

Natural Gradient Works Efficiently in Learning

Shun-ichi Amari

RIKEN Frontier Research Program, Saitama 351-01, Japan

When a parameter space has a certain underlying structure, the ordinary gradient of a function does not represent its steepest direction, but the natural gradient does. Information geometry is used for calculating the natural gradients in the parameter space of perceptrons, the space of matrices (for blind source separation), and the space of linear dynamical systems (for blind source deconvolution). The dynamical behavior of natural gradient online learning is analyzed and is proved to be Fisher efficient, implying that it has asymptotically the same performance as the optimal batch estimation of parameters. This suggests that the plateau phenomenon, which appears in the backpropagation learning algorithm of multilayer perceptrons, might disappear or might not be so serious when the natural gradient is used. An adaptive method of updating the learning rate is proposed and analyzed.

1 Introduction

The stochastic gradient method (Widrow, 1963; Amari, 1967; Tsypkin, 1973; Rumelhart, Hinton, & Williams, 1986) is a popular learning method in the general nonlinear optimization framework. The parameter space is not Euclidean but has a Riemannian metric structure in many cases. In these cases, the ordinary gradient does not give the steepest direction of a target function; rather, the steepest direction is given by the natural (or contravariant) gradient. The Riemannian metric structures are introduced by means of information geometry (Amari, 1985; Murray and Rice, 1993; Amari, 1997a; Amari, Kurata, & Nagoska, 1992). This article gives the natural gradients explicitly in the case of the space of perceptrons for neural learning, the space of matrices for blind source separation, and the space of linear dynamical systems for blind multichannel source deconvolution. This is an extended version of an earlier article (Amari, 1996), including new results.

How good is natural gradient learning compared to conventional gradient learning? The asymptotic behavior of online natural gradient learning is studied for this purpose. Training examples can be used only once in online learning when they appear. Therefore, the asymptotic performance of online learning cannot be better than the optimal batch procedure where all the examples can be reused again and again. However, we prove that natural gradient online learning gives the Fisher-efficient estimator in the sense

of asymptotic statistics when the loss function is differentiable, so that it is asymptotically equivalent to the optimal batch procedure (see also Amari, 1995; Oppen, 1996). When the loss function is nondifferentiable, the accuracy of asymptotic online learning is worse than batch learning by a factor of 2 (see, for example, Van den Broeck & Reimann, 1996). It was shown in Amari et al. (1992) that the dynamic behavior of natural gradient in the Boltzmann machine is excellent.

It is not easy to calculate the natural gradient explicitly in multilayer perceptrons. However, a preliminary analysis (Yang & Amari, 1997), by using a simple model, shows that the performance of natural gradient learning is remarkably good, and it is sometimes free from being trapped in plateaus, which give rise to slow convergence of the backpropagation learning method (Saad & Solla, 1995). This suggests that the Riemannian structure might eliminate such plateaus or might make them not so serious.

Online learning is flexible, because it can track slow fluctuations of the target. Such online dynamics were first analyzed in Amari (1967) and then by many researchers recently. Sompolinsky, Barkai, and Seung (1995), and Barkai, Seung, and Sompolinsky (1995) proposed an adaptive method of adjusting the learning rate (see also Amari, 1967). We generalize their idea and evaluate its performance based on the Riemannian metric of errors.

The article is organized as follows. The natural gradient is defined in section 2. Section 3 formulates the natural gradient in various problems of stochastic descent learning. Section 4 gives the statistical analysis of efficiency of online learning, and section 5 is devoted to the problem of adaptive changes in the learning rate. Calculations of the Riemannian metric and explicit forms of the natural gradients are given in sections 6, 7, and 8.

2 Natural Gradient

Let $S = \{\mathbf{w} \in \mathbf{R}^n\}$ be a parameter space on which a function $L(\mathbf{w})$ is defined. When S is a Euclidean space with an orthonormal coordinate system $\mathbf{w}$, the squared length of a small incremental vector $d\mathbf{w}$ connecting $\mathbf{w}$ and $\mathbf{w} + d\mathbf{w}$ is given by

$$|d\mathbf{w}|^2 = \sum_{i=1}^n (dw_i)^2,$$

where dw_i are the components of $d\mathbf{w}$. However, when the coordinate system is nonorthonormal, the squared length is given by the quadratic form

$$|d\mathbf{w}|^2 = \sum_{i,j} g_{ij}(\mathbf{w}) dw_i dw_j. \quad (2.1)$$

When S is a curved manifold, there is no orthonormal linear coordinates, and the length of $d\mathbf{w}$ is always written as in equation 2.1. Such a space is

a Riemannian space. We show in later sections that parameter spaces of neural networks have the Riemannian character. The $n \times n$ matrix $G = (g_{ij})$ is called the Riemannian metric tensor, and it depends in general on $\mathbf{w}$. It reduces to

$$g_{ij}(\mathbf{w}) = \delta_{ij} = \begin{cases} 1, & i = j, \\ 0, & i \neq j \end{cases}$$

in the Euclidean orthonormal case, so that G is the unit matrix I in this case.

The steepest descent direction of a function $L(\mathbf{w})$ at $\mathbf{w}$ is defined by the vector $d\mathbf{w}$ that minimizes $L(\mathbf{w} + d\mathbf{w})$ where $|d\mathbf{w}|$ has a fixed length, that is, under the constraint

$$|d\mathbf{w}|^2 = \varepsilon^2 \quad (2.2)$$

for a sufficiently small constant ε .

Theorem 1. *The steepest descent direction of $L(\mathbf{w})$ in a Riemannian space is given by*

$$-\tilde{\nabla}L(\mathbf{w}) = -G^{-1}(\mathbf{w})\nabla L(\mathbf{w}) \quad (2.3)$$

where $G^{-1} = (g^{ij})$ is the inverse of the metric $G = (g_{ij})$ and ∇L is the conventional gradient,

$$\nabla L(\mathbf{w}) = \left(\frac{\partial}{\partial w_1} L(\mathbf{w}), \dots, \frac{\partial}{\partial w_n} L(\mathbf{w}) \right)^T,$$

the superscript T denoting the transposition.

Proof. We put

$$d\mathbf{w} = \varepsilon \mathbf{a},$$

and search for the $\mathbf{a}$ that minimizes

$$L(\mathbf{w} + d\mathbf{w}) = L(\mathbf{w}) + \varepsilon \nabla L(\mathbf{w})^T \mathbf{a}$$

under the constraint

$$|\mathbf{a}|^2 = \sum g_{ij} a_i a_j = 1.$$

By the Lagrangean method, we have

$$\frac{\partial}{\partial a_i} \{ \nabla L(\mathbf{w})^T \mathbf{a} - \lambda \mathbf{a}^T G \mathbf{a} \} = 0.$$

This gives

$$\nabla L(\mathbf{w}) = 2\lambda G\mathbf{a}$$

or

$$\mathbf{a} = \frac{1}{2\lambda} G^{-1} \nabla L(\mathbf{w}),$$

where λ is determined from the constraint.

We call

$$\tilde{\nabla} L(\mathbf{w}) = G^{-1} \nabla L(\mathbf{w})$$

the natural gradient of L in the Riemannian space. Thus, $-\tilde{\nabla} L$ represents the steepest descent direction of L . (If we use the tensorial notation, this is nothing but the contravariant form of $-\nabla L$.) When the space is Euclidean and the coordinate system is orthonormal, we have

$$\tilde{\nabla} L = \nabla L. \quad (2.4)$$

This suggests the natural gradient descent algorithm of the form

$$\mathbf{w}_{t+1} = \mathbf{w}_t - \eta_t \tilde{\nabla} L(\mathbf{w}_t), \quad (2.5)$$

where η_t is the learning rate that determines the step size.

3 Natural Gradient Learning

Let us consider an information source that generates a sequence of independent random variables $z_1, z_2, \dots, z_t, \dots$, subject to the same probability distribution $q(z)$. The random signals z_t are processed by a processor (like a neural network) that has a set of adjustable parameters $\mathbf{w}$. Let $l(z, \mathbf{w})$ be a loss function when signal z is processed by the processor whose parameter is $\mathbf{w}$. Then the risk function or the average loss is

$$L(\mathbf{w}) = E[l(z, \mathbf{w})], \quad (3.1)$$

where E denotes the expectation with respect to z . Learning is a procedure to search for the optimal $\mathbf{w}^*$ that minimizes $L(\mathbf{w})$.

