



A dynamical-system analysis of the optimum s-gradient algorithm

Luc Pronzato, Henry P. Wynn, Anatoly A. Zhigljavsky

► To cite this version:

Luc Pronzato, Henry P. Wynn, Anatoly A. Zhigljavsky. A dynamical-system analysis of the optimum s-gradient algorithm. Luc Pronzato, Anatoly Zhigljavsky. Optimal Design and Related Areas in Optimization and Statistics, Springer Verlag, pp.39-80, 2009, Springer Optimization and its Applications, 10.1007/978-0-387-79936-0_3 . hal-00358755

HAL Id: hal-00358755

<https://hal.science/hal-00358755>

Submitted on 4 Feb 2009

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.

A dynamical-system analysis of the optimum s -gradient algorithm

L. Pronzato, H.P. Wynn, and A. Zhigljavsky

Summary. We study the asymptotic behaviour of Forsythe's s -optimum gradient algorithm for the minimization of a quadratic function in $\mathbb{R}^d$ using a renormalization that converts the algorithm into iterations applied to a probability measure. Bounds on the performance of the algorithm (rate of convergence) are obtained through optimum design theory and the limiting behaviour of the algorithm for $s = 2$ is investigated into details. Algorithms that switch periodically between $s = 1$ and $s = 2$ are shown to converge much faster than when s is fixed at 2.

1.1 Introduction

The asymptotic behavior of the steepest-descent algorithm (that is, the optimum 1-gradient method) for the minimization of a quadratic function in $\mathbb{R}^d$ is well-known, see Akaike (1959); Nocedal *et al.* (1998, 2002) and Chap. 7 of (Pronzato *et al.*, 2000). Any vector y of norm one with two nonzero components only is a fixed point for two iterations of the algorithm after a suitable renormalization. The main result is that, in the renormalized space, one typically observes convergence to a two-point limit set which lies in the space spanned by the eigenvectors corresponding to the smallest and largest eigenvalues of the matrix A of the quadratic function. The proof for bounded quadratic operators in Hilbert space is similar to the proof for $\mathbb{R}^d$ although more technical, see Pronzato *et al.* (2001, 2006). In both cases, the method consists of converting the renormalized algorithm into iterations applied to a measure ν_k supported on the spectrum of A . The additional technicalities arise from the fact that in the Hilbert space case the measure may be continuous. For $s = 1$, the well-known inequality of Kantorovich gives a bound on the rate of convergence of the algorithm, see Kantorovich and Akilov (1982) and (Luenberger, 1973, p. 151). However, the actual asymptotic rate of convergence, although satisfying the Kantorovich bound, depends on the starting point and is difficult to predict; a lower bound can be obtained (Pronzato *et al.*, 2001, 2006) from considerations on the stability of the fixed points for the attractor.

The situation is much more complicated for the optimum s -gradient algorithm with $s \geq 2$ and the paper extends the results presented in (Forsythe,

1968) in several directions. First, two different sequences are shown to be monotonically increasing along the trajectory followed by the algorithm (after a suitable renormalization) and a link with optimum design theory is established for the construction of upper bounds for these sequences. Second, the case $s = 2$ is investigated into details and a precise characterization of the limiting behavior of the renormalized algorithm is given. Finally, we show how switching periodically between the algorithms with respectively $s = 1$ and $s = 2$ drastically improves the rate of convergence. The resulting algorithm is shown to have superlinear convergence in $\mathbb{R}^3$ and we give some explanations for the fast convergence observed in simulations in $\mathbb{R}^d$ with d large: by switching periodically between algorithms one destroys the stability of the limiting behavior obtained when s is fixed (which is always associated with slow convergence).

The chapter is organized as follows. Sect. 1.2 presents the optimum s -gradient algorithm for the minimization of a quadratic function in $\mathbb{R}^d$, first in the original space and then, after a suitable renormalization, as a transformation applied to a probability measure. Rates of convergence are defined in the same section. The asymptotic behavior of the optimum s -gradient algorithm in $\mathbb{R}^d$ is considered in Sect. 1.3 where some of the properties established in (Forsythe, 1968) are recalled. The analysis for the case $s = 2$ is detailed in Sect. 1.4. Switching strategies that periodically alternate between $s = 1$ and $s = 2$ are considered in Sect. 1.5.

1.2 The optimum s -gradient algorithm for the minimization of a quadratic function

Let A be a real bounded self-adjoint (symmetric) operator in a real Hilbert space $\mathcal{H}$ with inner product (x, y) and norm given by $\|x\| = (x, x)^{1/2}$. We shall assume that A is positive, bounded below, and its spectral boundaries will be denoted by m and M :

$$m = \inf_{\|x\|=1} (Ax, x), \quad M = \sup_{\|x\|=1} (Ax, x),$$

with $0 < m < M < \infty$. The function f_0 to be minimized with respect to $t \in \mathcal{H}$ is the quadratic form

$$f_0(t) = \frac{1}{2}(At, t) - (t, y)$$

for some $y \in \mathcal{H}$, the minimum of which is located at $t^* = A^{-1}y$. By a translation of the origin, which corresponds to the definition of $x = t - t^*$ as the variable of interest, the minimization of f_0 becomes equivalent to that of f defined by

$$f(x) = \frac{1}{2}(Ax, x), \tag{1.1}$$

which is minimum at $x^* = 0$. The directional derivative of f at x in the direction u is

$$\nabla_u f(x) = (Ax, u).$$

The direction of steepest descent at x is $-g$, with $g = g(x) = Ax$ the gradient of f at x . The minimum of f along the line $\mathcal{L}_1(x) = \{x + \gamma Ax, \gamma \in \mathbb{R}\}$ is obtained for the optimum step-length

$$\gamma^* = -\frac{(g, g)}{(Ag, g)},$$

which corresponds to the usual steepest-descent algorithm. One iteration of the steepest-descent algorithm, or optimum 1-gradient method, is thus

$$x_{k+1} = x_k - \frac{(g_k, g_k)}{(Ag_k, g_k)} g_k, \quad (1.2)$$

with $g_k = Ax_k$ and x_0 some initial element in $\mathcal{H}$. For any integer $s \geq 1$, define the s -dimensional plane of steepest descent by

$$\mathcal{L}_s(x) = \left\{ x + \sum_{i=1}^s \gamma_i A^i x, \gamma_i \in \mathbb{R} \text{ for all } i \right\}.$$

In the optimum s -gradient method, x_{k+1} is chosen as the point in $\mathcal{L}_s(x_k)$ that minimizes f . When $\mathcal{H} = \mathbb{R}^d$, A is $d \times d$ symmetric positive-definite matrix with minimum and maximum eigenvalues respectively m and M , and x_{k+1} is uniquely defined provided that the d eigenvalues of A are all distinct. Also, in that case $\mathcal{L}_d(x_k) = \mathbb{R}^d$ and only the case $s \leq d$ is of interest. We shall give special attention to the case $s = 2$.

1.2.1 Updating rules

Similarly to (Pronzato *et al.*, 2001, 2006) and Chap. 2 of this volume, consider the renormalized gradient

$$z(x) = \frac{g(x)}{(g(x), g(x))^{1/2}},$$

so that $(z(x), z(x)) = 1$ and denote $z_k = z(x_k)$ for all k . Also define

$$\mu_j^k = (A^j z_k, z_k), \quad j \in \mathbb{Z}, \quad (1.3)$$

so that $\mu_0^k = 1$ for any k and the optimum step-length of the optimum 1-gradient at step k is $-1/\mu_1^k$, see (1.2). The optimum choice of the s γ_i 's in the optimum s -gradient can be obtained by direct minimization of f over $\mathcal{L}_s(x_k)$. A simpler construction follows from the observation that g_{k+1} , and thus z_{k+1} , must be orthogonal to $\mathcal{L}_s(x_k)$, and thus to $z_k, Az_k, \dots, A^{s-1}z_k$. The vector of

optimum step-lengths at step k , $\overrightarrow{\gamma}_k = (\gamma_1^k, \dots, \gamma_s^k)^\top$, is thus solution of the following system of s linear equations

$$\mathbf{M}_{s,1}^k \overrightarrow{\gamma}_k = -(1, \mu_1^k, \dots, \mu_{s-1}^k)^\top, \quad (1.4)$$

where $\mathbf{M}_{s,1}^k$ is the $s \times s$ (symmetric) matrix with element (i, j) given by $\{\mathbf{M}_{s,1}^k\}_{i,j} = \mu_{i+j-1}^k$.

The following remark will be important later on, when we shall compare the rates of convergence of different algorithms.

Remark 1. One may notice that one step of the optimum s -gradient method starting from some x in $\mathcal{H}$ corresponds to s successive steps of the conjugate gradient algorithm starting from the same x , see (Luenberger, 1973, p. 179). $\square$

The next remark shows the connection with optimum design of experiments, which will be further considered in Sect. 1.2.3 (see also Pronzato *et al.* (2005) where the connection is developed around the case of the steepest-descent algorithm).

Remark 2. Consider Least-Squares (LS) estimation in a regression model

$$\sum_{i=1}^s \gamma_i t^i = -1 + \varepsilon_i$$

with (ε_i) a sequence of i.i.d. errors with zero mean. Assume that the t_i 's are generated according to a probability (design) measure ξ . Then, the LS estimator of the parameters γ_i , $i = 1, \dots, s$ is

$$\hat{\gamma} = - \left[\int (t, t^2, \dots, t^s)^\top (t, t^2, \dots, t^s) \xi(dt) \right]^{-1} \int (t, t^2, \dots, t^s)^\top \xi(dt)$$

and coincides with $\overrightarrow{\gamma}_k$ when ξ is such that

$$\int t^{j+1} \xi(dt) = \mu_j^k, \quad j = 0, 1, 2, \dots$$

The information matrix $\mathbf{M}(\xi)$ for this LS estimation problem then coincides with $\mathbf{M}_{s,1}^k$. $\square$

Using (1.4), one iteration of the optimum s -gradient method thus gives

$$x_{k+1} = Q_s^k(A)x_k, \quad g_{k+1} = Q_s^k(A)g_k \quad (1.5)$$

where $Q_s^k(t)$ is the polynomial $Q_s^k(t) = 1 + \sum_{i=1}^s \gamma_i^k t^i$ with the γ_i^k solutions of (1.4). Note that the use of any other polynomial $P(t)$ of degree s or less, and such that $P(0) = 1$, yields a larger value for $f(x_{k+1})$. Using (1.4), we obtain

$$Q_s^k(t) = 1 - (1, \mu_1^k, \dots, \mu_{s-1}^k) [\mathbf{M}_{s,1}^k]^{-1} \begin{pmatrix} t \\ \vdots \\ t^s \end{pmatrix}$$

and direct calculations give

$$Q_s^k(t) = \frac{\begin{vmatrix} 1 & \mu_1^k & \dots & \mu_{s-1}^k & 1 \\ \mu_1^k & \mu_2^k & \dots & \mu_s^k & t \\ \vdots & \vdots & \dots & \vdots & \vdots \\ \mu_s^k & \mu_{s+1}^k & \dots & \mu_{2s-1}^k & t^s \end{vmatrix}}{|\mathbf{M}_{s,1}^k|} \quad (1.6)$$

where, for any square matrix $\mathbf{M}$, $|\mathbf{M}|$ denotes its determinant. The derivation of the updating rule for the normalized gradient z_k relies on the computation of the inner product (g_{k+1}, g_{k+1}) . From the orthogonality property of g_{k+1} to $g_k, Ag_k, \dots, A^{s-1}g_k$ we get

$$\begin{aligned} (g_{k+1}, g_{k+1}) &= (g_{k+1}, \gamma_s^k A^s g_k) = \gamma_s^k (Q_s^k(A) A^s g_k, g_k) \\ &= \frac{\begin{vmatrix} 1 & \mu_1^k & \dots & \mu_{s-1}^k & \mu_s^k \\ \mu_1^k & \mu_2^k & \dots & \mu_s^k & \mu_{s+1}^k \\ \vdots & \vdots & \dots & \vdots & \vdots \\ \mu_s^k & \mu_{s+1}^k & \dots & \mu_{2s-1}^k & \mu_{2s}^k \end{vmatrix}}{|\mathbf{M}_{s,1}^k|} \gamma_s^k (g_k, g_k), \end{aligned} \quad (1.7)$$

where γ_s^k , the coefficient of t^s in $Q_s^k(t)$, is given by

$$\gamma_s^k = \frac{\begin{vmatrix} 1 & \mu_1^k & \dots & \mu_{s-1}^k \\ \mu_1^k & \mu_2^k & \dots & \mu_s^k \\ \vdots & \vdots & \dots & \vdots \\ \mu_{s-1}^k & \mu_s^k & \dots & \mu_{2s-2}^k \end{vmatrix}}{|\mathbf{M}_{s,1}^k|}. \quad (1.8)$$

1.2.2 The optimum s -gradient algorithm as a sequence of transformations of a probability measure

When $\mathcal{H} = \mathbb{R}^d$, we can assume that A is already diagonalized, with eigenvalues $0 < m = \lambda_1 \leq \lambda_2 \leq \dots \leq \lambda_d = M$, and consider $[z_k]_i^2$, with $[z_k]_i$ the i -th component of z_k , as a mass on the eigenvalue λ_i (note that $\sum_{i=1}^d [z_k]_i^2 = \mu_0^k = 1$). Define the discrete probability measure ν_k supported on $(\lambda_1, \dots, \lambda_d)$ by $\nu_k(\lambda_i) = [z_k]_i^2$, so that its j -th moment is μ_j^k , $j \in \mathbb{Z}$, see (1.3).

Remark 3. When two eigenvalues λ_i and λ_j of A are equal, their masses $[z_k]_i^2$ and $[z_k]_j^2$ can be added since the updating rule is the same for the two components $[z_k]_i$ and $[z_k]_j$. Concerning the analysis of the rate of convergence of

the optimum s -gradient algorithm, confounding masses associated with equal eigenvalues simply amounts to reducing the dimension of $\mathcal{H}$ and we shall therefore assume that all eigenvalues of A are different when studying the evolution of ν_k . $\square$

In the general case where $\mathcal{H}$ is a Hilbert space, let E_λ denote the spectral family associated with the operator A ; we then define the measure ν_k by $\nu_k(d\lambda) = d(E_\lambda z_k, z_k)$, $m \leq \lambda \leq M$. In both cases $\mathcal{H} = \mathbb{R}^d$ and $\mathcal{H}$ a Hilbert space, we consider ν_k as the spectral measure of A at the iteration k of the algorithm, and write

$$\mu_j^k = \int t^j \nu_k(dt).$$

For any measure ν on the interval $[m, M]$, any $\alpha \in \mathbb{R}$ and any positive integer m define

$$\mathbf{M}_{m,\alpha}(\nu) = \int t^\alpha (1, t, t^2, \dots, t^m)^\top (1, t, t^2, \dots, t^m) \nu(dt). \quad (1.9)$$

For both $\mathcal{H} = \mathbb{R}^d$ and $\mathcal{H}$ a Hilbert space, the iteration on z_k can be written as

$$\begin{aligned} z_k \rightarrow z_{k+1} = T_z(z_k) &= \frac{g_{k+1}}{(g_{k+1}, g_{k+1})^{1/2}} \\ &= \frac{(g_k, g_k)^{1/2}}{(g_{k+1}, g_{k+1})^{1/2}} Q_s^k(A) z_k \\ &= \frac{|\mathbf{M}_{s,1}^k|}{|\mathbf{M}_{s,0}^k|^{1/2} |\mathbf{M}_{s-1,0}^k|^{1/2}} Q_s^k(A) z_k, \end{aligned} \quad (1.10)$$

with $\mathbf{M}_{s,1}^k = \mathbf{M}_{s-1,1}(\nu_k)$ and $\mathbf{M}_{m,0}^k = \mathbf{M}_{m,0}(\nu_k)$, i.e. the $(m+1) \times (m+1)$ matrix with element (i, j) given by $\{\mathbf{M}_{m,0}^k\}_{i,j} = \mu_{i+j-2}^k$. The iteration on z_k can be interpreted as a transformation of the measure ν_k

$$\nu_k \rightarrow \nu_{k+1} = T_\nu(\nu_k) \quad \text{with} \quad \nu_{k+1}(dx) = H_k(x) \nu_k(dx), \quad (1.11)$$

where, using (1.5, 1.6, 1.7) and (1.8), we have

$$\begin{aligned} H_k(x) &= \frac{[Q_s^k(x)]^2 |\mathbf{M}_{s,1}^k|^2}{|\mathbf{M}_{s,0}^k| |\mathbf{M}_{s-1,0}^k|} \\ &= (1, x, \dots, x^s) [\mathbf{M}_{s,0}^k]^{-1} \begin{pmatrix} 1 \\ x \\ \vdots \\ x^s \end{pmatrix} - (1, x, \dots, x^{s-1}) [\mathbf{M}_{s-1,0}^k]^{-1} \begin{pmatrix} 1 \\ x \\ \vdots \\ x^{s-1} \end{pmatrix}. \end{aligned} \quad (1.12)$$

As moment matrices, $\mathbf{M}_{s,1}^k$ and $\mathbf{M}_{m,0}^k$ are positive semi-definite and $|\mathbf{M}_{s,1}^k| \geq 0$ for any $s \geq 1$ (respectively $|\mathbf{M}_{m,0}^k| \geq 0$ for any $m \geq 0$), with equality if and only if ν_k is supported on strictly less than s points (respectively

on m points or less). Also note that from the construction of the polynomial $Q_s^k(t)$, see (1.5), we have

$$\int Q_s^k(t) t^i \nu_k(dt) = (Q_s^k(A)g_k, A^i g_k) = 0, \quad i = 0, \dots, s-1, \quad (1.13)$$

which can be interpreted as an orthogonality property between the polynomials $Q_s^k(t)$ and t^i for $i = 0, \dots, s-1$. From this we can easily deduce the following.

