

多层前馈感知器的 高阶序贯非线性 Kalman 滤波学习算法

邓志东 孙增圻

(清华大学计算机系·北京, 100084)

摘要: 本文提出了高阶序贯非线性增广 Kalman 滤波 (SEKF), 并将其应用于多层前馈感知器 (MLPs) 的学习问题. 文中给出了 MLPs 的 SEKF 算法, 得到了与 BP 算法类似的正向与 ψ 反向传播过程, 并且详细地推导了核心的量测 Jacobian 矩阵. 结合一非线性正弦函数, DEKF 和 SEKF 的仿真结果被进一步给出.

关键词: 多层前馈感知器; 学习算法; 非线性 Kalman 滤波; 高阶序贯估计

1 引言

Singhal 等^[1]基于非线性增广 Kalman 滤波 (EKF), 率先提出了 MLPs 的全局 EKF 学习算法 (GEKF). 为了降低复杂性与存储容量, Kollias 等^[2]和 Shah 等^[3]对网络每个节点之输入权组分别应用了独立的 EKF (IEKF). 而 Puskorius 等^[4]则进一步给出了将连接权进行分组, 忽略其耦合的所谓解耦的 EKF 学习算法 (DEKF). 由于连接权分组的各种可能性, 因而它更具一般性, 并可将 GEKF 和 IEKF 作为极限情况而包蕴其中.

不过上述各种方法仅采用了典型的 EKF, 而 MLPs 的学习问题显然是一个具有高度量测非线性 $h(\cdot)$ 的滤波问题, 仅沿预报值 $\bar{w}_k$ 的一阶 Taylor 级数展开可能丧失应有的精度. 为此本文提出了一种高阶序贯的 EKF (SEKF), 并将其应用于 MLPs 的快速学习问题. 仿真结果验证了本文方法的有效性.

2 MLPs 的 SEKF 快速学习算法

假定 MLPs 各层的神经元个数分别为 $n_p, p=1, 2, \dots, m$. 若记第 p 层的第 i 个神经元为 net_i^p , 第 k 步学习迭代时 net_i^p 的输入为 i_p^k , 输出为 o_i^p , 即有 $o_i^p = f(i_p^k)$, 这里 $f(x) = 1/(1 + e^{-x})$ 为 s 型激发函数, 且 $i_p = 1, 2, \dots, n_p, k=0, 1, \dots$.

又令 θ_i^p 为 net_i^p 的阈值, $\bar{w}_{i_p}^p \in R^{n_{p-1} \times 1}$ 为来自 $p-1$ 层各神经元的输入权向量, 相应构成的 $(n_{p-1}+1)$ 维增广输入权向量为 $w_{i_p}^{pT} = [\bar{w}_{i_p}^{pT} \theta_i^p]$.

不难确定, 上述 MLPs 共有 $M = \sum_{p=2}^m (n_{p-1}+1)n_p$ 个连接权和 $G = \sum_{p=2}^m n_p$ 个神经元 (不包括输入层), 因此若按神经元的输入权分组, 则可将 MLPs 的所有 M 个连接权分成如下 G 组或 G 个子系统

$$w_{i_p+1}^p = w_{i_p}^p + \Delta w_{i_p}^p. \quad (1a)$$

这里 $i_p = 1, 2, \dots, n_p, p=m, m-1, \dots, 2$.

若记 $\xi_k^i \triangleq 4w_k^i$, 则上式可写成

$$w_{k+1}^i = w_k^i + \xi_k^i. \quad (1b)$$

其中 $E\{\xi_k^i\} \triangleq 0$, $E\{\xi_k^i \xi_k^{jT}\} \triangleq Q_k \delta_{ij}$, $Q_k = Q_k^T \geq 0$.