The stochastic gradient descent learning method can be formulated in general as

$$\mathbf{w}_{t+1} = \mathbf{w}_t - \eta_t C(\mathbf{w}_t) \nabla l(z_t, \mathbf{w}_t), \quad (3.2)$$

where η_t is a learning rate that may depend on t and $C(\mathbf{w})$ is a suitably chosen positive definite matrix (see Amari, 1967). In the natural gradient online learning method, it is proposed to put $C(\mathbf{w})$ equal to $G^{-1}(\mathbf{w})$ when the Riemannian structure is defined. We give a number of examples to be studied in more detail.

3.1 Statistical Estimation of Probability Density Function. In the case of statistical estimation, we assume a statistical model $\{p(\mathbf{z}, \mathbf{w})\}$, and the problem is to obtain the probability distribution $p(\mathbf{z}, \hat{\mathbf{w}})$ that approximates the unknown density function $q(\mathbf{z})$ in the best way—that is, to estimate the true $\mathbf{w}$ or to obtain the optimal approximation $\mathbf{w}$ from the observed data. A typical loss function is

$$l(\mathbf{z}, \mathbf{w}) = -\log p(\mathbf{z}, \mathbf{w}). \quad (3.3)$$

The expected loss is then given by

$$\begin{aligned} L(\mathbf{w}) &= -E[\log p(\mathbf{z}, \mathbf{w})] \\ &= E_q \left[\log \frac{q(\mathbf{z})}{p(\mathbf{z}, \mathbf{w})} \right] + H_Z, \end{aligned}$$

where H_Z is the entropy of $q(\mathbf{z})$ not depending on $\mathbf{w}$. Hence, minimizing L is equivalent to minimizing the Kullback-Leibler divergence

$$D[q(\mathbf{z}) : p(\mathbf{z}, \mathbf{w})] = \int q(\mathbf{z}) \log \frac{q(\mathbf{z})}{p(\mathbf{z}, \mathbf{w})} d\mathbf{z} \quad (3.4)$$

of two probability distributions $q(\mathbf{z})$ and $p(\mathbf{z}, \mathbf{w})$. When the true distribution $q(\mathbf{z})$ is written as $q(\mathbf{z}) = p(\mathbf{z}, \mathbf{w}^*)$, this is equivalent to obtain the maximum likelihood estimator $\hat{\mathbf{w}}$.

The Riemannian structure of the parameter space of a statistical model is defined by the Fisher information (Rao, 1945; Amari, 1985)

$$g_{ij}(\mathbf{w}) = E \left[\frac{\partial \log p(\mathbf{x}, \mathbf{w})}{\partial w_i} \frac{\partial \log p(\mathbf{x}, \mathbf{w})}{\partial w_j} \right] \quad (3.5)$$

in the component form. This is the only invariant metric to be given to the statistical model (Chentsov, 1972; Campbell, 1985; Amari, 1985). The learning equation (see equation 3.2) gives a sequential estimator $\hat{\mathbf{w}}_t$.

3.2 Multilayer Neural Network. Let us consider a multilayer feedforward neural network specified by a vector parameter $\mathbf{w} = (w_1, \dots, w_n)^T \in \mathbf{R}^n$. The parameter $\mathbf{w}$ is composed of modifiable connection weights and thresholds. When input $\mathbf{x}$ is applied, the network processes it and calculates the outputs $\mathbf{f}(\mathbf{x}, \mathbf{w})$. The input $\mathbf{x}$ is subject to an unknown probability

distribution $q(\mathbf{x})$. Let us consider a teacher network that, by receiving $\mathbf{x}$, generates the corresponding output $\mathbf{y}$ subject to a conditional probability distribution $q(\mathbf{y} | \mathbf{x})$. The task is to obtain the optimal $\mathbf{w}^*$ from examples such that the student network approximates the behavior of the teacher.

Let us denote by $l(\mathbf{x}, \mathbf{w})$ a loss when input signal $\mathbf{x}$ is processed by a network having parameter $\mathbf{w}$. A typical loss is given,

$$l(\mathbf{x}, \mathbf{y}, \mathbf{w}) = \frac{1}{2} |\mathbf{y} - \mathbf{f}(\mathbf{x}, \mathbf{w})|^2, \quad (3.6)$$

where $\mathbf{y}$ is the output given by the teacher.

Let us consider a statistical model of neural networks such that its output $\mathbf{y}$ is given by a noisy version of $\mathbf{f}(\mathbf{x}, \mathbf{w})$,

$$\mathbf{y} = \mathbf{f}(\mathbf{x}, \mathbf{w}) + \mathbf{n}, \quad (3.7)$$

where $\mathbf{n}$ is a multivariate gaussian noise with zero mean and unit covariance matrix I . By putting $\mathbf{z} = (\mathbf{x}, \mathbf{y})$, which is an input-output pair, the model specifies the probability density of $\mathbf{z}$ as

$$p(\mathbf{z}, \mathbf{w}) = cq(\mathbf{x}) \exp \left\{ -\frac{1}{2} |\mathbf{y} - \mathbf{f}(\mathbf{x}, \mathbf{w})|^2 \right\}, \quad (3.8)$$

where c is a normalizing constant and the loss function (see equation 3.6) is rewritten as

$$l(\mathbf{z}, \mathbf{w}) = \text{const} + \log q(\mathbf{x}) - \log p(\mathbf{z}, \mathbf{w}). \quad (3.9)$$

Given a sequence of examples $(\mathbf{x}_1, \mathbf{y}_1), \dots, (\mathbf{x}_t, \mathbf{y}_t), \dots$, the natural gradient online learning algorithm is written as

$$\mathbf{w}_{t+1} = \mathbf{w}_t - \eta_t \tilde{\nabla} l(\mathbf{x}_t, \mathbf{y}_t, \mathbf{w}_t). \quad (3.10)$$

Information geometry (Amari, 1985) shows that the Riemannian structure is given to the parameter space of multilayer networks by the Fisher information matrix,

$$g_{ij}(\mathbf{w}) = E \left[\frac{\partial \log p(\mathbf{x}, \mathbf{y}; \mathbf{w})}{\partial w_i} \frac{\partial \log p(\mathbf{x}, \mathbf{y}; \mathbf{w})}{\partial w_j} \right]. \quad (3.11)$$

We will show how to calculate $G = (g_{ij})$ and its inverse in a later section.

3.3 Blind Separation of Sources. Let us consider m signal sources that produce m independent signals $s_i(t)$, $i = 1, \dots, m$, at discrete times $t = 1, 2, \dots$. We assume that $s_i(t)$ are independent at different times and that the

expectations of s_i are 0. Let $r(s)$ be the joint probability density function of s . Then it is written in the product form

$$r(s) = \prod_{i=1}^m r_i(s_i). \quad (3.12)$$

Consider the case where we cannot have direct access to the source signals $s(t)$ but we can observe their m instantaneous mixtures $x(t)$,

$$x(t) = As(t) \quad (3.13)$$

or

$$x_i(t) = \sum_{j=1}^m A_{ij}s_j(t),$$

where $A = (A_{ij})$ is an $m \times m$ nonsingular mixing matrix that does not depend on t , and $x = (x_1, \dots, x_m)^T$ is the observed mixtures.

Blind source separation is the problem of recovering the original signals $s(t)$, $t = 1, 2, \dots$ from the observed signals $x(t)$, $t = 1, 2, \dots$ (Jutten & Hérault, 1991). If we know A , this is trivial, because we have

$$s(t) = A^{-1}x(t).$$

The “blind” implies that we do not know the mixing matrix A and the probability distribution densities $r_i(s_i)$.

A typical algorithm to solve the problem is to transform $x(t)$ into

$$y(t) = W_t x(t), \quad (3.14)$$

where W_t is an estimate of A^{-1} . It is modified by the following learning equation:

$$W_{t+1} = W_t - \eta_t F(x_t, W_t). \quad (3.15)$$

Here, $F(x, W)$ is a special matrix function satisfying

$$E[F(x, W)] = 0 \quad (3.16)$$

for any density functions $r(s)$ in equation 3.12 when $W = A^{-1}$. For W_t of equation 3.15 to converge to A^{-1} , equation 3.16 is necessary but not sufficient, because the stability of the equilibrium is not considered here.