Theorem 1. *Assume that ν_k is supported on $s+1$ points at least. Then the polynomial $Q_s^k(t)$ defined by (1.5) has s roots in the open interval (m, M) .*

Proof. Let ζ_i , $i = 1, \dots, q-1$, denote the roots of $Q_s^k(t)$ in (m, M) . Suppose that $q-1 < s$. Consider the polynomial $T(t) = (-1)^{q-1} Q_s^k(m) \prod_{i=1}^{q-1} (t - \zeta_i)$, it satisfies $T(t)Q_s^k(t) > 0$ for all $t \in (m, M)$, $t \neq \zeta_i$. Therefore, $\int T(t)Q_s^k(t)\nu_k(dt) > 0$, which contradicts (1.13) since $T(t)$ has degree $q-1 \leq s-1$. ■

Remark 4. Theorem 1 implies that $Q_s^k(m)$ has the same sign as $Q_s^k(0) = (-1)^s |\mathbf{M}_{s-1,0}^k|/|\mathbf{M}_{s,1}^k|$, that is, $(-1)^s$. Similarly, $Q_s^k(M)$ has the same sign as $\lim_{t \rightarrow \infty} Q_s^k(t)$ which is positive. □

Using the orthogonality property (1.13), we can also prove the next two theorems concerning the support of ν_k , see Forsythe (1968).

Theorem 2. *Assume that the measure ν_k is supported on $s+1$ points at least. Then this is also true for the measure ν_{k+1} obtained through (1.11).*

Proof. We only need to consider the case when the support $\mathcal{S}_k$ of ν_k is finite, that is, when $\mathcal{H} = \mathbb{R}^d$. Suppose that ν_k is supported on n points, $n \geq s+1$. The determinants $|\mathbf{M}_{s,1}^k|$, $|\mathbf{M}_{s,0}^k|$ and $|\mathbf{M}_{s-1,1}^k|$ in (1.12) are thus strictly positive. Let q be the largest integer such that there exist $\lambda_{i_1} < \lambda_{i_2} < \dots < \lambda_{i_q}$ in $\mathcal{S}_k$ with $Q_s^k(\lambda_{i_j})Q_s^k(\lambda_{i_{j+1}}) < 0$, $j = 1, \dots, q-1$. We shall prove that $q \geq s+1$. From (1.11, 1.12) it implies that ν_k is supported on $s+1$ points at least.

From (1.13), $\int Q_s^k(t) \nu_k(dt) = 0$, so that there exist λ_{i_1} and λ_{i_2} in $\mathcal{S}_k$ with $Q_s^k(\lambda_{i_1})Q_s^k(\lambda_{i_2}) < 0$, therefore $q \geq 2$. Suppose that $q \leq s$. By construction $Q_s^k(\lambda_{i_1})$ is of the same sign as $Q_s^k(m)$ and we can construct q disjoint open intervals A_j , $j = 1, \dots, q$ such that $\lambda_{i_j} \in A_j$ and $Q_s^k(\lambda_i)Q_s^k(\lambda_{i_j}) \geq 0$ for all $\lambda_i \in \mathcal{S}_k \cap A_j$ with $\cup_{j=1}^q \bar{A}_j = [m, M]$, where $\bar{A}_j$ is the closure of A_j (notice that $Q_s^k(\lambda)$ may change sign in A_j but all the $Q_s^k(\lambda_i)$'s are of the same sign for $\lambda_i \in \mathcal{S}_k \cap A_j$). Consider the $q-1$ scalars ζ_i , $i = 1, \dots, q-1$, defined by the endpoints of the A_j 's, m and M excluded; they satisfy $\lambda_{i_1} < \zeta_1 < \lambda_{i_2} < \dots < \zeta_{q-1} < \lambda_{i_q}$. Form now the polynomial $T(t) = (-1)^{s+q-1} (t - \zeta_1) \times \dots \times (t - \zeta_{q-1})$, one can check that $Q_s^k(\lambda_i)T(\lambda_i) \geq 0$ for all λ_i in $\mathcal{S}_k$. Also, $Q_s^k(\lambda_{i_j})T(\lambda_{i_j}) > 0$ for $j = 1, \dots, q$. This implies $\sum_{i=1}^d T(\lambda_i)Q_s^k(\lambda_i)[z_k]_i^2 > 0$, which contradicts (1.13) since $T(t)$ has degree $q-1 \leq s-1$. ■

Corollary 1. *If x_0 is such that ν_0 is supported on $n_0 = s + 1$ points at least, then $(g_k, g_k) > 0$ for all $k \geq 1$. Also, the determinants of all $(m + 1) \times (m + 1)$ moment matrices $\mathbf{M}_{m,\alpha}(\nu_k)$, see (1.9), are strictly positive for all $k \geq 1$.*

Theorem 3. *Assume that ν_0 is supported on $n_0 = s + 1$ points. Then $\nu_{2k} = \nu_0$ for all k .*

Proof. It is enough to prove that g_{2k} is parallel to g_0 , and thus that g_2 is parallel to g_0 . Since the updating rule only concerns nonzero components, we may assume that $d = s + 1$. We have $g_1 = Q_s^0(A)g_0$, $g_2 = Q_s^1(A)g_1$, g_1 is orthogonal to $g_0, Ag_0, \dots, A^{s-1}g_0$, which are independent, and g_2 is orthogonal to g_1 . We can thus decompose g_2 with respect to the basis $g_0, Ag_0, \dots, A^{s-1}g_0$ as

$$g_2 = \sum_{i=0}^{s-1} \alpha_i A^i g_0.$$

Now, g_2 is orthogonal to Ag_1 , and thus

$$g_1^\top Ag_2 = 0 = \sum_{i=0}^{s-1} \alpha_i g_1^\top A^{i+1} g_0 = \alpha_{s-1} g_1^\top A^s g_0$$

with $g_1^\top A^s g_0 \neq 0$ since otherwise g_1 would be zero. Therefore, $\alpha_{s-1} = 0$. Similarly, g_2 is orthogonal to $A^2 g_1$, which gives

$$0 = \sum_{i=0}^{s-2} \alpha_i g_1^\top A^{i+2} g_0 = \alpha_{s-2} g_1^\top A^s g_0$$

so that $\alpha_{s-2} = 0$. Continuing like that up to $g_1^\top A^{s-1} g_2$ we obtain $\alpha_1 = \alpha_2 = \dots = \alpha_{s-1} = 0$ and $g_2 = \alpha_0 g_0$. Notice that $\alpha_0 > 0$ since $(A^{-1} g_2, g_0) = (A^{-1} g_1, g_1)$, see (1.18). ■

The transformation $z_k \rightarrow z_{k+1} = T_z(z_k)$ (respectively $\nu_k \rightarrow \nu_{k+1} = T_\nu(\nu_k)$) can be considered as defining a dynamical system with state $z_k \in \mathcal{H}$ at iteration k (respectively, $\nu_k \in \Pi$, the set of probability measures defined on the spectrum of A). One purpose of the paper is to investigate the limit set of the orbit of the system starting at z_0 or ν_0 . As it is classical in the study of stability of dynamical systems where Lyapunov functions often play a key role (through the Lyapunov Stability Theorem or Lasalle's Invariance Principle, see, e.g., (Elaydi, 2005, Chap. 4)), the presence of monotone sequences in the dynamics of the renormalized algorithm will be an important ingredient of the analysis. Theorem 3 shows that the behavior of the renormalized algorithm may be periodic with period 2. We shall see that this type of behavior is typical, although the structure of the attractor may be rather complicated.

1.2.3 Rates of convergence and monotone sequences

Rates of convergence

Consider the following rate of convergence of the algorithm at iteration k ,

$$r_k = \frac{f(x_{k+1})}{f(x_k)} = \frac{(Ax_{k+1}, x_{k+1})}{(Ax_k, x_k)} = \frac{(A^{-1}g_{k+1}, g_{k+1})}{(A^{-1}g_k, g_k)}. \quad (1.14)$$

From the orthogonality property of g_{k+1} we have

$$(A^{-1}g_{k+1}, g_{k+1}) = (A^{-1}Q_s^k(A)g_k, g_{k+1}) = (A^{-1}g_k, g_{k+1})$$

and thus, using (1.6),

$$r_k = \frac{|\mathbf{M}_{s,-1}^k|}{|\mathbf{M}_{s,1}^k| \mu_{-1}^k}$$

with $\mathbf{M}_{s,-1}^k = \mathbf{M}_{s,-1}(\nu_k)$, see (1.9), that is, the $(s+1) \times (s+1)$ moment matrix with element (i, j) given by $\{\mathbf{M}_{s,-1}^k\}_{i,j} = \mu_{i+j-3}^k$. Using the orthogonality property of g_{k+1} again, we can easily prove that the sequence of rates (r_k) is non-decreasing along the trajectory followed by the algorithm.

Theorem 4. *When x_0 is such that ν_0 is supported on $s+1$ points at least, the rate of convergence r_k defined by (1.14) is non-decreasing along the path followed by the optimum s -gradient algorithm. It also satisfies*

$$r_k \leq R_s^* = T_s^{-2} \left(\frac{\varrho + 1}{\varrho - 1} \right) \quad (1.15)$$

where $\varrho = M/m$ is the condition number of A and $T_s(t)$ is the s -th Chebyshev polynomial (normalized so that $\max_{t \in [-1, 1]} |T_s(t)| = 1$)

$$T_s(t) = \cos[s \arccos(t)] = \frac{(t + \sqrt{t^2 - 1})^s + (t - \sqrt{t^2 - 1})^s}{2}. \quad (1.16)$$

Moreover, the equality in (1.15) is obtained when ν_k is the measure ν_s^* defined by

$$\nu_s^*(y_0) = \nu_s^*(y_s) = 1/(2s), \quad \nu_s^*(y_j) = 1/s, \quad 1 \leq j \leq s-1, \quad (1.17)$$

where $y_j = (M+m)/2 + [\cos(j\pi/s)](M-m)/2$.

Proof. From Corollary 1, $(g_0, g_0) > 0$ implies $(g_k, g_k) > 0$ for all k and r_k is thus well defined. Straightforward manipulations give

$$\begin{aligned} (A^{-1}g_{k+1}, g_{k+1}) - (A^{-1}g_{k+2}, g_k) &= (A^{-1}[Q_s^k(A) - Q_s^{k+1}(A)]g_k, Q_s^k(A)g_k) \\ &= \left(A^{-1} \sum_{i=1}^s (\gamma_i^k - \gamma_i^{k+1}) A^i g_k, Q_s^k(A)g_k \right) = 0 \end{aligned} \quad (1.18)$$

where equality to zero follows from (1.13). Therefore, from the Cauchy–Schwarz inequality,

$$(A^{-1}g_{k+1}, g_{k+1})^2 = (A^{-1}g_{k+2}, g_k)^2 \leq (A^{-1}g_{k+2}, g_{k+2})(A^{-1}g_k, g_k)$$

and $r_k \leq r_{k+1}$, with equality if and only if $g_{k+2} = \alpha g_k$ for some $\alpha \in \mathbb{R}^+$ ($\alpha > 0$ since $(A^{-1}g_{k+2}, g_k) = (A^{-1}g_{k+1}, g_{k+1})$). This shows that r_k is non-decreasing.

The rate (1.14) can also be written as

$$r_k = [\mu_{-1}^k \{(\mathbf{M}_{s,-1}^k)^{-1}\}_{1,1}]^{-1}.$$

Define the measure $\bar{\nu}_k$ by $\bar{\nu}_k(dt) = \nu_k(dt)/(t\mu_{-1}^k)$ (so that $\int \bar{\nu}_k(dt) = 1$) and denote $\bar{\mathbf{M}}_{m,n}^k$ the matrix obtained by substituting $\bar{\nu}_k$ for ν_k in $\mathbf{M}_{m,n}^k$ for any n, m . Then, $\bar{\mathbf{M}}_{s,0}^k = \mathbf{M}_{s,-1}^k/\mu_{-1}^k$ and $r_k = [\{(\bar{\mathbf{M}}_{s,0}^k)^{-1}\}_{1,1}]^{-1}$. The maximum value for r_k is thus obtained for the D_s -optimal measure $\bar{\nu}_s^*$ on $[m, M]$ for the estimation θ_0 in the linear regression model $\eta(\theta, x) = \sum_{i=0}^s \theta_i x^i$ with i.i.d. errors, see, e.g., (Fedorov, 1972, p. 144) and (Silvey, 1980, p. 10) ($\bar{\nu}_s^*$ is also c -optimal for $c = (1, 0, \dots, 0)^\top$). This measure is uniquely defined, see Hoel and Levine (1964), (Sahm, 1998, p. 52): it is supported at the $s+1$ points $y_j = (M+m)/2 + [\cos(j\pi/s)](M-m)/2$, $j = 0, \dots, s$, and each y_j receives a weight proportional to α_j/y_j , with $\alpha_0 = \alpha_s = 1/2$ and $\alpha_j = 1$ for $j = 1, \dots, s-1$. Applying the transformation $\nu(dt) = t\bar{\nu}(dt)\mu_{-1}$ we obtain the measure ν_s^* given by (1.17). ■

Remark 5. Meinardus (1963) and Forsythe (1968) arrive at the result (1.15) by a different route. They write

$$r_k = \frac{(A^{-1}g_{k+1}, g_{k+1})}{(A^{-1}g_k, g_k)} = \frac{\int [Q_s^k(t)]^2 t^{-1} \nu_k(dt)}{\mu_{-1}^k}.$$

Since $Q_s^k(t)$ minimizes $f(x_{k+1})$, $r_k \leq (1/\mu_{-1}^k) \int P^2(t) t^{-1} \nu_k(dt)$ for any s -degree polynomial $P(t)$ such that $P(0) = 1$. Equivalently,

$$r_k \leq \frac{\int S^2(t) t^{-1} \nu_k(dt)}{S^2(0)\mu_{-1}^k}$$

for any s -degree polynomial $S(t)$. Take $S(t) = S^*(t) = T_s[(M+m-2t)/(M-m)]$, so that $S^2(t) \leq 1$ for $t \in [m, M]$, then r_k satisfies $r_k \leq [S^*(0)]^{-2} = R_s^*$ with R_s^* given by (1.15). □

Notice that $T_s[(\varrho+1)/(\varrho-1)] > 1$ in (1.15), so that we have the following.

Corollary 2. *If x_0 is such that ν_0 is supported on $s+1$ points at least, then the optimum s -gradient algorithm converges linearly to the optimum, that is,*

$$0 < c_1 = \frac{f(x_1)}{f(x_0)} \leq \frac{f(x_{k+1})}{f(x_k)} \leq R_s^* < 1, \quad \text{for all } k.$$

Moreover, the convergence slows down monotonically on the route to the optimum and the rate r_k given by (1.14) tends to a limit r_∞ .

The monotonicity of the sequence (r_k) , together with Theorem 3, has the following consequence.

Corollary 3. *Assume that ν_0 is supported on $n_0 = s + 1$ points. Then $r_{k+1} = r_k$ for all k .*

Other rates of convergence can be defined as

$$R_k(W) = \frac{(Wg_{k+1}, g_{k+1})}{(Wg_k, g_k)} \quad (1.19)$$

with W be a bounded positive self-adjoint operator in $\mathcal{H}$. However, the following theorem shows that all such rates are asymptotically equivalent, see Pronzato *et al.* (2006).

Theorem 5. *Let W be a bounded positive self-adjoint operator in $\mathcal{H}$, with bounds c and C such that $0 < c < C < \infty$ (when $\mathcal{H} = \mathbb{R}^d$, W is a $d \times d$ positive-definite matrix with minimum and maximum eigenvalues respectively c and C). Consider the rate of convergence defined by (1.19) if $\|g_k\| \neq 0$ and $R_k(W) = 1$ otherwise. Apply the optimum s -gradient algorithm (1.5), initialized at $g_0 = g(x_0)$, for the minimization of $f(x)$ given by (1.1). Then the limit*

$$R(W, x_0) = \lim_{n \rightarrow \infty} \left[\prod_{k=0}^{n-1} R_k(W) \right]^{1/n}$$

exists for all x_0 in $\mathcal{H}$ and $R(W, x_0) = R(x_0)$ does not depend on W . In particular,

$$R(W, x_0) = r_\infty = \lim_{n \rightarrow \infty} \left(\prod_{k=0}^{n-1} r_k \right)^{1/n} \quad (1.20)$$

with r_k defined by (1.14).

Proof. Assume that x_0 is such that for some $k \geq 0$, $\|g_{k+1}\| = 0$ with $\|g_i\| > 0$ for all $i \leq k$ (that is, $x_{k+1} = x^*$ and $x_i \neq x^*$ for $i \leq k$). This implies $R_k(W) = 0$ for any W , and therefore $R(W, x_0) = R(x_0) = 0$.

