.4 and to the continuity of the function  $x \mapsto q_k(x, x_k^*)$  (assumption  $\mathbf{L}^\varepsilon$ ). Besides, lemma 4.4.4 also insures that

$$\lim_{\varepsilon \rightarrow 0} \sup_{x \in E} \left\{ \frac{\int_E g_k^\varepsilon(x') q_k(x, x') dx'}{\int_E g_k^\varepsilon(x') p_k^{\varepsilon-}(x') dx'} \right\} \leq \frac{\sup_{(x, x') \in E^2} \{q_k(x, x')\}}{q_k(x_{k-1}^*, x_k^*)}.$$

Hence, for any constant  $\nu > 0$ ,

$$\frac{\int_E |a_k^{\varepsilon-}(x')| g_k^\varepsilon(x') dx'}{l_k^\varepsilon} \leq \frac{\sup_{(x, x') \in E^2} \{q_k(x, x')\}}{q_k(x_{k-1}^*, x_k^*)} A_{k-1} + \alpha_k + \alpha'_k M'_k + \alpha''_k (M'_k + M_k'^2) + \nu$$

for all sufficiently small  $\varepsilon$ . Finally, from (4.11) and the above inequality, for any  $\eta > 0$ , we have

$$\int_E |a_k^\varepsilon(x')| dx' \leq 2 \left( \frac{\sup_{(x, x') \in E^2} \{q_k(x, x')\}}{q_k(x_{k-1}^*, x_k^*)} A_{k-1} + \alpha_k + \alpha'_k M'_k + \alpha''_k (M'_k + M_k'^2) \right) + \beta_k + \eta$$

when  $\varepsilon$  is small enough.

At time  $k = 0$ , we simply have  $\hat{p}_0^\varepsilon(x) = p_0^\varepsilon(x)(1 + \varepsilon r_0^\varepsilon)$ , so that  $\hat{p}_0^\varepsilon(x) = p_0^\varepsilon(x) + \varepsilon a_0^\varepsilon(x)$ , with  $a_0^\varepsilon(x) = r_0^\varepsilon p_0^\varepsilon(x)$ . Then,  $\int_E |a_0^\varepsilon(x)| dx = |r_0^\varepsilon| \leq \beta_0$  and the recursion is initialized.  $\square$

## Conclusion

We propose in this chapter a recursive filtering algorithm for nonlinear models, exclusively based on the Laplace method. We are able to study the consistency of the approximated posterior density and the propagation of the approximation error over

time, under the restrictive assumption that, any time  $k$ , the likelihood function  $g_k^\varepsilon$  admits a unique global maximum when the observation noise intensity  $\varepsilon$  is small enough. This assumption is often not fulfilled in practice. For example, in bearings-only target tracking, the likelihood at time  $k$  attains its maximum over a linear subspace of the state space (the target state being its position and velocity), when both the target and the observer obey to a rectilinear uniform motion [Joannides and Le Gland (2002)]. The bearings-only tracking problem is addressed in chapter 8.

In the next chapter, we propose a way to overcome the restrictive assumption on the likelihood, which consists in considering that the state dynamics noise is of the same order  $\varepsilon$  as the observation noise.

## Chapter 5

# The Laplace method and Kalman filtering in dynamic models with small dynamics noise

In chapter 4, we have proposed a Bayesian filtering algorithm based uniquely on the Laplace method. The Laplace method is used to compute the integrals involved in the prediction step and in the update step. In order to study the asymptotic behaviour of this filter, we have considered a small noise observation model with the likelihood admitting an unique global maximum at each time step.

To overcome this unrealistic assumption on the likelihood, one can consider that the state dynamics model noise is of the same order as the observation noise, which is the choice we do in this chapter. For such models, we propose a novel filtering algorithm, where the prediction step is done like in the extended Kalman filter and the update step is performed thanks to the Laplace method. This algorithm is close but different from the Laplace Gaussian filter proposed by [Koyama et al. \(2010\)](#) and it is named the Kalman Laplace filter (KLF).

The small noise model we consider is defined in section 5.1, the KLF is presented in section 5.2 and the Laplace Gaussian filter of Koyama et al. is presented in section 5.3.

### 5.1 Model set-up

Let  $\varepsilon > 0$  be a (small) parameter. We consider the state-space model defined by the state process  $\{X_k\}_{k \geq 0}$  and the observation process  $\{Y_k\}_{k \geq 0}$ . The hidden state dynamics

is

$$X_k = F_k(X_{k-1}) + \sqrt{\varepsilon} V_k \quad (5.1)$$

in  $\mathbb{R}^d$ , where  $\{V_k\}_{k \geq 1}$  is a Gaussian white noise such that  $V_k \sim \mathcal{N}(0, \Sigma_k)$ , with invertible matrix  $\Sigma_k$ . The Markov kernel is

$$\mathbb{P}[X_k \in dx' | X_{k-1} = x] = Q_k^\varepsilon(x, dx')$$

and its density is

$$q_k^\varepsilon(x, x') = |(2\pi\varepsilon)^d \det \Sigma_k|^{-1/2} \exp \left( -\frac{1}{2\varepsilon} (x' - F_k(x))^T \Sigma_k^{-1} (x' - F_k(x)) \right),$$

with the initial condition  $X_0 \sim Q_0$ . The observation model is defined by the sequence of likelihood functions  $\{g_k^\varepsilon : \mathbb{R}^d \rightarrow \mathbb{R}^+\}_{k \geq 0}$ , parametrized by  $\varepsilon$ .  $\varepsilon$  is typically the observation noise intensity, as in chapter 4.

At time  $k$ , let  $\mu_k^\varepsilon$  and  $\mu_k^{\varepsilon-}$  be respectively the posterior and the predictor,  $p_k^\varepsilon(x)$  and  $p_k^{\varepsilon-}(x)$  their densities.

In the next section, we heuristically derive a filtering algorithm based on Kalman filtering and on the Laplace method. We suppose that the quantities involved in this algorithm exist, namely the maximum a posteriori (MAP) and the observed information matrix evaluated at the MAP.

## 5.2 The Kalman Laplace filter

Let  $k \geq 1$  and  $\hat{p}_{k-1}^\varepsilon(x)$  be a pointwise approximation of the posterior density at time  $k-1$ . We aim at deriving  $\hat{p}_k^\varepsilon(x)$ .

We first justify why it is very hard to use the Laplace method for the prediction step in this set-up, so that we proceed like in the extended Kalman filter instead.

The predictor density must verify

$$\hat{p}_k^{\varepsilon-}(x') \approx \int_E q_k^\varepsilon(x, x') \hat{p}_{k-1}^\varepsilon(x) dx. \quad (5.2)$$

Here, the small parameter  $\varepsilon$  intervenes in both the Markov kernel density and the posterior density at previous time  $k-1$ . This is different from equation (4.3) in chapter 4, where the Markov kernel does not depend on  $\varepsilon$ . Therefore, applying the Laplace

method on integral (5.2) yields

$$\begin{aligned} \hat{p}_k^{\varepsilon-}(x') &= (2\pi)^{d/2} \det \left[ -\frac{\partial^2 \log q_k^{\varepsilon}}{\partial x^2}(\hat{x}_k^{\varepsilon-}(x'), x') - (\log \hat{p}_{k-1}^{\varepsilon})''(\hat{x}_k^{\varepsilon-}(x')) \right]^{-1/2} \\ &\quad \times q_k^{\varepsilon}(\hat{x}_k^{\varepsilon-}(x'), x') \hat{p}_{k-1}^{\varepsilon}(\hat{x}_k^{\varepsilon-}(x')) \end{aligned}$$

where  $\hat{x}_k^{\varepsilon-}(x') = \operatorname{argmax}_{x \in E} \{q_k^{\varepsilon}(x, x') \hat{p}_{k-1}^{\varepsilon}(x)\}$ , which is a much more complicated expression than (4.4) in chapter 4. Then, from Bayes' rule, the approximated posterior density verifies

$$\hat{p}_k^{\varepsilon}(x') \approx \frac{g_k^{\varepsilon}(x') \hat{p}_k^{\varepsilon-}(x')}{\int_E g_k^{\varepsilon}(x') \hat{p}_k^{\varepsilon-}(x') dx'}. \quad (5.3)$$

Applying once again the Laplace method to compute the denominator of (5.3) yields

$$\int_E g_k^{\varepsilon}(x') \hat{p}_k^{\varepsilon-}(x') dx' \approx (2\pi)^{d/2} \det \left[ -(\log g_k^{\varepsilon})''(\hat{x}_k^{\varepsilon}) - (\log \hat{p}_k^{\varepsilon-})''(\hat{x}_k^{\varepsilon}) \right]^{-1/2} g_k^{\varepsilon}(\hat{x}_k^{\varepsilon}) \hat{p}_k^{\varepsilon-}(\hat{x}_k^{\varepsilon}) \quad (5.4)$$

where  $\hat{x}_k^{\varepsilon} = \operatorname{argmax}_{x' \in E} \{g_k^{\varepsilon}(x') \hat{p}_k^{\varepsilon-}(x')\}$ . The dependency of  $\hat{x}_k^{\varepsilon-}(x')$  in  $x'$  makes that maximization and differentiation are difficult to perform, so that there is no simple expression for  $\hat{x}_k^{\varepsilon}$  and  $(\log \hat{p}_k^{\varepsilon-})''(\hat{x}_k^{\varepsilon-})$ . It is consequently hard to compute (5.4).

These difficulties are overcome when computing integral (5.2) like in the extended Kalman filter. Let  $\varphi_{m,Q}$  denote the Gaussian measure with expectation  $m$  and covariance matrix  $Q$  and let

$$\hat{x}_{k-1}^{\varepsilon} = \operatorname{argmax}_{x \in E} \{\hat{p}_{k-1}^{\varepsilon}(x)\}$$

and

$$\hat{J}_{k-1}^{\varepsilon}(x) = -(\log \hat{p}_{k-1}^{\varepsilon})''(x).$$

Suppose that the posterior at time  $k-1$  is approximately Gaussian, with expectation  $\hat{m}_{k-1}^{\varepsilon}$  and covariance matrix  $\hat{P}_{k-1}^{\varepsilon}$ . Then, the integral (5.2) is approximately equal to

$$\int_E \exp \left( -\frac{1}{2\varepsilon} (x' - F_k(x))^T \Sigma_k^{-1} (x' - F_k(x)) - \frac{1}{2} (x - \hat{m}_{k-1}^{\varepsilon})^T [\hat{P}_{k-1}^{\varepsilon}]^{-1} (x - \hat{m}_{k-1}^{\varepsilon}) \right) dx$$

up to a normalizing constant. Therefore, we define

$$\hat{\mu}_k^{\varepsilon-} = \varphi_{\hat{m}_k^{\varepsilon-}, \hat{P}_k^{\varepsilon-}}$$

where

$$\hat{m}_k^{\varepsilon-} = F_k(\hat{m}_{k-1}^{\varepsilon})$$

and

$$\hat{P}_k^{\varepsilon-} = F_k'(\hat{m}_{k-1}^{\varepsilon})\hat{P}_{k-1}^{\varepsilon}F_k'(\hat{m}_{k-1}^{\varepsilon})^T + \varepsilon\Sigma_k$$

as the approximation of the predictor. This is identical to the prediction step in the extended Kalman filter algorithm (see [Arulampalam et al. (2002)], e.g.). To approximate the posterior density, we now use the Laplace method to compute the denominator of (5.3). It yields

$$\int_E g_k^{\varepsilon}(x')\hat{\mu}_k^{\varepsilon-}(dx') \approx (2\pi)^{d/2} \det \left[ \hat{J}_k^{\varepsilon}(\hat{x}_k^{\varepsilon}) \right]^{-1/2} g_k^{\varepsilon}(\hat{x}_k^{\varepsilon})\hat{p}_k^{\varepsilon-}(\hat{x}_k^{\varepsilon})$$

where

$$\hat{x}_k^{\varepsilon} = \operatorname{argmax}_{x \in E} \{g_k^{\varepsilon}(x)\hat{p}_k^{\varepsilon-}(x)\}$$

and

$$\hat{J}_k^{\varepsilon}(x) = -(\log g_k^{\varepsilon})''(x) + \left[ \hat{P}_k^{\varepsilon-} \right]^{-1}.$$

Hence, we define

$$\hat{p}_k^{\varepsilon}(x') = (2\pi)^{-d/2} \det \left[ \hat{J}_k^{\varepsilon}(\hat{x}_k^{\varepsilon}) \right]^{1/2} \frac{g_k^{\varepsilon}(x')\hat{p}_k^{\varepsilon-}(x')}{g_k^{\varepsilon}(\hat{x}_k^{\varepsilon})\hat{p}_k^{\varepsilon-}(\hat{x}_k^{\varepsilon})}$$

as the approximation of the posterior density.

With this approach, it is sufficient to require that  $g_k^{\varepsilon}\hat{p}_k^{\varepsilon-}$  has an unique global maximum, which is less restrictive than assuming that this is the case for  $g_k^{\varepsilon}$  only, as in chapter 4. Indeed, it is very likely that the product of the likelihood  $g_k^{\varepsilon}$  times the Gaussian prior density  $\hat{p}_k^{\varepsilon-}$  admits an unique global maximum, even though  $g_k^{\varepsilon}$  is multimodal.

Using the multidimensional Laplace approximations (3.2) and (3.3) of chapter 3, we let

$$\hat{m}_k^{\varepsilon} = \hat{x}_k^{\varepsilon} - \frac{1}{2}\hat{J}_k^{\varepsilon}(\hat{x}_k^{\varepsilon})^{-1}(\hat{J}_k^{\varepsilon})'(\hat{x}_k^{\varepsilon})^T \operatorname{vec} \left[ \hat{J}_k^{\varepsilon}(\hat{x}_k^{\varepsilon})^{-1} \right]$$

and

$$\begin{aligned}\hat{P}_k^\varepsilon &= \hat{J}_k^\varepsilon(\hat{x}_k^\varepsilon)^{-1} + \frac{1}{2} \hat{J}_k^\varepsilon(\hat{x}_k^\varepsilon)^{-1} (\hat{J}_k^\varepsilon)'(\hat{x}_k^\varepsilon)^T (\hat{J}_k^\varepsilon(\hat{x}_k^\varepsilon)^{-1} \otimes \hat{J}_k^\varepsilon(\hat{x}_k^\varepsilon)^{-1}) (\hat{J}_k^\varepsilon)'(\hat{x}_k^\varepsilon) \hat{J}_k^\varepsilon(\hat{x}_k^\varepsilon)^{-1} \\ &\quad + \frac{1}{2} (I_d \otimes \text{vec}(\hat{J}_k^\varepsilon(\hat{x}_k^\varepsilon)^{-1})^T (\hat{J}_k^\varepsilon)'(\hat{x}_k^\varepsilon)) (\hat{J}_k^\varepsilon(\hat{x}_k^\varepsilon)^{-1} \otimes \hat{J}_k^\varepsilon(\hat{x}_k^\varepsilon)^{-1}) (\hat{J}_k^\varepsilon)'(\hat{x}_k^\varepsilon) \hat{J}_k^\varepsilon(\hat{x}_k^\varepsilon)^{-1} \\ &\quad - \frac{1}{2} \hat{J}_k^\varepsilon(\hat{x}_k^\varepsilon)^{-1} (I_d \otimes \text{vec}(\hat{J}_k^\varepsilon(\hat{x}_k^\varepsilon)^{-1})^T) (\hat{J}_k^\varepsilon)''(\hat{x}_k^\varepsilon) \hat{J}_k^\varepsilon(\hat{x}_k^\varepsilon)^{-1}.\end{aligned}$$

The above computations can then be re-iterated.

The Kalman Laplace filter (KLF) is summarized in algorithm 6. Like in algorithm 5 in chapter 4, the computation of the approximated posterior density  $\hat{p}_k^\varepsilon$  is not necessary in the recursion. However, the MAP and the posterior moments ( $\hat{x}_k^\varepsilon$ ,  $\hat{m}_k^\varepsilon$  and  $\hat{P}_k^\varepsilon$ ) are recursively computed.

---

**Algorithm 6** Kalman Laplace filter.

---

- $k = 0$ .
    - Compute  $\hat{x}_0^\varepsilon = \underset{x \in E}{\text{argmax}} \{g_0^\varepsilon(x)q_0^\varepsilon(x)\}$ .
    - Compute  $\hat{J}_0^\varepsilon(\hat{x}_0^\varepsilon) = -(\log g_0^\varepsilon)''(\hat{x}_0^\varepsilon) - (\log q_0^\varepsilon)''(\hat{x}_0^\varepsilon)$ .
    - Compute  $\hat{p}_0^\varepsilon(x) = (2\pi)^{-d/2} \det [\hat{J}_0^\varepsilon(\hat{x}_0^\varepsilon)]^{1/2} \frac{g_0^\varepsilon(x)q_0^\varepsilon(x)}{g_0^\varepsilon(\hat{x}_0^\varepsilon)q_0^\varepsilon(\hat{x}_0^\varepsilon)}$ .
    - Compute  $\hat{m}_0^\varepsilon$  and  $\hat{P}_0^\varepsilon$ .
  - $k \geq 1$ .
    - Compute  $\hat{m}_k^{\varepsilon-} = F_k(\hat{m}_{k-1}^\varepsilon)$  and  $\hat{P}_k^{\varepsilon-} = F_k'(\hat{m}_{k-1}^\varepsilon) \hat{P}_{k-1}^\varepsilon F_k'(\hat{m}_{k-1}^\varepsilon)^T + \varepsilon \Sigma_k$ .
    - Compute  $\hat{p}_k^{\varepsilon-}(x) \propto \exp \left( -\frac{1}{2} (x - \hat{m}_k^{\varepsilon-})^T [\hat{P}_k^{\varepsilon-}]^{-1} (x - \hat{m}_k^{\varepsilon-}) \right)$ .
    - Compute  $\hat{x}_k^\varepsilon = \underset{x \in E}{\text{argmax}} \{g_k^\varepsilon(x) \hat{p}_k^{\varepsilon-}(x)\}$ .
    - Compute  $\hat{J}_k^\varepsilon(\hat{x}_k^\varepsilon) = -(\log g_k^\varepsilon)''(\hat{x}_k^\varepsilon) + [\hat{P}_k^{\varepsilon-}]^{-1}$ .
    - Compute  $\hat{p}_k^\varepsilon(x) = (2\pi)^{-d/2} \det [\hat{J}_k^\varepsilon(\hat{x}_k^\varepsilon)]^{1/2} \frac{g_k^\varepsilon(x) \hat{p}_k^{\varepsilon-}(x)}{g_k^\varepsilon(\hat{x}_k^\varepsilon) \hat{p}_k^{\varepsilon-}(\hat{x}_k^\varepsilon)}$ .
    - Compute  $\hat{m}_k^\varepsilon$  and  $\hat{P}_k^\varepsilon$ .
-

### 5.3 The Laplace Gaussian filter

The Laplace Gaussian filter (LGF) is a nonlinear filtering algorithm that has been proposed in [Koyama et al. (2010)]. As the KLF, the LGF is based on the Laplace method and on Gaussian approximation. It is presented in algorithm 7 below.

In the LGF, the posterior expectation is approximated thanks to the fully exponential Laplace approximation, see theorem 3.2.2 in chapter 3. Since this method is applicable to compute the expectation of positive functions only, Koyama et al. follow an idea introduced in [Tierney et al. (1989)], which consists in adding a large constant to the function to integrate and then subtract this constant from the approximation of the integral. That is, to approximate each component  $\mathbb{E}[X_{k,i}|Y_{0:k}]$  ( $i \in \{1, \dots, d\}$ ) of the column vector  $\mathbb{E}[X_k|Y_{0:k}]$ , they proceed as follows: they compute the fully exponential Laplace approximation of the expectation  $\mathbb{E}[X_{k,i} + c|Y_{0:k}] = \frac{\int_E (x_i + c) g_k^\varepsilon(x) \hat{p}_k^{\varepsilon,-}(x) dx}{\int_E g_k^\varepsilon(x) \hat{p}_k^{\varepsilon,-}(x) dx}$ , where  $c > 0$  is a large positive constant and  $x_i$  is the  $i$ th component of  $x \in \mathbb{R}^d$ , and then subtract  $c$  from this approximation. The constant  $c$  must be large enough so that  $X_{k,i} + c > 0$  with a (posterior) probability close to one. The  $d$  components of the column vector  $\mathbb{E}[X_k|Y_{0:k}]$  are approximated in the same manner.