Let $K(W)$ be an operator that maps a matrix to a matrix. Then

$$\tilde{F}(x, W) = K(W)F(x, W)$$

satisfies equation 3.16 when F does. The equilibrium of F and $\tilde{F}$ is the same, but their stability can be different. However, the natural gradient does not alter the stability of an equilibrium, because G^{-1} is positive-definite.

Let $l(\mathbf{x}, W)$ be a loss function whose expectation

$$L(W) = E[l(\mathbf{x}, W)]$$

is the target function minimized at $W = A^{-1}$. A typical function F is obtained by the gradient of l with respect to W ,

$$F(\mathbf{x}, W) = \nabla l(\mathbf{x}, W). \quad (3.17)$$

Such an F is also obtained by heuristic arguments. Amari and Cardoso (in press) gave the complete family of F satisfying equation 3.16 and elucidated the statistical efficiency of related algorithms.

From the statistical point of view, the problem is to estimate $W = A^{-1}$ from observed data $\mathbf{x}(1), \dots, \mathbf{x}(t)$. However, the probability density function of $\mathbf{x}$ is written as

$$p_X(\mathbf{x}; W, r) = |W| r(W\mathbf{x}), \quad (3.18)$$

which is specified not only by W to be estimated but also by an unknown function r of the form 3.12. Such a statistical model is said to be semiparametric and is a difficult problem to solve (Bickel, Klassen, Ritov, & Wellner, 1993), because it includes an unknown function of infinite degrees of freedom. However, we can apply the information-geometrical theory of estimating functions (Amari & Kawanabe, 1997) to this problem.

When F is given by the gradient of a loss function (see equation 3.17), where ∇ is the gradient $\partial/\partial W$ with respect to a matrix, the natural gradient is given by

$$\tilde{\nabla} l = G^{-1} \circ \nabla l. \quad (3.19)$$

Here, G is an operator transforming a matrix to a matrix so that it is an $m^2 \times m^2$ matrix. G is the metric given to the space $Gl(m)$ of all the nonsingular $m \times m$ matrices. We give its explicit form in a later section based on the Lie group structure. The inverse of G is also given explicitly. Another important problem is the stability of the equilibrium of the learning dynamics. This has recently been solved by using the Riemannian structure (Amari, Chen, & Chichocki, in press; see also Cardoso & Laheld, 1996). The superefficiency of some algorithms has been also proved in Amari (1997b) under certain conditions.

3.4 Blind Source Deconvolution. When the original signals $s(t)$ are mixed not only instantaneously but also with past signals as well, the prob-

lem is called blind source deconvolution or equalization. By introducing the time delay operator z^{-1} ,

$$z^{-1}s(t) = s(t-1), \quad (3.20)$$

we have a mixing matrix filter $\mathbf{A}$ denoted by

$$\mathbf{A}(z) = \sum_{k=0}^{\infty} \mathbf{A}_k z^{-k}, \quad (3.21)$$

where $\mathbf{A}_k$ are $m \times m$ matrices. The observed mixtures are

$$\mathbf{x}(t) = \mathbf{A}(z)\mathbf{s}(t) = \sum_k \mathbf{A}_k \mathbf{s}(t-k). \quad (3.22)$$

To recover the original independent sources, we use the finite impulse response model

$$\mathbf{W}(z) = \sum_{k=0}^d \mathbf{W}_k z^{-1} \quad (3.23)$$

of degree d . The original signals are recovered by

$$\mathbf{y}(t) = \mathbf{W}_t(z)\mathbf{x}(t), \quad (3.24)$$

where $\mathbf{W}_t$ is adaptively modified by

$$\mathbf{W}_{t+1}(z) = \mathbf{W}_t(z) - \eta_t \nabla l\{\mathbf{x}_t, \mathbf{x}_{t-1}, \dots, \mathbf{W}_t(z)\}. \quad (3.25)$$

Here, $l(\mathbf{x}_t, \mathbf{x}_{t-1}, \dots, \mathbf{W})$ is a loss function that includes some past signals. We can summarize the past signals into a current state variable in the on-line learning algorithm. Such a loss function is obtained by the maximum entropy method (Bell & Sejnowski, 1995), independent component analysis (Comon, 1994), or the statistical likelihood method.

In order to obtain the natural gradient learning algorithm

$$\mathbf{W}_{t+1}(z) = \mathbf{W}_t(z) - \eta_t \tilde{\nabla} l(\mathbf{x}_t, \mathbf{x}_{t-1}, \dots, \mathbf{W}_t),$$

we need to define the Riemannian metric in the space of all the matrix filters (multiterminal linear systems). Such a study was initiated by Amari (1987). It is possible to define G and to obtain G^{-1} explicitly (see section 8). A preliminary investigation into the performance of the natural gradient learning algorithm has been undertaken by Douglas, Chichocki, and Amari (1996) and Amari et al. (1997).

4 Natural Gradient Gives Fisher-Efficient Online Learning Algorithms

This section studies the accuracy of natural gradient learning from the statistical point of view. A statistical estimator that gives asymptotically the best result is said to be Fisher efficient. We prove that natural gradient learning attains Fisher efficiency.

Let us consider multilayer perceptrons as an example. We study the case of a realizable teacher, that is, the behavior of the teacher is given by $q(\mathbf{y} | \mathbf{x}) = p(\mathbf{y} | \mathbf{x}, \mathbf{w}^*)$. Let $D_T = \{(\mathbf{x}_1, \mathbf{y}_1), \dots, (\mathbf{x}_T, \mathbf{y}_T)\}$ be T -independent input-output examples generated by the teacher network having parameter $\mathbf{w}^*$. Then, minimizing the log loss,

$$l(\mathbf{x}, \mathbf{y}; \mathbf{w}) = -\log p(\mathbf{x}, \mathbf{y}; \mathbf{w}),$$

over the training data D_T is to obtain $\hat{\mathbf{w}}_T$ that minimizes the training error

$$L_{\text{train}}(\mathbf{w}) = \frac{1}{T} \sum_{t=1}^T l(\mathbf{x}_t, \mathbf{y}_t; \mathbf{w}). \quad (4.1)$$

This is equivalent to maximizing the likelihood $\prod_{t=1}^T p(\mathbf{x}_t, \mathbf{y}_t; \mathbf{w})$. Hence, $\hat{\mathbf{w}}_T$ is the maximum likelihood estimator. The Cramér-Rao theorem states that the expected squared error of an unbiased estimator satisfies

$$E[(\hat{\mathbf{w}}_T - \mathbf{w}^*)(\hat{\mathbf{w}}_T - \mathbf{w}^*)^T] \geq \frac{1}{T} G^{-1}, \quad (4.2)$$

where the inequality holds in the sense of positive definiteness of matrices. An estimator is said to be efficient or Fisher efficient when it satisfies equation 4.2 with equality for large T . The maximum likelihood estimator is Fisher efficient, implying that it is the best estimator attaining the Cramér-Rao bound asymptotically,

$$\lim_{T \rightarrow \infty} TE[(\hat{\mathbf{w}}_T - \mathbf{w}^*)(\hat{\mathbf{w}}_T - \mathbf{w}^*)^T] = G^{-1}, \quad (4.3)$$

where G^{-1} is the inverse of the Fisher information matrix $G = (g_{ij})$ defined by equation 3.11.

Examples $(\mathbf{x}_1, \mathbf{y}_1), (\mathbf{x}_2, \mathbf{y}_2) \dots$ are given one at a time in the case of online learning. Let $\tilde{\mathbf{w}}_t$ be an online estimator at time t . At the next time, $t + 1$, the estimator $\tilde{\mathbf{w}}_t$ is modified to give a new estimator $\tilde{\mathbf{w}}_{t+1}$ based on the current observation $(\mathbf{x}_t, \mathbf{y}_t)$. The old observations $(\mathbf{x}_1, \mathbf{y}_1), \dots, (\mathbf{x}_{t-1}, \mathbf{y}_{t-1})$ cannot be reused to obtain $\tilde{\mathbf{w}}_{t+1}$, so the learning rule is written as

$$\tilde{\mathbf{w}}_{t+1} = \mathbf{m}(\mathbf{x}_t, \mathbf{y}_t, \tilde{\mathbf{w}}_t).$$

The process $\{\tilde{\mathbf{w}}_t\}$ is Markovian. Whatever learning rule $\mathbf{m}$ is chosen, the behavior of the estimator $\tilde{\mathbf{w}}_t$ is never better than that of the optimal batch estimator $\hat{\mathbf{w}}_t$ because of this restriction. The gradient online learning rule

$$\tilde{\mathbf{w}}_{t+1} = \tilde{\mathbf{w}}_t - \eta_t C \frac{\partial l(\mathbf{x}_t, \mathbf{y}_t; \tilde{\mathbf{w}}_t)}{\partial \mathbf{w}},$$

was proposed where C is a positive-definite matrix, and its dynamical behavior was studied by Amari (1967) when the learning constant $\eta_t = \eta$ is fixed. Heskes and Kappen (1991) obtained similar results, which ignited research into online learning. When η_t satisfies some condition, say, $\eta_t = c/t$, for a positive constant c , the stochastic approximation guarantees that $\tilde{\mathbf{w}}_t$ is a consistent estimator converging to $\mathbf{w}^*$. However, it is not Fisher efficient in general.