Assume now that $\|g_k\| > 0$ for all k . Consider

$$V_n = \left[\prod_{k=0}^{n-1} R_k(W) \right]^{1/n} = \left[\prod_{k=0}^{n-1} \frac{(Wg_{k+1}, g_{k+1})}{(Wg_k, g_k)} \right]^{1/n} = \left[\frac{(Wg_n, g_n)}{(Wg_0, g_0)} \right]^{1/n}.$$

We have,

$$\forall z \in \mathcal{H}, \quad c\|z\|^2 \leq (Wz, z) \leq C\|z\|^2,$$

and thus

$$(c/C)^{1/n} \left[\frac{(g_n, g_n)}{(g_0, g_0)} \right]^{1/n} \leq V_n \leq (C/c)^{1/n} \left[\frac{(g_n, g_n)}{(g_0, g_0)} \right]^{1/n}.$$

Since $(c/C)^{1/n} \rightarrow 1$ and $(C/c)^{1/n} \rightarrow 1$ as $n \rightarrow \infty$, $\liminf_{n \rightarrow \infty} V_n$ and $\limsup_{n \rightarrow \infty} V_n$ do not depend on W . Taking $W = A^{-1}$ we get $R_k(W) = r_k$. The sequence (r_k) is not decreasing, and thus $\lim_{n \rightarrow \infty} V_n = r_\infty$ for any W . ■

For any fixed $\varrho = M/m$, the bound R_s^* given by (1.15) tends to zero as s tends to infinity whatever the dimension d when $\mathcal{H} = \mathbb{R}^d$, and also when $\mathcal{H}$ is a Hilbert space. However, since one step of the optimum s -gradient method corresponds to s successive steps of the conjugate gradient algorithm, see Remark 1, a normalized version of the convergence rate allowing comparison with classical steepest descent is $r_k^{1/s}$, which is bounded by

$$N_s^* = (R_s^*)^{1/s} = T_s^{-2/s} \left(\frac{\varrho + 1}{\varrho - 1} \right) \quad (1.21)$$

where $T_s(t)$ is the s -th Chebyshev polynomial, see (1.16). The quantity N_s^* is a decreasing function of s , see Fig. 1.1, but has a positive limit when s tends to infinity,

$$\lim_{s \rightarrow \infty} N_s^* = N_\infty^* = \frac{(\sqrt{\varrho} - 1)^2}{(\sqrt{\varrho} + 1)^2}. \quad (1.22)$$

Fig. 1.2 shows the evolution of N_∞^* as a function of the condition number ϱ .

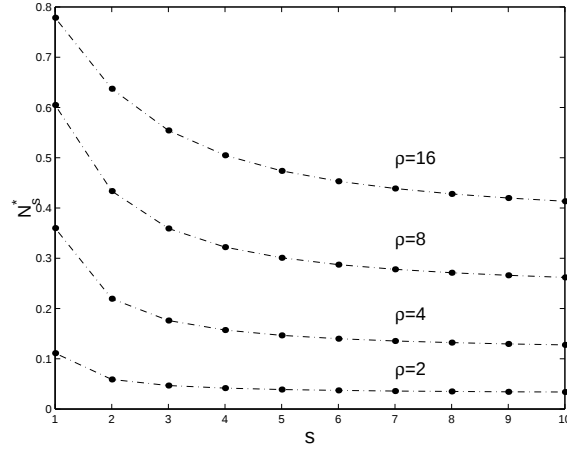


Fig. 1.1. Upper bounds N_s^* , see (1.21), as functions of s for different values of the condition number ϱ

A second monotone bounded sequence

Another quantity q_k also turns out to be non-decreasing along the trajectory followed by the algorithm, as shown in the following theorem.

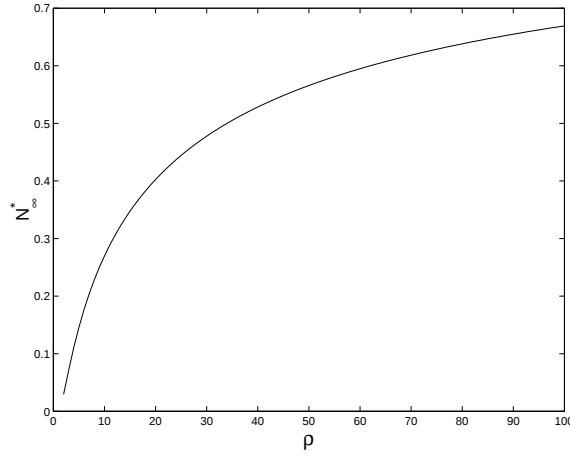


Fig. 1.2. Limiting value N_∞^* as a function of ρ

Theorem 6. When x_0 is such that ν_0 is supported on $s + 1$ points at least, the quantity q_k defined by

$$q_k = \frac{(g_{k+1}, g_{k+1})}{(\gamma_s^k)^2 (g_k, g_k)}, \quad (1.23)$$

with γ_s^k given by (1.8), is non-decreasing along the path followed by the optimum s -gradient algorithm. Moreover, it satisfies

$$q_k \leq q_s^* = \frac{(M - m)^{2s}}{2^{4s-2}}, \quad \text{for all } k, \quad (1.24)$$

where the equality is obtained when ν_k is the measure (1.17) of Theorem 4.

Proof. From Corollary 1, $(g_0, g_0) > 0$ implies $(g_k, g_k) > 0$ and $\gamma_s^k > 0$ for all k so that q_k is well defined. Using the same approach as for r_k in Theorem 4, we write

$$\begin{aligned} (g_{k+1}, g_{k+1})/\gamma_s^k - (g_{k+2}, g_{k+2})/\gamma_s^{k+1} &= (Q_s^k(A)g_k, g_{k+1})/\gamma_s^k \\ &\quad - (Q_s^{k+1}(A)g_k, g_{k+1})/\gamma_s^{k+1} \\ &= ([Q_s^k(A)/\gamma_s^k - Q_s^{k+1}(A)/\gamma_s^{k+1}]g_k, g_{k+1}) \\ &= \left(\left[(1/\gamma_s^k - 1/\gamma_s^{k+1})I + \sum_{i=1}^{s-1} (\gamma_i^k/\gamma_s^k - \gamma_i^{k+1}/\gamma_s^{k+1})A^i \right] g_k, g_{k+1} \right) \\ &= 0, \end{aligned}$$

where equality to zero follows from (1.13). The Cauchy–Schwarz inequality then implies

$$(g_{k+1}, g_{k+1})^2 / (\gamma_s^k)^2 = (g_{k+2}, g_k)^2 / (\gamma_s^{k+1})^2 \leq (g_{k+2}, g_{k+2})(g_k, g_k) / (\gamma_s^{k+1})^2$$

and $q_k \leq q_{k+1}$ with equality if and only if $g_{k+2} = \alpha g_k$ for some $\alpha \in \mathbb{R}^+$ ($\alpha > 0$) since $(g_{k+2}, g_k) = (g_{k+1}, g_{k+1})\gamma_s^{k+1}/\gamma_s^k$ and $\gamma_s^k > 0$ for all k , see Corollary 1). This shows that q_k is non-decreasing.

The determination of the probability measure that maximizes q_k is again related to optimal design theory. Using (1.7) and (1.8), we obtain

$$q_k = \frac{|\mathbf{M}_{s,0}^k|}{|\mathbf{M}_{s-1,0}^k|}.$$

Hence, using the inversion of a partitioned matrix, we can write

$$q_k = \mu_{2s}^k - (1, \mu_1^k, \dots, \mu_{s-1}^k)(\mathbf{M}_{s-1,0}^k)^{-1} \begin{pmatrix} 1 \\ \mu_1^k \\ \vdots \\ \mu_{s-1}^k \end{pmatrix} = \left(\{(\mathbf{M}_{s,0}^k)^{-1}\}_{s+1, s+1} \right)^{-1},$$

so that the maximization of q_k with respect to ν_k is equivalent to the determination of a D_s -optimum measure on $[m, M]$ for the estimation of θ_s in the linear regression model $\eta(\theta, x) = \sum_{i=0}^s \theta_i x^i$ with i.i.d. errors (or to the determination of a c -optimal measure on $[m, M]$ with $c = (0, \dots, 0, 1)^\top$). This measure is uniquely defined, see (Kiefer and Wolfowitz, 1959, p. 283): when the design interval is normalized to $[-1, 1]$, the optimum measure ξ^* is supported on $s+1$ points given by ± 1 and the $s-1$ zeros of the derivative of s -th Chebyshev polynomial $T_s(t)$ given by (1.16) and the weights are

$$\xi^*(-1) = \xi^*(1) = 1/(2s), \quad \xi^*(\cos[j\pi/s]) = 1/s, \quad 1 \leq j \leq s-1.$$

The transformation $t \in [-1, 1] \mapsto z = (M+m)/2 + t(M-m)/2 \in [m, M]$ gives the measure ν_s^* on $[m, M]$. The associated maximum value for q_k is q_s^* given by (1.24), see (Kiefer and Wolfowitz, 1959, p. 283). $\blacksquare$

The monotonicity of the sequence (q_k) , together with Theorem 3, implies the following analogue to Corollary 3.

Corollary 4. *Assume that ν_0 is supported on $n_0 = s+1$ points. Then $q_{k+1} = q_k$ for all k .*

As a non-decreasing and bounded sequence, (q_k) tends to a limit q_∞ . The existence of limiting values r_∞ and q_∞ will be essential for studying the limit points of the orbits (z_{2k}) or (ν_{2k}) in the next sections. In the rest of the paper we only consider the case where $\mathcal{H} = \mathbb{R}^d$ and assume that A is diagonalized with d distinct eigenvalues $0 < m = \lambda_1 < \lambda_2 < \dots < \lambda_d = M$.

1.3 Asymptotic behavior of the optimum s -gradient algorithm in $\mathbb{R}^d$

The situation is much more complex when $s \geq 2$ than for $s = 1$, and we shall recall some of the properties established in (Forsythe, 1968) for $\mathcal{H} = \mathbb{R}^d$. If x_0 is such that the initial measure ν_0 is supported on s points or less, the algorithm terminates in one step. In the rest of the section we thus suppose that ν_0 is supported on $n_0 \geq s + 1$ points. The algorithm then converges linearly to the optimum, see Corollary 2. The set $\mathcal{Z}(x_0)$ of limit points of the sequence of renormalized gradients z_{2k} satisfies the following.

Theorem 7. *If x_0 is such that ν_0 has $s + 1$ support points at least, the set $\mathcal{Z}(x_0)$ of limit points of the sequence of renormalized gradients z_{2k} , $k = 0, 1, 2, \dots$ is a closed connected subset of the d -dimensional unit sphere $\mathcal{S}_d$. Any y in $\mathcal{Z}(x_0)$ satisfies $T_z^2(y) = y$, where $T_z^2(y) = T_z[T_z(y)]$ with T_z defined by (1.10).*

Proof. Using (1.18) we get

$$\begin{aligned} r_{k+1} - r_k &= \frac{(A^{-1}g_{k+2}, g_{k+2})}{(A^{-1}g_{k+1}, g_{k+1})} \times \left[1 - \frac{(A^{-1}g_{k+2}, g_k)^2}{(A^{-1}g_{k+2}, g_{k+2})(A^{-1}g_k, g_k)} \right] \\ &= r_{k+1} \times \left[1 - \frac{(A^{-1}z_{k+2}, z_k)^2}{(A^{-1}z_{k+2}, z_{k+2})(A^{-1}z_k, z_k)} \right]. \end{aligned}$$

Since r_k is non-decreasing and bounded, see Theorem 4, r_k tends to a limit $r_\infty(x_0)$ and, using Cauchy-Schwarz inequality $\|z_{k+2} - z_k\| \rightarrow 0$. The set $\mathcal{Z}(x_0)$ of limit points of the sequence z_{2k} is thus a continuum on $\mathcal{S}_d$.

Take any $y \in \mathcal{Z}(x_0)$. There exists a subsequence (k_i) such that $z_{2k_i} \rightarrow y$ as $i \rightarrow \infty$, and $z_{2k_i+2} = T_z^2(z_{2k_i}) \rightarrow y$ since $\|z_{2k_i+2} - z_{2k_i}\| \rightarrow 0$. The continuity of T_z then implies $z_{2k_i+2} \rightarrow T_z^2(y)$ and thus $T_z^2(y) = y$. ■

Obviously, the sequence (z_{2k+1}) satisfies a similar property (consider ν_1 as a new initial measure ν_0), so that we only need to consider the sequence of even iterates.

Remark 6. Forsythe (1968) conjectures that the continuum for $\mathcal{Z}(x_0)$ is in fact always a single point. Although it is confirmed by numerical simulations, we are not aware of any proof of attraction to a single point. One may however think of $\mathcal{Z}(x_0)$ as the set of possible limit points for the sequence (z_{2k}) , leaving open the possibility for attraction to a particular point y^* in $\mathcal{Z}(x_0)$. Examples of sets $\mathcal{Z}(x_0)$ will be presented in Sect. 1.4. Note that we shall speak of attraction and attractors although the terms are somewhat inaccurate: starting from x'_0 such that $z'_0 = g(x'_0)/\|g(x'_0)\|$ is arbitrarily close to some y in $\mathcal{Z}(x_0)$ yields a limit set $\mathcal{Z}(x'_0)$ for the iterates z'_{2k} close to $\mathcal{Z}(x_0)$ but $\mathcal{Z}(x'_0) \neq \mathcal{Z}(x_0)$. □

Some coordinates $[y]_i$ of a given y in $\mathcal{Z}(x_0)$ may equal zero, $i \in \{1, \dots, d\}$. Define the asymptotic spectrum $\mathcal{S}(x_0, y)$ at $y \in \mathcal{Z}(x_0)$ as the set of eigenvalues λ_i such that $[y]_i \neq 0$ and let $n = n(x_0, y)$ be the number of points in $\mathcal{S}(x_0, y)$. We shall then say that $\mathcal{S}(x_0, y)$ is a n -point asymptotic spectrum. We know from Theorem 3 that if ν_0 is supported on exactly $n_0 = s + 1$ points then $n(x_0, y) = s + 1$ (and $\nu_{2k} = \nu_0$ for all k so that $\mathcal{Z}(x_0)$ is the singleton $\{z_0\}$). In the more general situation where ν_0 is supported on $n_0 \geq s + 1$ points, $n(x_0, y)$ satisfies the following, see Forsythe (1968).

Theorem 8. *Assume that ν_0 is supported on $n_0 > s + 1$ points. Then the number of points $n(x_0, y)$ in the asymptotic spectrum $\mathcal{S}(x_0, y)$ of any $y \in \mathcal{Z}(x_0)$ satisfies*

$$s + 1 \leq n(x_0, y) \leq 2s.$$

Proof. Take any y in $\mathcal{Z}(x_0)$, let n be the number of its nonzero components. To this y we associate a measure ν through $\nu(\lambda_i) = [y]_i^2$, $i = 1, \dots, d$ and we construct a polynomial $Q_s(t)$ from the moments of ν , see (1.6). Applying the transformation (1.11) to ν we get the measure ν' from which we construct the polynomial $Q'_s(t)$. The invariance property $T_z^2(y) = T_z[T_z(y)] = y$, with T_z defined by (1.10), implies $Q_s(\lambda_i)Q'_s(\lambda_i)[y]_i = c[y]_i$, $c > 0$, where $[y]_i$ is any nonzero component of y , $i = 1, \dots, n$. The equation $Q_s(t)Q'_s(t) = c > 0$ can have between 1 and $2s$ solutions in (m, M) . We know already from Theorem 2 that $n(x_0, y) \geq s + 1$. ■

The following Theorem shows that when m and M are support points of ν_0 , then the asymptotic spectrum of any $y \in \mathcal{Z}(x_0)$ also contains m and M .

Theorem 9. *Assume that ν_0 is supported on $n_0 \geq s + 1$ points and that $\nu_0(m) > 0$, $\nu_0(M) > 0$. Then $\liminf_{k \rightarrow \infty} \nu_k(m) > 0$ and $\liminf_{k \rightarrow \infty} \nu_k(M) > 0$.*

Proof. We only consider the case for M , the proof being similar for m .

First notice that from Theorem 1, all roots of the polynomials $Q_s^k(t)$ lie in the open interval (m, M) , so that $\nu_k(M) > 0$ for any k .

Suppose that $\liminf_{k \rightarrow \infty} \nu_k(M) = 0$. Then there exists a subsequence (k_i) such that z_{2k_i} tends to some y in $\mathcal{Z}(x_0)$ and $\lim_{i \rightarrow \infty} \nu_{2k_i}(M) = 0$. To this y we associate a measure ν as in the proof of Theorem 8 and construct a polynomial $Q_s(t)$ from the moments of ν , see (1.6). Since $\nu_{2k_i}(M) \rightarrow 0$, $\nu(M) = 0$. Let λ_j be the largest eigenvalue of A such that $\nu(\lambda_j) > 0$. Then, all zeros of $Q_s(t)$ lie in (m, λ_j) and the same is true for the polynomial $Q'_s(t)$ constructed from the measure $\nu' = T_\nu(\nu)$ obtained by the transformation (1.11). Hence, $Q_s(t)$ and $Q'_s(t)$ are increasing (and positive, see Remark 4) for t between λ_j and M , so that $Q_s(M)Q'_s(M) > Q_s(\lambda_j)Q'_s(\lambda_j)$. This implies by continuity

$$Q_s^{2k_i}(M)Q_s^{2k_i+1}(M) \geq cQ_s^{2k_i}(\lambda_j)Q_s^{2k_i+1}(\lambda_j)$$

for some $c > 1$ and all i larger than some i_0 . Therefore,

$$\frac{[g_{2k_i+2}]_d^2}{[g_{2k_i}]_d^2} \geq c^2 \frac{[g_{2k_i+2}]_j^2}{[g_{2k_i}]_j^2}, \quad i > i_0,$$

see (1.6), and thus

$$\frac{[z_{2k_i+2}]_d^2}{[z_{2k_i+2}]_j^2} \geq c^2 \frac{[z_{2k_i}]_d^2}{[z_{2k_i}]_j^2}, \quad i > i_0.$$

Since $[z_k]_d^2 > 0$ for all k and $[z_{2k_i}]_j^2 \rightarrow \nu(\lambda_j) > 0$ this implies $[z_{2k_i}]_d^2 \rightarrow \infty$ as $i \rightarrow \infty$, which is impossible. Therefore, $\liminf_{k \rightarrow \infty} \nu_k(M) > 0$. ■

The properties above explain the asymptotic behavior of the steepest-descent algorithm in $\mathbb{R}^d$: when $s = 1$ and ν_0 is supported on two points at least, including m and M , then $n(x_0, y) = 2$ for any y in $\mathcal{Z}(x_0)$ and m and M are in the asymptotic spectrum $S(x_0, y)$ of any $y \in \mathcal{Z}(x_0)$. Therefore, $S(x_0, y) = \{m, M\}$ for all $y \in \mathcal{Z}(x_0)$. Since $\mathcal{Z}(x_0)$ is a part of the unit sphere $\mathcal{S}_d$, $\|y\| = 1$ and there is only one degree of freedom. The limiting value r_∞ of r_k then defines the attractor uniquely and $\mathcal{Z}(x_0)$ is a singleton.