More precisely, let  $\hat{p}_k^\varepsilon(x) = g_k^\varepsilon(x) \hat{p}_k^{\varepsilon,-}(x)$  and  $\hat{p}_{k,i}^{\varepsilon,c}(x) = (x_i + c) g_k^\varepsilon(x) \hat{p}_k^{\varepsilon,-}(x)$ . Let  $\hat{x}_k^\varepsilon = \operatorname{argmax}_{x \in E} \{\hat{p}_k^\varepsilon(x)\}$  and  $\hat{x}_{k,i}^{\varepsilon,c} = \operatorname{argmax}_{x \in E} \{\hat{p}_{k,i}^{\varepsilon,c}(x)\}$ . Then,

$$\hat{m}_{k,i}^\varepsilon = \frac{\hat{p}_{k,i}^{\varepsilon,c}(\hat{x}_{k,i}^{\varepsilon,c})}{\hat{p}_k^\varepsilon(\hat{x}_k^\varepsilon)} \left( \frac{\det[-(\log \hat{p}_k^\varepsilon)''(\hat{x}_k^\varepsilon)]}{\det[-(\log \hat{p}_{k,i}^{\varepsilon,c})''(\hat{x}_{k,i}^{\varepsilon,c})]} \right)^{1/2} - c \quad (5.5)$$

and the approximation of the posterior expectation is  $\hat{m}_k^\varepsilon = \begin{pmatrix} \hat{m}_{k,1}^\varepsilon & \dots & \hat{m}_{k,d}^\varepsilon \end{pmatrix}^T$ .

Moreover, the approximation of the posterior covariance matrix is

$$\hat{P}_k^\varepsilon = -(\log \hat{p}_k^\varepsilon)''(\hat{x}_k^\varepsilon)^{-1} \quad (5.6)$$

in the LGF.

Lastly, the integral (5.7) defining the predictor approximation in the LGF is computed thanks to numerical approximation or to some asymptotic expansion (see [Koyama et al. (2010)] and its appendix for further details).

---

**Algorithm 7** Laplace Gaussian filter.

---

- $k = 0$ .

- Compute  $\hat{m}_0^\varepsilon$  and  $\hat{P}_0^\varepsilon$ .

- $k \geq 1$ .

- Compute

$$\begin{aligned} \hat{p}_k^{\varepsilon-}(x') &= |(2\pi)^d \det \hat{P}_{k-1}^\varepsilon|^{-1/2} \\ &\quad \times \int_E q_k^\varepsilon(x, x') \exp \left( -\frac{1}{2} (x - \hat{m}_{k-1}^\varepsilon)^T [\hat{P}_{k-1}^\varepsilon]^{-1} (x - \hat{m}_{k-1}^\varepsilon) \right) dx. \end{aligned} \quad (5.7)$$

- Compute  $\hat{m}_k^\varepsilon$  and  $\hat{P}_k^\varepsilon$  thanks to formulas (5.5) and (5.6).

---

## Conclusion

We propose in this chapter a recursive filtering algorithm, the Kalman Laplace filter (KLF), where the prediction step is performed as in the extended Kalman filter and the update step is done thanks to the Laplace method.

The KLF is close to the LGF proposed in [Koyama et al. (2010)]. The main differences are that, unlike the LGF, the KLF provides a pointwise non-Gaussian approximation of the posterior density and that the posterior expectation and covariance matrix are computed in a different manner.



## Part III

# Particle approximation and the Laplace method in filtering algorithms



# Chapter 6

## Importance sampling and the Laplace method in static models

Several authors have associated the Laplace method and Monte Carlo sampling for Bayesian estimation. In importance sampling, authors propose to design a Gaussian sampling distribution with the same mode and the same curvature around the mode as the distribution of interest [Pinheiro and Bates (1995); Kuk (1999); Strasser (2002); Jungbacker and Koopman (2007); Kleppe and Skaug (2012)]. That is, if  $p(x)$  is the density of the target distribution, the Gaussian distribution  $\mathcal{N}(\hat{x}, -[(\log p)''(\hat{x})]^{-1})$ , where  $\hat{x} = \operatorname{argmax}\{p(x)\}$ , is chosen as a sampling distribution. This strategy is called Laplace importance sampling in [Kuk (1999)] and [Kleppe and Skaug (2012)]. The Laplace method has also been applied in Gibbs sampling in [Guihenneuc-Jouyaux and Rousseau (2005)], in order to speed up computation.

The strategy we propose in this thesis consists in applying a deterministic affine transformation on the prior, this transformation being based on the Laplace method. The coefficients of the transformation are the multidimensional Laplace approximations of the posterior expectation and covariance matrix. The transformed distribution is taken as a sampling distribution in importance sampling. Simulation results show that the estimation is improved when sampling from this transformed distribution rather than from the prior.

The principle of our strategy is presented in section 6.1. We apply it on a simple triangulation problem in section 6.2.

## 6.1 Importance sampling based on the Laplace method

Consider a Bayesian model defined by the likelihood function  $g$  and the prior  $\eta$ , with density  $q$  on  $E \subset \mathbb{R}^d$ . According to Bayes' rule, the posterior is then

$$\mu(dx) = \frac{g(x)\eta(dx)}{\int_E g(x)\eta(dx)}.$$

We denote its density  $p$ .

The Bayesian estimators we aim at computing are the posterior expectation

$$m = \int_E x \mu(dx)$$

and the posterior covariance matrix

$$P = \int_E xx^T \mu(dx) - mm^T.$$

Importance sampling approximations of  $m$  and  $P$  are classically obtained by sampling from the prior and weighting according to the likelihood, i.e.

$$m^N = \sum_{i=1}^N w^i \xi^i \tag{6.1}$$

and

$$P^N = \sum_{i=1}^N w^i \xi^i (\xi^i)^T - m^N (m^N)^T \tag{6.2}$$

where  $(\xi^1, \dots, \xi^N)$  is an i.i.d. particles sample from  $\eta$  and  $w^i = \frac{g(\xi^i)}{\sum_{i=1}^N g(\xi^i)}$  for all  $i \in \{1, \dots, N\}$ .

Let

$$\hat{x} = \operatorname{argmax}_{x \in E} \{p(x)\} = \operatorname{argmax}_{x \in E} \{g(x)q(x)\}$$

be the maximum a posteriori (MAP) and

$$J(x) = -(\log p)''(x) = -(\log g)''(x) - (\log q)''(x)$$

be the observed information matrix.

It is assumed in the sequel that we can compute the Laplace approximations of

$m$  and  $P$ . That is, the MAP  $\hat{x}$  exists and the posterior density is sufficiently regular around the MAP. From chapter 3, these Laplace approximations are respectively

$$\hat{m} = \hat{x} - \frac{1}{2} \hat{J}^{-1}(\hat{J}')^T \text{vec}[\hat{J}^{-1}] \quad (6.3)$$

and

$$\begin{aligned} \hat{P} = & \hat{J}^{-1} + \frac{1}{2} \hat{J}^{-1}(\hat{J}')^T (\hat{J}^{-1} \otimes \hat{J}^{-1}) \hat{J}' \hat{J}^{-1} + \frac{1}{2} (I_d \otimes \text{vec}(\hat{J}^{-1})^T \hat{J}') (\hat{J}'^{-1} \otimes \hat{J}^{-1}) \hat{J}' \hat{J}^{-1} \\ & - \frac{1}{2} \hat{J}^{-1} (I_d \otimes \text{vec}(\hat{J}^{-1})^T) \hat{J}'' \hat{J}^{-1}, \end{aligned} \quad (6.4)$$

where  $\hat{J} = J(\hat{x})$ ,  $\hat{J}' = J'(\hat{x})$  and  $\hat{J}'' = J''(\hat{x})$

We propose to design a new sampling distribution  $\eta'$  using these approximations.  $\eta'$  is different from the prior and takes into account the likelihood. It is obtained by shifting and scaling the prior so that its expectation and covariance match the approximations given by the Laplace approximations (6.3) and (6.4).

More precisely, let  $m_q = \int x \eta(dx)$  and  $Q = \int x x^T \eta(dx) - m_q m_q^T$  the expectation and covariance of the prior. Then, we define

$$T(x) = \hat{P}^{1/2} Q^{-1/2} (x - m_q) + \hat{m} \quad (6.5)$$

as an affine transformation on  $\mathbb{R}^d$  whose coefficients are based on the Laplace approximations (6.3) and (6.4). The sampling distribution we propose to use is

$$\eta' = \eta \circ T^{-1}, \quad (6.6)$$

so that its density is

$$\begin{aligned} q'(x) &= q(T^{-1}(x)) |\det [(T^{-1})'(x)]| \\ &= q\left(Q^{1/2} \hat{P}^{-1/2} (x - \hat{m}) + m_q\right) \det Q^{1/2} \det \hat{P}^{-1/2}, \end{aligned}$$

since  $T^{-1}(x) = Q^{1/2} \hat{P}^{-1/2} (x - \hat{m}) + m_q$ .

Our goal is to improve the quality of the approximation of  $m$  and  $P$  when sampling from  $\eta'$  instead of  $\eta$ . We expect the particles sampled from  $\eta'$  to be more efficiently weighted since  $\eta'$  depends on the likelihood model, i.e. on the observations. In that

case, the importance sampling approximations of  $m$  and  $P$  are respectively

$$\hat{m}^N = \sum_{i=1}^N w^i \xi^i \quad (6.7)$$

and

$$\hat{P}^N = \sum_{i=1}^N w^i \xi^i (\xi^i)^T - \hat{m}^N (\hat{m}^N)^T \quad (6.8)$$

where  $(\xi^1, \dots, \xi^N)$  is an i.i.d. particles sample from  $\eta'$  (6.6) and  $w^i = \frac{g(\xi^i)q(\xi^i)/q'(\xi^i)}{\sum_{i=1}^N g(\xi^i)q(\xi^i)/q'(\xi^i)}$  for all  $i \in \{1, \dots, N\}$ .

## 6.2 Example: triangulation

To illustrate the efficiency of the method we propose, we consider in this section the same triangulation problem as in [Bui Quang et al. (2012)].

Let  $X = \begin{pmatrix} X_1 & X_2 \end{pmatrix}^T$  be the position of an object in the Cartesian plane.  $n$  sensors observe the azimuth of the target. Each sensor is located at  $\begin{pmatrix} s_{n,1} & s_{n,2} \end{pmatrix}^T$ , for  $k \in \{1, \dots, n\}$ . The  $n$  angular observations are delivered according to the model

$$Y_k = \arctan \frac{X_2 - s_{k,2}}{X_1 - s_{k,1}} + W_k$$

for  $k \in \{1, \dots, n\}$  where  $W_k$  is a Gaussian white noise,  $W_k \sim \mathcal{N}(0, \sigma^2)$ . Hence, the likelihood function is

$$g(x) = |\sqrt{2\pi}\sigma|^{-n} \exp \left( -\frac{1}{2\sigma^2} \sum_{k=0}^n \left( Y_k - \arctan \frac{x_2 - s_{k,2}}{x_1 - s_{k,1}} \right)^2 \right)$$

where  $x = \begin{pmatrix} x_1 & x_2 \end{pmatrix}^T$ .  $X$  follows the prior distribution  $\eta$ , which is set to a Gaussian distribution with expectation  $m_q$  and covariance matrix  $Q$ . We aim at estimating the position of the object given the observations. Notice that, as soon as  $n \geq 2$  and the sensors are located at different positions ( $s_{1,1} \neq s_{2,1}$  or  $s_{1,2} \neq s_{2,2}$ ), the posterior is unimodal so that the MAP can be computed. Besides, the model is sufficiently regular so that the Laplace formulas are computable as well.

We approximate the Bayesian estimators  $m = \mathbb{E}[X|Y_{1:n}]$  and  $P = \mathbb{V}[X|Y_{1:n}]$  thanks

to several methods, yielding different approximations:

- $m^N$  and  $P^N$ : importance sampling approximations (6.1) and (6.2) with  $\eta$  as the sampling distribution;
- $\hat{m}$  and  $\hat{P}$ : Laplace approximations (6.3) and (6.4);
- $\hat{m}^N$  and  $\hat{P}^N$ : importance sampling approximations (6.7) and (6.8) with  $\eta' = \eta \circ T^{-1}$  (6.6) as the sampling distribution, where  $T$  is the transformation (6.5) based on the Laplace approximations;
- $\hat{m}_0^N$  and  $\hat{P}_0^N$ : importance sampling approximations where the sampling distribution is  $\mathcal{N}(\hat{x}, J(\hat{x})^{-1})$ . This corresponds to the Laplace importance sampling technique of [Kuk (1999)] and [Kleppe and Skaug (2012)]. Since the prior is Gaussian in this section, this is equivalent to apply transformation  $T$  when replacing  $\hat{m}$  and  $\hat{P}$  by  $\hat{x}$  and  $J(\hat{x})^{-1}$  in (6.5).

To quantify the quality of the approximations, we use the classical root mean squared error (RMSE) criterion. For any approximation  $\tilde{m}$  of  $m$ , it is defined by

$$\text{RMSE}(\tilde{m}) = \mathbb{E} [(\tilde{m} - m)^T (\tilde{m} - m)]^{1/2},$$

and for any approximation  $\tilde{P}$  of  $P$ , we use

$$\text{RMSE}(\tilde{P}) = \mathbb{E} [\|\tilde{P} - P\|_F^2]^{1/2},$$

where  $\|\cdot\|_F$  is the Frobenius matrix norm. In both cases, the expectation is taken w.r.t. the distribution of the observations  $Y_{1:n}$ , and w.r.t. the sampling distribution when the approximation is obtained by importance sampling (i.e.,  $\eta$  or  $\eta'$ ).

The simulation parameters are set as follows. The number of sensors is  $n = 2$  with positions  $s_1 = \begin{pmatrix} 0 & 0 \end{pmatrix}^T$  and  $s_2 = \begin{pmatrix} 0 & 50 \end{pmatrix}^T$  (the units are, say, meters). The standard deviation of the observation noise is  $\sigma = 0.0524$  (radians, i.e.  $3^\circ$ ). The expectation and covariance matrix of the prior are respectively  $m_q = \begin{pmatrix} 2000 & 3000 \end{pmatrix}^T$  and  $Q = 1000^2 I_2$ .

Note that, to compute the RMSEs, one needs the exact values  $m$  and  $P$  of the posterior expectation and covariance matrix. Since they are unknown in this nonlinear example (which is why we use approximation methods), we compute them accurately using importance sampling with a very large number of particles ( $N = 10^6$ ).

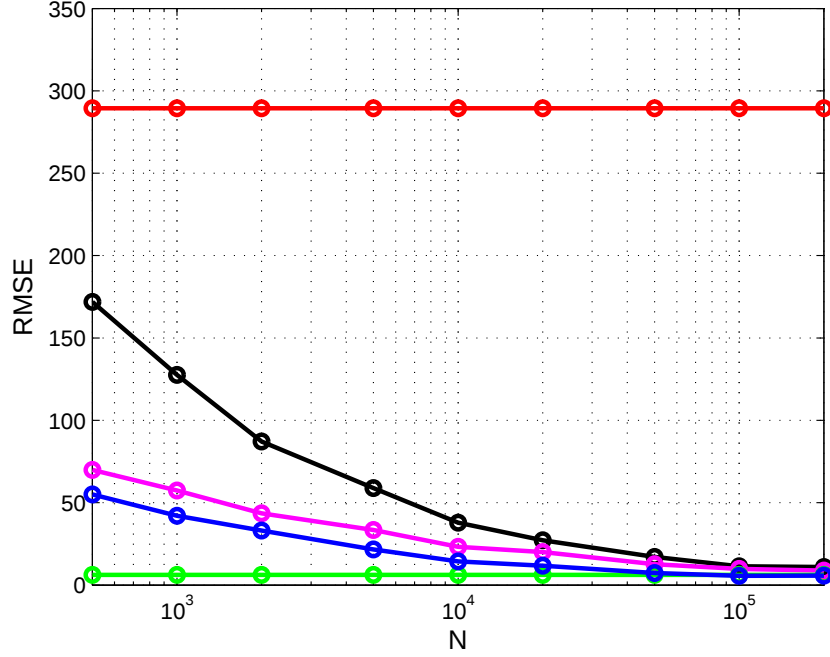


Figure 6.1: RMSE of the approximations of posterior expectation  $m$ . Black line:  $m^N$ , blue line  $\hat{m}_0^N$ , pink line:  $\hat{m}^N$ , green line:  $\hat{m}$ , red line:  $\hat{x}$ .

Figure 6.1 shows the quality of the different approximations of  $m$  when the number of particles  $N$  increases. The Laplace approximation  $\hat{m}$  (which does not depend on  $N$ ) is very accurate. When  $N$  increases, the RMSE of the importance sampling approximations  $m^N$ ,  $\hat{m}_0^N$  and  $\hat{m}^N$  goes to 0. The RMSE of  $\hat{m}^N$ , based on transformation  $T$ , converges faster to 0 than the RMSEs of  $m^N$  and  $\hat{m}_0^N$ .

We also show the RMSE of the MAP  $\hat{x}$  on figure 6.1. The MAP is required to compute the Laplace approximations  $\hat{m}$  (6.3) and  $\hat{P}$  (6.4). One can see that the difference between the MAP and the posterior expectation is large. This illustrates the asymmetry of the posterior density and the importance of the second term in formula (6.3). Indeed, this term takes into account the asymmetry of  $p$  at the MAP through  $J'(\hat{x}) = -(\log p)'''(\hat{x})$ , which involves third-order derivatives of the model. This asymmetry may explain why  $\hat{m}^N$  is a better approximation than  $\hat{m}_0^N$ , since the former depends on  $J'(\hat{x})$  and  $J''(\hat{x})$  through  $\hat{m}$  and  $\hat{P}$ , unlike the latter.

In figure 6.2, the RMSEs of the approximations of  $P$  are plotted. Once again, it can be seen that the Laplace approximation  $\hat{P}$  is very accurate, and that importance sampling based on the Laplace method yields a better approximation than importance

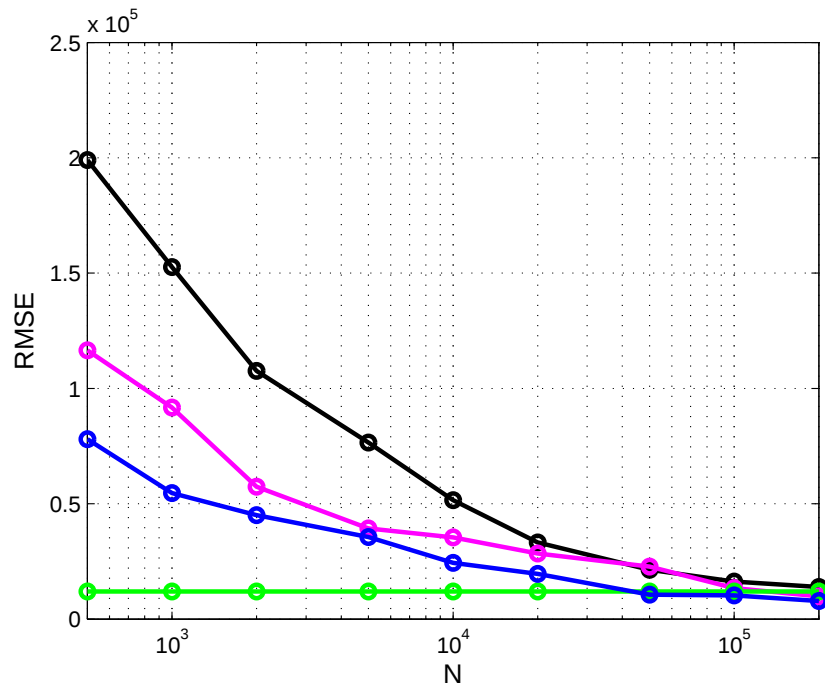


Figure 6.2: RMSE of the approximations of posterior covariance matrix  $P$ . Black line:  $P^N$ , blue line  $\hat{P}^N$ , pink line:  $\hat{P}_0^N$ , green line:  $\hat{P}$ .

sampling based on the prior.