There arises a question of whether there exists a learning rule that gives an efficient estimator. If it exists, the asymptotic behavior of online learning is equivalent to that of the best batch estimation method. This article answers the question affirmatively, by giving an efficient online learning rule (see Amari, 1995; see also Oppen, 1996).

Let us consider the natural gradient learning rule,

$$\tilde{\mathbf{w}}_{t+1} = \tilde{\mathbf{w}}_t - \frac{1}{t} \tilde{\nabla} l(\mathbf{x}_t, \mathbf{y}_t, \tilde{\mathbf{w}}_t). \quad (4.4)$$

Theorem 2. *Under the learning rule (see equation 4.4), the natural gradient online estimator $\tilde{\mathbf{w}}_t$ is Fisher efficient.*

Proof. Let us denote the covariance matrix of estimator $\tilde{\mathbf{w}}_t$ by

$$\tilde{V}_{t+1} = E[(\tilde{\mathbf{w}}_{t+1} - \mathbf{w}^*)(\tilde{\mathbf{w}}_{t+1} - \mathbf{w}^*)^T]. \quad (4.5)$$

This shows the expectation of the squared error. We expand

$$\begin{aligned} \frac{\partial l(\mathbf{x}_t, \mathbf{y}_t; \tilde{\mathbf{w}}_t)}{\partial \mathbf{w}} &= \frac{\partial l(\mathbf{x}_t, \mathbf{y}_t; \mathbf{w}^*)}{\partial \mathbf{w}} + \frac{\partial^2 l(\mathbf{x}_t, \mathbf{y}_t; \mathbf{w}^*)}{\partial \mathbf{w} \partial \mathbf{w}} (\tilde{\mathbf{w}}_t - \mathbf{w}^*) \\ &\quad + O(|\tilde{\mathbf{w}}_t - \mathbf{w}^*|^2). \end{aligned}$$

By subtracting $\mathbf{w}^*$ from the both sides of equation 4.4 and taking the expectation of the square of the both sides, we have

$$\tilde{V}_{t+1} = \tilde{V}_t - \frac{2}{t} \tilde{V}_t + \frac{1}{t^2} G^{-1} + O\left(\frac{1}{t^3}\right), \quad (4.6)$$

where we used

$$E\left[\frac{\partial l(\mathbf{x}_t, \mathbf{y}_t; \mathbf{w}^*)}{\partial \mathbf{w}}\right] = 0, \quad (4.7)$$

$$E \left[\frac{\partial^2 l(\mathbf{x}_t, \mathbf{y}_t; \mathbf{w}^*)}{\partial \mathbf{w} \partial \mathbf{w}} \right] = G(\mathbf{w}^*), \quad (4.8)$$

$$G(\tilde{\mathbf{w}}_t) = G(\mathbf{w}^*) + O\left(\frac{1}{t}\right),$$

because $\tilde{\mathbf{w}}_t$ converges to $\mathbf{w}^*$ as guaranteed by stochastic approximation under certain conditions (see Kushner & Clark, 1978). The solution of equation 4.6 is written asymptotically as

$$\tilde{\mathbf{V}}_t = \frac{1}{t} G^{-1} + O\left(\frac{1}{t^2}\right),$$

proving the theorem.

The theory can be extended to be applicable to the unrealizable teacher case, where

$$K(\mathbf{w}) = E \left[\frac{\partial^2}{\partial \mathbf{w} \partial \mathbf{w}} l(\mathbf{x}, \mathbf{y}; \mathbf{w}) \right] \quad (4.9)$$

should be used instead of $G(\mathbf{w})$ in order to obtain the same efficient result as the optimal batch procedure. This is locally equivalent to the Newton-Raphson method. The results can be stated in terms of the generalization error instead of the covariance of the estimator, and we can obtain more universal results (see Amari, 1993; Amari & Murata, 1993).

Remark. In the cases of blind source separation and deconvolution, the models are semiparametric, including the unknown function r (see equation 3.18). In such cases, the Cramér-Rao bound does not necessarily hold. Therefore, Theorem 2 does not hold in these cases. It holds when we can estimate the true r of the source probability density functions and use it to define the loss function $l(\mathbf{x}, \mathbf{W})$. Otherwise equation 4.8 does not hold. The stability of the true solution is not necessarily guaranteed either. Amari, Chen, & Cichocki (in press) have analyzed this situation and proposed a universal method of attaining the stability of the equilibrium solution.

5 Adaptive Learning Constant

The dynamical behavior of the learning rule (see equation 3.2) was studied in Amari (1967) when η_t is a small constant η . In this case, $\mathbf{w}_t$ fluctuates around the (local) optimal value $\mathbf{w}^*$ for large t . The expected value and variance of $\mathbf{w}_t$ was studied, and the trade-off between the convergence speed and accuracy of convergence was demonstrated.

When the current $\mathbf{w}_t$ is far from the optimal $\mathbf{w}^*$, it is desirable to use a relatively large η to accelerate the convergence. When it is close to $\mathbf{w}^*$, a

small η is preferred in order to eliminate fluctuations. An idea of an adaptive change of η was discussed in Amari (1967) and was called “learning of learning rules.”

Sompolinsky et al. (1995) (see also Barkai et al., 1995) proposed a rule of adaptive change of η_t , which is applicable to the pattern classification problem where the expected loss $L(\mathbf{w})$ is not differentiable at $\mathbf{w}^*$. This article generalizes their idea to a more general case where $L(\mathbf{w})$ is differentiable and analyzes its behavior by using the Riemannian structure.

We propose the following learning scheme:

$$\mathbf{w}_{t+1} = \mathbf{w}_t - \eta_t \tilde{\nabla} l(\mathbf{x}_t, \mathbf{y}_t; \hat{\mathbf{w}}_t) \quad (5.1)$$

$$\eta_{t+1} = \eta_t \exp\{\alpha[\beta l(\mathbf{x}_t, \mathbf{y}_t; \hat{\mathbf{w}}_t) - \eta_t]\}, \quad (5.2)$$

where α and β are constants. We also assume that the training data are generated by a realizable deterministic teacher and that $L(\mathbf{w}^*) = 0$ holds at the optimal value. (See Murata, Müller, Ziehe, and Amari (1996) for a more general case.) We try to analyze the dynamical behavior of learning by using the continuous version of the algorithm for the sake of simplicity,

$$\frac{d}{dt}\mathbf{w}_t = -\eta_t G^{-1}(\mathbf{w}_t) \frac{\partial}{\partial \mathbf{w}} l(\mathbf{x}_t, \mathbf{y}_t; \mathbf{w}_t), \quad (5.3)$$

$$\frac{d}{dt}\eta_t = \alpha \eta_t [\beta l(\mathbf{x}_t, \mathbf{y}_t; \mathbf{w}_t) - \eta_t]. \quad (5.4)$$

