



Neural Kalman filter

Gábor Szirtes, Barnabás Póczos, András Lőrincz*

*Department of Information Systems, Eötvös Loránd University, Pázmány Péter sétány 1/C,
1117 Budapest, Hungary*

Available online 15 December 2004

Abstract

Anticipating future events is a crucial function of the central nervous system and can be modelled by Kalman filter-like mechanisms, which are optimal for predicting *linear* dynamical systems. *Connectionist* representation of such mechanisms with Hebbian learning rules has not yet been derived. We show that the recursive prediction error method offers a solution that can be mapped onto the entorhinal–hippocampal loop in a biologically plausible way. Model predictions are provided.

© 2004 Elsevier B.V. All rights reserved.

Keywords: Kalman filter; Hippocampus; Entorhinal cortex; Neural prediction

1. Introduction

Linear dynamical systems (LDS) are widely applied tools in state estimation and control tasks. The so-called Kalman-filter recursion (KFR) makes the inference in LDS simple; the resulting estimations are unbiased and have minimized covariance. Here, an approximation of the KFR is provided (Section 2) by applying the recursive prediction error (RPE) method [5]. We show that the approximation (i) can be represented in neuronal form, (ii) is efficient and (iii) can be mapped (Section 3) onto the entorhinal–hippocampal (EC–HC) loop, the center of memory functions. Conclusions are drawn in Section 4.

*Corresponding author. Tel.: +36 1 4633515; fax: +36 1 4633490.

E-mail address: andras.lorincz@elte.hu (A. Lőrincz).

2. Approximated Kalman-filter recursion

Consider the following LDS:

$$\text{observation process : } \mathbf{y}^t = \mathbf{H}\mathbf{x}^t + \mathbf{n}^t, \quad (1)$$

$$\text{hidden process : } \mathbf{x}^{t+1} = \mathbf{F}\mathbf{x}^t + \mathbf{m}^t, \quad (2)$$

where variables $\mathbf{m}^t \propto \mathcal{N}(0, \Pi)$ and $\mathbf{n}^t \propto \mathcal{N}(0, \Sigma)$ are independent Gaussian noise processes. The task is to estimate the hidden variables $\mathbf{x}^t \in \mathbf{R}^n$ given the series of observations $\mathbf{y}^\tau \in \mathbf{R}^p$, $\tau \leq t$. For squared norm in the cost, the optimal solution was derived in [4]. The *prediction equation* is used to estimate $\mathbf{x}$ before the $(t+1)^{\text{th}}$ measurement:

$$\hat{\mathbf{x}}^{(t+1|t)} = \mathbf{F}\hat{\mathbf{x}}^{(t|t-1)} + \mathbf{K}^t(\mathbf{y}^t - \mathbf{H}\hat{\mathbf{x}}^{(t|t-1)}) = \mathbf{F}\hat{\mathbf{x}}^{(t|t)}, \quad (3)$$

where $\mathbf{K}^t$ is the ‘Kalman gain’, which can be computed by means of a priori and posteriori covariance matrices of $\hat{\mathbf{x}}^t$. Expression $\mathbf{e}^t = \mathbf{y}^t - \mathbf{H}\hat{\mathbf{x}}^{(t|t-1)}$ in Eq. (3) can be identified as the *reconstruction error*, because, for the noiseless case, $\mathbf{H}\hat{\mathbf{x}}^{(t|t-1)}$ should perfectly match the input. Kalman-gain balances error $\mathbf{e}^t$ and model-based prediction $\mathbf{F}\hat{\mathbf{x}}^{(t|t-1)}$ to optimize the estimation.

The first problem of the classical solution is that covariance matrices of measurement and observation noises (Π and Σ) are generally assumed to be known. The second problem is that to ensure dynamic adaptation of the Kalman gain, the algorithm requires the calculation of a matrix inversion, which is hard to interpret in neurobiological terms. In this section we derive an approximation of the Kalman gain, which eliminates these problems. Our approximation makes use of the RPE method. The resulting scheme is (i) local, (ii) well suited to ‘track’ the changing world and (iii) asymptotically optimal by construction under mild conditions [9]. Let $\mathbf{K}^t \mathbf{z} \approx \theta^t \cdot * \mathbf{K} \mathbf{z}$ denote an arbitrary parametrization of $\mathbf{K}^t$, where $*$ denotes elementwise multiplication. The RPE approximation of KFR using this arbitrary parametrization is as follows:

$$\hat{\mathbf{x}}_i^{t+1} = \sum_j F_{ij} \hat{\mathbf{x}}_j^t + \theta_i^t \sum_l K_{il} e_l^t \quad (4)$$

in which $\theta^t \in \mathbf{R}^p$, $\hat{\mathbf{x}}^{t+1} = \hat{\mathbf{x}}^{(t+1|t)}$. For simplicity, the notation of the error’s explicit dependence on θ is dropped. Let us suppose that a suboptimal matrix $\mathbf{K}(\theta^0)$ is given at time $t=0$. Our goal is to tune parameter θ^t in order to minimize $J_k(\theta_k) = \frac{1}{2} E[(e_k^t)^2]$ with respect to θ_k , where $E[\cdot]$ is the expectation operator and $e_k^t = \sum_l K_{kl} e_l^t$ is the transformed error. Stochastic gradient approximation provides the following update equations:

$$\theta_k^{t+1} = \theta_k^t + \alpha \sum_{lj} K_{kl} H_{lj} W_{jk} e_k^t, \quad (5)$$

$$W_{ik}^{t+1} = \sum_j F_{ij} W_{jk}^t \zeta_k - \theta_i^t \sum_{lj} K_{il} H_{lj} W_{jk}^t \zeta_k + \delta_{ik} e_k^t, \quad (6)$$

Download English Version:

<https://daneshyari.com/en/article/9653500>

Download Persian Version:

<https://daneshyari.com/article/9653500>

[Daneshyari.com](https://daneshyari.com)