Figure 6.3 shows the prior density  $q$ , the density of the shifted and scaled density  $q'$  and the posterior density  $p$ . It is clear that the application of the transformation  $T$  makes the sampling distribution closer to the posterior distribution.

## Conclusion

We propose in this chapter a novel way to use the Laplace method within an importance sampling scheme. We design a sampling distribution, taking into account the likelihood, by shifting and scaling the prior so that its expectation and covariance matrix match the corresponding Laplace approximations. This technique is different than the way the Laplace method and importance sampling have previously been associated in the literature.

Since the Laplace formulas are very accurate, one could argue that there is no point to use importance sampling. Although this may be true in the simple triangulation example we consider in this chapter, it is not hard to build a model where the Laplace

approximations alone give poor results (a multimodal model, e.g.). Besides, importance sampling approximation always outperforms deterministic Laplace approximation when the number of particles is large enough, for a fixed number of observations. The former is consistent as the number of particles tends to infinity, whereas the latter is consistent as the number of observations tends to infinity but biased when the number of observation is fixed.

In the next chapter, we adapt the idea introduced in this chapter to a dynamic model framework and we propose a novel particle filter which combines importance sampling and the Laplace method. In particular, the computation of the MAP and of the observed information matrix and its derivatives, which is difficult in a filtering framework, will be discussed.

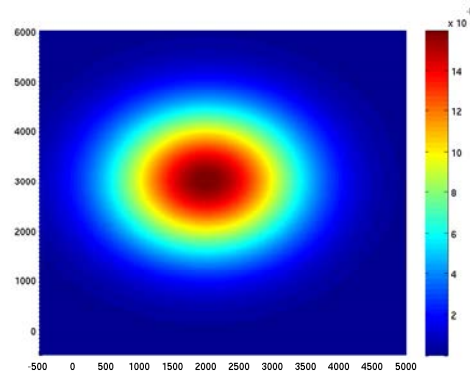
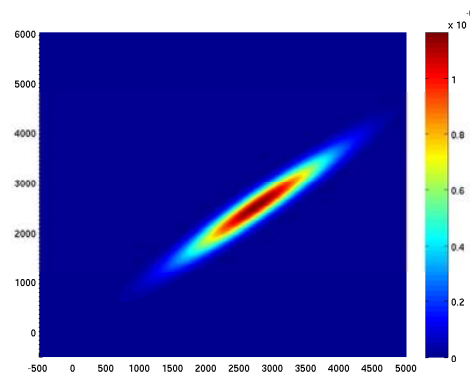
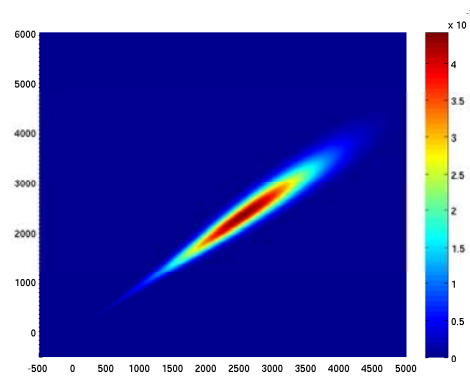
(a) Prior density  $q$ .(b) Transformed density  $q'$ .(c) Posterior density  $p$  (unnormalized).

Figure 6.3: Transformation of the prior density.



# Chapter 7

## Particle filtering and the Laplace method in dynamic models

The Laplace method has been used in state-space models in [Strasser (2002)], [Jungbacker and Koopman (2007)] and [Kleppe and Skaug (2012)] for the estimation of static parameters. In [Kleppe and Skaug (2012)], it is associated to sequential importance sampling, yielding the so-called Laplace accelerated sequential importance sampling (LASIS) algorithm. In these articles, the distribution of interest is approximated by a Gaussian distribution with the same mode and the same curvature around the mode, which is called its Laplace approximation and taken as a sampling distribution.

In this chapter, we associate sequential Monte Carlo and the Laplace methods for nonlinear Bayesian filtering. The idea of shifting and scaling the prior thanks to the multidimensional Laplace approximations of posterior moments, introduced in chapter 6, is applied in a dynamic and recursive framework. We propose a novel particle filtering algorithm based on this idea: the Laplace particle filter (LPF).

We define the model set-up (nonlinear state-space models) in section 7.1. The classical particle filters against which the LPF will be benchmarked in chapter 8 – namely, the sequential importance resampling (SIR) algorithm and the regularized particle filter (RPF) – are presented in section 7.2. Section 7.3 is devoted to the presentation of the LPF.

## 7.1 Model set-up

We consider a state-space model where the state process is nonlinear with additive Gaussian noise and the observation model is nonlinear.

The state space is an open subset  $E$  of  $\mathbb{R}^d$ . The state process is the Markov process  $\{X_k\}_{k \geq 0}$  defined by the nonlinear dynamics

$$X_k = F_k(X_{k-1}) + V_k,$$

for  $k \geq 1$ , where  $\{V_k\}_{k \geq 1}$  is a Gaussian white noise,  $V_k \sim \mathcal{N}(0, \Sigma_k)$  with invertible matrix  $\Sigma_k$ . The initial state distribution is  $X_0 \sim Q_0$ . The Markov kernel is denoted  $Q_k(x, dx')$  and its density  $q_k(x, x')$ .

The observation process is  $\{Y_k\}_{k \geq 0}$  and the likelihood function associated to the observation  $Y_k$  is denoted  $g_k$ .

For all  $k \geq 0$ , we denote  $\mu_k$  the conditional distribution of  $X_k$  given  $Y_{0:k}$  (the posterior) and  $p_k$  its density and we denote  $\mu_k^-$  the conditional distribution of  $X_k$  given  $Y_{0:k-1}$  (the predictor) and  $p_k^-$  its density (with the notation convention  $Y_{0:-1} = Y_0$ ).

## 7.2 Standard particle filtering and regularized particle filtering

For all  $k \geq 0$ , let  $\mu_k^N$  denote the particle approximation of the posterior and  $\mu_k^{-N}$  the particle approximation of the predictor.

### 7.2.1 Standard particle filtering

The most standard particle filter is algorithm 8. The particles are sampled from the Markov kernel  $Q_k$  and resampled at random times according to the effective sample size (ESS) criterion. We refer to this algorithm as the sequential importance resampling (SIR) algorithm, following the terminology of [Arulampalam et al. (2002)].

The approximation  $\mu_k^N$  of the posterior is

$$\mu_k^N = \sum_{i=1}^N w_k^i \delta_{\xi_k^i} \quad (7.1)$$

---

**Algorithm 8** Standard particle filter (SIR).

---

- $k = 0$ .
  - For  $i = 1, \dots, N$ , sample  $\xi_0^i \sim Q_0$ .
  - For  $i = 1, \dots, N$ , compute the weight

$$w_0^i \propto g_0(\xi_0^i).$$

- $k \geq 1$ .
  - ▷ If  $N_{\text{eff}} \geq N_{\text{th}}$ .
    - \* For  $i = 1, \dots, N$ , sample  $\xi_k^i \sim Q_k(\xi_{k-1}^i, dx)$ .
    - \* For  $i = 1, \dots, N$ , compute the weight

$$w_k^i \propto w_{k-1}^i g_k(\xi_k^i).$$

- ▷ If  $N_{\text{eff}} < N_{\text{th}}$ .
  - \* For  $i = 1, \dots, N$ , sample  $j_i \sim \sum_{i=1}^N w_{k-1}^i \delta_{\{i\}}$ .
  - \* For  $i = 1, \dots, N$ , sample  $\xi_k^i \sim Q_k(\xi_{k-1}^{j_i}, dx)$ .
  - \* For  $i = 1, \dots, N$ , compute the weight

$$w_k^i \propto g_k(\xi_k^i).$$


---

and the approximation  $\mu_k^{-N}$  of the predictor is

$$\mu_k^{-N} = \sum_{i=1}^N w_{k-1}^i \delta_{\xi_k^i}.$$

Besides, since the Markov kernel admits a density, the predictor admits a density as well, denoted  $p_k^-(x)$ . The predictor density can be approximated using the weighted particles, consistently as  $N \rightarrow \infty$ , by

$$p_k^{-N}(x) = \sum_{i=1}^N w_{k-1}^i q_k(\xi_{k-1}^i, x).$$

The posterior also admits a density  $p_k(x)$  which, however, cannot be straightforwardly approximated. A way to do so is to use kernel regularization.

### 7.2.2 Regularized particle filtering

The use of kernel regularization in particle filtering has been proposed in [Musso and Oudjane (1998); Le Gland et al. (1998); Oudjane and Musso (2000); Musso et al. (2001)].

Let  $K$  be a symmetric density and  $h$  a real positive parameter, called the bandwidth. A kernel is a function  $K_h$  such that  $K_h(x) = \frac{1}{h^d} K\left(\frac{x}{h}\right)$ . The kernel regularization of the posterior approximation  $\mu_k^N$  (7.1) is in the form of a finite mixture of shifted and rescaled densities,

$$p_k^{N,h}(x) = \sum_{i=1}^N w_k^i K_h(x - \xi_k^i). \quad (7.2)$$

In the regularized particle filter (RPF, algorithm 9 at the end of the chapter), regularization is performed after resampling by adding an independent noise to each resampled particle. The distribution of the regularized particles admits the density (7.2), which consistently approximates the posterior density  $p_k(x)$  as  $N \rightarrow \infty$  and  $h \rightarrow 0$ . See [Musso et al. (2001)] for further details, in particular how to efficiently choose  $h$ , which depends on the state space dimension  $d$  and the number of particles  $N$ .

An important feature of regularization is to improve particle approximation in the case of a small noise model. The performance of standard particle filtering (SIR) is often poor when the state dynamics noise or the observation noise are small [Oudjane and Musso (2000)]. In this case, bringing additional randomness to the particle population, thanks to an artificial algorithm noise, allows to cope with the lack of model noise and improves the performance of the algorithm. This improvement will be illustrated on simulations in chapter 8.

## 7.3 Laplace particle filtering

Let

$$J_k(x) = -(\log p_k)''(x) = -(\log g_k)''(x) - (\log p_k^-)''(x)$$

be the observed information matrix of the state  $x \in E$  and

$$J_k^N(x) = -(\log g_k)''(x) - (\log p_k^{-N})''(x)$$

---

**Algorithm 9** Regularized particle filter.

---

- $k = 0$ .

- For  $i = 1, \dots, N$ , sample  $\xi_0^i \sim Q_0$ .
- For  $i = 1, \dots, N$ , compute the weight

$$w_0^i \propto g_0(\xi_0^i).$$

- $k \geq 1$ .

▷ If  $N_{\text{eff}} \geq N_{\text{th}}$ .

- \* For  $i = 1, \dots, N$ , sample  $\xi_k^i \sim Q_k(\xi_{k-1}^i, dx)$ .
- \* For  $i = 1, \dots, N$ , compute the weight

$$w_k^i \propto w_{k-1}^i g_k(\xi_k^i).$$

▷ If  $N_{\text{eff}} < N_{\text{th}}$ .

- \* For  $i = 1, \dots, N$ , sample  $j_i \sim \sum_{i=1}^N w_{k-1}^i \delta_{\{i\}}$ .
- \* For  $i = 1, \dots, N$ , sample  $\xi_k^{j_i} \sim Q_k(\xi_{k-1}^{j_i}, dx)$ .
- \* Compute  $S_k^N = \frac{1}{N} \sum_{i=1}^N \xi_k^{j_i} (\xi_k^{j_i})^T - \frac{1}{N^2} \left( \sum_{i=1}^N \xi_k^{j_i} \right) \left( \sum_{i=1}^N \xi_k^{j_i} \right)^T$  the empirical covariance matrix of  $(\xi_k^{j_1}, \dots, \xi_k^{j_N})$ .
- \* For  $i = 1, \dots, N$ , sample  $Z^i \sim K(x)dx$ .
- \* For  $i = 1, \dots, N$ , compute  $\xi_k^i = \xi_k^{j_i} + h [S_k^N]^{1/2} Z^i$ .
- \* For  $i = 1, \dots, N$ , compute the weight

$$w_k^i \propto g_k(\xi_k^i).$$


---

its particle approximation. Let

$$\hat{x}_k = \operatorname{argmax}_{x \in E} \{g_k(x) p_k^-(x)\}$$

be the MAP and

$$\hat{x}_k^N = \operatorname{argmax}_{x \in E} \{g_k(x) p_k^{-N}(x)\}$$

its particle approximation.

### 7.3.1 Laplace approximations of the posterior moments

Thanks to the approximated density of the predictor  $p_k^{-N}(x)$  and to the likelihood function  $g_k(x)$ , one can compute the MAP  $\hat{x}_k^N$ , the observed information matrix  $J_k^N(x)$  and the derivatives of  $J_k^N$ . Hence, the quantities involved in the Laplace formulas for approximating the posterior expectation and covariance are available, so that these formulas are computable.

From the derivations in chapter 3, let

$$\begin{aligned}\hat{m}_k^N &= \hat{x}_k^N - \frac{1}{2}(\hat{J}_k^N)^{-1}(\hat{J}_k^N)^{T'} \text{vec}[(\hat{J}_k^N)^{-1}] \\ &\approx \mathbb{E}[X_k | Y_{0:k}]\end{aligned}\tag{7.3}$$

and

$$\begin{aligned}\hat{P}_k^N &= (\hat{J}_k^N)^{-1} + \frac{1}{2}(\hat{J}_k^N)^{-1}(\hat{J}_k^N)^{T'}((\hat{J}_k^N)^{-1} \otimes (\hat{J}_k^N)^{-1})(\hat{J}_k^N)'(\hat{J}_k^N)^{-1} \\ &\quad + \frac{1}{2}(I_d \otimes \text{vec}[(\hat{J}_k^N)^{-1}]^T)(\hat{J}_k^N)'((\hat{J}_k^N)^{-1} \otimes (\hat{J}_k^N)^{-1})(\hat{J}_k^N)'(\hat{J}_k^N)^{-1} \\ &\quad - \frac{1}{2}(\hat{J}_k^N)^{-1}(I_d \otimes \text{vec}[(\hat{J}_k^N)^{-1}]^T)(\hat{J}_k^N)''(\hat{J}_k^N)^{-1} \\ &\approx \mathbb{V}[X_k | Y_{0:k}],\end{aligned}\tag{7.4}$$

where  $\hat{J}_k^N = J_k^N(\hat{x}_k^N)$ ,  $(\hat{J}_k^N)' = (J_k^N)'(\hat{x}_k^N)$  and  $(\hat{J}_k^N)'' = (J_k^N)''(\hat{x}_k^N)$ , respectively be the Laplace approximations of the posterior expectation and covariance. The superscript  $N$  denotes the dependency of these approximations on the particles, through the particle approximation  $p_k^{-N}(x)$  of the predictor density.

### 7.3.2 Sampling particles according to the Laplace approximations

At time  $k$ , let  $(\xi_k^{-1}, \dots, \xi_k^{-N})$  be the particles sampled from the predictor approximation  $\mu_k^{-N}$  (with the convention  $\mu_0^- = Q_0$ ).

Following the idea introduced in chapter 6, we apply on these particles an affine transformation based on the Laplace approximations  $\hat{m}_k^N$  (7.3) and  $\hat{P}_k^N$  (7.4). The aim of this transformation is to relocate the particles  $(\xi_k^{-1}, \dots, \xi_k^{-N})$  so that their empirical mean and covariance matrix match  $\hat{m}_k^N$  and  $\hat{P}_k^N$ . The transformation to be applied on

each  $\xi_k^{-i}$  is thus

$$T_k^N(x) = \left[ \hat{P}_k^N \right]^{1/2} \left[ \bar{S}_k^{-N} \right]^{-1/2} (x - \bar{\xi}_k^{-N}) + \hat{m}_k^N, \quad (7.5)$$

where  $\bar{\xi}_k^{-N} = \frac{1}{N} \sum_{i=1}^N \xi_k^{-i}$  and  $\bar{S}_k^{-N} = \frac{1}{N} \sum_{i=1}^N (\xi_k^{-i} - \bar{\xi}_k^{-N})(\xi_k^{-i} - \bar{\xi}_k^{-N})^T$  are the empirical mean and covariance matrix of  $(\xi_k^{-1}, \dots, \xi_k^{-N})$ . Thus, the transformed particle sample  $(T_k^N(\xi_k^{-1}), \dots, T_k^N(\xi_k^{-N}))$  admit the Laplace approximations  $\hat{m}_k$  and  $\hat{P}_k$  as its first and second-order empirical moments.

In other words, we sample particles from the distribution

$$\mu_k'^{-N} = \mu_k^{-N} \circ (T_k^N)^{-1}$$

instead of sampling from  $\mu_k^{-N}$ . This new sampling distribution takes into account the observation  $Y_k$  through the Laplace approximations involved in the coefficients of transformation  $T_k^N$ . Note that  $\mu_k'^{-N}$  admits a density, which we denote  $p_k'^{-N}(x)$ ,

$$\begin{aligned} p_k'^{-N}(x) &= p_k^{-N}((T_k^N)^{-1}(x)) \left| \det \left[ ((T_k^N)^{-1})'(x) \right] \right| \\ &= p_k^{-N} \left( [\bar{S}_k^{-N}]^{1/2} [\hat{P}_k^N]^{-1/2} (x - \hat{m}_k^N) + \bar{\xi}_k^{-N} \right) \det [\bar{S}_k^{-N}]^{1/2} \det [\hat{P}_k^N]^{-1/2} \\ &\propto p_k^{-N} \left( [\bar{S}_k^{-N}]^{1/2} [\hat{P}_k^N]^{-1/2} (x - \hat{m}_k^N) + \bar{\xi}_k^{-N} \right). \end{aligned}$$

To perform nonlinear filtering using sequential Monte Carlo and the Laplace methods, the transformation  $T_k^N$  (7.5) is embedded in a particle filtering algorithm. We apply  $T_k^N$  in an adaptive way: when the ESS  $N_{\text{eff}}$  is below the threshold  $N_{\text{th}}$ , the Laplace formulas are computed and  $T_k^N$  is applied just after resampling. This strategy is implemented in algorithm 10.

### 7.3.3 Computational issues

**MAP computation.** The Laplace formulas used in transformation  $T_k^N$  involve the MAP,  $\hat{x}_k^N = \underset{x \in E}{\operatorname{argmax}} \{g_k(x)p_k^{-N}\}$ , which can be difficult to compute, especially when the state space dimension is large.

However, in Bayesian filtering it is often the case that the likelihood  $g_k$  depends on the state vector  $x$  only through a subvector  $x^o$ . The subvector  $x^o$  is said to be observable, whereas the remaining subvector  $x^n$  is said to be nonobservable. Let  $x =$

$(x^o, x^n) \in E \subset \mathbb{R}^d$  such that  $g_k(x) = g_k(x^o)$ . Notice that

$$\begin{aligned} \max_{x \in E} \{g_k(x)p_k^{-N}(x)\} &= \max_{x^o, x^n} \{g_k(x^o)p_k^{-N}(x^o, x^n)\} \\ &= \max_{x^o} \{g_k(x^o) \max_{x^n} \{p_k^{-N}(x^o, x^n)\}\} \\ &= \max_{x^o} \{g_k(x^o)p_k^{-N}(x^o, \hat{x}^n(x^o))\} \end{aligned}$$

where  $\hat{x}^n(x^o) = \operatorname{argmax}_{x^n} \{p_k^{-N}(x^o, x^n)\}$  for any  $x^o$ . Therefore, the particle approximation of the MAP can be expressed as

$$\hat{x}_k^N = \operatorname{argmax}_{x \in E} \{g_k(x)p_k^{-N}(x)\} = (\hat{x}^o, \hat{x}^n(\hat{x}^o)) \quad (7.6)$$

where  $\hat{x}^o = \operatorname{argmax}_{x^o} \{g_k(x^o)p_k^{-N}(x^o, \hat{x}^n(x^o))\}$ . If  $\hat{x}^n(x^o)$  can be computed analytically, one can simplify optimization by maximizing  $g_k(x^o)p_k^{-N}(x^o, \hat{x}^n(x^o))$  over the subspace  $\{x^o : (x^o, x^n) \in E\}$  of the state space corresponding to the observable part of the state.