In order to show the dynamical behavior of $(\mathbf{w}_t, \eta_t)$, we use the averaged version of equations 5.3 and 5.4 with respect to the current input-output pair $(\mathbf{x}_t, \mathbf{y}_t)$. The averaged learning equation (Amari, 1967, 1977) is written as

$$\frac{d}{dt}\mathbf{w}_t = -\eta_t G^{-1}(\mathbf{w}_t) \left\langle \frac{\partial}{\partial \mathbf{w}} l(\mathbf{x}, \mathbf{y}; \mathbf{w}_t) \right\rangle, \quad (5.5)$$

$$\frac{d}{dt}\eta_t = \alpha \eta_t \{\beta \langle l(\mathbf{x}, \mathbf{y}; \mathbf{w}_t) \rangle - \eta_t\}, \quad (5.6)$$

where $\langle \rangle$ denotes the average over the current $(\mathbf{x}, \mathbf{y})$. We also use the asymptotic evaluations

$$\begin{aligned} \left\langle \frac{\partial}{\partial \mathbf{w}} l(\mathbf{x}, \mathbf{y}; \mathbf{w}_t) \right\rangle &= \left\langle \frac{\partial}{\partial \mathbf{w}} l(\mathbf{x}, \mathbf{y}; \mathbf{w}^*) \right\rangle + \left\langle \frac{\partial^2}{\partial \mathbf{w} \partial \mathbf{w}} l(\mathbf{x}, \mathbf{y}; \mathbf{w}^*) (\mathbf{w}_t - \mathbf{w}^*) \right\rangle \\ &= G^*(\mathbf{w}_t - \mathbf{w}^*), \\ \langle l(\mathbf{x}, \mathbf{y}; \mathbf{w}_t) \rangle &= \frac{1}{2} (\mathbf{w}_t - \mathbf{w}^*)^T G^* (\mathbf{w}_t - \mathbf{w}^*), \end{aligned}$$

where $G^* = G(\mathbf{w}^*)$ and we used $L(\mathbf{w}^*) = 0$. We then have

$$\frac{d}{dt}\mathbf{w}_t = -\eta_t (\mathbf{w}_t - \mathbf{w}^*), \quad (5.7)$$

$$\frac{d}{dt}\eta_t = \alpha\eta_t \left\{ \frac{\beta}{2}(\mathbf{w}_t - \mathbf{w}^*)^T G^*(\mathbf{w}_t - \mathbf{w}^*) - \eta_t \right\}. \quad (5.8)$$

Now we introduce the squared error variable,

$$e_t = \frac{1}{2}(\mathbf{w}_t - \mathbf{w}^*)^T G^*(\mathbf{w}_t - \mathbf{w}^*), \quad (5.9)$$

where e_t is the Riemannian magnitude of $\mathbf{w}_t - \mathbf{w}^*$. It is easy to show

$$\frac{d}{dt}e_t = -2\eta_t e_t, \quad (5.10)$$

$$\frac{d}{dt}\eta_t = \alpha\beta\eta_t e_t - \alpha\eta_t^2. \quad (5.11)$$

The behavior of equations 5.10 and 5.11 is interesting. The origin $(0, 0)$ is its attractor. However, the basin of attraction has a boundary of fractal structure. Anyway, starting from an adequate initial value, it has the solution of the form

$$e_t = \frac{a}{t},$$

$$\eta_t = \frac{b}{t}.$$

The coefficients a and b are determined from

$$a = 2ab$$

$$b = -\alpha\beta ab + \alpha b^2.$$

This gives

$$b = \frac{1}{2},$$

$$a = \frac{1}{\beta} \left(\frac{1}{2} - \frac{1}{\alpha} \right), \quad \alpha > 2.$$

This proves the $1/t$ convergence rate of the generalization error, that is, the optimal order for any estimator $\hat{\mathbf{w}}_t$ converging to $\mathbf{w}^*$. The adaptive η_t shows a nice characteristic when the target teacher is slowly fluctuating or changes suddenly.

6 Natural Gradient in the Space of Perceptrons

The Riemannian metric and its inverse are calculated in this section to obtain the natural gradient explicitly. We begin with an analog simple perceptron whose input-output behavior is given by

$$y = f(\mathbf{w} \cdot \mathbf{x}) + n, \quad (6.1)$$

where n is a gaussian noise subject to $N(0, \sigma^2)$ and

$$f(u) = \frac{1 - e^{-u}}{1 + e^{-u}}. \quad (6.2)$$

The conditional probability density of y when $\mathbf{x}$ is applied is

$$p(y | \mathbf{x}; \mathbf{w}) = \frac{1}{\sqrt{2\pi}\sigma} \exp \left\{ -\frac{1}{2\sigma^2} [y - f(\mathbf{w} \cdot \mathbf{x})]^2 \right\}. \quad (6.3)$$

The distribution $q(\mathbf{x})$ of inputs $\mathbf{x}$ is assumed to be the normal distribution $N(0, I)$. The joint distribution of $(\mathbf{x}, y)$ is

$$p(y, \mathbf{x}; \mathbf{w}) = q(\mathbf{x})p(y | \mathbf{x}; \mathbf{w}).$$

In order to calculate the metric G of equation 3.11 explicitly, let us put

$$\mathbf{w}^2 = |\mathbf{w}|^2 = \sum w_i^2 \quad (6.4)$$

where $|\mathbf{w}|$ is the Euclidean norm. We then have the following theorem.

Theorem 3. *The Fisher information metric is*

$$G(\mathbf{w}) = \mathbf{w}^2 c_1(\mathbf{w}) I + \{c_2(\mathbf{w}) - c_1(\mathbf{w})\} \mathbf{w} \mathbf{w}^T, \quad (6.5)$$

where $c_1(\mathbf{w})$ and $c_2(\mathbf{w})$ are given by

$$\begin{aligned} c_1(\mathbf{w}) &= \frac{1}{4\sqrt{2\pi}\sigma^2\mathbf{w}^2} \int \{f^2(\mathbf{w}\varepsilon) - 1\}^2 \exp \left\{ -\frac{1}{2}\varepsilon^2 \right\} d\varepsilon, \\ c_2(\mathbf{w}) &= \frac{1}{4\sqrt{2\pi}\sigma^2\mathbf{w}^2} \int \{f^2(\mathbf{w}\varepsilon) - 1\}^2 \varepsilon^2 \exp \left\{ -\frac{1}{2}\varepsilon^2 \right\} d\varepsilon. \end{aligned}$$

Proof. We have

$$\log p(y, \mathbf{x}; \mathbf{w}) = \log q(\mathbf{x}) - \log(\sqrt{2\pi}\sigma) - \frac{1}{2\sigma^2} [y - f(\mathbf{w} \cdot \mathbf{x})]^2.$$

Hence,

$$\begin{aligned} \frac{\partial}{\partial w_i} \log p(y, \mathbf{x}; \mathbf{w}) &= \frac{1}{\sigma^2} [y - f(\mathbf{w} \cdot \mathbf{x})] f'(\mathbf{w} \cdot \mathbf{x}) x_i \\ &= \frac{1}{\sigma^2} n f'(\mathbf{w} \cdot \mathbf{x}) x_i. \end{aligned}$$

The Fisher information matrix is given by

$$\begin{aligned} g_{ij}(\mathbf{w}) &= E \left[\frac{\partial}{\partial w_i} \log p \frac{\partial}{\partial w_j} \log p \right] \\ &= \frac{1}{\sigma^2} E[\{f'(\mathbf{w} \cdot \mathbf{x})\}^2 x_i x_j], \end{aligned}$$

where $E[n^2] = \sigma^2$ is taken into account. This can be written, in the vector-matrix form, as

$$G(\mathbf{w}) = \frac{1}{\sigma^2} E[(f')^2 \mathbf{x} \mathbf{x}^T].$$

In order to show equation 6.5, we calculate the quadratic form $\mathbf{r}^T G(\mathbf{w}) \mathbf{r}$ for arbitrary $\mathbf{r}$. When $\mathbf{r} = \mathbf{w}$,

$$\mathbf{w}^T G \mathbf{w} = \frac{1}{\sigma^2} E[\{f'(\mathbf{w} \cdot \mathbf{x})\}^2 (\mathbf{w} \cdot \mathbf{x})^2].$$

Since $u = \mathbf{w} \cdot \mathbf{x}$ is subject to $N(0, w^2)$, we put $u = w\varepsilon$, where ε is subject to $N(0, 1)$. Noting that

$$f'(u) = \frac{1}{2} \{1 - f^2(u)\},$$

we have,

$$\mathbf{w}^T G(\mathbf{w}) \mathbf{w} = \frac{w^2}{4\sqrt{2\pi}\sigma^2} \int \varepsilon^2 \{f^2(w\varepsilon) - 1\}^2 \exp\left\{-\frac{\varepsilon^2}{2}\right\} d\varepsilon,$$

which confirms equation 6.5 when $\mathbf{r} = \mathbf{w}$. We next put $\mathbf{r} = \mathbf{v}$, where $\mathbf{v}$ is an arbitrary unit vector orthogonal to $\mathbf{w}$ (in the Euclidean sense). We then have

$$\mathbf{v}^T G(\mathbf{w}) \mathbf{v} = \frac{1}{4\sigma^2} E[\{f^2(\mathbf{w} \cdot \mathbf{x}) - 1\}^2 (\mathbf{v} \cdot \mathbf{x})^2].$$

Since $u = \mathbf{w} \cdot \mathbf{x}$ and $v = \mathbf{v} \cdot \mathbf{x}$ are independent, and v is subject to $N(0, 1)$, we have

$$\begin{aligned} \mathbf{v}^T G(\mathbf{w}) \mathbf{v} &= \frac{1}{4\sigma^2} E[(\mathbf{v} \cdot \mathbf{x})^2] E[\{f^2(\mathbf{w} \cdot \mathbf{x}) - 1\}^2] \\ &= \frac{1}{4\sqrt{2\pi}\sigma^2} \int \{f^2(w\varepsilon) - 1\}^2 \exp\left\{-\frac{\varepsilon^2}{2}\right\} d\varepsilon. \end{aligned}$$

Since $G(\mathbf{w})$ in equation 6.5 is determined by the quadratic forms for n -independent $\mathbf{w}$ and $\mathbf{v}$'s, this proves equation 6.5.