In the case where s is even, Forsythe (1968) gives examples of invariant measures ν_0 satisfying $\nu_{k+2} = \nu_k$ and supported on $2q$ points with $s + 1 < 2q \leq 2s$, or supported on $2q + 1$ points with $s + 1 \leq 2q + 1 < 2s$. The nature of the sets $\mathcal{Z}(x_0)$ and $S(x_0, y)$ is investigated more deeply in the next section for the case $s = 2$.

1.4 The optimum 2-gradient algorithm in $\mathbb{R}^d$

In all the section we omit the index k in the moments μ_j^k and matrices $\mathbf{M}_{m,n}^k$. The polynomial $Q_2^k(t)$ defined by (1.6) is then

$$Q_2^k(t) = \frac{\begin{vmatrix} 1 & \mu_1 & 1 \\ \mu_1 & \mu_2 & t \\ \mu_2 & \mu_3 & t^2 \end{vmatrix}}{|\mathbf{M}_{2,1}|} = \frac{\begin{vmatrix} 1 & \mu_1 & 1 \\ \mu_1 & \mu_2 & t \\ \mu_2 & \mu_3 & t^2 \end{vmatrix}}{\begin{vmatrix} \mu_1 & \mu_2 \\ \mu_2 & \mu_3 \end{vmatrix}}$$

and the function $H_k(x)$, see (1.12), is given by

$$\begin{aligned} H_k(x) &= \frac{\begin{vmatrix} 1 & \mu_1 & 1 \\ \mu_1 & \mu_2 & x \\ \mu_2 & \mu_3 & x^2 \end{vmatrix}^2}{\begin{vmatrix} 1 & \mu_1 \\ \mu_1 & \mu_2 \end{vmatrix} \begin{vmatrix} 1 & \mu_1 & \mu_2 \\ \mu_1 & \mu_2 & \mu_3 \\ \mu_2 & \mu_3 & \mu_4 \end{vmatrix}} \\ &= (1 \ x \ x^2) \begin{pmatrix} 1 & \mu_1 & \mu_2 \\ \mu_1 & \mu_2 & \mu_3 \\ \mu_2 & \mu_3 & \mu_4 \end{pmatrix}^{-1} \begin{pmatrix} 1 \\ x \\ x^2 \end{pmatrix} - (1 \ x) \begin{pmatrix} 1 & \mu_1 \\ \mu_1 & \mu_2 \end{pmatrix}^{-1} \begin{pmatrix} 1 \\ x \end{pmatrix}. \end{aligned} \tag{1.25}$$

The monotone sequences (r_k) and (q_k) of Sect. 1.2.3, see (1.14, 1.23), are given by

$$r_k = \frac{\begin{vmatrix} \mu_{-1} & 1 & \mu_1 \\ 1 & \mu_1 & \mu_2 \\ \mu_1 & \mu_2 & \mu_3 \end{vmatrix}}{\mu_{-1} \begin{vmatrix} \mu_1 & \mu_2 \\ \mu_2 & \mu_3 \end{vmatrix}}, \quad q_k = \frac{\begin{vmatrix} 1 & \mu_1 & \mu_2 \\ \mu_1 & \mu_2 & \mu_3 \\ \mu_2 & \mu_3 & \mu_4 \end{vmatrix}}{\begin{vmatrix} 1 & \mu_1 \\ \mu_1 & \mu_2 \end{vmatrix}}. \quad (1.26)$$

When ν_0 is supported on three points, $\nu_{k+2} = \nu_0$ for all k from Theorem 3 and, when ν_0 is supported on less than three points, the algorithm converges in one iteration. In the rest of the section we thus assume that ν_0 is supported on more than three points. Without any loss of generality, we may take d as the number of components in the support of ν_0 and m and M respectively as the minimum and maximum values of these components.

1.4.1 A characterization of limit points through the transformation $\nu_k \rightarrow \nu_{k+1}$

From Theorem 8, the number of components $n(x_0, y)$ of the asymptotic spectrum $\mathcal{S}(x_0, y)$ of any $y \in \mathcal{Z}(x_0)$ satisfies $3 \leq n(x_0, y) \leq 4$ and, from Theorem 9, $\mathcal{S}(x_0, y)$ always contains m and M .

Consider the function

$$\bar{Q}_2^k(t) = \frac{Q_2^k(t) |\mathbf{M}_{2,1}|}{|\mathbf{M}_{1,0}|^{1/2} |\mathbf{M}_{2,0}|^{1/2}} = \frac{\begin{vmatrix} 1 & \mu_1 & 1 \\ \mu_1 & \mu_2 & t \\ \mu_2 & \mu_3 & t^2 \end{vmatrix}}{\begin{vmatrix} 1 & \mu_1 \\ \mu_1 & \mu_2 \end{vmatrix}^{1/2} \begin{vmatrix} 1 & \mu_1 & \mu_2 \\ \mu_1 & \mu_2 & \mu_3 \\ \mu_2 & \mu_3 & \mu_4 \end{vmatrix}^{1/2}}. \quad (1.27)$$

It satisfies $[\bar{Q}_2^k(t)]^2 = H_k(t)$, $z_{k+1} = \bar{Q}_2^k(A)z_k$, see (1.10), and can be considered as a normalized version of $Q_2^k(t)$; $z_{k+2} = z_k$ is equivalent to $\bar{Q}_2^{k+1}(A)\bar{Q}_2^k(A)z_k = z_k$, that is $\bar{Q}_2^k(\lambda_i)\bar{Q}_2^{k+1}(\lambda_i) = 1$ for all i 's such that $[z_k]_i \neq 0$. $\bar{Q}_2^k(t)$ and $\bar{Q}_2^{k+1}(t)$ are second order polynomials in t with two zeros in (m, M) , see Theorem 1, and we write

$$\bar{Q}_2^k(t) = \alpha_k t^2 + \beta_k t + \omega_k, \quad \bar{Q}_2^{k+1}(t) = \alpha_{k+1} t^2 + \beta_{k+1} t + \omega_{k+1}.$$

From the expressions (1.26, 1.27), $\alpha_k = 1/\sqrt{q_k}$ for any k , so that both α_k and α_{k+1} tend to some limit $1/\sqrt{q_\infty}$, see Theorem 6. From (1.27) and (1.7, 1.8),

$$\omega_k^2 = \frac{(g_k, g_k)}{(g_{k+1}, g_{k+1})}, \quad \omega_{k+1}^2 = \frac{(g_{k+1}, g_{k+1})}{(g_{k+2}, g_{k+2})}.$$

Since $\|z_{2k+2} - z_{2k}\|$ tends to zero, see Theorem 7, Theorem 5 implies that $\omega_k \omega_{k+1}$ tend to $1/r_\infty$ as k tends to infinity.

Three-point asymptotic spectra

Assume that x_0 is such that ν_0 has more than three support points. To any y in $\mathcal{Z}(x_0)$ we associate the measure ν defined by $\nu(\lambda_i) = [y]_i^2$, $i = 1, \dots, d$ and denote $Q_2(t)$ the polynomial obtained through (1.6) from the moments of ν . Denote ν' the iterate of ν through T_ν ; to ν' we associate $Q'_2(t)$ and write

$$Q_2(t) = \alpha t^2 + \beta t + \omega, \quad Q'_2(t) = \alpha' t^2 + \beta' t + \omega', \quad (1.28)$$

where the coefficients satisfy $\alpha = \alpha' = 1/\sqrt{q_\infty}$ and $\omega\omega' = 1/r_\infty$.

Suppose that $n(x_0, y) = 3$, with $\mathcal{S}(x_0, y) = \{m, \lambda_j, M\}$, where λ_j is some eigenvalue of A in (m, M) . We thus have

$$Q_2(m)Q'_2(m) = Q_2(\lambda_j)Q'_2(\lambda_j) = Q_2(M)Q'_2(M) = 1,$$

so that $Q_2(t)$ and $Q'_2(t)$ are uniquely defined, in the sense that the number of solutions in (β, β', ω) is finite. (ω is a root of a 6-th degree polynomial equation, with one value for β and β' associated with each root. There is always one solution at least: any measure supported on m, λ_j, M is invariant, so that at least two roots exist for ω . The numerical solution of a series of examples shows that only two roots exist, which renders the product $Q_2(t)Q'_2(t)$ unique, due to the possible permutation between (β, ω) and (β', ω') .) Fig. 1.3 presents a plot of the function $Q_2(t)Q'_2(t)$ when ν gives respectively the weights $1/4, 1/4$ and $1/2$ to the points $m = 1$, $\lambda_j = 4/3$ and $M = 2$ (which gives $Q_2(t)Q'_2(t) = (81t^2 - 249t + 176)(12t^2 - 33t + 22)/8$).

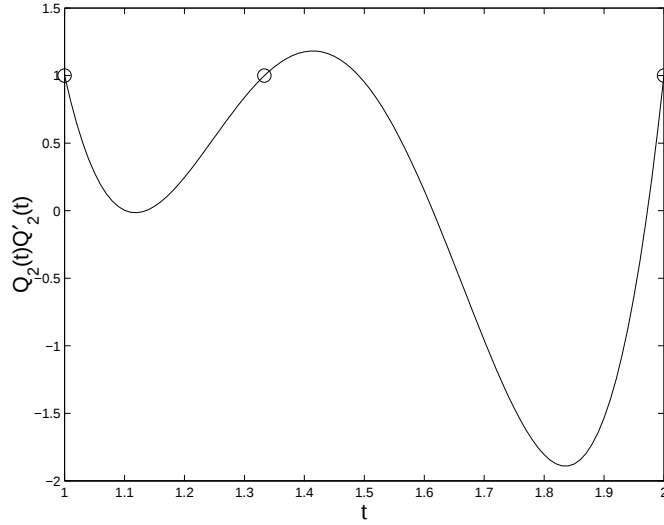


Fig. 1.3. $Q_2(t)Q'_2(t)$ when $\nu(1) = 1/4$, $\nu(4/3) = 1/4$ and $\nu(2) = 1/2$

The orthogonality property (1.13) for $i = 0$ gives $\int Q_2(t) \nu(dt) = 0$, that is,

$$\alpha\mu_2 + \beta\mu_1 + \omega = 0 \quad (1.29)$$

where $\mu_i = p_1 m^i + p_2 \lambda_j^i + (1 - p_1 - p_2)M^i$, $i = 1, 2, \dots$ with $p_1 = [y]_1^2$ and $p_2 = [y]_j^2$. Since ν has three support points, its moments can be expressed as linear combinations of μ_1 and μ_2 through the equations

$$\int t^i (t - m)(t - \lambda_j)(t - M) \nu(dt) = 0, \quad i \in \mathbb{Z}.$$

Using (1.27) and (1.29) μ_1 and μ_2 can thus be determined from two coefficients of $Q_2(t)$ only. After calculation, the Jacobians J_1, J_2 of the transformations $(\mu_1, \mu_2) \rightarrow (\alpha, \beta)$ and $(\mu_1, \mu_2) \rightarrow (\alpha, \omega)$ are found to be equal to

$$J_1 = \frac{Mm + m\lambda_j + \lambda_j M - 2\mu_1(m + \lambda_j + M) + \mu_1^2 + 2\mu_2}{2|\mathbf{M}_{2,0}||\mathbf{M}_{1,0}|},$$

$$J_2 = \frac{\mu_1^2(m + \lambda_j + M) - 2\mu_1\mu_2 - m\lambda_j M}{2|\mathbf{M}_{2,0}||\mathbf{M}_{1,0}|}.$$

The only measure $\tilde{\nu}$ supported on $\{m, \lambda_j, M\}$ for which $J_1 = J_2 = 0$ is given by $\mu_1 = \lambda_j$, $\mu_2 = [\lambda_j(m + \lambda_j + M) - mM]/2$, or equivalently $\tilde{\nu}(m) = (M - \lambda_j)/[2(M - m)]$, $\tilde{\nu}(\lambda_j) = 1/2$. The solution for μ_1, μ_2 (and thus for ν supported at m, λ_j, M) associated with a given polynomial $Q_2(t)$ through the pair r_∞, q_∞ is thus locally unique and there is no continuum for three-point asymptotic spectra, $\mathcal{Z}(x_0)$ is a singleton. As Fig. 1.3 illustrates, the existence of a continuum would require the presence of an eigenvalue λ^* in the spectrum of A to which some weight could be transferred from ν . This is only possible if $Q_2(\lambda^*)Q_2'(\lambda^*) = 1$, so that λ^* is uniquely defined ($\lambda^* = 161/108$ in Fig. 1.3). When this happens, it corresponds to a four-point asymptotic spectrum, a situation considered next.

Four-point asymptotic spectra

We consider the same setup as above with now $n(x_0, y) = 4$ and $\mathcal{S}(x_0, y) = \{m, \lambda_j, \lambda_k, M\}$, where $\lambda_j < \lambda_k$ are two eigenvalues of A in (m, M) . We thus have

$$Q_2(m)Q_2'(m) = Q_2(\lambda_j)Q_2'(\lambda_j) = Q_2(\lambda_k)Q_2'(\lambda_k) = Q_2(M)Q_2'(M) = 1,$$

where $Q_2(t), Q_2'(t)$ are given by (1.28) and satisfy $\alpha = \alpha' = 1/\sqrt{q_\infty}$, $\omega\omega' = 1/r_\infty$. The system of equations in $(\alpha, \beta, \omega, \alpha', \beta', \omega')$ is over-determined, which implies the existence of a relation between r_∞ and q_∞ . As it is the case for three-point asymptotic spectra, to a given value for q_∞ corresponds a unique polynomial $Q_2(t)$ (up to the permutation with $Q_2'(t)$). The measures ν associated with $Q_2(t)$ can be characterized by their three moments μ_1, μ_2, μ_3 ,

which can be obtained from the values of α, β, ω once μ_4 has been expressed as a function of μ_1, μ_2 and μ_3 using

$$\int (t - m)(t - \lambda_j)(t - \lambda_k)(t - M) \nu(dt) = 0.$$

Consider the Jacobian J of the transformation $(\mu_1, \mu_2, \mu_3) \rightarrow (\alpha, \beta, \omega)$. Setting J to zero defines a two-dimensional manifold in the space of moments μ_1, μ_2, μ_3 , or equivalently in the space of weights p_1, p_2, p_3 with $p_1 = \nu(m)$, $p_2 = \nu(\lambda_j)$, $p_3 = \nu(\lambda_k)$ (and $\nu(M) = 1 - p_1 - p_2 - p_3$). By setting some value to the limit r_∞ or q_∞ one removes one degree of freedom and the manifold becomes one-dimensional. Since $[y]_1^2 = p_1$, $[y]_j^2 = p_2$, $[y]_k^2 = p_3$, $[y]_d^2 = 1 - p_1 - p_2 - p_3$, the other components being zero, this also characterizes the limit set $\mathcal{Z}(x_0)$.

Let $x_1 < x_2$ (respectively $x'_1 < x'_2$) denote the two zeros of $\bar{Q}_2(t)$ (respectively $\bar{Q}'_2(t)$). Suppose that $\lambda_j < x_1$. Then $\bar{Q}_2(\lambda_j) > 0$ and $\bar{Q}_2(\lambda_j)\bar{Q}'_2(\lambda_j) = 1$ implies $\bar{Q}'_2(\lambda_j) > 0$ and thus $\lambda_j < x'_1$. But then, $\bar{Q}_2(\lambda_j) < \bar{Q}_2(m)$ and $\bar{Q}'_2(\lambda_j) < \bar{Q}'_2(m)$ which contradicts $\bar{Q}_2(\lambda_j)\bar{Q}'_2(\lambda_j) = \bar{Q}_2(m)\bar{Q}'_2(m) = 1$. Therefore, $\lambda_j > x_1$, and similarly $x_2 > \lambda_k$, that is

$$m < x_1 < \lambda_j < \lambda_k < x_2 < M. \quad (1.30)$$

Denote

$$\begin{aligned} S &= x_1 + x_2 = -\frac{\beta}{\alpha}, \quad P = x_1 x_2 = \frac{\omega}{\alpha}, \\ S_\lambda &= \lambda_j + \lambda_k, \quad P_\lambda = \lambda_j \lambda_k, \\ S_m &= m + M, \quad P_m = m M, \end{aligned}$$

and

$$E = (S_\lambda - S)[S(S - S_m) + (P_m - P)] - (P_\lambda - P)(S - S_m). \quad (1.31)$$

One can easily check by direct calculation that

$$J = \frac{|\mathbf{M}_{1,0}|^2}{2\alpha|\mathbf{M}_{2,0}|^3} E$$

so that the set of limit points y in $\mathcal{Z}(x_0)$ with $n(x_0, y) = 4$ is characterized by $E = 0$. In the next section we investigate the form of the corresponding manifold into more details in the case where the spectrum $\mathcal{S}(x_0, y)$ is symmetric with respect to $c = (m + M)/2$.