**Computation of the information matrix and its derivatives.** Laplace approximations also involve the observed information matrix and its derivatives.  $J_k^N$ ,  $(J_k^N)'$  and  $(J_k^N)''$  need to be evaluated only at the MAP. Thus, once the MAP is available,  $J_k^N(\hat{x}_k^N)$ ,  $(J_k^N)'(\hat{x}_k^N)$  and  $(J_k^N)''(\hat{x}_k^N)$  can be computed by numerical differentiation around  $\hat{x}_k^N$  (see the supplementary material of [Koyama et al. (2010)]).

**Computation of the predictor density.** The weighting step

$$w_k^i \propto \frac{g_k(\xi_k^i)p_k^{-N}(\xi_k^i)}{p_k^{-N}((T_k^N)^{-1}(\xi_k^i))} = \frac{g_k(\xi_k^i)p_k^{-N}(\xi_k^i)}{p_k^{-N}(\xi_k^{-i})}$$

is difficult to perform in practice because of the form of the approximated predictor density involved at the numerator,

$$p_k^{-N}(x) = \sum_{i=1}^N w_{k-1}^i q_k(\xi_{k-1}^i, x) \propto \sum_{i=1}^N w_{k-1}^i \exp \left( -\frac{1}{2}(x - F_k(\xi_{k-1}^i))^T \Sigma_k^{-1}(x - F_k(\xi_{k-1}^i)) \right). \quad (7.7)$$

This density is a mixture of Gaussian densities that may be concentrated around their maxima, if the covariance  $\Sigma_k$  of the state dynamics noise is small, e.g. In that case, it is likely that, once transformation  $T_k^N$  has been applied, many particles  $\xi_k^i$  lie between

the modes  $F_k(\xi_{k-1}^i)$  of mixture (7.7). Then,  $p_k^{-N}(\xi_k^i) \approx 0$ , so that  $w_k^i \approx 0$ . A way to cope with this phenomenon is to use kernel regularization to smooth (7.7), which then becomes

$$p_k^{-N,h}(x) = \frac{1}{N} \sum_{i=1}^N w_{k-1}^i K_h(x - \xi_k^{-i}),$$

where  $\xi_k^{-1}, \dots, \xi_k^{-N}$  are the particles after resampling. The weights are then computed as

$$w_k^i \propto \frac{g_k(\xi_k^i) p_k^{-N,h}(\xi_k^i)}{p_k^{-N,h}(\xi_k^{-i})}. \quad (7.8)$$

The kernel approximation of the predictor is consistent as  $N \rightarrow \infty$  and  $h \rightarrow 0$ . However, kernel regularization requires to tune the bandwidth parameter  $h$ , which depends on  $d$  and  $N$ , and demands a lot of computational power.

### 7.3.4 Gaussian approximation of the predictor

We now propose a drastic way to simplify the computations involved in the Laplace approximations and the weighting step, which will be demonstrated practically efficient in the numerical experiments in chapter 8.

The simplification relies on the Gaussian approximation of the predictor. Let  $\tilde{\mu}_k^{-N} = \varphi_{m_k^{-N}, P_k^{-N}}$  be a Gaussian distribution approximating  $\mu_k^{-N}$ , with expectation  $m_k^{-N}$  and covariance matrix  $P_k^{-N}$ , and let  $\tilde{p}_k^{-N}$  be its density.

From equation (7.6), using  $\tilde{p}_k^{-N}$  instead of  $p_k^{-N}$ , the MAP is computed as

$$\hat{x}_k^N = \operatorname{argmax}_{x \in E} \{g_k(x) \tilde{p}_k^{-N}(x)\} = (\hat{x}^o, \hat{x}^n(\hat{x}^o)),$$

where  $\hat{x}^n(x^o) = \operatorname{argmax}_{x^n} \{\tilde{p}_k^{-N}(x^o, x^n)\}$ . In this case,  $\hat{x}^n(x^o)$  is the conditional expectation

$$\frac{\int_E x^n \tilde{p}_k^{-N}(x^o, x^n) dx^n}{\int_E \tilde{p}_k^{-N}(x^o, x^n) dx^n},$$

which can be computed exactly since  $\tilde{p}_k^{-N}$  is a Gaussian density (see remark 7.3.1 below).

**Remark 7.3.1.** Suppose that the expectation and the covariance matrix of the Gaussian approximation of the predictor are respectively in the form of  $m_k^{-N} = \begin{pmatrix} m_k^{-N,o} & m_k^{-N,n} \end{pmatrix}^T$

and  $P_k^{-N} = \begin{pmatrix} P_k^{-N,o} & P_k^{-N,on} \\ P_k^{-N,no} & P_k^{-N,n} \end{pmatrix}$ , where the upperscript o denotes the observable part and

the superscript  $n$  denotes the nonobservable part. Then, straightforward calculations yield  $\hat{x}^n(x^o) = m_k^{-N,n} + P_k^{-N,no} [P_k^{-N,o}]^{-1} (x^o - m_k^{-N,o})$ . Thus, the remaining optimization problem is  $\max_{x^o} \{g_k(x^o) \tilde{p}_k^{-N}(x^o, \hat{x}^n(x^o))\}$ , which is a lower dimension problem.

Moreover, the Gaussian approximation of the predictor allows to compute more easily the observed information matrix and its derivatives. Indeed,  $(\log \tilde{p}_k^{-N})''(x) = [P_k^{-N}]^{-1}$  for all  $x \in E$ , so that  $(\log \tilde{p}_k^{-N})'''(x)$  and  $(\log \tilde{p}_k^{-N})^{(4)}(x)$  (which are involved in  $(J_k^N)'(x)$  and  $(J_k^N)''(x)$ ) are zero. Therefore, the difficulty is reduced to the computation of  $(\log \tilde{g}_k)''(\hat{x}_k^N)$ ,  $(\log \tilde{g}_k)'''(\hat{x}_k^N)$  and  $(\log \tilde{g}_k)^{(4)}(\hat{x}_k^N)$ . On the other hand, when working with the consistent density approximation  $p_k^{-N}(x) = \sum_{i=1}^N w_{k-1}^i q_k(\xi_{k-1}^i, x)$  rather than with the Gaussian approximation  $\tilde{p}_k^{-N}$ , one must compute

$$\begin{aligned} & (\log p_k^{-N})''(x) \\ &= -\frac{1}{\left(\sum_{i=1}^N w_{k-1}^i q_k(\xi_{k-1}^i, x)\right)^2} \left(\sum_{i=1}^N w_{k-1}^i \frac{\partial q_k}{\partial x}(\xi_{k-1}^i, x)\right)^T \left(\sum_{i=1}^N w_{k-1}^i \frac{\partial q_k}{\partial x}(\xi_{k-1}^i, x)\right) \\ & \quad + \frac{1}{\sum_{i=1}^N w_{k-1}^i q_k(\xi_{k-1}^i, x)} \sum_{i=1}^N w_{k-1}^i \frac{\partial^2 q_k}{\partial x^2}(\xi_{k-1}^i, x) \end{aligned}$$

and its higher order derivatives, which is cumbersome.

Lastly, the Gaussian approximation can be used to compute the importance weight as

$$w_k^i \propto \frac{g_k(\xi_k^i) \tilde{p}_k^{-N}(\xi_k^i)}{\tilde{p}_k^{-N}(\xi_k^{-i})},$$

which is simpler than using kernel regularization as in (7.8).

However, the price to pay when approximating the predictor by a Gaussian is the loss of consistency of the approximation as  $N \rightarrow \infty$  (if one uses  $p_k^{-N}(x)$ ) or  $N \rightarrow \infty$  and  $h \rightarrow 0$  (if one uses  $p_k^{-N,h}(x)$ ). In practice, this drawback is often not prohibitive, as it will be illustrated in the simulation experiments in chapter 8. The applicability of this approximation is restrained to models where the predictor density has a pronounced main mode, that is the predictor is rather concentrated around its maximum value.

We detail now two methods to obtain  $\tilde{p}_k^{-N}(x)$ , i.e. to compute its expectation  $m_k^{-N}$  and its covariance  $P_k^{-N}$ . The first one is to use the particles propagated through the Markov kernel and to compute empirical moments. Let  $\xi_k^{t-i} \sim Q_k(\xi_{k-1}^i, dx)$  for  $i \in$

$\{1, \dots, N\}$ . Then,

$$m_k^{-N} = \sum_{i=1}^N w_{k-1}^i \xi_k'^{-i} \quad (7.9)$$

and

$$P_k^{-N} = \sum_{i=1}^N w_{k-1}^i (\xi_k'^{-i} - m_k^{-N})(\xi_k'^{-i} - m_k^{-N})^T \quad (7.10)$$

are respectively consistent approximations of the true expectation and covariance of  $\mu_k^-$ . Another way to obtain these moments is

$$m_k^{-N} = F_k(\tilde{m}_{k-1}^N) \quad (7.11)$$

and

$$P_k^{-N} = F'_k(\tilde{m}_{k-1}^N) \tilde{P}_{k-1}^N F'_k(\tilde{m}_{k-1}^N)^T + \Sigma_k, \quad (7.12)$$

like in the extended Kalman filter, where  $\tilde{m}_{k-1}^N$  and  $\tilde{P}_{k-1}^N$  are some approximations of the posterior expectation and covariance at time  $k-1$ .  $\tilde{m}_{k-1}^N$  and  $\tilde{P}_{k-1}^N$  can be Laplace approximations ( $\tilde{m}_{k-1}^N = \hat{m}_{k-1}^N$  and  $\tilde{P}_{k-1}^N = \hat{P}_{k-1}^N$ ) or particle approximations ( $\tilde{m}_{k-1}^N = \sum_{i=1}^N w_{k-1}^i \xi_{k-1}^i$  and  $\tilde{P}_{k-1}^N = \sum_{i=1}^N w_{k-1}^i (\xi_{k-1}^i - \tilde{m}_{k-1}^N)(\xi_{k-1}^i - \tilde{m}_{k-1}^N)^T$ ). Unlike the first method, the second method does not provide consistent approximations of the moments of  $\mu_k^-$ . (Note that  $m_k^{-N}$  and  $P_k^{-N}$  can also be computed using the unscented transformation, like in the unscented Kalman filter [Julier and Uhlmann (2004)].)

Finally, the sought Gaussian density, to be used in the Laplace approximations, is

$$\tilde{p}_k^{-N}(x) = |(2\pi)^d \det P_k^{-N}|^{-1/2} \exp \left( -\frac{1}{2} (x - m_k^{-N})^T [P_k^{-N}]^{-1} (x - m_k^{-N}) \right).$$

The algorithm embedding Laplace approximations with a Gaussian approximation of the predictor is algorithm 11. It is called the Laplace particle filter (LPF).

## Conclusion

In this chapter, we propose a novel particle filtering algorithm, called the Laplace particle filter (LPF), which associates sequential Monte Carlo sampling and the Laplace method. The principle of this algorithm is to apply a transformation on the particles when the ESS is too small, this transformation being based on the Laplace approximations for the posterior expectation and covariance matrix. The objective is to relocate

the particles so that their distribution is closer to the optimal posterior distribution, in order to improve weighting and thus avoid weight degeneracy.

This approach is a way to design a new sampling distribution that matches the posterior (approximate) expectation and covariance. However, it is sometimes not necessary that a good sampling distribution must have its first- and second-order moments close to those of the distribution of interest. For example, a sampling distribution with heavier tails than the posterior can be preferred, in which case Student's  $t$ -distribution is a good candidate [Evans and Swartz (1995)]. That is, if  $\hat{m}_k^N$  and  $\hat{P}_k^N$  are the Laplace approximations at time  $k$ , then the particles can be sampled as

$$\xi_k^{-i} = \hat{m}_k^N + \left[ \hat{P}_k^N \right]^{1/2} T_i, \quad (7.13)$$

for  $i \in \{1, \dots, N\}$ , where  $T_i \sim t_{\text{df}}$  ( $t_{\text{df}}$  being Student's  $t$ -distribution in  $\mathbb{R}^d$  with  $\text{df}$  degrees of freedom).

Formulas (7.3) and (7.4) look long and complicated, but what is actually complicated is:

- the computation of the MAP,
- the evaluation of the observed information matrix and its derivatives at the MAP.

However, computing the MAP can be drastically simplified thanks to the Gaussian approximation of the predictor (see remark 7.3.1). Besides, when the likelihood model is known well enough so that the derivation of  $(\log g_k)''(x)$ ,  $(\log g_k)'''(x)$  and  $(\log g_k)^{(4)}(x)$  can be done offline, the (online) evaluation of the observed information matrix and its derivatives at the MAP is straightforward. Then, if the values of  $\hat{x}_k$ ,  $J_k^N(\hat{x}_k)$ ,  $(J_k^N)'(\hat{x}_k)$ ,  $(J_k^N)''(\hat{x}_k)$  and  $(J_k^N)^{-1}(\hat{x}_k)$  are respectively contained in the variables

```
map; J; dJ; d2J; Jinv;
```

then the Laplace approximations  $\hat{m}_k^N$  and  $\hat{P}_k^N$  are computed in MATLAB as:

```
expectation_laplace = map - 0.5*inv(J)*dJ'*invJ(:);
covariance_laplace = invJ + .5*invJ*dJ'*kron(invJ,invJ)*dJ*invJ...
+ .5*kron(eye(d),invJ(:)')*dJ*kron(invJ,invJ)*dJ*invJ...
- .5*invJ*kron(eye(d),invJ(:)')*d2J*invJ;
```

which corresponds to formulas (7.3) and (7.4). The LPF is applied on practical nonlinear filtering problems in the next chapter.

---

**Algorithm 10**


---

- $k = 0$ .
  - For  $i = 1, \dots, N$ , sample  $\xi_0^{-i} \sim Q_0$ .
  - Compute  $\hat{x}_0 = \operatorname{argmax}_{x \in E} \{g_0(x)q_0(x)\}$ .
  - Compute  $J_0(\hat{x}_0) = -(\log g_0)''(\hat{x}_0) - (\log q_0)''(\hat{x}_0)$ ,  $(J_0)'(\hat{x}_0)$  and  $(J_0)''(\hat{x}_0)$ .
  - Compute  $\hat{m}_0$  and  $\hat{P}_0$  according to (7.3) and (7.4).
  - For  $i = 1, \dots, N$ , compute  $\xi_0^i = T_0(\xi_0^{-i})$  according to (7.5).
  - For  $i = 1, \dots, N$ , compute the weight

$$w_0^i \propto \frac{g_0(\xi_0^i)q_0(\xi_0^i)}{q_0(\xi_0^{-i})}.$$

- $k \geq 1$ .
  - If  $N_{\text{eff}} \geq N_{\text{th}}$ .
    - \* For  $i = 1, \dots, N$ , sample  $\xi_k^i \sim Q_k(\xi_{k-1}^i, dx)$ .
    - \* For  $i = 1, \dots, N$ , compute the weight

$$w_k^i \propto w_{k-1}^i g_k(\xi_k^i).$$

- If  $N_{\text{eff}} < N_{\text{th}}$ .
  - \* For  $i = 1, \dots, N$ , sample  $j_i \sim \sum_{i=1}^N w_{k-1}^i \delta_{\{i\}}$ .
  - \* For  $i = 1, \dots, N$ , sample  $\xi_k^{i-} \sim Q_k(\xi_{k-1}^{j_i}, dx)$ .
  - \* Compute  $\hat{x}_k^N = \operatorname{argmax}_{x \in E} \{g_k(x)p_k^{-N}(x)\}$ .
  - \* Compute  $J_k^N(\hat{x}_k^N) = -(\log g_k)''(\hat{x}_k^N) - (\log p_k^{-N})''(\hat{x}_k^N)$ ,  $(J_k^N)'(\hat{x}_k^N)$  and  $(J_k^N)''(\hat{x}_k^N)$ .
  - \* Compute  $\hat{m}_k^N$  and  $\hat{P}_k^N$  according to (7.3) and (7.4).
  - \* For  $i = 1, \dots, N$ , compute  $\xi_k^i = T_k^N(\xi_k^{i-})$  according to (7.5).
  - \* For  $i = 1, \dots, N$ , compute the weight

$$w_k^i \propto \frac{g_k(\xi_k^i)p_k^{-N}(\xi_k^i)}{p_k^{-N}(\xi_k^{i-})}.$$


---

---

**Algorithm 11** Laplace particle filter.

---

- $k = 0$ .
  - For  $i = 1, \dots, N$ , sample  $\xi_0^{-i} \sim Q_0$ .
  - Compute  $\hat{x}_0 = \operatorname{argmax}_{x \in E} \{g_0(x)q_0(x)\}$ .
  - Compute  $J_0(\hat{x}_0) = -(\log g_0)''(\hat{x}_0) - (\log q_0)''(\hat{x}_0)$ ,  $(J_0)'(\hat{x}_0)$  and  $(J_0)''(\hat{x}_0)$ .
  - Compute  $\hat{m}_0$  and  $\hat{P}_0$  according to (7.3) and (7.4).
  - For  $i = 1, \dots, N$ , compute  $\xi_0^i = T_0(\xi_0^{-i})$  according to (7.5).
  - For  $i = 1, \dots, N$ , compute the weight

$$w_0^i \propto \frac{g_0(\xi_0^i)q_0(\xi_0^i)}{q_0(\xi_0^{-i})}.$$

- $k \geq 1$ .

▷ If  $N_{\text{eff}} \geq N_{\text{th}}$ .

- \* For  $i = 1, \dots, N$ , sample  $\xi_k^i \sim Q_k(\xi_{k-1}^i, dx)$ .
- \* For  $i = 1, \dots, N$ , compute the weight

$$w_k^i \propto w_{k-1}^i g_k(\xi_k^i).$$

▷ If  $N_{\text{eff}} < N_{\text{th}}$ .

- \* Compute  $m_k^{-N}$  and  $P_k^{-N}$  according to (7.9) and (7.10) (or (7.11) and (7.12)).
- \* For  $i = 1, \dots, N$ , sample  $\xi_k^{-i} \sim \mathcal{N}(m_k^{-N}, P_k^{-N})$ .
- \* Compute  $\hat{x}_k^N = \operatorname{argmax}_{x \in E} \{g_k(x)\tilde{p}_k^{-N}(x)\}$ , where  $\tilde{p}_k^{-N}$  is the density of the distribution  $\mathcal{N}(m_k^{-N}, P_k^{-N})$ .
- \* Compute  $J_k^N(\hat{x}_k^N) = -(\log g_k)''(\hat{x}_k^N) - (\log \tilde{p}_k^{-N})''(\hat{x}_k^N)$ ,  $(J_k^N)'(\hat{x}_k^N)$  and  $(J_k^N)''(\hat{x}_k^N)$ .
- \* Compute  $\hat{m}_k^N$  and  $\hat{P}_k^N$  according to (7.3) and (7.4).
- \* For  $i = 1, \dots, N$ , compute  $\xi_k^i = T_k^N(\xi_k^{-i})$  according to (7.5).
- \* For  $i = 1, \dots, N$ , compute the weight

$$w_k^i \propto \frac{g_k(\xi_k^i)\tilde{p}_k^{-N}(\xi_k^i)}{\tilde{p}_k^{-N}(\xi_k^{-i})}.$$


---

# Chapter 8

## Simulation experiments

We apply in this chapter the algorithms proposed in the thesis and compare them to existing algorithms. The Laplace particle filter (chapter 7) and the Kalman Laplace filter (chapter 5) are used on three real-life Bayesian filtering problems: bearings-only target tracking, ballistic target tracking, neural decoding. We process simulated observations rather than real observations. Although less realistic, this allows us to make the model parameters vary and to study the performance of the algorithms depending on these parameters. We particularly focus on difficult situations where the model noise is small (see chapter 2).

We have not tuned the simulation parameters to optimize the performance of the algorithms. The ESS threshold under which the particles are resampled was set at  $N_{\text{th}} = 2/3N$  in all the particle filters. As mentioned in chapter 7, the computation of the MAP and of the observed information matrix and its derivatives are the tricky part of the LPF. In the simulations in this chapter, optimization to compute the MAP has been performed w.r.t. to the observable state vector, which has lower dimension than the whole state vector (see remark 7.3.1 in chapter 7), using the "fminsearch" MATLAB function. The observable state vector corresponds to the target position in the bearings-only target tracking and the ballistic target tracking applications. Besides, details regarding the differentiation of the log-likelihood, in order to compute the observed information matrix and its derivatives, are provided in appendix A.