To obtain the natural gradient, it is necessary to have an explicit form of G^{-1} . We can calculate $G^{-1}(\mathbf{w})$ explicitly in the perceptron case.

Theorem 4. The inverse of the Fisher information metric is

$$G^{-1}(\mathbf{w}) = \frac{1}{w^2 c_1(\mathbf{w})} I + \frac{1}{w^4} \left(\frac{1}{c_2(\mathbf{w})} - \frac{1}{c_1(\mathbf{w})} \right) \mathbf{w} \mathbf{w}^T. \quad (6.6)$$

This can easily be proved by direct calculation of GG^{-1} . The natural gradient learning equation (3.10) is then given by

$$\begin{aligned} \mathbf{w}_{t+1} = & \mathbf{w}_t + \eta_t \{y_t - f(\mathbf{w}_t \cdot \mathbf{x}_t)\} f'(\mathbf{w}_t \cdot \mathbf{x}_t) \\ & \left[\frac{1}{w_t^2 c_1(\mathbf{w}_t)} \mathbf{x}_t + \frac{1}{w_t^4} \left(\frac{1}{c_2(\mathbf{w}_t)} - \frac{1}{c_1(\mathbf{w}_t)} \right) (\mathbf{w}_t \cdot \mathbf{x}_t) \mathbf{w}_t \right]. \end{aligned} \quad (6.7)$$

We now show some other geometrical characteristics of the parameter space of perceptrons. The volume V_n of the manifold of simple perceptrons is measured by

$$V_n = \int \sqrt{|G(\mathbf{w})|} d\mathbf{w} \quad (6.8)$$

where $|G(\mathbf{w})|$ is the determinant of $G = (g_{ij})$, which represents the volume density by the Riemannian metric. It is interesting to see that the manifold of perceptrons has a finite volume.

Bayesian statistics considers that $\mathbf{w}$ is randomly chosen subject to a prior distribution $\pi(\mathbf{w})$. A choice of $\pi(\mathbf{w})$ is the Jeffrey prior or noninformative prior given by

$$\pi(\mathbf{w}) = \frac{1}{V_n} \sqrt{|G(\mathbf{w})|}. \quad (6.9)$$

The Jeffrey prior is calculated as follows.

Theorem 5. The Jeffrey prior and the volume of the manifold are given, respectively, by

$$\sqrt{|G(\mathbf{w})|} = \frac{w}{V_n} \sqrt{c_2(\mathbf{w}) \{c_1(\mathbf{w})\}^{n-1}}, \quad (6.10)$$

$$V_n = a_{n-1} \int \sqrt{c_2(\mathbf{w}) \{c_1(\mathbf{w})\}^{n-1}} w^n d\mathbf{w}, \quad (6.11)$$

respectively, where a_{n-1} is the area of the unit $(n-1)$ -sphere.

The Fisher metric G can also be calculated for multilayer perceptrons. Let us consider a multilayer perceptron having m hidden units with sigmoidal activation functions and a linear output unit. The input-output relation is

$$y = \sum v_i f(\mathbf{w}_i \cdot \mathbf{x}) + n,$$

or the conditional probability is

$$p(y | \mathbf{x}; \mathbf{v}, \mathbf{w}_1, \dots, \mathbf{w}_m) = c \exp \left[-\frac{1}{2} \{y - \sum v_i f(\mathbf{w}_i \cdot \mathbf{x})\}^2 \right]. \quad (6.12)$$

The total parameter $\mathbf{w}$ consist of $\{\mathbf{v}, \mathbf{w}_1, \dots, \mathbf{w}_m\}$. Let us calculate the Fisher information matrix G . It consists of $m+1$ blocks corresponding to these $\mathbf{w}_i$'s and $\mathbf{v}$.

From

$$\frac{\partial}{\partial \mathbf{w}_i} \log p(y | \mathbf{x}; \mathbf{w}) = n v_i f'(\mathbf{w}_i \cdot \mathbf{x}) \mathbf{x},$$

we easily obtain the block submatrix corresponding to $\mathbf{w}_i$ as

$$\begin{aligned} E \left[\frac{\partial}{\partial \mathbf{w}_i} \log p \frac{\partial}{\partial \mathbf{w}_i} \log p \right] &= \frac{1}{\sigma^4} E[n^2] v_i^2 E[\{f'(\mathbf{w}_i \cdot \mathbf{x})\}^2 \mathbf{x} \mathbf{x}^T] \\ &= \frac{1}{\sigma^2} v_i^2 E[\{f'(\mathbf{w}_i \cdot \mathbf{x})\}^2 \mathbf{x} \mathbf{x}^T]. \end{aligned}$$

This is exactly the same as the simple perceptron case except for a factor of $(v_i)^2$. For the off-diagonal block, we have

$$E \left[\frac{\partial}{\partial \mathbf{w}_i} \log p \frac{\partial}{\partial \mathbf{w}_j} \log p \right] = \frac{1}{\sigma^2} v_i v_j E[f'(\mathbf{w}_i \cdot \mathbf{x}) f'(\mathbf{w}_j \cdot \mathbf{x}) \mathbf{x} \mathbf{x}^T].$$

In this case, we have the following form,

$$G_{\mathbf{w}_i \mathbf{w}_j} = c_{ij} I + d_{ii} \mathbf{w}_i \mathbf{w}_i^T + d_{ij} \mathbf{w}_i \mathbf{w}_j^T + d_{ji} \mathbf{w}_j \mathbf{w}_i^T + d_{jj} \mathbf{w}_j \mathbf{w}_j^T, \quad (6.13)$$

where the coefficients c_{ij} and d_{ij} 's are calculated explicitly by similar methods.

The $\mathbf{v}$ block and $\mathbf{v}$ and $\mathbf{w}_i$ block are also calculated similarly. However, the inversion of G is not easy except for simple cases. It requires inversion of a $2(m+1)$ dimensional matrix. However, this is much better than the direct inversion of the original $(n+1)m$ -dimensional matrix of G . Yang and Amari (1997) performed a preliminary study on the performance of the natural gradient learning algorithm for a simple multilayer perceptron. The result shows that natural gradient learning might be free from the plateau phenomenon. Once the learning trajectory is trapped in a plateau, it takes a long time to get out of it.

7 Natural Gradient in the Space of Matrices and Blind Source Separation

We now define a Riemannian structure to the space of all the $m \times m$ nonsingular matrices, which forms a Lie group denoted by $Gl(m)$, for the purpose of introducing the natural gradient learning rule to the blind source separation problem. Let dW be a small deviation of a matrix from W to $W + dW$. The tangent space T_W of $Gl(m)$ at W is a linear space spanned by all such small deviations dW_{ij} 's and is called the Lie algebra.

We need to introduce an inner product at W by defining the squared norm of dW

$$ds^2 = \langle dW, dW \rangle_W = \|dW\|^2.$$

By multiplying W^{-1} from the right, W is mapped to $WW^{-1} = I$, the unit matrix, and $W + dW$ is mapped to $(W + dW)W^{-1} = I + dX$, where

$$dX = dWW^{-1}. \quad (7.1)$$

This shows that a deviation dW at W is equivalent to the deviation dX at I by the correspondence given by multiplication of W^{-1} . The Lie group invariance requires that the metric is kept invariant under this correspondence, that is, the inner product of dW at W is equal to the inner product of dWY at WY for any Y ,

$$\langle dW, dW \rangle_W = \langle dWY, dWY \rangle_{WY}. \quad (7.2)$$

When $Y = W^{-1}$, $WY = I$. This principle was used to derive the natural gradient in Amari, Cichocki, and Yang (1996); see also Yang and Amari (1997) for detail. Here we give its analysis by using dX .