Four-point symmetric asymptotic spectra

When the spectrum is symmetric with respect to $c = (m + M)/2$, $S_\lambda = S_m$ so that the equation $E = 0$, with E given by (1.31), becomes

$$(S - S_m)[S(S - S_m) + (P_m - P) + (P_\lambda - P)] = 0.$$

This defines a two-dimensional manifold with two branches: $\mathcal{M}_1$ defined by $S = S_m$ and $\mathcal{M}_2$ defined by $S(S - S_m) + (P_m - P) + (P_\lambda - P) = 0$. The manifolds $\mathcal{M}_1$ and $\mathcal{M}_2$ only depend on the spectrum $\{m, \lambda_j, \lambda_k, M\}$. Note that on the branch $\mathcal{M}_1$ we have $(x_1 + x_2)/2 = (m + M)/2 = c$ so that $\bar{Q}_2(t)$ is symmetric with respect to c . One may also notice that $\bar{Q}_2(t)$ symmetric with respect to c implies that the spectrum $\mathcal{S}(x_0, y)$ is symmetric with respect to c when $E = 0$. Indeed, $\bar{Q}_2(t)$ symmetric implies $S = S_m$, (1.30) then implies $P \neq P_m$, so that $E = 0$ implies $S_\lambda = S = S_m$.

The branch $\mathcal{M}_1$ can be parameterized in P , and the values of r_∞ , q_∞ satisfy

$$r_\infty = \frac{(P_m - P)(P - P_\lambda)}{P(P_m + P_\lambda - P)},$$

$$q_\infty = (P_m - P)(P - P_\lambda).$$

Both r_∞ and q_∞ are maximum for $P = (P_\lambda + P_m)/2$. On $\mathcal{M}_1$, p_1, p_2, p_3 satisfy the following

$$p_1 = \frac{[p_2(M - \lambda_j)(\lambda_j - \lambda_k) - (P - P_m)](P - P_\lambda)}{(P - P_m)(M - m)(M - \lambda_j)}, \quad (1.32)$$

$$p_3 = \frac{P - P_m}{(M - \lambda_j)(\lambda_j - m)} - p_2, \quad (1.33)$$

where the value of P is fixed by r_∞ or q_∞ , and the one dimensional manifold for p_1, p_2, p_3 is a linear segment in $\mathbb{R}^3$.

We parameterize the branch $\mathcal{M}_2$ in S , and obtain

$$r_\infty = \frac{(m + \lambda_j - S)(M + \lambda_j - S)(m + \lambda_k - S)(M + \lambda_k - S)}{[(\lambda_k - S)(\lambda_j - S) + P_m][(\lambda_k - S)(\lambda_j - S) + P_m + 2S_m(S_m - S)]},$$

$$q_\infty = \frac{(m + \lambda_j - S)(M + \lambda_j - S)(m + \lambda_k - S)(M + \lambda_k - S)}{4}.$$

Both r_∞ and q_∞ are maximum for $S = S_m = m + M$. Hence, for each branch the maximum value for r_∞ and q_∞ is obtained on the intersection $\mathcal{M}_1 \cap \mathcal{M}_2$ where $S = S_m$ and $P = (P_m + P_\lambda)/2$. On $\mathcal{M}_2$, p_1, p_2, p_3 satisfy the following

$$p_2 = \frac{\lambda_k + m - S}{\lambda_k - \lambda_j} \left[\frac{\lambda_k + M - S}{2(M - \lambda_j)} - p_1 \frac{M - m}{\lambda_j + M - S} \right],$$

$$p_3 = \frac{\lambda_j + m - S}{\lambda_k - \lambda_j} \left[-\frac{\lambda_j + M - S}{2(\lambda_j - m)} + p_1 \frac{M - m}{\lambda_k + M - S} \right],$$

where now the value of S is fixed by r_∞ or q_∞ ; the one dimensional manifold for p_1, p_2, p_3 is again a linear segment in $\mathbb{R}^3$.

Figs. 1.4 and 1.5 present the two manifolds $\mathcal{M}_1$ and $\mathcal{M}_2$ in the space (p_1, p_2, p_3) when $m = 1$, $\lambda_j = 4/3$, $\lambda_k = 5/3$ and $M = 2$. The line segment

C, C' on Fig. 1.4 corresponds to symmetric distributions for which $p_2 = p_3$. On both figures, when starting the algorithm at x_0 such that the point with coordinates $([z_0]_1^2, [z_0]_2^2, [z_0]_3^2)$ is in A , to the even iterates z_{2k} correspond points that evolve along the line segment A, B , and to the odd iterates z_{2k+1} corresponds the line A', B' . The initial z_0 is chosen such that A is close to the manifold $\mathcal{M}_1$ in Fig. 1.4 and to the manifold $\mathcal{M}_2$ in Fig. 1.5. In both cases the limit set $\mathcal{Z}(x_0)$ is a singleton $\{y\}$ with $n(x_0, y) = 3$: $[y]_3 = 0$ in Fig. 1.4 and $[y]_2 = 0$ in Fig. 1.5.

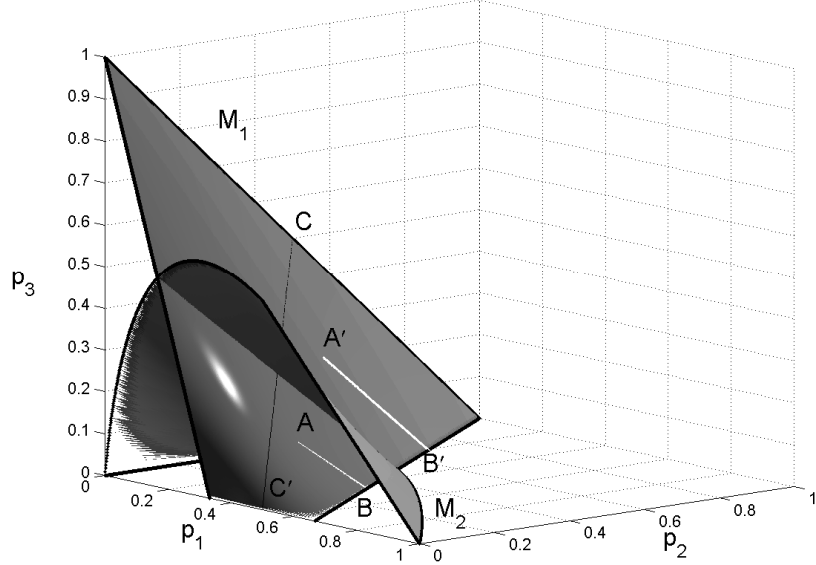


Fig. 1.4. Sequence of iterates close to the manifold $\mathcal{M}_1$: A, B for z_{2k} , A', B' for z_{2k+1} ; the line segment C, C' corresponds to symmetric distributions on $\mathcal{M}_1$ ($p_2 = p_3$)

1.4.2 A characterization of limit points through monotone sequences

Assume that x_0 is such that ν_0 has more than three support points. The limit point for the orbit (z_k) are such that $r_{k+1} = r_k$. To any y in $\mathcal{Z}(x_0)$ we associate the measure ν defined by $\nu(\lambda_i) = [y]_i^2$, $i = 1, \dots, d$ and then $\nu' = T_\nu(\nu)$ with T_ν given by (1.11); with ν and ν' we associate respectively $r(\nu)$ and $r(\nu')$, which are defined from their moments by (1.26). Then, $y \in \mathcal{Z}(x_0)$ implies $T_\nu[T_\nu(\nu)] = \nu$ and

$$\Delta(\nu) = r(\nu') - r(\nu) = 0.$$

We thus investigate the nature of the sets of measures satisfying $\Delta(\nu) = 0$.

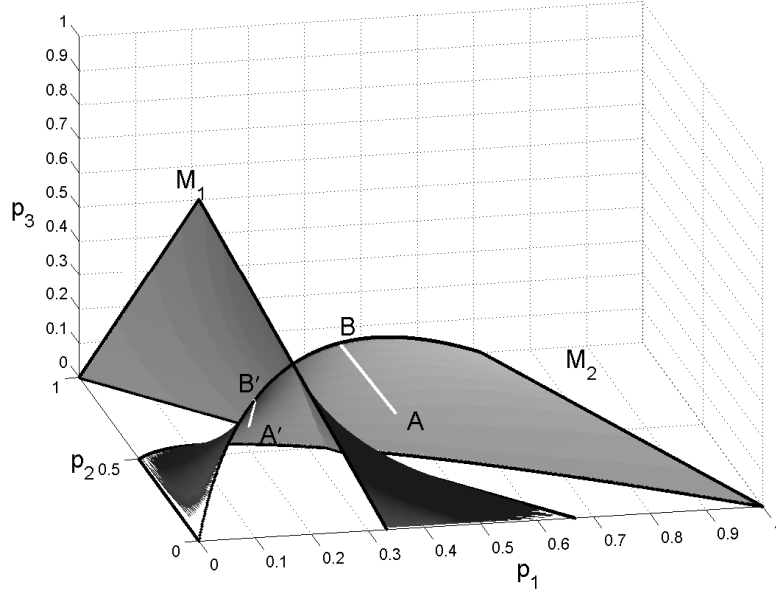


Fig. 1.5. Sequence of iterates close to the manifold $\mathcal{M}_2$: A, B for z_{2k} , A', B' for z_{2k+1}

Distributions that are symmetric with respect to μ_1

When ν is *symmetric* with respect to μ_1 direct calculation gives

$$\Delta(\nu) = \frac{|\mathbf{M}_{4,-1}| |\mathbf{M}_{2,1}|}{\mu_1 |\mathbf{M}_{3,0}| |\mathbf{M}_{2,0}|},$$

which is zero for any four-point distribution. Any four-point distribution ν that is symmetric with respect to μ_1 is thus invariant in two iterations, that is $T_\nu[T_\nu(\nu)] = \nu$. The expression above for $\Delta(\nu)$ is not valid when ν is not symmetric with respect to μ_1 , a situation considered below.

General situation

Direct (but lengthy) calculations give

$$\Delta(\nu) = \frac{N}{D}$$

with

$$D = \mu_{-1} |\mathbf{M}_{2,1}| |\mathbf{M}_{2,-1}| \left[|\mathbf{M}_{1,0}|^2 (|\mathbf{M}_{1,-1}| |\mathbf{M}_{4,1}| - \mu_1 |\mathbf{M}_{4,-1}|) \right. \\ \left. + |\mathbf{M}_{2,0}|^2 (\mu_{-1} |\mathbf{M}_{3,1}| - |\mathbf{M}_{3,-1}|) - 2 |\mathbf{M}_{1,0}| |\mathbf{M}_{2,0}| |\mathbf{M}_{3,0}| \right] \quad (1.34)$$

$$N = |\mathbf{M}_{4,-1}| |\mathbf{M}_{1,0}|^2 \left[\mu_1 (\mu_{-1} |\mathbf{M}_{2,1}| - |\mathbf{M}_{2,-1}|)^2 - \mu_{-1} |\mathbf{M}_{2,1}|^2 \right] \\ + a^2 |\mathbf{M}_{4,1}| + b^2 |\mathbf{M}_{3,-1}| - 2ab |\mathbf{M}_{3,0}| \quad (1.35)$$

with $a = |\mathbf{M}_{1,0}||\mathbf{M}_{2,-1}|$ and $b = (\mu_{-1}|\mathbf{M}_{2,1}| - |\mathbf{M}_{2,-1}|)|\mathbf{M}_{2,0}|$. The determinants that are involved satisfy some special identities

$$\begin{aligned} |\mathbf{M}_{1,0}|^2 &= |\mathbf{M}_{1,-1}||\mathbf{M}_{2,1}| - \mu_1|\mathbf{M}_{2,-1}| \\ |\mathbf{M}_{2,0}|^2 &= |\mathbf{M}_{2,-1}||\mathbf{M}_{3,1}| - |\mathbf{M}_{2,1}||\mathbf{M}_{3,-1}| \\ |\mathbf{M}_{3,0}|^2 &= |\mathbf{M}_{3,-1}||\mathbf{M}_{4,1}| - |\mathbf{M}_{3,1}||\mathbf{M}_{4,-1}| \end{aligned} \quad (1.36)$$

and

$$\begin{aligned} \mu_{-1}|\mathbf{M}_{1,1}| - |\mathbf{M}_{1,-1}| &= 1, \quad |\mathbf{M}_{1,-1}||\mathbf{M}_{1,1}| - \mu_1|\mathbf{M}_{1,-1}| = 0 \\ (|\mathbf{M}_{1,-1}||\mathbf{M}_{2,1}| - \mu_1|\mathbf{M}_{2,-1}|)(\mu_{-1}|\mathbf{M}_{1,1}| - |\mathbf{M}_{1,-1}|) &= |\mathbf{M}_{1,0}|^2 \\ (|\mathbf{M}_{1,-1}||\mathbf{M}_{3,1}| - \mu_1|\mathbf{M}_{3,-1}|)(\mu_{-1}|\mathbf{M}_{2,1}| - |\mathbf{M}_{2,-1}|) &= |\mathbf{M}_{2,0}|^2 \\ &\quad + (|\mathbf{M}_{1,-1}||\mathbf{M}_{2,1}| - \mu_1|\mathbf{M}_{2,-1}|)(\mu_{-1}|\mathbf{M}_{3,1}| - |\mathbf{M}_{3,-1}|) \\ (|\mathbf{M}_{1,-1}||\mathbf{M}_{4,1}| - \mu_1|\mathbf{M}_{4,-1}|)(\mu_{-1}|\mathbf{M}_{3,1}| - |\mathbf{M}_{3,-1}|) &= |\mathbf{M}_{3,0}|^2 \\ &\quad + (|\mathbf{M}_{1,-1}||\mathbf{M}_{3,1}| - \mu_1|\mathbf{M}_{3,-1}|)(\mu_{-1}|\mathbf{M}_{4,1}| - |\mathbf{M}_{4,-1}|) \end{aligned}$$

where all the terms inside the brackets are non-negative. Using these identities we obtain

$$\begin{aligned} D &= \mu_{-1}|\mathbf{M}_{2,1}||\mathbf{M}_{2,-1}| \left\{ \left[|\mathbf{M}_{1,0}|(|\mathbf{M}_{1,-1}||\mathbf{M}_{4,1}| - \mu_1|\mathbf{M}_{4,-1}|)^{1/2} \right. \right. \\ &\quad \left. \left. + |\mathbf{M}_{2,0}|(\mu_{-1}|\mathbf{M}_{3,1}| - |\mathbf{M}_{3,-1}|)^{1/2} \right]^2 \right. \\ &\quad \left. + 2|\mathbf{M}_{1,0}||\mathbf{M}_{2,0}| \right. \\ &\quad \left. \times \frac{(|\mathbf{M}_{1,-1}||\mathbf{M}_{3,1}| - \mu_1|\mathbf{M}_{3,-1}|)(\mu_{-1}|\mathbf{M}_{4,1}| - |\mathbf{M}_{4,-1}|)}{(|\mathbf{M}_{1,-1}||\mathbf{M}_{4,1}| - \mu_1|\mathbf{M}_{4,-1}|)^{1/2}(\mu_{-1}|\mathbf{M}_{3,1}| - |\mathbf{M}_{3,-1}|)^{1/2} + |\mathbf{M}_{3,0}|} \right\} \end{aligned}$$

and thus $D > 0$ when ν has three support points or more.

We also get

$$a^2|\mathbf{M}_{4,1}| + b^2|\mathbf{M}_{3,-1}| - 2ab|\mathbf{M}_{3,0}| \geq |\mathbf{M}_{1,0}|^2|\mathbf{M}_{2,1}|^2 \frac{|\mathbf{M}_{3,1}||\mathbf{M}_{4,-1}|}{|\mathbf{M}_{3,-1}|}$$

which gives

$$\begin{aligned} N &\geq \frac{|\mathbf{M}_{1,0}|^2|\mathbf{M}_{4,-1}|}{|\mathbf{M}_{3,-1}|} [|\mathbf{M}_{1,0}|^2(\mu_{-1}|\mathbf{M}_{2,1}| - |\mathbf{M}_{2,-1}|)|\mathbf{M}_{3,-1}| + |\mathbf{M}_{2,0}|^3|\mathbf{M}_{2,-1}|] \\ &\geq 0. \end{aligned}$$

Now, $y \in \mathcal{Z}(x_0)$ implies $N = 0$. Since $|\mathbf{M}_{2,-1}| > 0$ when ν has three support points or more, $N = 0$ implies $|\mathbf{M}_{4,-1}| = 0$, that is, ν has three or four support points only, and we recover the result of Theorem 8.

Setting $|\mathbf{M}_{4,-1}| = 0$ in (1.35) we obtain that $y \in \mathcal{Z}(x_0)$ implies

$$a^2|\mathbf{M}_{4,1}| - 2ab|\mathbf{M}_{3,0}| + b^2|\mathbf{M}_{3,-1}| = 0. \quad (1.37)$$

The condition (1.37) is satisfied for any three-point distribution. For a four point-distribution it is equivalent to b/a being the (double) root of the following quadratic equation in t , $|\mathbf{M}_{4,1}| - 2t|\mathbf{M}_{3,0}| + |\mathbf{M}_{3,-1}|t^2 = 0$. This condition can be written as

$$\delta(\nu) = (\mu_{-1}|\mathbf{M}_{2,1}| - |\mathbf{M}_{2,-1}|)|\mathbf{M}_{2,0}||\mathbf{M}_{3,-1}| - |\mathbf{M}_{1,0}||\mathbf{M}_{3,0}||\mathbf{M}_{2,-1}| = 0. \quad (1.38)$$

(One may notice that when ν is symmetric with respect to μ_1 , $\delta(\nu)$ becomes $\delta(\nu) = [|\mathbf{M}_{1,0}||\mathbf{M}_{2,0}|^2/|\mathbf{M}_{3,0}|] |\mathbf{M}_{4,-1}|$ which is equal to zero for a four-point distribution. We thus recover the result of Sect. 1.4.2.)