We describe in section 8.1 how we assess the performance of the algorithms, in terms of accuracy and robustness to divergence. The simulation experiments and their results for bearings-only target tracking, ballistic target tracking and neural decoding are respectively presented in sections 8.2, 8.3 and 8.4.

## 8.1 Performance evaluation in Bayesian filtering

We describe in this section how to assess the performance of filtering algorithms in simulation experiments.

### 8.1.1 Simulation set-up

As we consider exclusively simulated problems, it is possible to run the filtering algorithm as many times as we want. Let  $n$  be the number of discrete-time steps of the state-space model and let  $R$  be the number of times we run the algorithm.

To each run  $r \in \{0, \dots, R\}$  corresponds a "true" state sequence  $X_{0:n}^{\text{true},r}$ . The initial true state of each sequence is sampled from the initial state prior distribution,

$$X_0^{\text{true},r} \sim Q_0(dx)$$

independently for  $r \in \{0, \dots, R\}$ .  $X_0^{\text{true},r}$  is then deterministically propagated according to a noise-free state dynamics in the form of  $X_k = F_k(X_{k-1})$ , for  $k \in \{1, \dots, n\}$ .  $Q_0(dx)$  is chosen such that it has a high level of uncertainty (its variance is large); thus the initial true states  $X_0^{\text{true},1}, \dots, X_0^{\text{true},R}$  are very different among themselves and so are the true state sequences  $X_{0:n}^{\text{true},1}, \dots, X_{0:n}^{\text{true},R}$ .

From the  $R$  true state sequences,  $R$  observation sequences are independently simulated thanks to the likelihood model,

$$Y_{0:n}^r \sim \prod_{k=0}^n \mathbb{P}[Y_k \in dy_k | X_k^{\text{true},r}],$$

for  $r \in \{1, \dots, R\}$ . Using each of the  $R$  observation sequences, a sequence of estimators of the hidden state is computed thanks to the filtering algorithm. Let

$$\left\{ \hat{X}_k^r(Y_{1:k}^r) \right\}_{0 \leq k \leq n}$$

be the  $r$ th output of the algorithm, where the observations  $Y_{0:n}^r$  are processed. In filtering, the estimator  $\hat{X}_k^r(Y_{1:k}^r)$  depends only on the past observations  $Y_{1:k}^r$ .

After  $R$  runs of the algorithm, we end up with  $R$  realizations of trajectories of the state estimator process. Indicators of statistical performance can be computed thanks to these multiple realizations. In all the simulation experiments of this chapter, the

number of runs is set to  $R = 500$ .

### 8.1.2 Accuracy

We drop momentarily the time index. We define the mean squared error (MSE) between the hidden state  $X$  and an estimator  $\hat{X}(Y)$  as

$$\text{MSE}(\hat{X}(Y)) = \mathbb{E} \left[ \left| \hat{X}(Y) - X \right|^2 \right].$$

(Notice that this definition is different than the one used in chapter 6.) It can be decomposed as

$$\mathbb{E} \left[ \left| \hat{X}(Y) - X \right|^2 \right] = \mathbb{E} \left[ \left| \hat{X}(Y) - \mathbb{E}[X|Y] \right|^2 \right] + \mathbb{E} \left[ \left| \mathbb{E}[X|Y] - X \right|^2 \right], \quad (8.1)$$

where the first term is the error due to the algorithm which provides  $\hat{X}(Y)$ , and the second term is intrinsic to the model (see figure 8.1). From decomposition (8.1), it can be seen that the estimator  $\hat{X}(Y) = \mathbb{E}[X|Y]$  minimizes the algorithm error. In general, the posterior expectation  $\mathbb{E}[X|Y]$  is unknown.

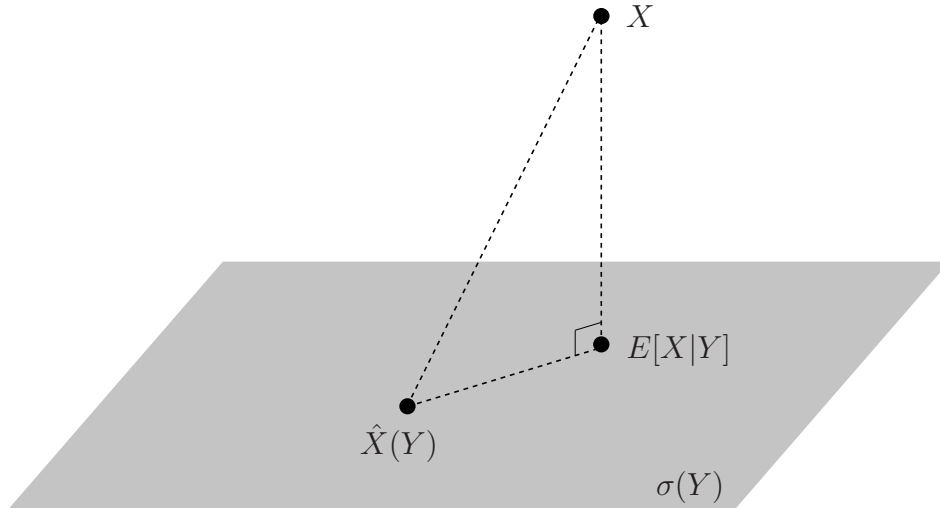


Figure 8.1: decomposition of the statistical error.  $\sigma(Y)$  denotes the  $\sigma$ -field generated by observation  $Y$ .

Under technical conditions which are satisfied in most cases, its MSE is bounded from below by the posterior Cramér-Rao bound. Let  $p(x|Y)$  be the posterior density,

which is supposed to be sufficiently regular w.r.t. the state variable  $x$ . The posterior Cramér-Rao bound is the inverse of the Fisher information matrix,

$$\mathbb{E} \left[ \left| \hat{X}(Y) - X \right|^2 \right] \geq J^{-1} \quad (8.2)$$

where the Fisher information matrix is

$$J = -\mathbb{E} [(\log p)''(X|Y)].$$

See [Van Trees (2001)] for the derivation of the posterior Cramér-Rao bound.

Thus, if the Fisher information matrix exists, from (8.1) and (8.2) we have that

$$\text{MSE}(\hat{X}(Y)) = \mathbb{E} \left[ \left| \hat{X}(Y) - X \right|^2 \right] \geq \mathbb{E} [|\mathbb{E}[X|Y] - X|^2] \geq J^{-1}$$

for any estimator  $\hat{X}(Y)$ , where the second inequality holds since  $\mathbb{E}[X|Y]$  is an unbiased estimator of  $X$ .

Let us re-introduce the time index  $k$ . The MSE of the estimator  $\hat{X}_k(Y_{1:k})$  of  $X_k$  is an expectation w.r.t. to the joint distribution of  $(X_k, Y_{1:k})$ . In practice, it is difficult to integrate w.r.t. this distribution so that the true MSE is generally not computable. In this chapter, to assess the accuracy of the estimators, we will approximate it as

$$\text{MSE}(\hat{X}_k(Y_{1:k})) = \mathbb{E} \left[ \left| \hat{X}_k(Y_{1:k}) - X_k \right|^2 \right] \approx \frac{1}{R} \sum_{i=1}^R \left| \hat{X}_k^r(Y_{1:k}^r) - X_k^{\text{true},r} \right|^2$$

(recall  $R$  is the number of runs).

Note that, in state-space models, the Fisher information matrix  $J_k = -\mathbb{E} [(\log p_k)''(X_k|Y_{1:k})]$  can be computed recursively using the algorithm proposed in [Tichavský et al. (1998)].

### 8.1.3 Robustness to divergence

In nonlinear filtering, it is a crucial issue to avoid divergence [Gustafsson (2010)]. We define here the criterion that is used in this chapter to decide whether or not a run of the filtering algorithm is divergent.

Our divergence criterion is based on a Gaussian approximation of the posterior at final time  $n$ . In each algorithm we consider, the hidden state is estimated by an approximation of the posterior expectation, i.e.  $\hat{X}_n^r(Y_{0:n}^r) \approx \mathbb{E}[X_n|Y_{0:n}]$ . Each algorithm

also delivers an approximation of the posterior covariance matrix, denoted  $\hat{P}_n^r(Y_{0:n}^r)$ , i.e.  $\hat{P}_n^r(Y_{0:n}^r) \approx \mathbb{V}[X_n|Y_{0:n}]$ . We say that the  $r$ th run of the algorithm is divergent if its final output is such that the final true state  $X_n^{\text{true},r}$  does not belong to the 99% confidence ellipsoid associated with the Gaussian distribution with expectation  $\hat{X}_n^r(Y_{0:n}^r)$  and covariance  $\hat{P}_n^r(Y_{0:n}^r)$ . That is, if

$$\left(\hat{X}_n^r(Y_{0:n}^r) - X_n^{\text{true},r}\right)^T \left[\hat{P}_n^r(Y_{0:n}^r)\right]^{-1} \left(\hat{X}_n^r(Y_{0:n}^r) - X_n^{\text{true},r}\right) > \chi_{d,0.99}^2,$$

where  $\chi_{d,0.99}^2$  is the 0.99-order quantile of the  $\chi^2$  distribution with  $d$  degrees of freedom ( $d$  being the state dimension), then the  $r$ th run is said to be divergent.

In a nutshell, a run is divergent when the final true state lies on the tail of the Gaussian approximation of the posterior delivered by the algorithm.

## 8.2 Bearings-only target tracking

We first consider the problem of tracking an object in the plane with angle-only observations [Ristic et al. (2001); Ristic and Arulampalam (2003); Musso et al. (2011)]. Here, we are interested in the case where the observations are delivered by a maneuvering sensor.

### 8.2.1 Model

A target in the Cartesian plane moves according to a rectilinear uniform motion. A moving sensor measures the azimuth of the target w.r.t. a reference direction. The sensor must maneuver to insure the observability of the target over time [Le Cadre and Jauffret (1997)].

The observations  $\{Y_k\}_{k \geq 0}$  are delivered every  $\Delta$  seconds. At time  $k\Delta$ , we denote  $s_k = (s_{k,1} \ s_{k,2})^T$  the sensor position,  $(X_{k,1} \ X_{k,2})^T$  the target position and  $(\dot{X}_{k,1} \ \dot{X}_{k,2})^T$  the target velocity. The state vector at time  $k\Delta$  is then  $X_k = (X_{k,1} \ \dot{X}_{k,1} \ X_{k,2} \ \dot{X}_{k,2})^T$  and the observation model is

$$Y_k = \arctan \frac{X_{k,2} - s_{k,2}}{X_{k,1} - s_{k,1}} + W_k$$

where  $\{W_k\}_{k \geq 0}$  is a Gaussian white noise with  $W_k \sim \mathcal{N}(0, \sigma^2)$ .

The initial state of the target is modeled a priori by a Gaussian distribution with expectation  $m_0$  and covariance matrix  $\Sigma_0$ .

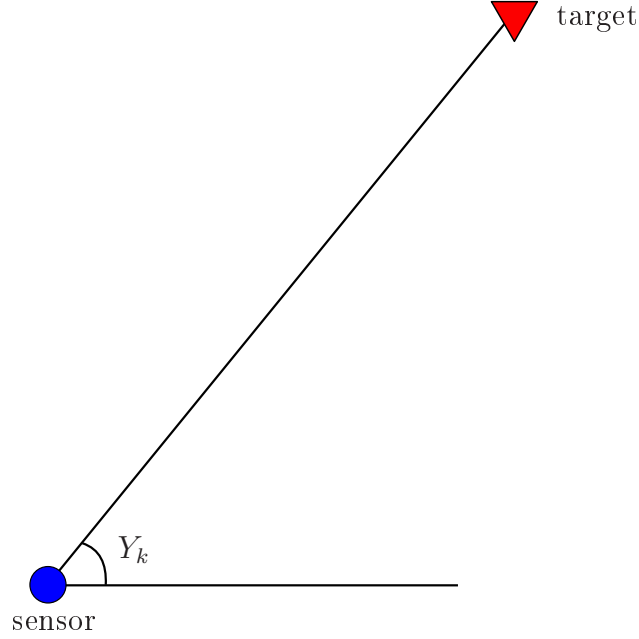


Figure 8.2: Bearings-only observation.

The state dynamics is a Markov process,

$$X_k = FX_{k-1} + V_k,$$

where  $\{V_k\}_{k \geq 0}$  is a Gaussian white noise, with  $V_k \sim \mathcal{N}(0, \Sigma_k)$ , and where

$$F = \begin{pmatrix} 1 & \Delta & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & \Delta \\ 0 & 0 & 0 & 1 \end{pmatrix}.$$

For  $k \geq 1$ , the state noise covariance matrix is

$$\Sigma_k = s \begin{pmatrix} \frac{\Delta^3}{3} & \frac{\Delta^2}{2} & 0 & 0 \\ \frac{\Delta^2}{2} & \Delta & 0 & 0 \\ 0 & 0 & \frac{\Delta^3}{3} & \frac{\Delta^2}{2} \\ 0 & 0 & \frac{\Delta^2}{2} & \Delta \end{pmatrix}. \quad (8.3)$$

### 8.2.2 Simulation parameters

We consider two simulation scenarii, corresponding to different trajectories of the sensor and the target. In scenario 1, the state dynamics is noise-free and the sensor maneuvers continuously. In scenario 2, there is a state dynamics noise and the sensor maneuvers abruptly. These scenarii are shown on figure 8.3.

**Scenario 1.** The expectation of the initial target position is  $\begin{pmatrix} 4000 & 4000 \end{pmatrix}^T$  (in meters) and its initial velocity is  $\begin{pmatrix} 7/\sqrt{2} & 7/\sqrt{2} \end{pmatrix}^T$  (in meters per second). The initial sensor position is  $\begin{pmatrix} 0 & 0 \end{pmatrix}^T$  and its initial course is  $\pi/4$  radians. The sensor moves at speed 15 m/s (like the target) and its course rotates at speed  $-\pi/600$  radians per second. In this scenario, we suppose that the state dynamics is noise-free, i.e.  $s = 0$  in (8.3). A noise-free state dynamics is challenging in particle filtering. Indeed, the generally chosen sampling distribution is the distribution of the state process noise, like in the SIR algorithm (algorithm 8, see chapter 7). When only  $X_0$  is random and the dynamics is noise-free, the particles are sampled initially once according to  $Q_0(dx)$  and then are deterministically propagated. This situation may lead to severe weight degeneracy. The use of kernel regularization is mandatory in this case, as it will be demonstrated in the experimental results below. The lack of model noise must be balanced by an artificial algorithm noise.

**Scenario 2.** The expectation of initial target position is  $\begin{pmatrix} 4000 & 4000 \end{pmatrix}^T$  and its initial velocity is  $\begin{pmatrix} 7 & 0 \end{pmatrix}^T$ . The initial sensor position is  $\begin{pmatrix} 0 & 0 \end{pmatrix}^T$ . Its velocity is  $\begin{pmatrix} 7 & 0 \end{pmatrix}^T$  during the first half of the simulation. At time step  $\lfloor n/2 \rfloor$  (recall  $n$  denote the number of time iterations), the sensor makes a  $2\pi/3$  radians rotation and pursues its course at constant velocity  $\begin{pmatrix} -7/2 & 7\sqrt{3}/2 \end{pmatrix}^T$ . The state dynamics noise covariance (8.3) is such that  $s = 0.1 \text{ m}^2/\text{s}^3$ .

In both scenarii, the period between two observations is  $\Delta = 1 \text{ s}$ . The number of time iterations is  $n = 121$ , so that the duration of the scenario is 120 s. The initial state covariance matrix is  $\Sigma_0 = \text{diag}(1000^2, 2^2, 1000^2, 2^2)$  (in  $\text{m}^2, (\text{m/s})^2, \text{m}^2, (\text{m/s})^2$ ). The observation noise standard deviation  $\sigma$  is set at different values, from  $\sigma = 0.01^\circ$  to  $1^\circ$ .

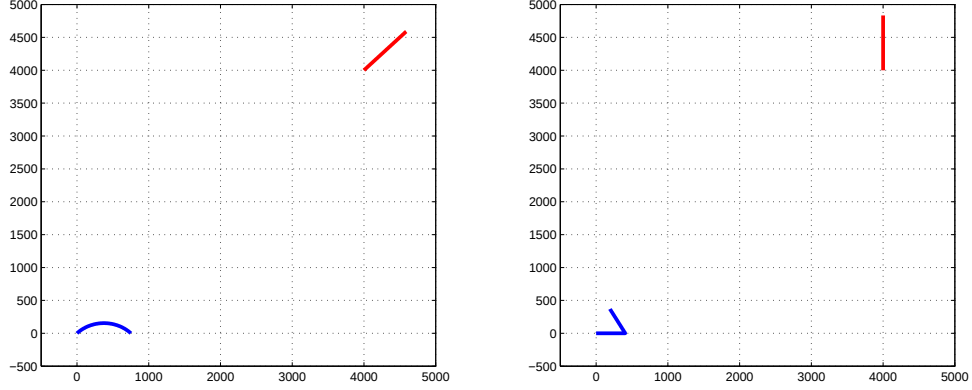


Figure 8.3: scenario 1 (left) and scenario 2 (right); red line: target, blue line: sensor.

### 8.2.3 Results

The KLF appears to be totally inadequate for the bearings-only tracking problem, since it always diverges in our simulations.

We then compare the three other algorithms: the SIR algorithm, the RPF and the LPF. Non-divergence rates are computed for scenario 1 (table 8.1) and scenario 2 (table 8.2), for different values of the observation noise standard deviation  $\sigma$  and of the number of particles  $N$ . We can first notice that the SIR algorithm behaves extremely poorly in scenario 1, which is due to the lack of state dynamics noise. The RPF is more robust to divergence than the SIR algorithm in the two scenarios, but its performance is degraded when  $\sigma$  is small. The LPF appears to be the best algorithm as it almost never diverges, regardless of  $\sigma$ , even with a small number of particles.

When we further reduce the number of particles, the LPF remains very robust to divergence. For  $\sigma = 0.1^\circ$ , the non-divergence rate of the LPF is 86% when  $N = 300$  and 44% when  $N = 100$  for scenario 1, and 97% when  $N = 300$  and 76% when  $N = 100$  for scenario 2. Comparatively, the non-divergence rate of the RPF is 1% when  $N = 300$  and 0% when  $N = 100$  for scenario 1, and 22% when  $N = 300$  and 4% when  $N = 100$  for scenario 2.

The evolution in time of the RMSE of the estimated position and speed are shown for scenario 1 (figure 8.4) and scenario 2 (figure 8.5), for the RPF and the LPF, with  $\sigma = 0.1^\circ$  and  $N = 3000$  and using non-divergent runs. Both algorithms exhibit a similar accuracy in scenario 1. (Notice that, due to the large amount of divergence of the RPF in scenario 1, one computes the RMSE of a biased estimator when considering only

|          | SIR        | RPF | LPF | SIR        | RPF | LPF |
|----------|------------|-----|-----|------------|-----|-----|
| 0.01°    | 00         | 00  | 97  | 00         | 01  | 98  |
| 0.05°    | 00         | 05  | 96  | 00         | 29  | 98  |
| 0.1°     | 00         | 17  | 96  | 00         | 52  | 98  |
| 0.5°     | 00         | 67  | 97  | 04         | 88  | 98  |
| 1°       | 05         | 81  | 97  | 35         | 94  | 97  |
| $\sigma$ | $N = 1000$ |     |     | $N = 3000$ |     |     |

Table 8.1: non-divergence rates (in percentage) for scenario 1 (without state noise).

|          | SIR        | RPF | LPF | SIR        | RPF | LPF |
|----------|------------|-----|-----|------------|-----|-----|
| 0.01°    | 4          | 12  | 99  | 5          | 29  | 100 |
| 0.05°    | 7          | 42  | 100 | 15         | 73  | 100 |
| 0.1°     | 12         | 58  | 100 | 33         | 83  | 100 |
| 0.5°     | 49         | 89  | 99  | 78         | 96  | 100 |
| 1°       | 69         | 93  | 100 | 89         | 98  | 99  |
| $\sigma$ | $N = 1000$ |     |     | $N = 3000$ |     |     |

Table 8.2: non-divergence rates (in percentage) for scenario 2 (with state noise).

non-divergent runs. This explains why the RMSE of the RPF can be slightly below the posterior Cramér-Rao bound.) In scenario 2, the LPF is noticeably better than the RPF, especially for the estimation of the speed.

