

Global Optimization via SPSA

John L. Maryak and Daniel C. Chin

The Johns Hopkins University, Applied Physics Laboratory
11100 Johns Hopkins Road, Laurel, Maryland 20723-6099

john.maryak@jhuapl.edu and daniel.chin@jhuapl.edu

Phone (240) 228-4959 and (240) 228-4956

FAX (240) 228-1093

ABSTRACT

A desire with iterative optimization techniques is that the algorithm reaches the global optimum rather than get stranded at a local optimum value. In this paper, we examine the theoretical and numerical global convergence properties of a certain “gradient free” stochastic approximation algorithm called “SPSA,” that has performed well in complex optimization problems. We establish two theorems on the global convergence of SPSA. The first provides conditions under which SPSA will converge in probability to a global optimum using the well-known method of injected noise. The injected noise prevents the algorithm from converging prematurely to a local optimum point. In the second theorem, we show that, under different conditions, “basic” SPSA *without injected noise* can achieve convergence in probability to a global optimum. This occurs because of the noise *effectively* (and automatically) introduced into the algorithm by the special form of the SPSA gradient approximation. This global convergence without injected noise can have important benefits in the setup (tuning) and performance (rate of convergence) of the algorithm. The discussion is supported by numerical studies showing favorable comparisons of SPSA to simulated annealing and genetic algorithms.

KEYWORDS: *Stochastic Optimization, Global Convergence, Stochastic Approximation, Simultaneous Perturbation Stochastic Approximation (SPSA), Recursive Annealing*

1. INTRODUCTION

A problem of great practical importance is the problem of stochastic optimization, which may be stated as the problem of finding a minimum point, $q^* \in R^p$, of a real-valued function $L(q)$, called the “loss function,” that is observed in the presence of noise. Many approaches have been devised for numerous applications over the long history of this problem. A common desire in many applications is that the algorithm reaches the global minimum rather than get stranded at a local minimum value. In this paper, we consider the popular stochastic optimization technique of stochastic approximation (SA), in particular, the form that may be called “gradient-free” SA. This refers to the case where the gradient, $g(q) = \nabla L(q) / \nabla q$, of the loss function is not readily available or not directly measured (even with noise). This is a common occurrence, for example, in complex systems where the exact functional relationship between the loss function value and the parameters, q , is not known and the loss function is evaluated by measurements on the system (or by other means, such as simulation). In such cases, one uses instead an approximation to $g(q)$ (the well-known form of SA called the Kiefer-Wolfowitz type is an example).

Report Documentation Page				Form Approved OMB No. 0704-0188	
Public reporting burden for the collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to a penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number.					
1. REPORT DATE AUG 2002		2. REPORT TYPE		3. DATES COVERED 00-00-2002 to 00-00-2002	
4. TITLE AND SUBTITLE Global Optimization via SPSA				5a. CONTRACT NUMBER	
				5b. GRANT NUMBER	
				5c. PROGRAM ELEMENT NUMBER	
6. AUTHOR(S)				5d. PROJECT NUMBER	
				5e. TASK NUMBER	
				5f. WORK UNIT NUMBER	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) Johns Hopkins University, Applied Physics Laboratory, 11100 Johns Hopkins Road, Laurel, MD, 20723-6099				8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES)				10. SPONSOR/MONITOR'S ACRONYM(S)	
				11. SPONSOR/MONITOR'S REPORT NUMBER(S)	
12. DISTRIBUTION/AVAILABILITY STATEMENT Approved for public release; distribution unlimited					
13. SUPPLEMENTARY NOTES Proceedings of the 2002 Performance Metrics for Intelligent Systems Workshop (PerMIS '02), Gaithersburg, MD on August 13-15, 2002					
14. ABSTRACT see report					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT Same as Report (SAR)	18. NUMBER OF PAGES 13	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

The usual form of this type of SA recursion is:

$$\hat{q}_{k+1} = \hat{q}_k - a_k \hat{g}_k(\hat{q}_k), \quad (1)$$

where $\hat{g}_k(q)$ is an approximation (at the k^{th} step of the recursion) of the gradient $g(q)$, and $\{a_k\}$ is a sequence of positive scalars that decreases to zero (in the standard implementation) and satisfies other properties. This form of SA has been extensively studied, and is known to converge to a local minimum of the loss function under various conditions.