To summarize, $y \in \mathcal{Z}(x_0)$ implies that ν is supported on three or four points and satisfies (1.38). In the case of a four-point distribution supported on $m, \lambda_j, \lambda_k, M$, after expressing μ_{-1}, μ_4, μ_5 and μ_6 as functions of μ_1, μ_2, μ_3 though

$$\int t^i(t-m)(t-\lambda_j)(t-\lambda_k)(t-M)\nu(dt) = 0, \quad i \in \mathbb{Z},$$

we obtain $\delta(\nu) = KJ$ with

$$K = 2 \frac{\alpha |\mathbf{M}_{2,0}|^3 |\mathbf{M}_{1,0}|}{m \lambda_j \lambda_k M} p_1 p_2 p_3 (1 - p_1 - p_2 - p_3) \\ \times (\lambda_k - \lambda_j)^2 (\lambda_k - m)^2 (M - \lambda_k)^2 (\lambda_j - m)^2 (M - \lambda_j)^2 (M - m)^2 > 0,$$

where p_1, p_2, p_3, α and J are defined as in Sect. 1.4.1. Since $K > 0$, the attractor is equivalently defined by $J = 0$, which is precisely the situation considered in Sect. 1.4.1.

1.4.3 Stability

Not all three or four-point asymptotic spectra considered in Sect. 1.4.1 correspond to stable attractors. Although ν associated with some vector y on the unit sphere $\mathcal{S}_d$ may be invariant by two iterations of (1.11), a measure ν_k arbitrarily close to ν (that is, associated with a renormalized gradient z_k close to y) may lead to an iterate ν_{k+2} far from ν_k . The situation can be explained from the example of a three-point distribution considered in Fig. 1.3 of Sect. 1.4.1. The measure ν is invariant in two iteration of T_ν given by (1.11). Take a measure $\nu_k = (1 - \kappa)\nu + \kappa\nu'$ where $0 < \kappa < 1$ and ν' is a measure on (m, M) that puts some positive weight to some point λ^* in the intervals $(4/3, 161/108)$ or $(1.7039, 1.9337)$. Then, for κ small enough the function $Q_2(t)Q'_2(t)$ obtained for ν' is similar to that plotted for ν on Fig. 1.3, and the weight of λ^* will increase in two iterations since $|Q_2(\lambda^*)Q'_2(\lambda^*)| > 1$. The invariant measure ν is thus an unstable fixed point for T_ν^2 when the spectrum of A contains some eigenvalues in $(4/3, 161/108) \cup (1.7039, 1.9337)$. The analysis is thus similar to that in (Pronzato *et al.*, 2001, 2006) for the steepest-descent algorithm ($s = 1$), even if the precise derivation of stability regions (in terms of the weights that two-step invariant measures give to their three or four support points), for a given spectrum for A , is much more difficult for $s = 2$.

1.4.4 Open questions and extension to $s > 2$

For $s = 2$, as shown in Sect. 1.4.1, if x_0 is such that z_0 is exactly on one of the manifolds $\mathcal{M}_1$ or $\mathcal{M}_2$, then by construction $z_{2k} = z_0$ for all k . Although it might be possible that choosing z_0 close enough to $\mathcal{M}_1$ (respectively $\mathcal{M}_2$) would force z_{2k} to converge to a limit point before reaching the plane $p_3 = 0$ (respectively $p_2 = 0$), the monotonicity of the trajectory along the line segment A, B indicates that there is no continuum and the limit set $\mathcal{Z}(x_0)$ is a single point. This was conjectured by Forsythe (1968) and is still an open question. We conjecture that additionally for almost all initial points the trajectory always attracts to a three-point spectrum (though the attraction may take a large number of iterations when z_k is very close to $\mathcal{M}_1 \cup \mathcal{M}_2$ for some k). One might think of using the asymptotic equivalence of rates of convergence, as stated in Theorem 5, to prove this conjecture. However, numerical calculations show that for any point on the manifold $\mathcal{M}_1$ defined in Sect. 1.4.1, the product of the rates at two consecutive iterations satisfies $R_k(W)R_{k+1}(W) = r_\infty^2$ for any positive-definite matrix W (so that $R_k(W)R_{k+1}(W)$ does not depend on p_2 on the linear segment defined by r_∞ on $\mathcal{M}_1$, see (1.32, 1.33), and all points on this segment can thus be considered as asymptotically equivalent).

Extending the approach of Sect. 1.4.2 for the characterization of limit points to the case $s > 2$ seems rather difficult, and the method used in Sect. 1.4.1 is more promising. A function $\bar{Q}_s^k(t)$ can be defined similarly to (1.27), leading to two s -degree polynomials $Q_s(t)$, $Q'_s(t)$, see (1.28), each of them having s roots in (m, M) . Let α and α' denote the coefficients of terms of highest degree in $Q_s(t)$ and $Q'_s(t)$ respectively, then $\alpha = \alpha' = 1/\sqrt{q_\infty}$. Also, let ω and ω' denote the constant terms in $Q_s(t)$ and $Q'_s(t)$; Theorem 5 implies that $\omega\omega' = 1/r_\infty$. Let n be the number of components in the asymptotic spectrum $\mathcal{S}(x_0, y)$, with $s+1 \leq n \leq 2s$ from Theorem 8; we thus have $2(s+1)$ coefficients to determine, with $n+3$ equations:

$$\alpha = \alpha' = 1/\sqrt{q_\infty}, \quad \omega\omega' = 1/r_\infty \quad \text{and} \quad Q_s(\lambda_i)Q'_s(\lambda_i) = 1, \quad i = 1, \dots, n$$

where the λ_i 's are the eigenvalues of A in $\mathcal{S}(x_0, y)$ (including m and M , see Theorem 9). We can then demonstrate that the functions $Q_s(t)$ and $Q'_s(t)$ are uniquely defined when $n = 2s$ and $n = 2s-1$ (which are the only possible cases when $s = 2$). When $n = 2s$, the system is over-determined, as it is the case in Sect. 1.4.1 for four-point asymptotic spectra when $s = 2$. The limit set $\mathcal{Z}(x_0)$ corresponds to measures ν for which the weights p_i , $i = 1, \dots, n-1$ belong to a $n-2$ -dimensional manifold. By setting some value to r_∞ or q_∞ the manifold becomes $(n-3)$ -dimensional. When $n = 2s-1$, ν can be characterized by its $2s-2$ first moments, which cannot be determined uniquely when $s > 2$. Therefore, although $Q_s(t)$ and $Q'_s(t)$ are always uniquely defined for $n = 2s$ and $n = 2s-1$, the possibility of a continuum for $\mathcal{Z}(x_0)$ still exists. The situation is even more complex when $s+1 \leq n \leq 2s-2$ and $s > 2$.

1.5 Switching algorithms

The bound $N_2^* = (R_2^*)^{1/2}$ on the convergence rate of the optimum 2-gradient algorithm, see (1.21), is smaller than the bound $N_1^* = R_1^*$, which indicates a slower convergence for the latter. Let ϵ denote a required precision on the squared norm of the gradient g_k , then the number of gradient evaluations needed to reach the precision ϵ is bounded by $\log(\epsilon)/\log(N_2^*)$ for the optimum 2-gradient algorithm and by $\log(\epsilon)/\log(N_1^*)$ for the steepest-descent algorithm. To compare the number of gradient evaluations for the two algorithms we thus compute the ratio $L_{1/2}^* = \log(N_1^*)/\log(N_2^*)$. The evolution of $L_{1/2}^*$ as a function of ϱ is presented in Fig. 1.6 in solid line. The improvement of the optimum 2-gradient over steepest descent is small for small ϱ but $L_{1/2}^*$ tends to $1/2$ as ϱ tends to infinity. (More generally, the ratio $L_{1/s}^* = \log(N_1^*)/\log(N_s^*)$ tends to $1/s$ as $\varrho \rightarrow \infty$.) The ratio $L_{1/\infty}^* = \log(N_1^*)/\log(N_\infty^*)$, where N_∞^* is defined in (1.22), is also presented in Fig. 1.6. It tends to zero as $1/\sqrt{\varrho}$ when ϱ tends to ∞ .

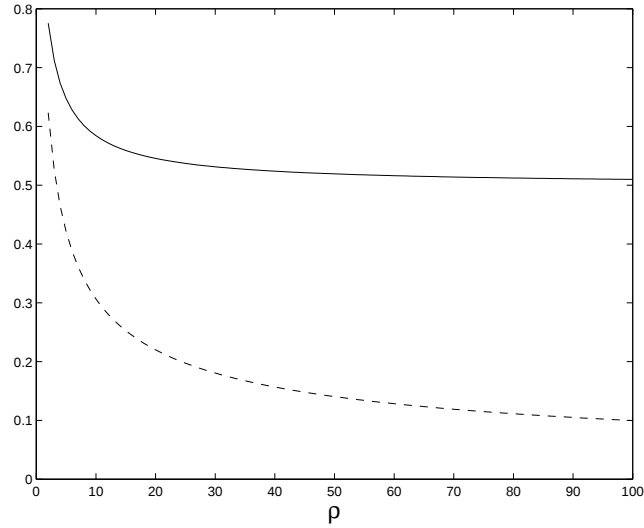


Fig. 1.6. Ratios $L_{1/2}^*$ (solid line) and $L_{1/\infty}^*$ (dashed line) as functions of ϱ

The slow convergence of steepest descent, or, more generally, of the optimum s -gradient algorithm, is partly due to the existence of a measure ν_s^* associated with a large value for the rate of convergence R_s^* , see Theorem 4, but mainly to the fact that ν_s^* is supported on $s+1$ points and is thus invariant in two steps of the algorithm, that is, $T_\nu[T_\nu(\nu_s^*)] = \nu_s^*$, see Theorem 3 (in fact, ν_s^* is even invariant in one step, i.e. $T_\nu(\nu_s^*) = \nu_s^*$).

Switching between algorithms may then be seen as an attractive option to destroy the stability of this worst-case behavior. The rest of the paper is devoted to the analysis of the performance obtained when the two algorithms for $s = 1$ and $s = 2$ are combined, in the sense that the resulting algorithm switches between steepest descent and optimum 2-gradient. Notice that no measure exists that is invariant in two iterations for both the steepest descent and the optimum 2-gradient algorithms (this is no longer true for larger values of s , for instance, a symmetric four-point distribution is invariant in two iterations for $s = 2$, see Sect. 1.4.2, and $s = 3$, see Theorem 3).

1.5.1 Superlinear convergence in $\mathbb{R}^3$

We suppose that $d = 3$, A is already diagonalized with eigenvalues $m < \lambda < M = \varrho m$, and that x_0 is such that ν_0 puts a positive weight on each of them. We denote $p_k = \nu_k(m)$, $t_k = \nu_k(\lambda)$ (so that $\nu_k(M) = 1 - p_k - t_k$). Consider the following algorithm.

Algorithm A

- Step 0: Fix n , the total number iterations allowed, choose ϵ , a small positive number; go to Step 1.
- Step 1 ($s = 1$): Use steepest descent from $k = 0$ to k^* , the first value of k such that $q_{k+1} - q_k < \epsilon$, where q_k is given by (1.23) with $s = 1$; go to Step 2.
- Step 2 ($s = 2$): Use the optimum 2-gradient algorithm for iterations k^* to n .

Notice that since q_k is non-decreasing and bounded, see Theorem 6, switching will always occur for some finite k^* . Since for the steepest-descent algorithm the measure ν_k converges to a two-point measure supported at m and M , after switching the optimum 2-gradient algorithm is in a position where its convergence is fast provided that k^* is large enough for t_{k^*} to be small (it would converge in one iteration if the measure were supported on exactly two points, i.e. if t_{k^*} were zero). Moreover, the 3-point measure ν_{k^*} is invariant in two iterations of optimum 2-gradient, that is, $\nu_{k^*+2j} = \nu_{k^*}$ for any j , so that this fast convergence is maintained from iterations k^* to n . This can be formulated more precisely as follows.

Theorem 10. *When $\epsilon = \epsilon^*(n) = C \log(n)/n$ in Algorithm A, with C an arbitrary positive constant, the global rate of convergence*

$$R_n(x_0) = \left[\frac{f(x_n)}{f(x_0)} \right]^{1/n} = \left(\prod_{k=0}^{n-1} r_k \right)^{1/n}, \quad (1.39)$$

with r_k given by (1.14), satisfies

$$\limsup_{n \rightarrow \infty} \frac{\log[R_n(x_0)]}{\log(n)} < -\frac{1}{2}.$$

Proof. We know from Theorem 6 that the sequence (q_k) is non-decreasing and bounded. For steepest descent $s = 1$ and the bound (1.24) is $q_1^* = (M - m)^2/4$. Since $q_0 > 0$, it implies $k^* < (M - m)^2/(4\epsilon)$. Also, direct calculation gives $q_{k+1} - q_k = |\mathbf{M}_{2,0}^k|/q_k^2 = |\mathbf{M}_{2,0}^k|/|\mathbf{M}_{1,0}^k|^2$ and, using (1.36),

$$q_{k+1} - q_k = \frac{|\mathbf{M}_{2,0}^k|}{|\mathbf{M}_{2,1}^k|} \frac{1}{|\mathbf{M}_{1,-1}^k| - \mu_1 \frac{|\mathbf{M}_{2,-1}^k|}{|\mathbf{M}_{2,1}^k|}} > \frac{|\mathbf{M}_{2,0}^k|}{|\mathbf{M}_{2,1}^k|} \frac{1}{|\mathbf{M}_{1,-1}^k|},$$

so that $q_{k^*+1} - q_{k^*} < \epsilon$ implies

$$\frac{|\mathbf{M}_{2,0}^{k^*}|}{|\mathbf{M}_{2,1}^{k^*}|} < \epsilon |\mathbf{M}_{1,-1}^{k^*}|. \quad (1.40)$$

For steepest descent, $r_k = |\mathbf{M}_{1,-1}^k|/(\mu_{-1}^k |\mathbf{M}_{1,1}^k|) = 1 - 1/(\mu_1^k \mu_{-1}^k) < R_1^*$, so that $|\mathbf{M}_{1,-1}^{k^*}| = \mu_1^{k^*} \mu_{-1}^{k^*} - 1 < R_1^*/(1 - R_1^*) = (M - m)^2/(4mM)$ and (1.40) gives

$$\frac{|\mathbf{M}_{2,0}^{k^*}|}{|\mathbf{M}_{2,1}^{k^*}|} < \epsilon \frac{(M - m)^2}{4mM}. \quad (1.41)$$

The first iteration of the optimum 2-gradient algorithm has the rate

$$r_{k^*} = \frac{f(x_{k^*+1})}{f(x_{k^*})} = \frac{|\mathbf{M}_{2,-1}^{k^*}|}{\mu_{-1}^{k^*} |\mathbf{M}_{2,1}^{k^*}|} = \frac{|\mathbf{M}_{2,0}^{k^*}|}{mM\lambda\mu_{-1}^{k^*} |\mathbf{M}_{2,1}^{k^*}|},$$

and using $\mu_{-1}^{k^*} > 1/M$ and (1.41) we get $r_{k^*} < B\epsilon$ with $B = (M - m)^2/[4M\lambda m^2]$. Since $d = 3$, $\nu_{k^*+2j} = \nu_{k^*}$ and $r_{k^*+2j} = r_{k^*}$ for $j = 1, 2, 3 \dots$

Now, for each iteration of steepest descent we bound r_k by R_1^* ; for the optimum 2-gradient we use $r_k < B\epsilon$ for $k = k^* + 2j$ and $r_k < R_2^*$ for $k = k^* + 2j + 1$, $j = 1, 2, 3 \dots$. We have

$$\log[R_n(x_0)] = \frac{1}{n} \left\{ \log \left[\prod_{k=0}^{k^*-1} r_k \right] + \log \left[\prod_{k=k^*}^{n-1} r_k \right] \right\}.$$

Since $R_2^* < R_1^*$, $B\epsilon < R_2^*$ for ϵ small enough, and $k^* < \bar{k} = (M - m)^2/(4\epsilon)$ we can write

$$\log[R_n(x_0)] < L_n(\epsilon) = \frac{1}{n} \left\{ \bar{k} \log R_1^* + (n - \bar{k}) \frac{1}{2} [\log(B\epsilon) + \log(R_2^*)] \right\}.$$

Taking $\epsilon = C \log(n)/n$ and letting n tend to infinity, we obtain

$$\lim_{n \rightarrow \infty} L_n / \log(n) = -1/2.$$

■

Algorithm A requires to fix the number n of iterations *a priori* and to choose ϵ as a function of n . The next algorithm does not require any such prior choice and uses alternatively a fixed number of iterations of steepest descent and optimum 2-gradient.

Algorithm B

- Step 1 ($s = 1$): Use steepest descent for $m_1 \geq 1$ iterations; go to Step 2.
 Step 2 ($s = 2$): Use the optimum 2-gradient algorithm for $2m_2$ iterations, $m_2 \geq 1$; return to Step 1.

Its performance satisfies the following.