We define the inner product at I by

$$\langle dX, dX \rangle_I = \sum_{i,j} (dX_{ij})^2 = \text{tr}(dX^T dX). \quad (7.3)$$

We then have the Riemannian metric structure at W as

$$\langle dW, dW \rangle_W = \text{tr}\{(W^{-1})^T dW^T dW W^{-1}\}. \quad (7.4)$$

We can write the metric tensor G in the component form. It is a quantity having four indices $G_{ij,kl}(W)$ such that

$$\begin{aligned} ds^2 &= \sum G_{ij,kl}(W) dW_{ij} dW_{kl}, \\ G_{ij,kl}(W) &= \sum_m \delta_{ik} W_{jm}^{-1} W_{lm}^{-1}, \end{aligned} \quad (7.5)$$

where W_{jm}^{-1} are the components of W^{-1} . While it may not appear to be straightforward to obtain the explicit form of G^{-1} and natural gradient $\tilde{\nabla}L$, in fact it can be calculated as shown below.

Theorem 6. *The natural gradient in the matrix space is given by*

$$\tilde{\nabla}L = (\nabla L)W^TW. \quad (7.6)$$

Proof. The metric is Euclidean at I , so that both $G(I)$ and its inverse, $G^{-1}(I)$, are the identity. Therefore, by mapping dW at W to dX at I , the natural gradient learning rule in terms of dX is written as

$$\frac{dX}{dt} = -\eta_t G^{-1}(I) \frac{\partial L}{\partial X} = -\eta_t \frac{\partial L}{\partial X}, \quad (7.7)$$

where the continuous time version is used. We have from equation 7.1

$$\frac{dX}{dt} = \frac{dW}{dt} W^{-1}. \quad (7.8)$$

The gradient $\partial L / \partial X$ is calculated as

$$\frac{\partial L}{\partial X} = \frac{\partial L(W)}{\partial W} \left(\frac{\partial W^T}{\partial X} \right) = \frac{\partial L}{\partial W} W^T.$$

Therefore, the natural gradient learning rule is

$$\frac{dW}{dt} = -\eta_t \frac{\partial L}{\partial W} W^T W,$$

which proves equation 7.6.

The $dX = dW W^{-1}$ forms a basis of the tangent space at W , but this is not integrable; that is, we cannot find any matrix function $X = X(W)$ that satisfies equation 7.1. Such a basis is called a nonholonomic basis. This is a locally defined basis but is convenient for our purpose. Let us calculate the natural gradient explicitly. To this end, we put

$$l(x, W) = -\log \det |W| - \sum_{i=1}^n \log f_i(y_i), \quad (7.9)$$

where $y = Wx$ and $f_i(y_i)$ is an adequate probability distribution. The expected loss is

$$L(W) = E[l(x, W)],$$

which represents the entropy of the output $\mathbf{y}$ after a componentwise non-linear transformation (Nadal & Parga, 1994; Bell & Sejnowski, 1995). The independent component analysis or the mutual information criterion also gives a similar loss function (Comon, 1994; Amari et al., 1996; see also Oja & Karhunen, 1995). When f_i is the true probability density function of the i th source, $l(\mathbf{x}, W)$ is the negative of the log likelihood.

The natural gradient of l is calculated as follows. We calculate the differential

$$dl = l(\mathbf{x}, W + dW) - l(\mathbf{x}, W) = -d \log \det |W| - \sum d \log f_i(y_i)$$

due to change dW . Then,

$$\begin{aligned} d \log \det |W| &= \log \det |W + dW| - \log \det |W| \\ &= \log \det |(W + dW)W^{-1}| = \log(\det |I + dX|) \\ &= \text{tr} dX. \end{aligned}$$

Similarly, from $d\mathbf{y} = dW\mathbf{x}$,

$$\begin{aligned} \sum d \log f_i(y_i) &= -\varphi(\mathbf{y})^T dW\mathbf{x} \\ &= -\varphi(\mathbf{y})^T dX\mathbf{y}, \end{aligned}$$

where $\varphi(\mathbf{y})$ is the column vector

$$\begin{aligned} \varphi(\mathbf{y}) &= [\varphi_1(y_1), \dots, \varphi_m(y_m)], \\ \varphi_i(y_i) &= -\frac{d}{dy} \log f_i(y_i). \end{aligned} \tag{7.10}$$

This gives $\partial L / \partial X$, and the natural gradient learning equation is

$$\frac{dW}{dt} = \eta_t (I - \varphi(\mathbf{y})^T \mathbf{y}) W. \tag{7.11}$$

The efficiency of this equation is studied from the statistical and information geometrical point of view (Amari & Kawanabe, 1997; Amari & Cardoso, in press). We further calculate the Hessian by using the natural frame dX ,

$$d^2 l = \mathbf{y}^T dX^T \dot{\varphi}(\mathbf{y}) dX \mathbf{y} + \varphi(\mathbf{y})^T dX dX \mathbf{y}, \tag{7.12}$$

where $\dot{\varphi}(\mathbf{y})$ is the diagonal matrix with diagonal entries $d\varphi_i(y_i)/dy_i$. Its expectation can be explicitly calculated (Amari et al., in press). The Hessian is decomposed into diagonal elements and two-by-two diagonal blocks (see also Cardoso & Laheld, 1996). Hence, the stability of the above learning rule is easily checked. Thus, in terms of dX , we can solve the two fundamental problems: the efficiency and the stability of learning algorithms of blind source separation (Amari & Cardoso, in press; Amari et al., in press).

8 Natural Gradient in Systems Space

The problem is how to define the Riemannian structure in the parameter space $\{W(z)\}$ of systems, where z is the time-shift operator. This was given in Amari (1987) from the point of view of information geometry (Amari, 1985, 1997a; Murray & Rice, 1993). We show here only ideas (see Douglas et al., 1996; Amari, Douglas, Cichocki, & Yang, 1997, for preliminary studies).

In the case of multiterminal deconvolution, a typical loss function l is given by

$$l = -\log \det |W_0| - \sum_i \int p\{y_i; W(z)\} \log f_i(y_i) dy_i, \quad (8.1)$$

where $p\{y_i; W(z)\}$ is the marginal distribution of $y(t)$ which is derived from the past sequence of $x(t)$ by matrix convolution $W(z)$ of equation 3.24. This type of loss function is obtained from maximization of entropy, independent component analysis, or maximum likelihood.

The gradient of l is given by

$$\nabla_m l = -(W_0^{-1})^T \delta_{0m} + \varphi(y_t) x^T(t-m), \quad (8.2)$$

where

$$\nabla_m = \frac{\partial}{\partial W_m},$$

and

$$\nabla l = \sum_{m=0}^d (\nabla_m l) z^{-m}. \quad (8.3)$$

In order to calculate the natural gradient, we need to define the Riemannian metric G in the manifold of linear systems. The geometrical theory of the manifold of linear systems by Amari (1987) defines the Riemannian metric and a pair of dual affine connections in the space of linear systems.

Let

$$dW(z) = \sum_m dW_m z^{-m} \quad (8.4)$$

be a small deviation of $W(z)$. We postulate that the inner product $\langle dW(z), dW(z) \rangle$ is invariant under the operation of any matrix filter $Y(z)$,

$$\langle dW(z), dW(z) \rangle_{W(z)} = \langle dW(z)Y(z), dW(z)Y(z) \rangle_{WY}, \quad (8.5)$$

where $\mathbf{Y}(z)$ is any system matrix. If we put

$$\mathbf{Y}(z) = \{\mathbf{W}(z)\}^{-1},$$

which is a general system not necessarily belonging to FIR,

$$\mathbf{W}(z)\{\mathbf{W}(z)\}^{-1} = \mathbf{I}(z),$$

which is the identity system

$$\mathbf{I}(z) = I$$

not including any z^{-m} terms. The tangent vector $d\mathbf{W}(z)$ is mapped to

$$dX(z) = d\mathbf{W}(z)\{\mathbf{W}(z)\}^{-1}. \quad (8.6)$$

The inner product at I is defined by

$$\langle dX(z), dX(z) \rangle_I = \sum_{m,ij} (dX_{m,ij})^2, \quad (8.7)$$

where $dX_{m,ij}$ are the elements of matrix dX_m .