Several authors (e.g., Chin (1994), Gelfand and Mitter (1991), Kushner (1987), and Styblinski and Tang (1990)) have examined the problem of *global* optimization using various forms of gradient-free SA. The usual version of this algorithm is based on using the standard “finite difference” gradient approximation for $\hat{g}_k(q)$. It is known that carefully injecting noise into the recursion based on this standard gradient can result in an algorithm that converges (in some sense) to the global minimum. For a discussion of the conditions, results, and proofs, see, e.g., Fang et al. (1997), Gelfand and Mitter (1991), and Kushner (1987). The amplitude of the injected noise is decreased over time (a process called “annealing”), so that the algorithm can finally converge when it reaches the neighborhood of the global minimum point.

A somewhat different version of SA is obtained by using a “simultaneous perturbation” gradient approximation, as described in Spall (1992) for multivariable ($p > 1$) problems. The gradient approximation in simultaneous-perturbation SA (SPSA) is much faster to compute than the finite-difference approximation in multivariable problems. More significantly, using SPSA often results in a recursion that is much more economical, in terms of loss-function evaluations, than the standard version of SA. The loss function evaluations can be the most expensive part of an optimization, especially if computing the loss function requires making measurements on the physical system. Several studies (e.g., Spall (1992), Chin (1997)) have shown SPSA to be very effective in complex optimization problems. A considerable body of theory has been developed for SPSA (Spall (1992), Chin (1997), Dippon and Renz (1997), Spall (2000), and the references therein), but, because of the special form of its gradient approximation, existing theory on global convergence of standard SA algorithms is not directly applicable to SPSA. In Section 2 of this paper, we present a theorem showing that SPSA can achieve global convergence (in probability) by the technique of injecting noise. The “convergence in probability” results of our Theorem 1 (Section 2) and Theorem 2 (Section 3) are standard types of global convergence results. Several authors have shown or discussed global convergence in probability or in distribution (Chiang *et al.* (1987), Gelfand and Mitter (1991), Gelfand and Mitter (1993), Geman and Geman (1984), Fang *et al.* (1997), Hajek (1988), Kushner (1987), Yakowitz *et al.* (2000), and Yin (1999)). Stronger “almost sure” global convergence results seem only to be available by using generally infeasible exhaustive search (Dippon and Fabian (1994)) or random search methods (Yakowitz (1993)), or for cases of optimization in a discrete (q -) space (Alrefaei and Andradottir (1999)).

The method of injection of noise into the recursions has proven useful, but naturally results in a relative slowing of the rate of convergence of the algorithm (e.g., Yin (1999)) due to the continued injection of noise when the recursion is near a global solution. In addition, the implementation of the extra noise terms adds to the complexity of setting up the algorithm. In Section 3, we present a theorem showing that, under different (more demanding) conditions, the basic version of SPSA can perform as a

global optimizer *without* the need for injected noise. Section 4 contains numerical studies demonstrating SPSA's performance compared to two other popular strategies for global optimization, namely, simulated annealing and genetic algorithms; and Section 5 is a summary. The Appendix provides some technical details.

2. SPSA WITH INJECTED NOISE AS A GLOBAL OPTIMIZER

Our first theorem applies to the following algorithm, which is the basic SPSA recursion indicated in equation (1), modified by the addition of extra noise terms:

$$\hat{q}_{k+1} = \hat{q}_k - a_k \hat{g}_k(\hat{q}_k) + q_k w_k, \quad (2)$$

where $w_k \in R^p$ is i.i.d. $N(0, I)$ injected noise, $a_k = a/k$, $q_k^2 = q/k \log \log k$, $a > 0$, $q > 0$, and $\hat{g}_k(\bullet)$ is the "simultaneous perturbation" gradient defined as follows: (3)

$$\hat{g}_k(q) \equiv (2c_k \Delta_k)^{-1} [L(q + c_k \Delta_k) - L(q - c_k \Delta_k) + e_k^{(+)} - e_k^{(-)}],$$

where $c_k, e_k^{(\pm)}$ are scalars, $\Delta_k \in R^p$, and the inverse of a vector is defined to be the vector of inverses. This gradient definition follows that given in Spall (1992). The e_k terms represent (unknown) additive noise that may contaminate the loss function observation, c_k and Δ_{kl} are parameters of the algorithm, the c_k sequence decreases to zero, and the Δ_{kl} components are chosen randomly according to the conditions in Spall (1992), usually (but not necessarily) from the Bernoulli (± 1) distribution. (Uniformly or normally distributed perturbations are *not* allowed by the regularity conditions.)