Theorem 11. *For any choice of m_1 and m_2 in Algorithm B, the global rate (1.39) satisfies $R_n(x_0) \rightarrow 0$ as the number n of iterations tends to infinity.*

Proof. Denote $k_j = (j-1)(m_1 + 2m_2) + m_1$, $j = 1, 2, \dots$, the iteration number for the j -th switching from steepest descent to optimum 2-gradient. Notice that $\nu_{k_j-1} = \nu_{k_j+2m_2-1}$ since any three-point measure is invariant in two steps of the optimum 2-gradient algorithm. The repeated use of Step 2, with $2m_2$ iterations each time, has thus no influence on the behavior of the steepest-descent iterations used in Step 1. Therefore, $\epsilon_j = q_{k_j} - q_{k_j-1}$ tends to zero as j increases. Using the same arguments and the same notation as in the proof of Theorem 10 we thus get $r_{k_j} < B\epsilon_j$ for the first of the $2m_2$ iterations of the optimum 2-gradient algorithm, with $B = (M - m)^2/[4M\lambda m^2]$. For large n , we write

$$j = \left\lfloor \frac{n}{m_1 + 2m_2} \right\rfloor$$

$n' = n - j(m_1 + 2m_2) < (m_1 + 2m_2)$. For the last n' iterations we bound r_k by R_1^* ; for steepest-descent iterations we use $r_k < R_1^*$; at the j -th call of Step 2 we use $r_{k_j} < B\epsilon_j$ for the first iteration of optimum 2-gradient and $r_k < R_2^*$ for the subsequent ones. This yields the bound

$$\begin{aligned} \log[R_n(x_0)] &< \frac{1}{j(m_1 + 2m_2) + n'} \\ &\times \left\{ n' \log(R_1^*) + jm_1 \log(R_1^*) + j(2m_2 - 1) \log(R_2^*) + \sum_{i=1}^j \log(B\epsilon_i) \right\}. \end{aligned}$$

Finally, we use the concavity of the logarithm and write

$$\frac{\sum_{i=1}^j \log(B\epsilon_i)}{j} < \log \left(B \frac{\sum_{i=1}^j \epsilon_i}{j} \right) = \log \left(B \frac{q_{k_j} - q_{m_1-1}}{j} \right) < \log \left(B \frac{q_1^*}{j} \right)$$

with $q_1^* = (M - m)^2/4$, see (1.24). Therefore, $\log[R_n(x_0)] \rightarrow -\infty$ and $R_n(x_0) \rightarrow 0$ as $n \rightarrow \infty$. $\blacksquare$

Fig. 1.7 presents a typical evolution of $\log(r_k)$ and $\log(R_k)$, with R_k the global rate of convergence (1.39), as functions of the iteration number k in Algorithm B when $m_1 = m_2 = 1$ (z_0 is a random point on the unit sphere $\mathcal{S}_3$ and A has the eigenvalues 1, $3/2$ and 2). The rate of convergence r_k of the steepest-descent iterations is slightly increasing, but this is compensated by the regular decrease of the rate for the pairs of optimum 2-gradient iterations and the global rate R_k decreases towards zero.

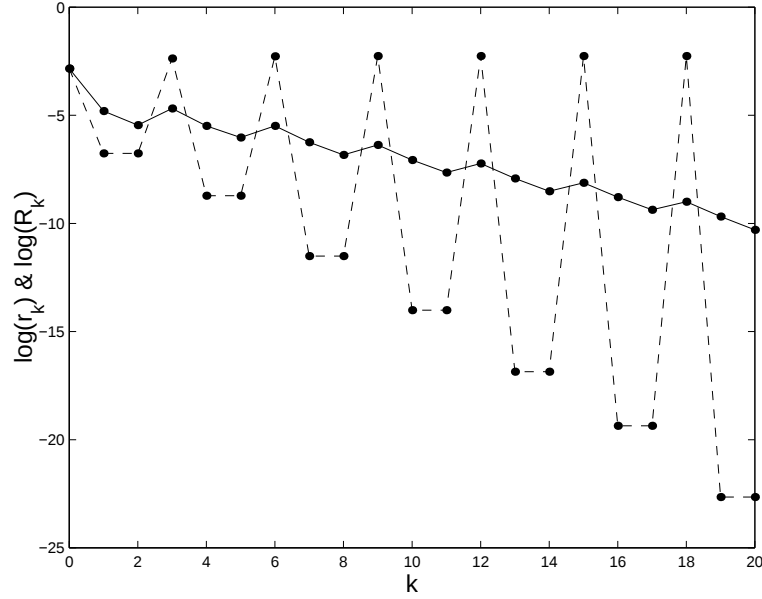


Fig. 1.7. Typical evolution of the logarithms of the rate of convergence r_k , in dashed line, and of the global rate R_k defined by (1.39), in solid line, as functions of k in Algorithm B

Remark 7. In Theorems 10 and 11, n counts the number of iterations. Let n now denote the number of gradient evaluations (remember that one iteration of the optimum 2-gradient algorithm corresponds to two steps of the conjugate gradient method, see Remark 1, and thus requires two gradient evaluations), and define

$$N_n(x_0) = \left[\frac{f[x(n)]}{f(x_0)} \right]^{1/n} \quad (1.42)$$

with $x(n)$ the value of x generated by the algorithm after n gradient evaluations. Following the same lines as in the proof of Theorem 10 we get for Algorithm A

$$\begin{aligned} \log[N_n(x_0)] &= \frac{1}{n} \left\{ \log \left[\prod_{k=0}^{k^*-1} r_k \right] + \log \left[\prod_{k=k^*}^{n-1} \sqrt{r_k} \right] \right\} \\ &< L'_n(\epsilon) = \frac{1}{n} \left\{ \bar{k} \log R_1^* + (n - \bar{k}) \frac{1}{4} [\log(B\epsilon) + \log(R_2^*)] \right\}. \end{aligned}$$

Taking $\epsilon = C \log(n)/n$ and letting n tend to infinity, we then obtain $\lim_{n \rightarrow \infty} L'_n / \log(n) = -1/4$.

Similarly to Theorem 11 we also have $N_n(x_0) \rightarrow 0$ as $n \rightarrow \infty$ in Algorithm B. $\square$

The importance of the results in Theorems 10 and 11 should not be overemphasized. After all, the optimum 3-gradient converges in one iteration in $\mathbb{R}^3$! It shows, however, that an algorithm with fast convergence can be obtained from the combination of two algorithms with rather poor performance, which opens a promising route for further developments. As a first step in this direction, next section shows that the combination used in Algorithm B still has a good performance in dimensions $d > 3$, with a behavior totally different from the regular one observed in $\mathbb{R}^3$.

1.5.2 Switching algorithms in $\mathbb{R}^d$, $d > 3$

We suppose that A is diagonalized with eigenvalues $m = \lambda_1 < \lambda_2 < \dots < \lambda_d = M = \varrho m$ and that x_0 is such that ν_0 has $n_0 > 3$ support points. The behavior of Algorithm B is then totally different from the case $d = 3$, where convergence is superlinear. Numerical simulations (see Sect. 1.5.2) indicate that the convergence is then only linear, although faster than for the optimum 2-gradient algorithm for suitable choices of m_1 and m_2 . A simple interpretation is as follows. Steepest descent tends to force ν_k to be supported on m and M only. If m_1 is large, when switching to optimum 2-gradient, say at iteration $k_j = j(m_1 + 2m_2) + m_1$, the first iteration has then a very small rate r_{k_j} . Contrary to the case $d = 3$, $\nu_{k_j+2} \neq \nu_{k_j}$, so that the rate r_k quickly deteriorates as k increases and ν_k converges to a measure with three or four support points. However, when switching back to steepest descent at iteration $(j+1)(m_1 + 2m_2) = k_j + 2m_2$, the rate is much better than the bound R_1^* since $\nu_{k_j+2m_2}$ is far from a two-point measure. This alternation of phases where ν_k converges towards a two-point measure and then to a three or four-point measure renders the behavior of the algorithm hardly predictable (Sect. 1.5.2 shows that a direct worst-case analysis is doomed to failure). On the other hand, each switching forces ν_k to jump to regions where convergence is fast. The main interest of switching is thus to prevent the renormalized gradient z_k from approaching its limit set where convergence is slow (since r_k is non-decreasing), and we shall see in Sect. 1.5.2 that choosing $m_1 = 1$ and $1 \leq m_2 \leq 5$ in Algorithm B is suitable.

The limits of a worst-case analysis

One of the simplest construction for a switching algorithm is as follows.

Algorithm C

Use the optimum 2-gradient algorithm if its rate of convergence is smaller than some value $R < R_2^*$ and steepest descent otherwise.

When the state of the algorithm is given by the measure ν_k , denote $r_k = r(\nu_k)$ (respectively $r'_k = r'(\nu_k)$) the rate of convergence (1.14) if a steepest-descent (respectively an optimum 2-gradient) iteration is used. Despite the simplicity of its construction, the performance of Algorithm C resists to a worst-case analysis when one tries to bound the rate of convergence at each iteration of the algorithm. Indeed, one can easily check that the measures ν_s^* associated with the worst rates R_s^* for $s = 1, 2$, see (1.17), satisfy $r'(\nu_1^*) = 0$ and $r(\nu_2^*) = \sqrt{r'(\nu_2^*)} = \sqrt{R_2^*} = N_2^* > R_2^*$, see (1.21). This implies that when the state of the algorithm is given by the measure ν_2^* the rate of convergence equals R_2^* for an optimum 2-gradient iteration and is larger than R_2^* for a steepest-descent iteration; it thus ruins any hope to improve the performance of optimum 2-gradient at each iteration. The situation is not better when measuring the performance per gradient evaluation: the rate of convergence then equals N_2^* for both a steepest-descent and an optimum 2-gradient iteration for the measure ν_2^* . Therefore, the only possibility for obtaining an improvement over optimum 2-gradient is in the long run behavior of the algorithm: when iterations with a slow rate of convergence occur they are compensated by a fast rate at some other iterations. This phenomenon is difficult to analyse since it requires to study several consecutive iterations of steepest descent and/or optimum 2-gradient, which is still an open issue. Some encouraging simulations results are presented in Sect. 1.5.2.

Although improving the value R_2^* for the rate of convergence is doomed to failure, it is instructive to investigate the possible choices for R in Algorithm C.

Consider a steepest-descent iteration. Define $L_k = \mu_{-1}^k \mu_1^k$, so that $L_k = 1/(1 - r_k)$. Since r_k is non-decreasing, L_k is non-decreasing too, and bounded by $1/(1 - R_1^*) = (\varrho + 1)^2/(4\varrho)$. Direct calculation gives

$$L_{k+1} - L_k = \frac{\mu_1^k |\mathbf{M}_{2,-1}^k|}{|\mathbf{M}_{1,0}^k|^2}.$$

When the optimum 2-gradient is used, the rate of convergence r'_k satisfies

$$r'_k = \frac{|\mathbf{M}_{2,-1}^k|}{\mu_{-1}^k |\mathbf{M}_{2,1}^k|},$$

so that using (1.36)

$$\begin{aligned}
L_{k+1} - L_k &= \mu_1^k \mu_{-1}^k r'_k \frac{|\mathbf{M}_{2,1}^k|}{|\mathbf{M}_{1,0}^k|^2} = \mu_1^k \mu_{-1}^k r'_k \frac{1}{|\mathbf{M}_{1,-1}^k| - \mu_1^k \frac{|\mathbf{M}_{2,-1}^k|}{|\mathbf{M}_{2,1}^k|}} \\
&> \frac{\mu_1^k \mu_{-1}^k r'_k}{|\mathbf{M}_{1,-1}^k|} = \frac{r'_k}{r_k}.
\end{aligned} \tag{1.43}$$

Now,

$$L_{k+1} - L_k = \frac{r_{k+1} - r_k}{(1 - r_{k+1})(1 - r_k)} < \frac{r_{k+1} - r_k}{(1 - R_1^*)^2}$$

and $r_k < R_1^*$ so that $r_{k+1} - r_k < F r'_k$ where

$$F = \frac{(1 - R_1^*)^2}{R_1^*} = \frac{16\varrho^2}{(\varrho - 1)^2(\varrho + 1)^2}$$

(and $F < 1$ if $\varrho > 2 + \sqrt{5}$). In Algorithm C, by construction the rate of convergence is less than R when optimum 2-gradient is used. When steepest descent is used, it means that $r'_k \geq R$, and therefore

$$r_k < r_{k+1} - FR < \bar{r}_1(R) = R_1^* - FR. \tag{1.44}$$

Choosing R such that the bounds on the rates coincide for steepest-descent and optimum 2-gradient iterations, that is, such that $\bar{r}_1(R) = R$, gives

$$R = \bar{R} = \frac{(\varrho - 1)^4}{\varrho^4 + 14\varrho^2 + 1}, \tag{1.45}$$

which is larger than R_2^* for any $\varrho > 1$ and therefore cannot be used in Algorithm C. Fig. 1.8 presents R_1^* (dotted line), R_2^* (dashed line), $\bar{R}$ (dash-dotted line) and $\bar{r}_1(R_2^*)$, $\bar{r}_1(R_2^*/2)$ (solid lines, top for $\bar{r}_1(R_2^*/2)$), as functions of ϱ . The trade-off value $\bar{R}$ is clearly larger than R_2^* , taking $R = R_2^*$ gives a bound $\bar{r}_1(R_2^*)$ already close to R_1^* and getting even closer as R decreases.

The situation can be slightly improved by constructing a better bound than (1.44). We write

$$L_{k+1} - L_k = \frac{r_{k+1} - r_k}{(1 - r_{k+1})(1 - r_k)} < \frac{R_1^* - r_k}{(1 - R_1^*)(1 - r_k)}$$

so that for $r'_k \geq R$ (1.43) gives

$$(R_1^* - r_k)r_k > R(1 - R_1^*)(1 - r_k). \tag{1.46}$$

The quadratic equation $(R_1^* - x)x = R(1 - R_1^*)(1 - x)$ has two roots $r_2(R) < \bar{r}_2(R)$ if and only if

$$R > \bar{R}_\varrho = \frac{2 - R_1^* + 2\sqrt{1 - R_1^*}}{1 - R_1^*}$$

or

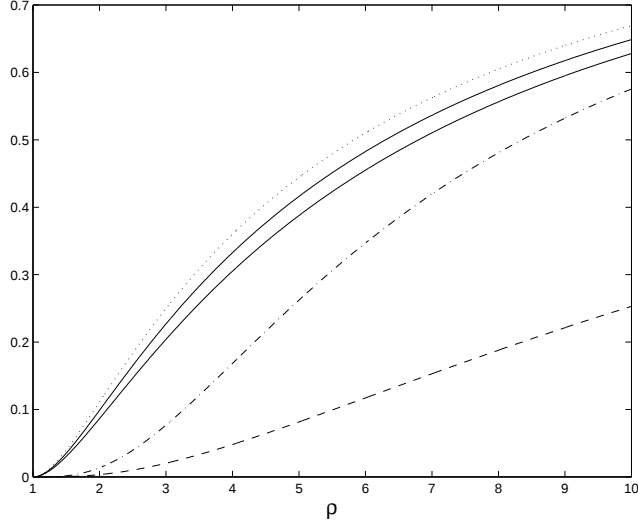


Fig. 1.8. Dotted line: R_1^* , dashed line: R_2^* , dash-dotted line: $\bar{R}$ defined by (1.45), solid lines: $\bar{r}_1(R_2^*) < \bar{r}_1(R_2^*/2)$, with $\bar{r}_1(R)$ defined by (1.44), as functions of ϱ

$$R < \underline{R}_\varrho = \frac{2 - R_1^* - 2\sqrt{1 - R_1^*}}{1 - R_1^*} = \frac{(\varrho + 1 - 2\sqrt{\varrho})^2}{4\varrho},$$

(and they are then positive since their product equals $R(1 - R_1^*)$ and their sum is $R_1^* + R(1 - R_1^*)$). One may easily check that $R_2^* < \underline{R}_\varrho$ and $\bar{R}_\varrho > 4$ for any $\varrho > 1$. Choosing $R \leq R_2^*$ thus ensures that the two roots $r_2(R), \bar{r}_2(R)$ exist and, from (1.46), r_k satisfies $r_2(R) < r_k < \bar{r}_2(R)$. We have $\bar{r}_2(R) < \bar{r}_1(R) = R_1^* - FR$, which thus improves (1.44) (notice that the equation $\bar{r}_2(R) = R$ has now no solution in R). Fig. 1.9 presents R_1^* (dotted line), $\sqrt{R_2^*}$ (dashed line), $\bar{r}_1(R_2^*) < \bar{r}_1(R_2^*/2)$ (solid lines) and $\bar{r}_2(R_2^*), \bar{r}_1(R_2^*/2)$ (dash-dotted lines, top for $\bar{r}_2(R_2^*/2)$), as functions of ϱ . The improvement of $\bar{r}_2(R)$ over $\bar{r}_1(R)$ is clear for R near R_2^* (although $\bar{r}_2(R_2^*)$ remains quite close to R_1^*) but is negligible for $R = R_2^*/2$.

Remark 8. The bounds $\bar{r}_1(R)$ and $\bar{r}_2(R)$ are rather pessimistic, as evidenced by Fig. 1.9: the constraint $r'_k \geq R_2^*$ implies that ν_k is the measure ν_2^* given by (1.17) in Theorem 4, and therefore $r_k = \sqrt{R_2^*}$. The gap between $\sqrt{R_2^*}$ and the lower curves in solid line and dash-dotted line illustrates the pessimism of $\bar{r}_1(R_2^*)$ and $\bar{r}_2(R_2^*)$ respectively.