进一步地, 我们有

$$o_k^m = h^m(w_k^i). \quad (2a)$$

即 MLPs 输出层的实际输出 o_k^m 是所有输入权向量 w_k^i 的非线性函数, 其中 o_k^m 与 $h^m(\cdot)$ 分别表示由 n_m 个 o_k^m 与 $h^m(\cdot)$ 所组成的向量.

显然, 若将训练样本集的期望输出 y_k 作为量测值 z_k , 则上述 MLPs 的学习问题即转化为典型的具有量测非线性性的 Kalman 滤波问题了.

为了在不致增加过多计算量的前提下, 尽量降低量测非线性 $h^m(\cdot)$ 的影响, 我们可考虑将经典 EKF 中 $h^m(\cdot)$ 沿预报值 $\bar{x}_k$ 的展开, 改为沿滤波值 $\hat{x}_k$ 的展开. 由于 $\hat{x}_k$ 更接近于真值 x_k , 因而通过此种沿 $\hat{x}_k$ 的重新线性化, 可望以较小的计算复杂性, 获得较低的量测线性化误差.

利用 Taylor 级数, 记 $\eta_k^m \triangleq \text{HOT}$ (二阶以上的高阶项), 将式 (2a) 沿邻域 λ_j 展开

$$o_k^m(\lambda_j) = h^m(\lambda_j) + H_k^{iT}(\lambda_j)(w_k^i - \lambda_j) + \eta_k^m \quad (2b)$$

其中 $E\{\eta_k^m\} \triangleq 0$, $E\{\eta_k^m \eta_j^{mT}\} \triangleq R_k \delta_{kj}$, $R_k = R_k^T > 0$, 且将

$$H_k^{iT}(\lambda_j) = \left. \frac{\partial h^m}{\partial w_k^i} \right|_{w_k^i = \lambda_j} \quad (2c)$$

称为 $n_m \times (n_{i-1} + 1)$ 维的量测 Jacobian 矩阵, $j = 0, 1, \dots, l-1$.

相应的高阶序贯 SEKF 学习算法为

$$\lambda_{j+1} = \lambda_j + K_k(\lambda_j)[y_k - h(\lambda_j)], \quad j = 0, 1, \dots, l-1, \quad (3a)$$

$$\hat{w}_k = \lambda_l, \quad \lambda_0 = \bar{w}_k, \quad (3b)$$

$$\bar{w}_k = \hat{w}_{k-1}, \quad (3c)$$

$$K_k(\lambda_j) = M_k H_k(\lambda_j) A_k(\lambda_j), \quad (3d)$$

$$A_k(\lambda_j) = [H_k^T(\lambda_j) M_k H_k(\lambda_j) + R_k]^{-1}, \quad (3e)$$

$$M_k = P_{k-1} + Q_{k-1}, \quad (3f)$$

$$P_k = [I - K_k(\lambda_0) H_k^T(\lambda_0)] M_k, \quad (3g)$$

$$\hat{w}_0 = w_0, \quad P_0 = \alpha^2 I. \quad (3h)$$

其中初始权值 $w_0 \in \mathbb{R}^{(n_{i-1}+1) \times 1}$ 可与 BP 算法类似取自均匀分布, $Q_k \in \mathbb{R}^{(n_{i-1}+1) \times (n_{i-1}+1)}$ 可取为 $10^{-6} \sim 10^{-2}$ 的对角矩阵以避免 P_k 负定或奇异, $R_k \in \mathbb{R}^{n_m \times n_m}$ 可取为等于或略小于 1 的对角矩阵, 且 α^2 一般可取为 $10^2 \sim 10^6$.

需要指出的是, 为了避免记号上的繁琐, 上式中的各项均没有特别注明上标 i . 事实上, 上述 SEKF 是按分组(解耦)后的 G 个子系统独立进行学习滤波的.

现在问题的关键已归结为如何构造式 (2c) 的量测 Jacobian 矩阵了.

与 BP 算法的 δ 反向传播过程类似, 利用求偏导数的链式法则, 我们这里也将有所谓的 ψ 反向传播过程.