The average resampling rate over time is plotted in figure 8.6 for the RPF and the LPF and for  $\sigma = 0.1^\circ$  and  $N = 3000$ . One can observe that it decreases over time, as the target becomes more and more observable [Le Cadre and Jauffret (1997)]. The resampling rate of the LPF is slightly below that of the RPF.

Concerning the computational time, a non-divergent run of the LPF requires approximately 1.8 times more time than a non-divergent run of the RPF.

### 8.3 Ballistic target tracking during atmospheric reentry

The second problem we consider is to track a ballistic object falling into the atmosphere, along a vertical line. This is the same problem as in [Ristic et al. (2003)], which is challenging due to the nonlinear nature of the object dynamics in this context.

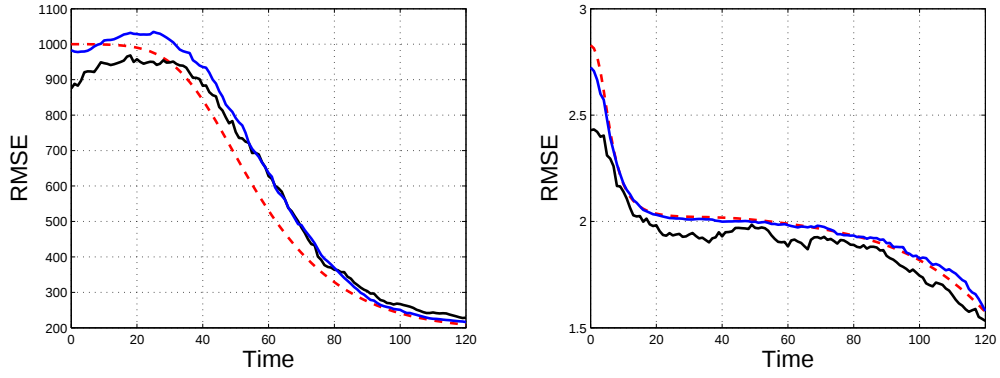


Figure 8.4: RMSE of the estimated position (left) and speed (right) in scenario 1 for  $\sigma = 0.1^\circ$  and  $N = 3000$  (non-divergent runs); black line: RPF, blue line: LPF, red dashed line: posterior Cramér-Rao bound.

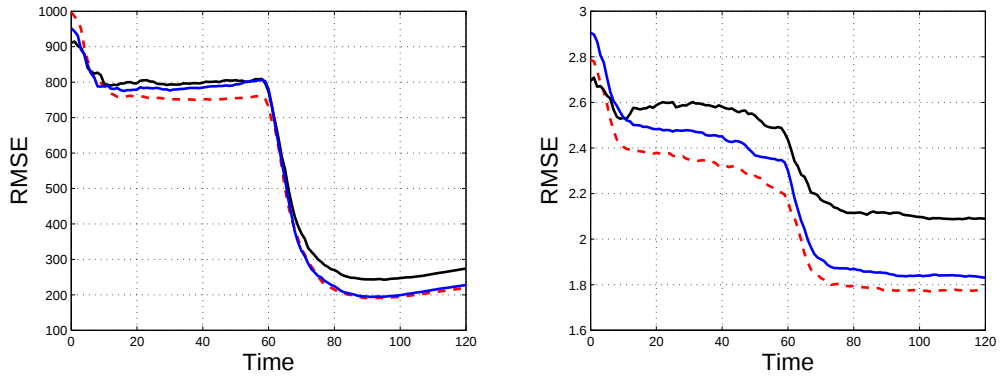


Figure 8.5: RMSE of the estimated position (left) and speed (right) in scenario 2 for  $\sigma = 0.1^\circ$  and  $N = 3000$  (non-divergent runs); black line: RPF, blue line: LPF, red dashed line: posterior Cramér-Rao bound.

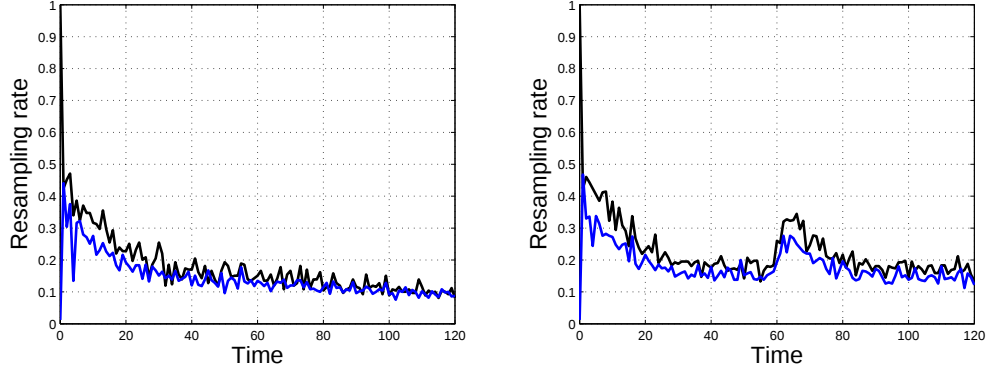


Figure 8.6: resampling rate over time for  $\sigma = 0.1^\circ$  and  $N = 3000$  (non-divergent runs); left: scenario 1 (without state noise), right: scenario 2 (with state noise); black line: RPF, blue line: LPF.

### 8.3.1 Model

At time  $t$ , let  $h_t$  be the altitude of the object,  $v_t$  its velocity and  $\beta_t$  its ballistic coefficient.  $\beta_t$  depends on the object mass, shape and cross-sectional area at time  $t$ . The main forces that act on a ballistic object falling into the atmosphere are weight and drag. We suppose the other forces are neglectible. Newton's second law then yields

$$\dot{v}_t = -\frac{\rho(h_t)gv_t^2}{2\beta_t} + g$$

where  $\rho(h_t)$  is the air density at altitude  $h_t$  and  $g = 9.81 \text{ m/s}^2$  is the weight acceleration.  $\rho$  is defined by  $\rho(h_t) = c_1 e^{-c_2 h_t}$ , where  $c_1 = 1.227$  and  $c_2 = 1.093 \cdot 10^{-4}$  when  $h_t < 9144$ , and  $c_1 = 1.754$  and  $c_2 = 1.49 \cdot 10^{-4}$  when  $h_t \geq 9144$  [Farina et al. (2002)]. The ballistic coefficient is supposed to be constant, i.e.

$$\dot{\beta}_t = 0,$$

which is generally the case for objects with a high super-sonic speed [Farina et al. (2002); Suwantong et al. (2012)].

Let  $X_t = \begin{pmatrix} h_t & v_t & \beta_t \end{pmatrix}^T$  be the state vector at time  $t$ , taking values in  $\mathbb{R}_+^* \times \mathbb{R} \times \mathbb{R}_+^*$ . A sensor measures the altitude of the object every  $\Delta$  seconds. The discrete-time observation model is

$$Y_k = HX_k + W_k$$

for  $k \geq 0$ , where

$$H = \begin{pmatrix} 1 & 0 & 0 \end{pmatrix}$$

and  $\{W_k\}_{k \geq 0}$  is a Gaussian white noise with  $W_k \sim \mathcal{N}(0, \sigma^2)$ .

After discretization, the state dynamics is

$$X_k = F_k(X_{k-1}) + V_k$$

for  $k \geq 1$ , where

$$F_k(X_{k-1}) = F_{\text{dyn}}X_{k-1} - \begin{pmatrix} 0 & \Delta & 0 \end{pmatrix} \left( \frac{\rho(X_{1,k-1})gX_{2,k-1}^2}{2X_{3,k-1}} - g \right),$$

with

$$F_{\text{dyn}} = \begin{pmatrix} 1 & -\Delta & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix}.$$

$X_{j,k}$ , for  $j \in \{1, 2, 3\}$ , denotes the  $j$ th component of the state vector  $X_k$ . The state noise process  $\{V_k\}_{k \geq 0}$  takes into account the model imperfections, such that unaccounted forces acting on the object or small variations of the ballistic coefficient. It is a Gaussian white noise,  $V_k \sim \mathcal{N}(0, \Sigma_k)$  where

$$\Sigma_k = \begin{pmatrix} s\frac{\Delta^3}{3} & s\frac{\Delta^2}{2} & 0 \\ s\frac{\Delta^2}{2} & s\Delta & 0 \\ 0 & 0 & s\beta\Delta \end{pmatrix}$$

when  $k \geq 1$ . The joint distribution of the initial altitude  $X_{1,0}$  and velocity  $X_{2,0}$  is a Gaussian with expectation  $m_0$  and  $2 \times 2$  covariance matrix  $\Sigma_0$ . The initial ballistic coefficient is independent from  $X_{1,0}$  and  $X_{2,0}$  and obeys to a beta distribution scaled on the support  $[\beta^-, \beta^+]$ , with parameters  $\lambda_1, \lambda_2$ . This distribution has the density

$$f(\beta_0) = \frac{\Gamma(\lambda_1 + \lambda_2)}{(\beta^+ - \beta^-)\Gamma(\lambda_1)\Gamma(\lambda_2)} \left( \frac{\beta_0 - \beta^-}{\beta^+ - \beta^-} \right)^{\lambda_1-1} \left( 1 - \frac{\beta_0 - \beta^-}{\beta^+ - \beta^-} \right)^{\lambda_2-1},$$

where  $\Gamma$  is Euler's gamma function. This choice of initial distribution on the ballistic coefficient is made in [Ristic et al. (2003)] since it is suitable for modeling the initial prior knowledge on the target.

This simple ballistic target tracking scenario is outlined in figure 8.7.

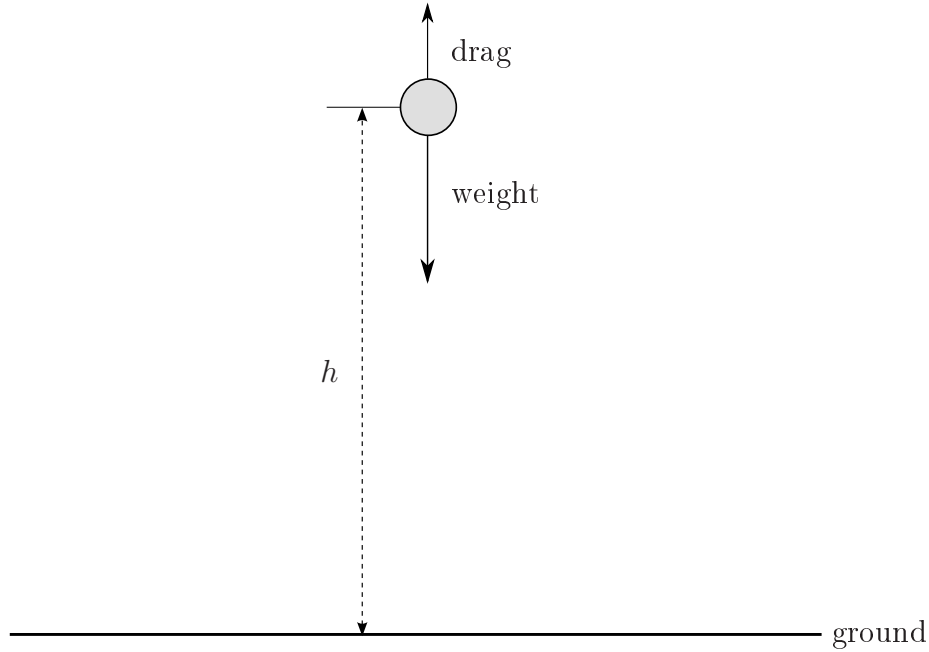


Figure 8.7: Ballistic object tracking scenario.

### 8.3.2 Simulation parameters

The period between two observations is  $\Delta = 0.1$  s and the duration of the object fall is 20 s, so that the number of discrete time iterations is  $n = 201$ . The observation noise standard deviation is set at different values, from  $\sigma = 10$  m to  $\sigma = 200$  m.

The expectation of the initial altitude and velocity is  $m_0 = \begin{pmatrix} 6.1 \cdot 10^4 & 3 \cdot 10^3 \end{pmatrix}^T$  (in m and m/s) and their covariance matrix is  $Q_0 = \text{diag}(10^4, 5 \cdot 10^2)$  (in  $\text{m}^2$  and  $(\text{m/s})^2$ ). The parameters of the initial ballistic coefficient prior distribution are  $\beta^- = 10000$ ,  $\beta^+ = 63000$  (in  $\text{kg}/(\text{ms}^2)$ )(which defines the range of admissible values for  $\beta_0$ ), and  $\lambda_1 = \lambda_2 = 1.1$ .

The parameters of the state noise covariance matrix are set at  $s = s_\beta = 5 \text{ m}^2/\text{s}^3$ .

### 8.3.3 Results

The compared algorithms are the extended Kalman filter (EKF), the SIR algorithm, the RPF and the LPF. Note that the EKF is identical to the KLF for the model considered in this section, since the observation function is linear.

Non-divergence rates are shown in table 8.3, for different values of  $\sigma$  and  $N$ . The SIR algorithm is totally unadapted to this tracking problem since it diverges for every

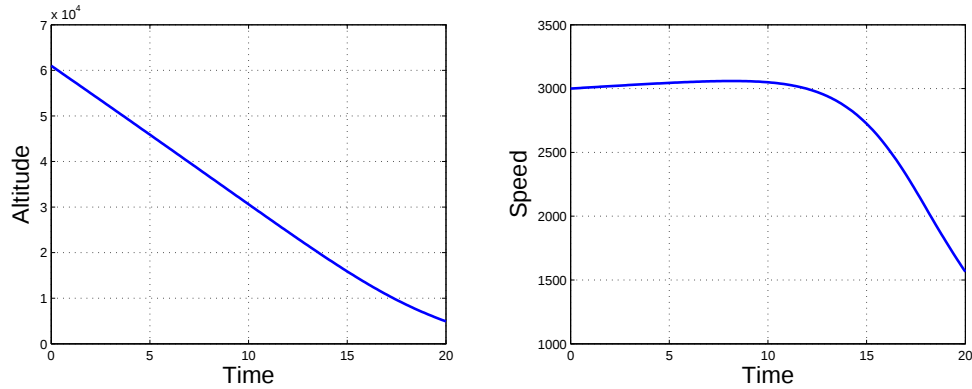


Figure 8.8: evolution in time of the altitude (in m, left) and speed (in m/s, right) of the falling object.

run. The performance of RPF degrades as  $\sigma$  becomes small. On the other hand, the EKF diverges less than the RPF when  $\sigma$  is small and more when  $\sigma$  is large. The LPF is the most robust to divergence for all the values of  $\sigma$ .

The LPF is also very robust to divergence when the number of particles is very small in this application. For  $\sigma = 100$  m, the non-divergence rate of the LPF is 96% when  $N = 300$  and 78% when  $N = 100$ . On the other hand, the non-divergence rate of the RPF is 9% when  $N = 300$  and 0% when  $N = 100$ .

|          | SIR        | RPF | LPF | SIR        | RPF | LPF | EKF |
|----------|------------|-----|-----|------------|-----|-----|-----|
| 10 m     | 00         | 01  | 100 | 00         | 11  | 100 | 89  |
| 20 m     | 00         | 06  | 100 | 00         | 24  | 100 | 81  |
| 100 m    | 00         | 40  | 97  | 00         | 75  | 98  | 50  |
| 200 m    | 00         | 59  | 97  | 00         | 84  | 96  | 42  |
| $\sigma$ | $N = 1000$ |     |     | $N = 3000$ |     |     |     |

Table 8.3: non-divergence rates (in percentage).

The RMSE of the estimated position and speed are plotted in figure 8.9 and the RMSE of the ballistic coefficient is plotted in figure 8.10, for  $\sigma = 100$  m and  $N = 3000$ , using non-divergent runs. The EKF, the RPF and the LPF have approximately the same accuracy when they do not diverge.

Figure 8.11 shows the resampling rate over time for the RPF and the LPF, for  $\sigma = 100$  m and  $N = 3000$ . The resampling rate of the LPF is a little lower than that of the RPF at the beginning of the scenario.

The RPF and the LPF are computationally as fast as each other for non-divergent

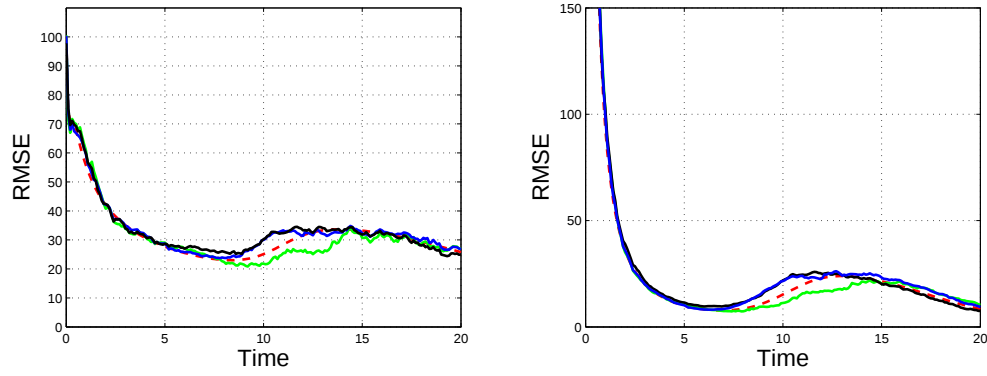


Figure 8.9: RMSE of the estimated position (left) and speed (right) for  $\sigma = 100$  m and  $N = 3000$  (non-divergent runs); black line: RPF, blue line: LPF, green line: EKF, red dashed line: posterior Cramér-Rao bound.

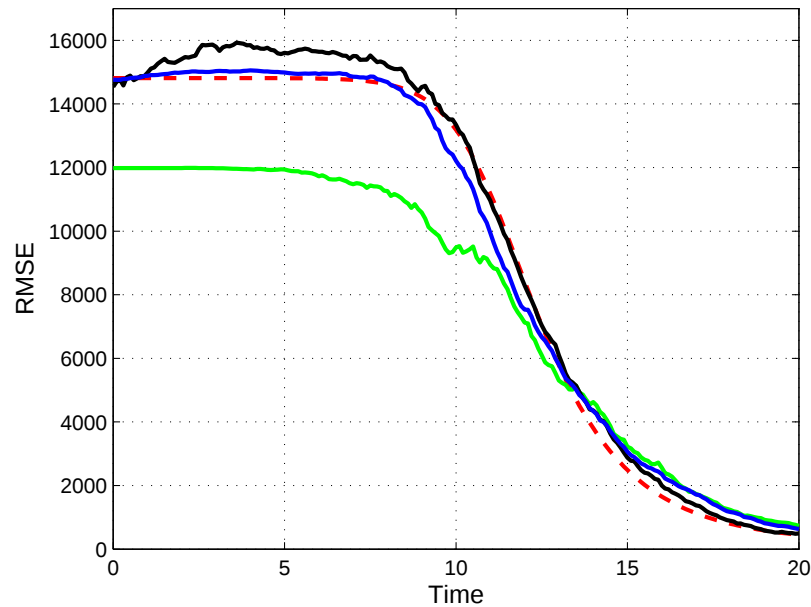


Figure 8.10: RMSE of the estimated ballistic coefficient for  $\sigma = 100$  m and  $N = 3000$  (non-divergent runs); black line: RPF, blue line: LPF, green line: EKF, red dashed line: posterior Cramér-Rao bound.