An exact bound could be obtained, at least numerically, for any value of $R \in [0, R_2^*]$ since the maximization of the rate r_k of the steepest-descent algorithm under the constraint that the rate r'_k of the optimum 2-gradient algorithm satisfies $r'_k \geq R$ for some $R \in [0, R_2^*]$ corresponds to a D_s -optimum design problem under constraint. Indeed, similarly to the proof of Theorem 4, the maximum value for r_k is obtained for the measure $\bar{\nu}$ supported on $[m, M]$

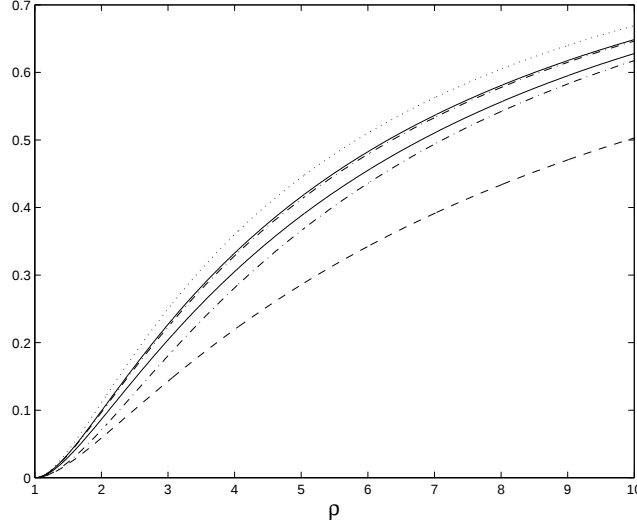


Fig. 1.9. Dotted line: R_1^* , dashed line: $\sqrt{R_2^*}$, solid lines: $\bar{r}_1(R_2^*) < \bar{r}_1(R_2^*/2)$, with $\bar{r}_1(R)$ defined by (1.44), dash-dotted lines: $\bar{r}_2(R_2^*) < \bar{r}_2(R_2^*/2)$, as functions of ρ

that minimizes the variance of the estimator of θ_0 in the regression model $\theta_0 + \theta_1 x$ with i.i.d. errors (a convex function of $\bar{\nu}$) under the restriction that the variance of the estimator of θ_0 in the model $\theta_0 + \theta_1 x + \theta_2 x^2$ is smaller than $1/R$ (which defines a convex constraint on $\bar{\nu}$). Numerical algorithms for solving such convex design problems can be constructed following e.g. the ideas presented in (Molchanov and Zuyev, 2001). $\square$

As an attempt to consider several consecutive iterations in order to circumvent the limits of the worst-case analysis above, one may consider the following algorithm.

- Step 1 ($s = 1$): Use steepest descent while $L_{k+1} - L_k \geq \epsilon$ for some $\epsilon > 0$, with $L_k = 1/(1 - r_k) = \mu_{-1}^k \mu_1^k$. When $L_{k+1} - L_k < \epsilon$, go to Step 2.
Step 2 ($s = 2$): Use the optimum 2-gradient algorithm while the rate of convergence r' is smaller than αR_2^* for some α , $\epsilon R_1^*/R_2^* < \alpha < 1$. When $r' > \alpha R_2^*$, return to Step 1.

The idea of the algorithm is that the rate of convergence of the first iteration of Step 2 is very good when ϵ is small enough. Indeed, when switching from $s = 1$ to $s = 2$, $L_{k+1} - L_k < \epsilon$ and (1.43) imply $r'_k < \epsilon r_k < \epsilon R_1^* < \alpha R_2^*$ (note that this switching necessarily occurs since L_k is not decreasing and bounded by $L^* = 1/(1 - R_1^*)$). We did not manage to improve the results above, however. The reason is that a worst-case analysis leads to consider cycles where the optimum 2-gradient algorithm is used for one iteration only, with a comeback to a sequence of n steepest-descent iterations with associated rates bounded by $1 - 1/[L^* - (n - 1)\epsilon], 1 - 1/[L^* - (n - 2)\epsilon], \dots, 1 - 1/[L^* - \epsilon], 1 - 1/L^*$.

The logarithm of the global rate of convergence over such a cycle of $n + 1$ iterations can then be bounded by

$$\begin{aligned}
\log R_{\max} &= \left(\frac{1}{n+1}\right) \left[\sum_{i=0}^n \log \left(1 - \frac{1}{L^* - i\epsilon}\right) + \log(R_1^* \epsilon) \right] \\
&< \left(\frac{n}{n+1}\right) \log \left[\frac{1}{n} \sum_{i=0}^n \left(1 - \frac{1}{L^* - i\epsilon}\right) \right] + \frac{\log(R_1^* \epsilon)}{n+1} \\
&< \left(\frac{n}{n+1}\right) \log \left[1 - \frac{1}{n(L^*)^2} \sum_{i=0}^n (L^* + i\epsilon) \right] + \frac{\log(R_1^* \epsilon)}{n+1} \\
&= \left(\frac{n}{n+1}\right) \log \left[1 - \frac{1}{L^*} - \frac{(n-1)\epsilon}{2(L^*)^2} \right] + \frac{\log(R_1^* \epsilon)}{n+1}
\end{aligned}$$

which should be maximized with respect to n (to bound the worst-case cycle) and then minimized with respect to ϵ to optimize the bound. A careful analysis shows that the optimum is always attained for ϵ as large as possible and n small. The situation is then similar to that considered for Algorithm C: when $n = 1$ one should choose ϵ such that both the steepest-descent and the optimum 2-gradient iterations have a rate of convergence smaller than R_2^* , which is impossible.

Table 1.1. Global rates R_{100} , N_{188} and their logarithms for Algorithm B ($m_1 = 1$, $m_2 = 4$), averaged over 1000 random problems in $\mathbb{R}^{1000}$ with $\varrho = 100$, together with their standard deviations, minimum and maximum values over the 1000 problems

	mean	std. deviation	minimum	maximum
R_{100}	0.5538	0.0157	0.5047	0.6138
$\log(R_{100})$	-0.5914	0.0284	-0.6838	-0.4881
N_{188}	0.7394	0.0105	0.7061	0.7781
$\log(N_{188})$	-0.3020	0.0142	-0.3481	-0.2508

Some simulation results

We apply Algorithm B with $m_1 = 1$, $m_2 = 4$ to a series of 1000 problems in $\mathbb{R}^d$ with $d = 1000$. For each problem, the eigenvalues of A are randomly generated with the uniform distribution in $[1, \varrho]$ and the initial renormalized gradient z_0 is also randomly generated, with the uniform distribution on the unit sphere $\mathcal{S}_{1000}$.

The algorithm is run for 100 iterations (which means 12 steepest-descent iterations and 88 iterations of the optimum 2-gradient algorithm, and thus 188 gradient evaluations). The results in terms of global rates R_n , see (1.39)

Table 1.2. Global rates R_{100} , N_{200} and their logarithms for the optimum 2-gradient algorithm, averaged over 1000 random problems in $\mathbb{R}^{1000}$ with $\varrho = 100$, together with their standard deviations, minimum and maximum values over the 1000 problems, and theoretical maxima

	mean	std. deviation	minimum	maximum	theoretical max.
R_{100}	0.8199	0.0101	0.7766	0.8399	$R_2^* \simeq 0.8548$
$\log(R_{100})$	-0.1986	0.0123	-0.2528	-0.1745	$\log(R_2^*) \simeq -0.1569$
N_{200}	0.9055	0.0056	0.8812	0.9164	$N_2^* \simeq 0.9245$
$\log(N_{200})$	-0.0993	0.0062	-0.1264	-0.0873	$\log(N_2^*) \simeq -0.0785$

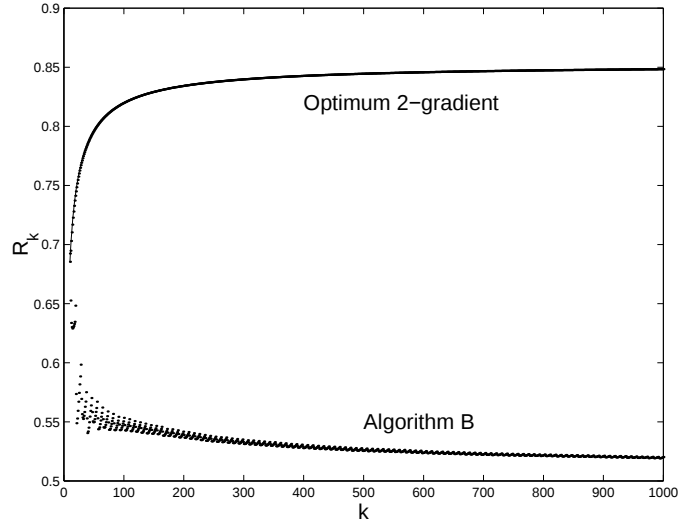


Fig. 1.10. Global rates R_k from iteration 1 to iteration k , (averaged over 1000 random problems) for the optimum 2-gradient and Algorithm B with $m_1 = 1$, $m_2 = 4$ as functions of k

and N_n , see (1.42), are summarized in Table 1.1 for the case $\varrho = 100$. For comparison, $R_1^* \simeq 0.9608$ for steepest descent, $R_2^* \simeq 0.8548$ and $N_2^* \simeq 0.9245$ for the optimum 2-gradient. The rate of convergence of the steepest-descent algorithm is known to be always close to its maximum value, see, e.g., Pronzato *et al.* (2001, 2006). Table 1.2 indicates that this is true also for the optimum 2-gradient algorithm: in that table, Algorithm B is run with $m_1 = 0$, that is, all iterations correspond to the optimum 2-gradient algorithm. One may notice that on average Algorithm B requires $0.3020/0.0993 \simeq 3$ times less gradient evaluations than the optimum 2-gradient algorithm to reach a given precision on the squared norm of the gradient. Even if one considers the very pessimistic situation that corresponds to the worst performance for Algorithm B and the best one for the optimum 2-gradient algorithm, the ratio is $0.2508/0.1264 \simeq$

2. To obtain similar performance for $\varrho = 100$ with an optimum s -gradient algorithm, that is, $N_s^* < 0.78$, one must take $s \geq 9$.

Tables 1.3 and 1.4 give the same information as Tables 1.1 and 1.2 respectively, but for the case when the algorithm is run for 1000 iterations (which means 2000 gradient evaluations for the optimum 2-gradient algorithm and 1888 for Algorithm B with $m_1 = 1$ and $m_2 = 4$). The performance of the optimum 2-gradient algorithm are worse in Table 1.4 than in Table 1.2 (which comes as no surprise since the rate of convergence of the algorithm is non-decreasing), but those of Algorithm B are better when the number of iterations increases. This is confirmed by Fig. 1.10 that shows the global rates R_k from iteration 1 to iteration k , see (1.39), averaged over 1000 random problems, for the optimum 2-gradient and Algorithm B as functions of k . Fig. 1.11 presents the rate of convergence r_k of both algorithms, averaged over the 1000 random problems, as a function of k . The regular increase of r_k is clear for the optimum 2-gradient algorithm, whereas Algorithm B exhibits a rather specific pattern: the dots above the full line correspond to the steepest-descent iterations and those below to optimum 2-gradient iterations, the (averaged) rates of which tend to follow $m_2 = 4$ rather well identified pairs of trajectories.

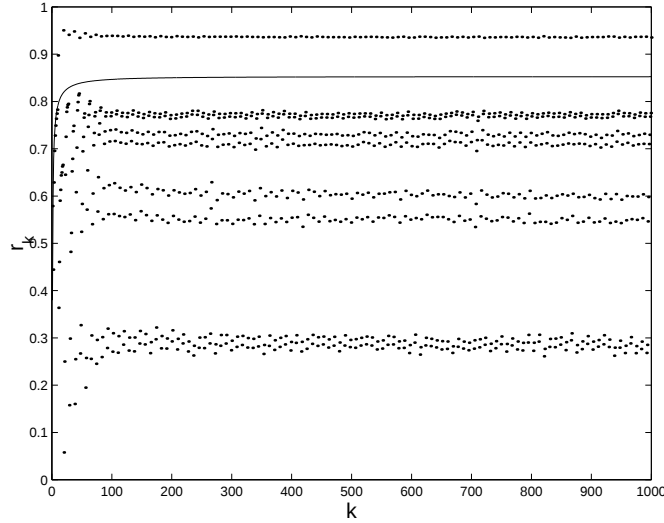


Fig. 1.11. Rate r_k at iteration k , averaged over 1000 random problems, for the optimum 2-gradient (solid line) and Algorithm B with $m_1 = 1$, $m_2 = 4$ (dots) as a function of k

Table 1.5 gives the same information as Table 1.1 but for the case $\varrho = 1000$ (which gives $R_2^* \simeq 0.9842$ and $N_2^* \simeq 0.9920$). To obtain $R_s^* < 0.93$ for $\varrho = 1000$ with an optimum s -gradient algorithms one must take $s \geq 5$ and to get $N_s^* < 0.956$ one must take $s \geq 13$.

Table 1.3. Global rates R_{100} , N_{188} and their logarithms for Algorithm B ($m_1 = 1$, $m_2 = 4$), averaged over 1000 random problems in $\mathbb{R}^{1000}$ with $\varrho = 100$, together with their standard deviations, minimum and maximum values over the 1000 problems

	mean	std. deviation	minimum	maximum
R_{1000}	0.5203	0.0079	0.4953	0.5467
$\log(R_{1000})$	-0.6535	0.0151	-0.7026	-0.6039
N_{1888}	0.7184	0.0055	0.7008	0.7365
$\log(N_{1898})$	-0.3307	0.0076	-0.3555	-0.3059

Table 1.4. Global rates R_{1000} , N_{2000} and their logarithms for the optimum 2-gradient algorithm, averaged over 1000 random problems in $\mathbb{R}^{1000}$ with $\varrho = 100$, together with their standard deviations, minimum and maximum values over the 1000 problems, and theoretical maxima

	mean	std. deviation	minimum	maximum	theoretical max.
R_{1000}	0.8484	0.0018	0.8403	0.8521	$R_2^* \simeq 0.8548$
$\log(R_{1000})$	-0.1644	0.0022	-0.1740	-0.1601	$\log(R_2^*) \simeq -0.1569$
N_{2000}	0.9211	0.0010	0.9167	0.9231	$N_2^* \simeq 0.9245$
$\log(N_{2000})$	-0.0822	0.0011	-0.0870	-0.0800	$\log(N_2^*) \simeq -0.0785$

Table 1.5. Global rates R_{100} , N_{188} and their logarithms for Algorithm B ($m_1 = 1$, $m_2 = 4$), averaged over 1000 random problems in $\mathbb{R}^{1000}$ with $\varrho = 1000$, together with their standard deviations, minimum and maximum values over the 1000 problems

	mean	std. deviation	minimum	maximum
R_{100}	0.8724	0.0182	0.8042	0.9154
$\log(R_{100})$	-0.1368	0.0209	-0.2179	-0.0884
N_{188}	0.9320	0.0099	0.8940	0.9554
$\log(N_{188})$	-0.0705	0.0106	-0.1121	-0.0457

References

- Akaike, H. (1959). On a successive transformation of probability distribution and its application to the analysis of the optimum gradient method. *Ann. Inst. Statist. Math. Tokyo*, **11**, 1–16.
- Elyadi, S. (2005). *An Introduction to Difference Equations*. Springer, Berlin. Third Edition.
- Fedorov, V. (1972). *Theory of Optimal Experiments*. Academic Press, New York.
- Forsythe, G. (1968). On the asymptotic directions of the s -dimensional optimum gradient method. *Numerische Mathematik*, **11**, 57–76.

- Hoel, P. and Levine, A. (1964). Optimal spacing and weighting in polynomial prediction. *Annals of Math. Stat.*, **35**(4), 1553–1560.
- Kantorovich, L. and Akilov, G. (1982). *Functional Analysis*. Pergamon Press, London. Second edition.
- Kiefer, J. and Wolfowitz, J. (1959). Optimum designs in regression problems. *Annals of Math. Stat.*, **30**, 271–294.
- Luenberger, D. (1973). *Introduction to Linear and Nonlinear Programming*. Addison-Wesley, Reading, Massachusetts.
- Meinardus, G. (1963). Über eine Verallgemeinerung einer Ungleichung von L.V. Kantorowitsch. *Numerische Mathematik*, **5**, 14–23.
- Molchanov, I. and Zuyev, S. (2001). Variational calculus in the space of measures and optimal design. In A. Atkinson, B. Bogacka, and A. Zhigljavsky, editors, *Optimum Design 2000*, chapter 8, pages 79–90. Kluwer, Dordrecht.
- Nocedal, J., Sartenaer, A., and Zhu, C. (1998). On the accuracy of nonlinear optimization algorithms. Technical Report Nov. 1998, ECE Department, Northwestern Univ., Evanston, IL 60208.
- Nocedal, J., Sartenaer, A., and Zhu, C. (2002). On the behavior of the gradient norm in the steepest descent method. *Computational Optimization and Applications*, **22**, 5–35.
- Pronzato, L., Wynn, H., and Zhigljavsky, A. (2000). *Dynamical Search*. Chapman & Hall/CRC, Boca Raton.
- Pronzato, L., Wynn, H., and Zhigljavsky, A. (2001). Renormalised steepest descent in Hilbert space converges to a two-point attractor. *Acta Applicandae Mathematicae*, **67**, 1–18.
- Pronzato, L., Wynn, H., and Zhigljavsky, A. (2005). Kantorovich-type inequalities for operators via D-optimal design theory. *Linear Algebra and its Applications (Special Issue on Linear Algebra and Statistics)*, **410**, 160–169.
- Pronzato, L., Wynn, H., and Zhigljavsky, A. (2006). Asymptotic behaviour of a family of gradient algorithms in $\mathbb{R}^d$ and Hilbert spaces. *Mathematical Programming*, **A107**, 409–438.
- Sahm, M. (1998). Optimal designs for estimating individual coefficients in polynomial regression. Ph.D. Thesis, Fakultät für Mathematik, Ruhr Universität Bochum.
- Silvey, S. (1980). *Optimal Design*. Chapman & Hall, London.