令 $\psi_k^{i,i} \triangleq \partial o_k^i / \partial w_k^i$, 由于

$$\frac{\partial h^n(w_p^i)}{\partial w_p^i} = \frac{\partial o_k^n}{\partial v_p^i} \frac{\partial v_p^i}{\partial w_p^i} = \psi_k^{i,n} o_k^{i,n-1}. \quad (4)$$

其中 $o_k^{i,n-1}$ 为由第 $p-1$ 层的 $(n_{p-1}+1)$ 个神经元(包括阈值)之输出 $o_k^{i,n-1}$ 所组成的输出向量, $i_p=1, 2, \dots, n_p, p=m-1, m-2, \dots, 2$, 从而

当 $p=m$ 时

$$\psi_k^{i,n,i_p} = \frac{\partial o_k^{i,n}}{\partial v_k^{i,n}} = f'(v_k^{i,n}) = o_k^{i,n}(1 - o_k^{i,n}), \quad (5)$$

当 $p \neq m$ 时

$$\begin{aligned} \psi_k^{i,n,i_p} &= \frac{\partial o_k^{i,n}}{\partial v_k^{i_p}} = \frac{\partial o_k^{i,n}}{\partial v_k^{i_p+1}} \frac{\partial v_k^{i_p+1}}{\partial o_k^{i_p}} \frac{\partial o_k^{i_p}}{\partial v_k^{i_p}} \\ &= \psi_k^{i,n,i_p+1} w_{k,i_p+1}^{i_p} f'(v_k^{i_p}) \\ &= o_k^{i_p}(1 - o_k^{i_p}) w_{k,i_p+1}^{i_p} \psi_k^{i,n,i_p+1}. \end{aligned} \quad (6)$$

故可构成如下 $n_m \times (n_{p-1}+1)$ 维的量测 Jacobian 矩阵

$$H_k^{i_p T} = \psi_k^{i,n,i_p T} o_k^{i,n-1}. \quad (7)$$

其中 ψ_k^{i,n,i_p} 为由 n_m 个 ψ_k^{i,n,i_p} 组成的 ψ 向量。

因此上述算法的基本思想是:在给出随机初始权后,首先正向传播得到各神经元的实际输出,然后利用式(5),(6)进行 ψ 反向传播,并据此构造量测 Jacobian 矩阵,最后即可利用式(3)的 SEKF 得到新的输入权估计值(按 G 个子系统分别进行),如此不断循环上述过程,直到满足精度为止。

3 仿真举例

以带常数项的正弦函数为例,考虑训练样本集为 $x_i = (1/20)i, y_{di} = (2/5)\sin(2\pi x_i) + 1/2, i=0, 1, \dots, 19$ 。

取三层前馈网络,其中输入层、隐层和输出层的神经元个数分别为 1, 10, 1。

为了进行比较,本文同时研究了 DEKF 和 SEKF 两种学习算法的情形。

仿真中,我们利用了用 C 语言自行研制开发的神经网络自动生成工具(THN²-tool),式(3)中各参数的选择为: $P_0^i = 10^6 I_i, Q_k^i = 10^{-6} I_i, R_k^i = 0.9999 I_i, w_0^i$ 取自 $(-0.5, 0.5)$ 的均匀分布, $i=1, 2, \dots, 11$, 即有 11 个子系统,其中第 1 个子系统(输出层)为 11 维,其他子系统(隐层)均为 2 维,这里 I_i 为第 i 个子系统的相应维单位矩阵。