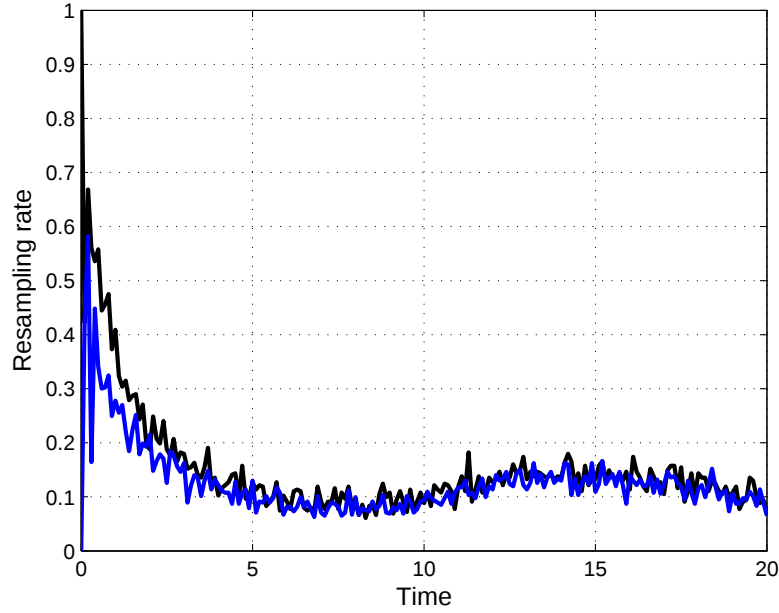


Figure 8.11: resampling rate over time for  $\sigma = 100$  m and  $N = 3000$  (non-divergent runs); black line: RPF, blue line: LPF.

runs in this application. The EKF is approximately 48 times faster than both of them when  $N = 3000$ .

## 8.4 Neural decoding

We consider a third simulation experiment which differs from classical target tracking. In this section, the observation model is not in the form of a nonlinear function of the hidden state corrupted by an additive Gaussian noise (i.e.,  $Y_k = H_k(X_k) + W_k$ ) as previously. The problem we are interested in here is called neural decoding and comes from neurosciences. It consists in extracting information from the brain activity, which is measured by the electrical activity of a set of neurons, thanks to statistical signal processing methods. In the recent years, neural decoding has thus been addressed as a nonlinear Bayesian filtering problem [Gao et al. (2002); Brockwell et al. (2004); Wu et al. (2006); Brockwell et al. (2007); Koyama et al. (2010)].

### 8.4.1 Model

The state–space model we consider here is similar to that in [Brockwell et al. (2007)] and [Koyama et al. (2010)].

The electrical activity of a set of neurons from an animal (e.g., a monkey) is recorded. The observation of one neuron’s electrical activity consists in a count of spikes, each spike being a fast change of voltage. The hidden state describes the motion of the animal (e.g., the hand motion) while its neural activity is monitored. Here, it is the 3-dimensional Cartesian position and corresponding velocity of the animal’s hand in some reference coordinates system, so that the state space dimension is  $d = 6$ .

The number of neurons monitored is  $M$ . The discrete–time observation process is  $\{Y_k\}_{k \geq 0}$  with  $Y_k = (Y_{1,k}, \dots, Y_{M,k})$ ,  $Y_{j,k}$  being the count of spikes fired by neuron  $j \in \{1, \dots, M\}$  in the short time interval  $[k\Delta, (k+1)\Delta]$ . Thus,  $Y_k$  takes values in  $\mathbb{N}^M$ . The distribution of  $Y_{j,k}$  is Poisson with intensity  $\lambda_j(X_k)\Delta$ ,

$$Y_{j,k} \sim \mathcal{P}(\lambda_j(X_k)\Delta),$$

$X_k$  denoting the hidden state at time  $k$ . The activity of each neuron is supposed to be independent. The likelihood function associated with observation  $Y_k$  is then

$$g_k(x) = \prod_{j=1}^M g_{j,k}(x),$$

where

$$g_{j,k}(x) = e^{-\lambda_j(x)\Delta} \frac{(\lambda_j(x)\Delta)^{Y_{j,k}}}{Y_{j,k}!}$$

is the likelihood associated with observation  $Y_{j,k}$  from neuron  $j$ . The relation between the intensity parameter of the Poisson observation model and the hidden state is

$$\lambda_j(x) = \exp(\alpha_j + \beta_j^T x)$$

for all  $x \in \mathbb{R}^d$ , where  $\alpha_i \in \mathbb{R}$  and  $\beta_i \in \mathbb{R}^d$ .

The state vector is

$$X_k = \begin{pmatrix} X_{1,k} & \dot{X}_{1,k} & X_{2,k} & \dot{X}_{2,k} & X_{3,k} & \dot{X}_{3,k} \end{pmatrix}^T$$

where  $(X_{1,k} \ X_{2,k} \ X_{3,k})^T$  is the hand position vector at time  $k\Delta$  and  $(\dot{X}_{1,k} \ \dot{X}_{2,k} \ \dot{X}_{3,k})^T$  is the average hand velocity during the time bin  $[k\Delta, (k+1)\Delta]$ . The state process is modeled as a linear Markov chain,

$$X_k = FX_{k-1} + V_k$$

where  $\{V_k\}_{k \geq 0}$  is a Gaussian white noise,  $V_k \sim \mathcal{N}(0, \Sigma_k)$ . The state transition matrix is

$$F = \begin{pmatrix} 1 & \Delta & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & \Delta & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & \Delta \\ 0 & 0 & 0 & 0 & 0 & 1 \end{pmatrix}$$

and the state noise covariance matrix is

$$\Sigma_k = s \begin{pmatrix} \frac{\Delta^3}{3} & \frac{\Delta^2}{2} & 0 & 0 & 0 & 0 \\ \frac{\Delta^2}{2} & \Delta & 0 & 0 & 0 & 0 \\ 0 & 0 & \frac{\Delta^3}{3} & \frac{\Delta^2}{2} & 0 & 0 \\ 0 & 0 & \frac{\Delta^2}{2} & \Delta & 0 & 0 \\ 0 & 0 & 0 & 0 & \frac{\Delta^3}{3} & \frac{\Delta^2}{2} \\ 0 & 0 & 0 & 0 & \frac{\Delta^2}{2} & \Delta \end{pmatrix}$$

for  $k \geq 1$ . The initial state distribution is Gaussian with expectation  $m_0$  and covariance  $\Sigma_0$ .

### 8.4.2 Simulation parameters

The duration of the time bins is  $\Delta = 30 \cdot 10^{-3}$  s and the duration of the experiment is 1.5 s, so that the number of discrete time steps is  $n = 51$ .

The number of monitored neurons is  $M = 100$ . The likelihood parameters do not have a clear "physical sense" in this application, unlike in classical target tracking. The  $\alpha_j$ 's and  $\beta_j$ 's ( $j \in \{1, \dots, M\}$ ) are somehow features of the  $M$  monitored neurons. In [Brockwell et al. (2007)], the likelihood parameters are learned thanks to a MCMC algorithm that is run before filtering. Here, we do not perform this learning step. We rather randomly draw the parameters  $\alpha_j$  and  $\beta_j$  once, before the 500 runs of the scenario,

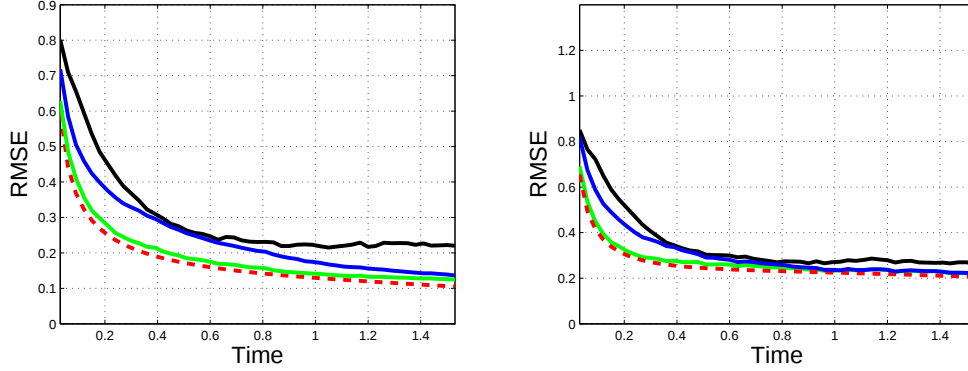


Figure 8.12: RMSE of the estimated position (left) and speed (right) (non-divergent runs); black line: RPF, blue line: LPF, green line: KLF, red dashed line: posterior Cramér-Rao bound.

like in the simulation study in [Koyama et al. (2010)]: the  $\alpha_j$ 's are independently sampled from the distribution  $\mathcal{N}(2.5, 1)$  and the  $\beta_j$ 's are independently sampled from the uniform distribution on the unit sphere in  $\mathbb{R}^6$ .

The parameter in the state noise covariance matrix is  $s = 0.05^2/\Delta \simeq 0.0833 \text{ m}^2/\text{s}^3$ . The parameters of the initial state distribution are  $m_0 = 0$  and  $\Sigma_0 = I_d$ .

### 8.4.3 Results

We compare the SIR algorithm, the RPF, the LPF and the KLF. The number of particles is set to  $N = 3000$  (recall that the KLF does not involve particles).

The non-divergence rates of the SIR algorithm, the RPF, the LPF and the KLF are respectively: 10%, 95%, 83% and 96%. The RPF and the KLF are thus more robust to divergence than the LPF for this model. The SIR algorithm is not adapted.

The RMSE of the estimated position and speed are plotted in figure 8.12. The better algorithm in terms of RMSE is the KLF. The LPF is more accurate than the RPF.

We have observed that initializing the LPF as a RPF allows to improve its robustness to divergence in this application. We set it as follows: from  $k = 0$  to 9 (i.e., for 0.3 s), the algorithm is a classical RPF, then at  $k = 10$  the algorithm switches to a LPF and it remains in this regime till the end of the scenario at time  $k = n = 51$ . The non-divergence rate of the LPF with this initialization increases to 94%. Its RMSE is shown in figure 8.13 where it is called the delayed-LPF.

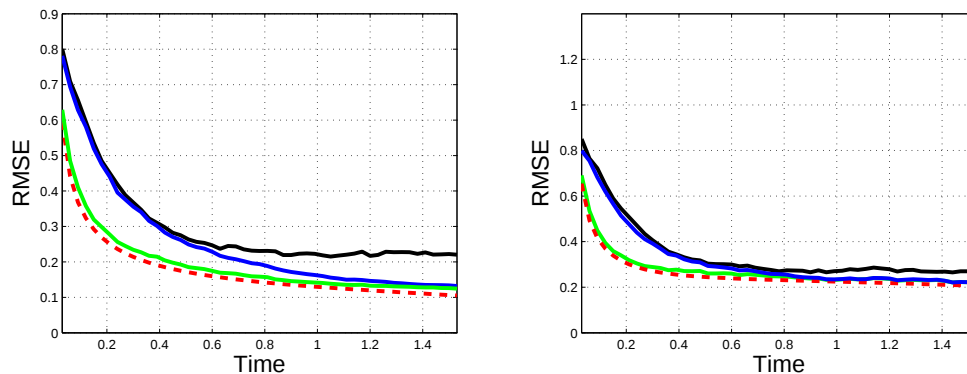


Figure 8.13: RMSE of the estimated position (left) and speed (right) (non-divergent runs); black line: RPF, blue line: delayed-LPF, green line: KLF, red dashed line: posterior Cramér-Rao bound.

The resampling rate for the RPF and the LPF is plotted in figure 8.14. The resampling rate of the LPF is approximately half that of the RPF in this application.

The computational times of the LPF and the KLF are respectively approximately 5.5 times and 12 times that of the RPF for non-divergent runs. Note that, in this problem, the optimization step for the computation of the MAP in the LPF and the KLF is performed over the whole (6-dimensional) state space. Indeed, the likelihood function depends here on the whole state vector, so that maximization cannot be performed over a lower dimension subspace (see remark 7.3.1 in chapter 7). Also note that, in the KLF, the computation of the MAP is performed at each time step, whereas it is performed only when  $N_{\text{eff}} < N$  in the LPF. This explains why the KLF is slower than the LPF although it does not involve particles.

## Conclusion

The simulation experiments in this chapter have demonstrated the potential of the LPF and the KLF. The classical SIR algorithm is ineffective in the three filtering problems we have considered.

In the bearings-only target tracking (section 8.2) and the ballistic target tracking (section 8.3) applications, the LPF is much more robust to divergence than the RPF. In the neural decoding application (section 8.4), the LPF is less robust to divergence than the RPF, but its robustness can be improved thanks to a different initialization. Regarding accuracy, the LPF error attains the Cramér-Rao lower bound. The good

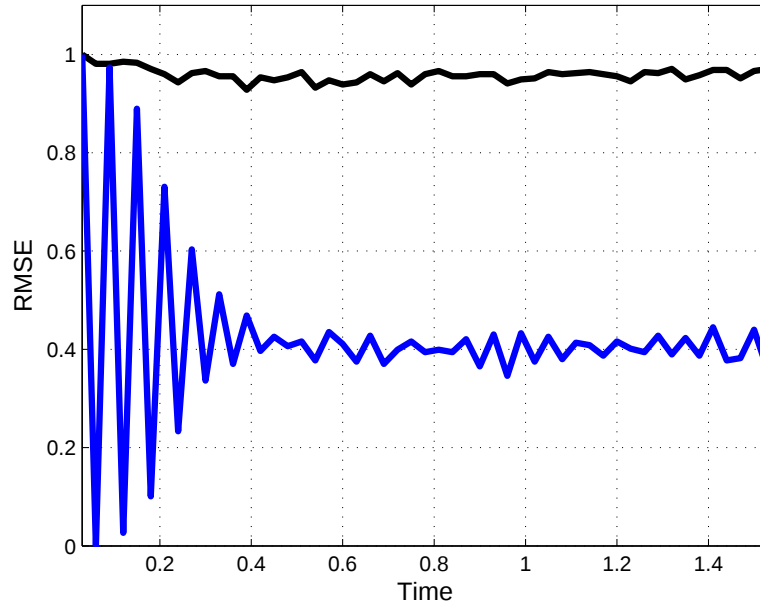


Figure 8.14: resampling rate over time; black line: RPF, blue line: LPF.

performance of the LPF is particularly pronounced in small noise models, which are generally difficult to handle in particle filtering, but it is not restricted to them.

The KLF is not adapted to the bearings-only target tracking problem and it is equivalent to the EKF in the ballistic target tracking application (because the observation model is linear and Gaussian). It is however extremely efficient in the neural decoding problem, both in terms of robustness to divergence and accuracy, although it is slower than particle filters.



# Conclusion

Nonlinear filtering is a Bayesian state estimation problem in state–space models, also known as hidden Markov models. When the model is nonlinear, approximation methods must be used in order to compute at each time step the conditional distribution of the hidden state given the past observations (i.e., the posterior).

Particle filtering is a sequential Monte Carlo method that recursively approximates the Bayesian filter as a sequence of weighted sum of Dirac measures. The approximation is consistent as the number of particles increases, which makes it a popular filtering algorithm. However, particle filtering algorithms suffer in practice from some drawbacks. In certain situations, the phenomenon of weight degeneracy is observed, which causes the impoverishment of the particle approximation. This phenomenon typically occurs when the model is highly informative (when the state dynamics noise or the observation noise are small, e.g.) or highly dimensional.

The Laplace method is an approximation technique for multidimensional integral which is widely used for nonlinear Bayesian estimation in static models. When the model is identifiable, it yields consistent approximations (at a fast rate) when the observation sample size increases or the observation noise intensity decreases, i.e. when the information brought by the likelihood model increases. Thus, the Laplace method is accurate precisely in a situation where particle approximation is poor. It is consequently a natural idea to combine both techniques.

The main contributions of this work are the following.

- **The use of the Laplace method within particle filtering.** The resulting algorithm, called the Laplace particle filter (LPF), is tested in practice on tracking problems and is demonstrated to be efficient, particularly when the system noise is small (chapters 6, 7, 8).
- **The study of an approximation method of the Bayesian filter based on the Laplace method only (without particles).** The propagation of the

approximation error over time is analyzed when the observation noise intensity tends to zero. This study was led under a strong assumption on the likelihood model that allows to apply the Laplace method at the prediction step and the update step (chapter 4).

The other contributions of the thesis are the following.

- **The analysis of the behaviour of importance sampling in difficult situations (high information or high dimensional model).** This analysis is done in an asymptotic framework that allows to apply the Laplace method. The Laplace method is used to compute the asymptotic variance of the importance weights, which describes the quality of importance sampling for a given model (chapter 2).
- **The derivation of multidimensional Laplace approximations formulas for the posterior expectation and covariance matrix.** These formulas were derived in the one-dimensional case by Tierney et al. (1989), we provide their multidimensional versions (which are used in practice in the LPF) using matrix calculus (chapter 3).
- **A recursive nonlinear filter based on the extended Kalman filter and the Laplace method.** The prediction step is performed like in the extended Kalman filter and the update step is done thanks to the Laplace method. This algorithm is called the Kalman Laplace filter (KLF) (chapter 5).
- **The comparative experimental analysis of the KLF and the LPF in the context of engineering applications.** The three problems we have considered are: bearings-only target tracking, ballistic target tracking during atmospheric reentry, neural decoding (chapter 8).

To end this dissertation, we state below some perspectives to extend and generalize our work.

- **The analysis of the consistency of the KLF and of the propagation of its approximation error over time.** The asymptotic framework is when the state dynamics noise and the observation noise tends to zero (at the same speed), so that the noise intensity  $\varepsilon$  is taken as the expansion parameter in the Laplace method.

- **The asymptotic study of the LPF.** It would be interesting to derive error bounds for the approximated Bayesian filter as  $N \rightarrow \infty$  (which improves particles approximation) and  $\varepsilon \rightarrow 0$  (which degrades particles approximation and improves Laplace approximation).
- **The adaptation of the LPF to models where the posterior is highly multimodal.** This may be done by adding a clustering step to the algorithm (like in [Murangira et al. (2011)], e.g.) and apply the Laplace method into each cluster.
- **Investigate if the LPF is efficient in high dimensional models.** These models typically arise in data assimilation and particle filters are notoriously inefficient in this application (see [Snyder et al. (2008)], e.g.).
- **The adaptation of the LPF to smoothing problems.**



# Appendix A

## Differentiation of the log-likelihood

In this appendix, we provide the derivatives of the log-likelihood which are needed to compute Laplace approximations. Differentiation is made thanks to the matrix calculus rules presented in [Magnus (2010)] and [Fackler (2005)].

### A.1 Bearings-only target tracking

In the bearings-only target tracking application, the likelihood function is in the form of

$$g(x) = |2\pi\sigma^2|^{-1/2} \exp\left(-\frac{1}{2\sigma^2}|Y - H(x)|^2\right),$$

where the nonlinear observation function is

$$H(x) = \arctan \frac{x_2}{x_1}$$

and  $x = \begin{pmatrix} x_1 & \dot{x}_1 & x_2 & \dot{x}_2 \end{pmatrix}^T \in E$  is the state vector.

The first, second and third-order derivatives of  $h$  are respectively

$$H'(x) = \frac{1}{x_1^2 + x_2^2} \begin{pmatrix} -x_2 & 0 & x_1 & 0 \end{pmatrix},$$

$$H''(x) = \frac{1}{(x_1^2 + x_2^2)^2} \begin{pmatrix} 2x_1x_2 & 0 & x_2^2 - x_1^2 & 0 \\ 0 & 0 & 0 & 0 \\ x_2^2 - x_1^2 & 0 & -2x_1x_2 & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix},$$

and

$$H'''(x) = \frac{2}{(x_1^2 + x_2^2)^3} \begin{pmatrix} -3x_1^2x_2 + x_2^3 & 0 & x_1^3 - 3x_1x_2^2 & 0 \\ 0 & 0 & 0 & 0 \\ x_1^3 - 3x_1x_2^2 & 0 & 3x_1^2x_2 - x_2^3 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ x_1^3 - 3x_1x_2^2 & 0 & 3x_1^2x_2 - x_2^3 & 0 \\ 0 & 0 & 0 & 0 \\ 3x_1^2x_2 - x_2^3 & 0 & -x_1^3 + 3x_1x_2^2 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix}.$$

Besides, the second and third-order derivatives of the log-likelihood are respectively

$$(\log g)''(x) = \frac{1}{\sigma^2}(y - H(x))H''(x) - \frac{1}{\sigma^2}H'(x)^T H'(x),$$

and

$$\begin{aligned} (\log g)'''(x) &= \frac{1}{\sigma^2} (-\text{vec}[H''(x)]H'(x) + (y - H(x)) \otimes H'''(x)) \\ &\quad - \frac{1}{\sigma^2} ((H'(x)^T \otimes I_d)H''(x) + (I_d \otimes H'(x)^T)H''(x)). \end{aligned}$$

$(\log g)^{(4)}(x)$  (which is a  $64 \times 6$  matrix) has been computed by numerical differentiation.