The natural gradient

$$\tilde{\nabla}l = G^{-1} \circ \nabla l$$

of the manifold of systems is given as follows.

Theorem 7. *The natural gradient of the manifold of systems is given by*

$$\tilde{\nabla}l = \nabla l(z) \mathbf{W}^T(z^{-1}) \mathbf{W}(z), \quad (8.8)$$

where operator z^{-1} should be operated adequately.

The proof is omitted. It should be remarked that $\tilde{\nabla}l$ does not belong to the class of FIR systems, nor does it satisfy the causality condition either. Hence, in order to obtain an online learning algorithm, we need to introduce time delay to map it to the space of causal FIR systems. This article shows only the principles involved; details will be published in a separate article by Amari, Douglas, and Cichocki.

9 Conclusions

This article introduces the Riemannian structures to the parameter spaces of multilayer perceptrons, blind source separation, and blind source deconvolution by means of information geometry. The natural gradient learning method is then introduced and is shown to be statistically efficient. This implies that optimal online learning is as efficient as optimal batch learning when the Fisher information matrix exists. It is also suggested that natural gradient learning might be easier to get out of plateaus than conventional stochastic gradient learning.

Acknowledgments

I thank A. Cichocki, A. Back, and H. Yang at RIKEN Frontier Research Program for their discussions.

References

- Amari, S. (1967). Theory of adaptive pattern classifiers. *IEEE Trans., EC-16*(3), 299–307.
- Amari, S. (1977). Neural theory of association and concept-formation. *Biological Cybernetics*, 26, 175–185.
- Amari, S. (1985). *Differential-geometrical methods in statistics*. Lecture Notes in Statistics 28. New York: Springer-Verlag.
- Amari, S. (1987). Differential geometry of a parametric family of invertible linear systems—Riemannian metric, dual affine connections and divergence. *Mathematical Systems Theory*, 20, 53–82.
- Amari, S. (1993). Universal theorem on learning curves. *Neural Networks*, 6, 161–166.
- Amari, S. (1995). Learning and statistical inference. In M. A. Arbib (Ed.), *Handbook of brain theory and neural networks* (pp. 522–526). Cambridge, MA: MIT Press.
- Amari, S. (1996). Neural learning in structured parameter spaces—Natural Riemannian gradient. In M. C. Mozer, M. I. Jordan, & Th. Petsche (Eds.), *Advances in neural processing systems*, 9. Cambridge, MA: MIT Press.
- Amari, S. (1997a). Information geometry. *Contemporary Mathematics*, 203, 81–95.
- Amari, S. (1997b). *Superefficiency in blind source separation*. Unpublished manuscript.
- Amari, S., & Cardoso, J. F. (In press). Blind source separation—Semi-parametric statistical approach. *IEEE Trans. on Signal Processing*.
- Amari, S., Chen, T.-P., & Cichocki, A. (In press). Stability analysis of learning algorithms for blind source separation. *Neural Networks*.
- Amari, S., Cichocki, A., & Yang, H. H. (1996). A new learning algorithm for blind signal separation, in *NIPS'95*, vol. 8, Cambridge, MA: MIT Press.
- Amari, S., Douglas, S. C., Cichocki, A., & Yang, H. H. (1997). Multichannel blind deconvolution and equalization using the natural gradient. *Signal Processing*

- Advance in Wireless Communication Workshop*, Paris.
- Amari, S., & Kawanabe, M. (1997). Information geometry of estimating functions in semiparametric statistical models, *Bernoulli*, 3, 29–54.
- Amari, S., Kurata, K., & Nagaoka, H. (1992). Information geometry of Boltzmann machines. *IEEE Trans. on Neural Networks*, 3, 260–271.
- Amari, S., & Murata, N. (1993). Statistical theory of learning curves under entropic loss criterion. *Neural Computation*, 5, 140–153.
- Barkai, N., Seung, H. S., & Sompolinsky, H. (1995). Local and global convergence of on-line learning. *Phys. Rev. Lett.*, 75, 1415–1418.
- Bell, A. J., & Sejnowski, T. J. (1995). An information-maximization approach to blind separation and blind deconvolution. *Neural Computation*, 7, 1129–1159.
- Bickel, P. J., Klassen, C. A. J., Ritov, Y., & Wellner, J. A. (1993). *Efficient and adaptive estimation for semiparametric models*. Baltimore: Johns Hopkins University Press.
- Campbell, L. L. (1985). The relation between information theory and the differential-geometric approach to statistics. *Information Sciences*, 35, 199–210.
- Cardoso, J. F., & Laheld, B. (1996). Equivariant adaptive source separation. *IEEE Trans. on Signal Processing*, 44, 3017–3030.
- Chentsov, N. N. (1972). *Statistical decision rules and optimal inference* (in Russian). Moscow: Nauka [translated in English (1982), Rhode Island: AMS].
- Comon, P. (1994). Independent component analysis, a new concept? *Signal Processing*, 36, 287–314.
- Douglas, S. C., Cichocki, A., & Amari, S. (1996). Fast convergence filtered regressor algorithms for blind equalization. *Electronics Letters*, 32, 2114–2115.
- Heskes, T., & Kappen, B. (1991). Learning process in neural networks. *Physical Review*, A44, 2718–2762.
- Jutten, C., & Héroult, J. (1991). Blind separation of sources, an adaptive algorithm based on neuromimetic architecture. *Signal Processing*, 24(1), 1–31.
- Kushner, H. J., & Clark, D. S. (1978). *Stochastic approximation methods for constrained and unconstrained systems*. Berlin: Springer-Verlag.
- Murata, N., & Müller, K. R., Ziehe, A., & Amari, S. (1996). Adaptive on-line learning in changing environments. In M. C. Mozer, M. I. Jordan, & Th. Petsche (Eds.), *Advances in neural processing systems*, 9. Cambridge, MA: MIT Press.
- Murray, M. K., & Rice, J. W. (1993). *Differential geometry and statistics*. New York: Chapman & Hall.
- Nadal, J. P. & Parga, N. (1994). Nonlinear neurons in the low noise limit—A factorial code maximizes information transfer. *Network*, 5, 561–581.
- Oja, E., & Karhunen, J. (1995). Signal separation by nonlinear Hebbian learning. In M. Palaniswami et al. (Eds.), *Computational intelligence—A dynamic systems perspective* (pp. 83–97). New York: IEEE Press.
- Opper, M. (1996). Online versus offline learning from random examples: General results. *Phys. Rev. Lett.*, 77, 4671–4674.
- Rao, C. R. (1945). Information and accuracy attainable in the estimation of statistical parameters. *Bulletin of the Calcutta Mathematical Society*, 37, 81–91.
- Rumelhart, D. E., Hinton, G. E., & Williams, R. J. (1986). Learning internal representations by error propagation. *Parallel Distributed Processing* (Vol. 1, pp. 318–362). Cambridge, MA: MIT Press.

- Saad, D., & Solla, S. A. (1995). On-line learning in soft committee machines. *Phys. Rev. E*, 52, 4225–4243.
- Sompolinsky, H., Barkai, N., & Seung, H. S. (1995). On-line learning of dichotomies: Algorithms and learning curves. In J.-H. Oh et al. (Eds.), *Neural networks: The statistical mechanics perspective* (pp. 105–130). Proceedings of the CTP-PBSRI Joint Workshop on Theoretical Physics. Singapore: World Scientific.
- Tsybkin, Ya. Z. (1973). *Foundation of the theory of learning systems*. New York: Academic Press.
- Van den Broeck, C., & Reimann, P. (1996). Unsupervised learning by examples: On-line versus off-line. *Phys. Rev. Lett.*, 76, 2188–2191.
- Widrow, B. (1963). *A statistical theory of adaptation*. Oxford: Pergamon Press.
- Yang, H. H., & Amari, S. (1997). *Application of natural gradient in training multilayer perceptrons*. Unpublished manuscript.
- Yang, H. H., & Amari, S. (In press). Adaptive on-line learning algorithms for blind separation—Maximum entropy and minimal mutual information. *Neural Computation*.

Received January 24, 1997; accepted May 20, 1997.