选择学习迭代步数 $k=200$, 图 1 给出了利用 DEKF($l=1$)和 SEKF($l=3$)的误差曲线。

从图中可以明显看出,SEKF 较之 DEKF 在滤波精度方面有进一步的改善,即在相同的精度控制指标下,SEKF 所需训练次数将比 DEKF 更少,而计算复杂性却

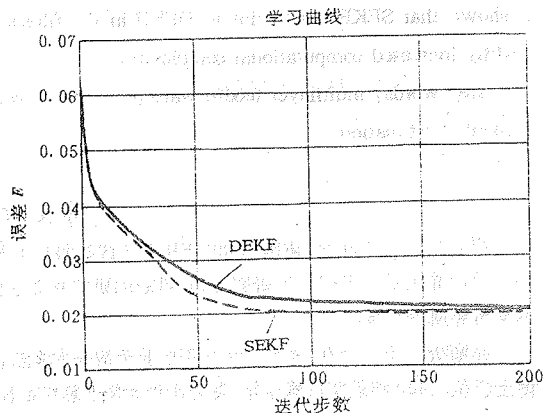


图 1 DEKF 与 SEKF 的收敛性

只增加了约 1.25 倍。

4 结 论

综上所述,我们还可以得到如下几点结论:

- 1) 当式(3)中 $l=1$ 时,上述 SEKF 学习算法退化为文[4]的 DEKF 学习算法;
- 2) 由于 MLPs 的输出层的神经元个数 n_m 通常较小,因此式(3e)的 $A_k^l \in R^{n_m \times n_m}$ 只是一个低阶矩阵,特别地,当 $n_m=1$ 时, A_k^l 仅是一个标量,相应的求逆计算量实际是很小的;
- 3) 式(5),(6)的 ψ 反向传播过程较之 BP 算法不存在 Σ 项,这是因为 BP 算法是对标量 E 求导的,而这里则是对 n_m 维向量 o_m^k 求导的,其实这一项已反映在式(3a)中了。

参 考 文 献

- [1] Singhal, S. and Wu, L. . Training Multilayer Perceptrons with the Extended Kalman Algorithm. In Advances in Neural Information Processing System 1, (Touretzky, D. S. , Ed.), San Mateo, CA;Morgan Kaufmann, 1989, 133—140
- [2] Kollias, S. and Anastassiou, D. . An Adaptive Least Squares Algorithm for the Efficient Training of Artificial Neural Network. IEEE Trans. on Circuits and Systems, 1989, 36(8);1092—1101
- [3] Shah, S. A. and Palmieri, F. . MEKA——A Fast, Local Algorithm for Training Feedforward Neural Networks. In International Joint Conference on Neural Networks, New York, 1990, III;41—45
- [4] Puskorius, G. V. and Feldkamp, L. A. . Decoupled Extended Kalman Filter Training of Feedforward Layered Networks. IJCNN-91, 1991, 1;771—777

Learning Algorithm of Multilayer Feedforward Perceptron with Higher Order Sequential Nonlinear Kalman Filter

DENG Zhidong and SUN Zengqi

(Department of Computer, Tsinghua University • Beijing, 100084; PRC)

Abstract: In this paper a higher order sequential nonlinear extended Kalman filter (SEKF), based on the DEKF algorithm given by ref. [4], is proposed and applied to the learning problem of MLPs. The simulation result is shown that SEKF is superior to DEKF in the filtering accuracy and the required learning number except the slightly increased computational complexity.

Key words: multilayer feedforward perceptron; learning algorithm; nonlinear Kalman filtering; higher order sequential estimation

本文作者简介

邓志东 1966年生. 讲师. 1986年毕业于成都科技大学计算机系. 1991年在哈尔滨工业大学获博士学位. 现在清华大学计算机系从事博士后研究工作. 目前的研究兴趣主要涉及神经网络控制、学习控制、混沌非线性以及飞行器的制导与导航等领域。

孙增圻 1943年生. 教授. 1966年毕业于清华大学自动控制系, 留校任教. 1979年赴瑞典进修, 1981年在瑞典获博士学位. 1990年赴美进修半年. 现为清华大学计算机系教授. 曾从事最优控制、控制系统CAD、计算机控制理论等方面的研究. 目前主要研究兴趣为智能控制、机器人控制和仿真以及神经网络控制等领域。