## A.2 Ballistic target tracking during atmospheric reentry

In the reentry target tracking application, the observation function is linear and the likelihood function is in the form of

$$g(x) = |2\pi\sigma^2|^{-1/2} \exp\left(-\frac{1}{2\sigma^2}|Y - Hx|^2\right),$$

where

$$H = \begin{pmatrix} 1 & 0 & 0 \end{pmatrix}.$$

The Hessian of the log-likelihood is

$$(\log g)''(x) = -\frac{1}{\sigma^2} H^T H.$$

The third and fourth-order derivatives are zero,

$$(\log g)'''(x) = 0$$

and

$$(\log g)^{(4)}(x) = 0.$$

## A.3 Neural decoding

In the neural decoding application, the likelihood function is in the form of

$$g(x) = \prod_{j=1}^M g_j(x),$$

where

$$g_j(x) = e^{-\lambda_j(x)\Delta} \frac{(\lambda_j(x)\Delta)^{Y_j}}{Y_j!}$$

and

$$\lambda_j(x) = e^{\alpha_j + \beta_j^T x}.$$

The log-likelihood and its derivatives verifies

$$(\log g)^{(l)} = \sum_{j=1}^M (\log g_j)^{(l)}$$

for all  $l \geq 0$ . The second, third and fourth-order derivatives are respectively

$$(\log g_j)''(x) = -\Delta e^{\alpha_j + \beta_j^T x} \beta_j \beta_j^T,$$

$$(\log g_j)'''(x) = -\Delta e^{\alpha_j + \beta_j^T x} \text{vec}[\beta_j \beta_j^T] \beta_j^T,$$

and

$$(\log g_j)^{(4)}(x) = -\Delta e^{\alpha_j + \beta_j^T x} \text{vec} [\text{vec}[\beta_j \beta_j^T] \beta_j^T] \beta_j^T.$$

# Bibliography

- Arulampalam, M. S., Maskell, S., Gordon, N., and Clapp, T. (2002). A tutorial on particle filters for online nonlinear/non-Gaussian Bayesian tracking. *IEEE Transactions on Signal Processing*, 50(2):174–188.
- Bengtsson, T., Bickel, P., and Li, B. (2008). Curse-of-dimensionality revisited: Collapse of the particle filter in very large scale systems. *Probability and Statistics: Essays in Honor of David A. Freedman*, 2:316–334.
- Bengtsson, T., Snyder, C., and Nychka, D. (2003). Toward a nonlinear ensemble filter for high dimensional systems. *Journal of Geophysical Research*, 108(D24):8775.
- Berliner, L. M. and Wickle, C. K. (2007). Approximate importance sampling Monte Carlo for data assimilation. *Physica D*, 230:37–49.
- Beskos, A., Crisan, D., and Jasra, A. (2011). On the Stability of Sequential Monte Carlo Methods in High Dimensions. *arXiv preprint, arXiv:1103.3965*.
- Bickel, P., Li, B., and Bengtsson, T. (2008). Sharp failure rates for the bootstrap particle filter in high dimensions. *Pushing the Limits of Contemporary Statistics: Contributions in Honor of Jayanta K. Ghosh*, 3:318–329.
- Brockwell, A. E., Kass, R. E., and Schwartz, A. B. (2007). Statistical Signal Processing and the Motor Cortex. *Proceedings of the IEEE*, 95(5):881–898.
- Brockwell, A. E., Rojas, A. L., and Kass, R. E. (2004). Recursive Bayesian Decoding of Motor Cortical Signals by Particle Filtering. *Journal of Neurophysiology*, 91:1899–1907.
- Bui Quang, P., Musso, C., and Le Gland, F. (2010). An insight into the issue of dimensionality in particle filtering. In *Proceedings of the 13th International Conference on Information Fusion*, Edinburgh, UK, July 26–29.

- Bui Quang, P., Musso, C., and Le Gland, F. (2012). Multidimensional Laplace Formulas for Nonlinear Bayesian Estimation. In *Proceedings of the 20th European Signal Processing Conference (EUSIPCO 2012)*, Bucharest, Romania, August 27–31.
- Cappé, O., Godsill, S., and Moulines, E. (2007). An Overview of Existing Methods and Recent Advances in Sequential Monte Carlo. *Proceedings of the IEEE*, 95(5):899–924.
- Chopin, N. (2004). Central Limit Theorem for sequential Monte Carlo methods and its application to Bayesian inference. *The Annals of Statistics*, 32(6):2385–2411.
- Chopin, N. and Pelgrin, F. (2004). Bayesian inference and state number determination for hidden Markov models: an application to the information content of the yield curve about inflation. *Journal of Econometrics*, 123(2):327–344.
- Chorin, A. J. and Krause, P. (2004). Dimensional reduction for a Bayesian filter. *Proceedings of the National Academy of Sciences*, 101(42):15013–15017.
- Crisan, D. and Doucet, A. (2002). A survey of convergence results on particle filtering methods for practitioners. *IEEE Transactions on Signal Processing*, 50(3):736–746.
- Daum, F., Co, R., and Huang, J. (2003). Curse of Dimensionality and Particle Filters. In *Proceedings of the IEEE Aerospace Conference*, Big Sky, MT, USA, March 2003.
- Del Moral, P. (2004). *Feynman–Kac Formulae: Genealogical and Interacting Particle Systems with Applications*. Springer.
- Dieudonné, J. (1960). *Foundations of Modern Analysis*. Academic Press.
- Doucet, A., Godsill, S., and Andrieu, C. (2000). On sequential Monte Carlo sampling methods for Bayesian filtering. *Statistics and Computing*, 10(3):197–208.
- Eddy, S. R. (1996). Hidden Markov models. *Current Opinion in Structural Biology*, 6:361–365.
- Efron, B. and Hinkley, D. V. (1978). Assessing the accuracy of the maximum likelihood estimator: Observed versus expected Fisher information. *Biometrika*, 65(3):457–482.
- Evans, M. J. and Swartz, T. (1995). Methods for Approximating Integrals in Statistics with Special Emphasis on Bayesian Integration Problems. *Statistical Science*, 10(3):254–272.

- Fackler, P. L. (2005). Notes on Matrix Calculus. Technical report, North Carolina State University.
- Farina, A., Ristic, B., and Benvenuti, D. (2002). Tracking a Ballistic Target: Comparison of Several Nonlinear Filters. *IEEE Transactions on Aerospace and Electronic Systems*, 38(3):854–867.
- Fink, D. (1997). A compendium of conjugate priors. [url: <http://citeseerx.ist.psu.edu/viewdoc/download?doi=10.1.1.157.5540&rep=rep1&type=pdf>].
- Gao, Y., Black, M. J., Bienenstock, E., Shoham, S., and Donoghue, J. P. (2002). Probabilistic inference of hand motion from neural activity in motor cortex. In Dietterich, T. G., Becker, S., and Ghahramani, Z., editors, *Advances in Neural Information Processing Systems 14: Proceedings of the 2002 Conference*, volume 1. MIT Press.
- Ghosh, J. K. (1994). *Higher Order Asymptotics*. NSF-CBMS Regional Conference Series in Probability and Statistics. Institute of Mathematical Statistics.
- Gordon, N. J., Salmond, D. J., and Smith, A. F. M. (1993). Novel approach to nonlinear/non-Gaussian Bayesian state estimation. *IEE Proceedings-F*, 140(2):107–113.
- Guihenneuc-Jouyaux, C. and Rousseau, J. (2005). Laplace expansions in Markov chain Monte Carlo algorithms. *Journal of Computational and Graphical Statistics*, 14(1):75–94.
- Gustafsson, F. (2010). Particle Filter Theory and Practice with Positioning Applications. *IEEE Aerospace and Electronic Systems Magazine*, 27(7):53–81.
- Holmström, L. and Klemelä, J. (1992). Asymptotic Bounds for the Expected  $L^1$  Error of a Multivariate Kernel Density Estimator. *Journal of Multivariate Analysis*, 42:245–266.
- Hue, C., Le Cadre, J., and Pérez, P. (2002). Tracking multiple objects with particle filtering. *IEEE Transactions on Aerospace and Electronic Systems*, 38(3):791–812.
- Joannides, M. and Le Gland, F. (2002). Small noise asymptotics of the Bayesian estimator in nonidentifiable models. *Statistical Inference for Stochastic Processes*, 5(1):95–130.

- Johnson, R. A. (1970). Asymptotic Expansions Associated with Posterior Distributions. *The Annals of Mathematical Statistics*, 41(3):851–864.
- Juang, B. H. and Rabiner, L. R. (1991). Hidden Markov models for speech recognition. *Technometrics*, 33(3):251–272.
- Julier, S. and Uhlmann, J. (2004). Unscented filtering and nonlinear estimation. *Proceedings of the IEEE*, 92(3):401–422.
- Jungbacker, B. and Koopman, S. J. (2007). Monte Carlo estimation for nonlinear non-Gaussian state space models. *Biometrika*, 94(4):827–839.
- Kalman, R. E. (1960). A New Approach to Linear Filtering and Prediction Problems. *Journal of Basic Engineering*, 82(1):35–45.
- Kass, R. E., Tierney, L., and Kadane, J. B. (1988). Asymptotics in Bayesian Computation. In DeGroot, M. H., Lindley, D. V., and Smith, A. F. M., editors, *Bayesian Statistics 3*, pages 261–278. Oxford University Press.
- Kass, R. E., Tierney, L., and Kadane, J. B. (1990). The validity of posterior expansions based on Laplace’s method. In Geisser, S., Hodges, J. S., Press, S. J., and Zellner, A., editors, *Bayesian and Likelihood Methods in Statistics and Econometrics: essays in honor of George A. Barnard*, pages 473–488. North-Holland.
- Kass, R. E., Tierney, L., and Kadane, J. B. (1991). Laplace’s method in Bayesian analysis. In Flournoy, N. and Tsutakawa, R. K., editors, *Contemporary Mathematics: Statistical Multiple Integration*, volume 115, pages 89–99. American Mathematical Society.
- Keziou, A. (2003). Dual representation of  $\phi$ -divergences and applications. *Comptes rendus de l’Académie des sciences. Série 1, Mathématique*, 336:857–862.
- Kleppe, T. S. and Skaug, H. J. (2012). Fitting general stochastic volatility models using Laplace accelerated sequential importance sampling. *Computational Statistics & Data Analysis*, 56(11):3105–3119.
- Kong, A. (1992). A Note on Importance Sampling using Standardized Weights. Technical Report 348, University of Chicago, Dept. of Statistics.

- Koyama, S., Pérez-Bolde, L. C., Shalizi, C. R., and Kass, R. E. (2010). Approximate Methods for State-Space Models. *Journal of the American Statistical Association*, 105(489):170–180.
- Kuk, A. Y. C. (1999). Laplace Importance Sampling for Generalized Linear Mixed Models. *Journal of Statistical Computation and Simulation*, 63(2):146–158.
- Le Cadre, J.-P. and Jauffret, C. (1997). Discrete-Time Observability and Estimability Analysis for Bearings-Only Target Motion Analysis. *IEEE Transactions on Aerospace and Electronic Systems*, 33(1):178–201.
- Le Gland, F. (2012). Filtrage Bayésien et Approximation Particulaire. Cours de l’École Nationale Supérieure de Techniques Avancées [url: <http://www.irisa.fr/aspi/legland/ensta/cours.pdf>].
- Le Gland, F., Musso, C., and Oudjane, N. (1998). An analysis of regularized interacting particle methods in nonlinear filtering. In *Proceedings of the 3rd IEEE European Workshop on Computer-Intensive Methods in Control and Signal Processing*, Prague, Czech Republic, September 7–9.
- Magnus, J. R. (2010). On the concept of matrix derivative. *Journal of Multivariate Analysis*, 101:2200–2206.
- Murangira, A., Musso, C., Dahia, K., and Allard, J.-M. (2011). Robust regularized particle filter for terrain navigation. In *Proceedings of the 14th International Conference on Information Fusion*, Chicago, IL, USA, July 5–8.
- Musso, C. (1993). *Méthodes rapides d’estimation en trajectographie par mesures d’azimuts*. PhD thesis, Université Joseph Fourier Grenoble I.
- Musso, C., Bui Quang, P., and Le Gland, F. (2011). Introducing the Laplace approximation in particle filtering. In *Proceedings of the 14th International Conference on Information Fusion*, Chicago, IL, USA, July 5–8.
- Musso, C. and Oudjane, N. (1998). Regularisation Schemes for Branching Particle Systems as a Numerical Solving Method of the Nonlinear Filtering Problem. In *Proceedings of the Irish Signals Systems Conference*, Dublin, Ireland, June.

- Musso, C. and Oudjane, N. (2005). Data reduction for particle filters. In *Proceedings of the 4th International Symposium on Image and Signal Processing and Analysis*, Zagreb, Croatia, September 15–17.
- Musso, C., Oudjane, N., and Le Gland, F. (2001). Improving regularised particle filters. In Doucet, A., de Freitas, N., and Gordon, N., editors, *Sequential Monte Carlo methods in practice*. Springer.
- Oudjane, N. and Musso, C. (2000). Progressive correction for regularized particle filters. In *Proceedings of the 3rd International Conference on Information Fusion*, Paris, France, July 10–13.
- Pinheiro, J. C. and Bates, D. M. (1995). Approximations to the Loglikelihood Function in the Nonlinear Mixed Effects Model. *Journal of Computational and Graphical Statistics*, 4:12–35.
- Ristic, B. and Arulampalam, M. S. (2003). Tracking a manoeuvring target using angle-only measurements: algorithms and performance. *Signal Processing*, 83:1223–1238.
- Ristic, B., Arulampalam, S., and Gordon, N. (2004). *Beyond the Kalman filter: Particle filters for tracking applications*. Artech House.
- Ristic, B., Arulampalam, S., and Musso, C. (2001). The influence of communication bandwidth on target tracking with angle only measurements from two platforms. *Signal Processing*, 81(9):1801–1811.
- Ristic, B., Farina, A., Benvenuti, D., and Arulampalam, M. S. (2003). Performance bounds and comparison of nonlinear filters for tracking a ballistic object on re-entry. *IEEE Proceedings–Radar Sonar Navigation*, 150(2):65–70.
- Rossi, A. and Gallo, G. M. (2006). Volatility estimation via hidden Markov models. *Journal of Empirical Finance*, 13(2):203–230.
- Rubinstein, R. Y. (1981). *Simulation and the Monte Carlo method*. John Wiley & Sons.
- Schervish, M. J. (1995). *Theory of Statistics*. Springer.
- Shun, Z. and McCullagh, R. (1995). Laplace Approximation of High Dimensional Integrals. *Journal of the Royal Statistical Society. Series B*, 57(4):749–760.

- Snyder, C., Bengtsson, T., Bickel, P., and Anderson, J. (2008). Obstacles to high-dimensional particle filtering. *Monthly Weather Review*, 136(12):4629–4640.
- Strasser, H. (2002). Automatic Differentiation to Facilitate Maximum Likelihood Estimation in Nonlinear Random Effects Models. *Journal of Computational and Graphical Statistics*, 11(2):458–470.
- Suwantong, R., Bertrand, S., Dumur, D., and Beauvois, D. (2012). Space debris trajectory estimation during atmospheric reentry using moving horizon estimator. In *Proceedings of the 51st IEEE Conference on Decision and Control*, Maui, HI, USA, December 10–13.
- Tichavský, P., Muravchik, C. H., and Nehorai, A. (1998). Posterior Cramer-Rao bounds for discrete-time nonlinear filtering. *IEEE Transactions on Signal Processing*, 46(5):1386–1396.
- Tierney, L. and Kadane, J. B. (1986). Accurate Approximations for Posterior Moments and Marginal Densities. *Journal of the American Statistical Association*, 81(393):82–86.
- Tierney, L., Kass, R. E., and Kadane, J. B. (1989). Fully Exponential Laplace Approximations to Expectations and Variances of Nonpositive Functions. *Journal of the American Statistical Association*, 84(407):710–716.
- van Leeuwen, P. J. (2009). Particle Filtering in Geophysical Systems. *Monthly Weather Review*, 137(12):4089–4114.
- Van Trees, H. L. (2001). *Detection, Estimation, and Modulation Theory*. John Wiley & Sons.
- Wu, W., Gao, Y., Bienenstock, E., Donoghue, J. P., and Black, M. J. (2006). Bayesian Population Decoding of Motor Cortical Activity Using a Kalman Filter. *Neural Computation*, 18:80–118.





**Abstract:** The thesis deals with the contribution of the Laplace method to the approximation of the Bayesian filter in hidden Markov models with continuous state-space, i.e. in a sequential framework, with target tracking as the main application domain. Originally, the Laplace method is an asymptotic method used to compute integrals, i.e. in a static framework, valid in theory as soon as the function to be integrated exhibits an increasingly dominating maximum point, which brings the essential contribution to the integral. The two main contributions of the thesis are the following. Firstly, we have combined the Laplace method and particle filters: indeed, it is well-known that sequential Monte Carlo methods based on importance sampling are inefficient when the weighting function (here, the likelihood function) is too much spatially localized, e.g. when the variance of the observation noise is too small, whereas this is precisely the situation where the Laplace method is efficient and theoretically justified, hence the natural idea of combining the two approaches. We thus propose an algorithm associating the Laplace method and particle filtering, called the Laplace particle filter. Secondly, we have analyzed the approximation of the Bayesian filter based on the Laplace method only (i.e. without any generation of random samples): the objective has been to control the propagation of the approximation error from one time step to the next time step, in an appropriate asymptotic framework, e.g. when the variance of the observation noise goes to zero, or when the variances of the model noise and of the observation noise jointly go (with the same rate) to zero, or more generally when the information contained in the system goes to infinity, with an interpretation in terms of identifiability.

**Keywords:** Bayesian statistics, filtering, Monte Carlo method, asymptotic expansions, stochastic approximation, tracking.

**Résumé :** La thèse porte sur l'apport de la méthode de Laplace pour l'approximation du filtre bayésien dans des modèles de Markov cachés généraux, c'est-à-dire dans un cadre séquentiel, avec comme domaine d'application privilégié la poursuite de cibles mobiles. A la base, la méthode de Laplace est une méthode asymptotique pour le calcul d'intégrales, c'est-à-dire dans un cadre statique, valide en théorie dès que la fonction à intégrer présente un maximum de plus en plus significatif, lequel apporte la contribution essentielle au résultat. En pratique, cette méthode donne des résultats souvent très précis même en dehors de ce cadre de validité théorique. Les deux contributions principales de la thèse sont les suivantes. Premièrement, nous avons utilisé la méthode de Laplace en complément du filtrage particulaire : on sait en effet que les méthodes de Monte Carlo séquentielles basées sur l'échantillonnage pondéré sont mises en difficulté quand la fonction de pondération (ici la fonction de vraisemblance) est trop localisée, par exemple quand la variance du bruit d'observation est trop faible, or c'est précisément là le domaine où la méthode de Laplace est efficace et justifiée théoriquement, d'où l'idée naturelle de combiner les deux points de vue. Nous proposons ainsi un algorithme associant la méthode de Laplace et le filtrage particulaire, appelé le Laplace particle filter. Deuxièmement, nous avons analysé l'approximation du filtre bayésien grâce à la méthode de Laplace seulement (c'est-à-dire sans génération d'échantillons aléatoires) : il s'agit ici de contrôler la propagation de l'erreur d'approximation d'un pas de temps au pas de temps suivant, dans un cadre asymptotique approprié, par exemple quand le bruit d'observation tend vers zéro, ou quand le bruit d'état et le bruit d'observation tendent conjointement (et à la même vitesse) vers zéro, ou plus généralement quand l'information contenue dans le système tend vers l'infini, avec une interprétation en terme d'identifiabilité.

**Mots-clés :** statistique bayésienne, filtrage, méthode de Monte Carlo, développements asymptotiques, approximation stochastique, trajectographie.