In this Section, we will refer to Gelfand and Mitter (1991) as GM91. Our theorem on global convergence of SPSA using injected noise is based on a result in GM91. In order to state the theorem, we need to develop some notation, starting with the definition of a key probability measure, p^h , used in hypothesis H8 below. Define for any $h > 0$: $dp^h(q)/dq = \exp(-2L(q)/h^2)/Z^h$, where

$Z^h = \int_{R^p} \exp(-2L(q)/h^2) dq$. Next, define an important constant, C_0 , for convergence theory as follows

(GM91). For $t \in R$ and $u_1, u_2 \in R^p$, let

$$I(t, u_1, u_2) = \inf_f \frac{1}{2} \int_0^t |df(s)/ds + g(f(s))|^2 ds,$$

where the *inf* is taken over all absolutely continuous functions $f: R \rightarrow R^p$ such that $f(0) = u_1$ and $f(t) = u_2$, and $|\bullet|$ is the Euclidean norm. Let $V(u_1, u_2) = \lim_{t \rightarrow \infty} I(t, u_1, u_2)$, and $S_0 = \{q | g(q) = 0\}$. Then

$$C_0 \equiv \frac{3}{2} \sup_{u_1, u_2 \in S_0} (V(u_1, u_2) - 2L(u_2)).$$

We will also need the following definition of *tightness*. If K is a compact subset of R^p and $\{X_k\}$ is a sequence of random p -dimensional vectors, then $\{X_k\}$ is tight in K if $X_0 \in K$ and for any $\epsilon > 0$, there exists a compact subset $K_\epsilon \subset R^p$ such that $P(X_k \in K_\epsilon) > 1 - \epsilon, \forall k > 0$. Finally, let $z_k^* \equiv \hat{g}_k(\hat{q}_k) - g(\hat{q}_k)$ and let superscript prime ($'$) denote transpose.

The following are the hypotheses used in Theorem 1.

H1.. Let $\Delta_k \in R^p$ be a vector of p mutually independent mean-zero random variables

$\{\Delta_{k1}, \Delta_{k2}, \dots, \Delta_{kp}\}'$ such that $\{\Delta_k\}$ is a mutually independent sequence that is also independent of

the sequences $\{\hat{q}_1, \dots, \hat{q}_{k-1}\}$, $\{e_1^{(\pm)}, \dots, e_{k-1}^{(\pm)}\}$, and $\{w_1, \dots, w_{k-1}\}$, and such that Δ_{ki} is symmetrically distributed about zero, $|\Delta_{ki}| \leq a_1 < \infty$ a.s. and $E|\Delta_{ki}^{-2}| \leq a_2 < \infty$, a.s. $\forall i, k$.

H2.. Let $e_k^{(+)}$ and $e_k^{(-)}$ represent random measurement noise terms that satisfy $E_k(e_k^{(+)} - e_k^{(-)}) = 0$ a.s. $\forall k$, where E_k denotes the conditional expectation given $\mathcal{I}_k \equiv$ the sigma algebra induced by $\{\hat{q}_0, w_1, \dots, w_{k-1}, z_1^*, \dots, z_{k-1}^*\}$. The $\{e_k^{(\pm)}\}$ sequences are not assumed independent. Assume that $E_k[(e_k^{(\pm)})^2] \leq a_3 < \infty$ a.s. $\forall k$.

H3. $L(q)$ is a thrice continuously differentiable map from R^p into R^1 ; $L(q)$ attains the minimum value of zero; as $|q| \rightarrow \infty$, we have $L(q) \rightarrow \infty$ and $|g(q)| \rightarrow \infty$; $\inf(|g(q)|^2 - \text{Lap}(L(q))) > -\infty$ (Lap here is the Laplacian, i.e., the sum of the second derivatives of $L(q)$ with respect to each of its components); $L^{(3)}(q) \equiv \partial^3 L(q) / \partial q' \partial q' \partial q'$ exists continuously with individual elements satisfying $|L_{i_1 i_2 i_3}^{(3)}(q)| \leq a_5 < \infty$.

H4. The algorithm parameters have the form $a_k = a/k$, $c_k = c/k^g$, for $k = 1, 2, \dots$, where $a, c > 0$, $q/a > C_0$, and $g \in [1/6, 1/2)$.

H5. $[(4p-4)/(4p-3)]^{1/2} < \liminf_{|q| \rightarrow \infty} (g(q)'q / (|g(q)| |q|))$.

H6. $E_k(L(\hat{q}_k \pm c_k \Delta_k))^2 \leq a_4 < \infty$ a.s. $\forall k$.

H7. Let w_k be an i.i.d. $N(0, I)$ sequence, independent of the sequences $\{\hat{q}_1, \dots, \hat{q}_{k-1}\}$, $\{e_1^{(\pm)}, \dots, e_{k-1}^{(\pm)}\}$, and $\{\Delta_1, \dots, \Delta_{k-1}\}$.

H8. For any $h > 0$, $z^h < \infty$; p^h has a unique weak limit p as $h \rightarrow 0$.

H9. There exists a compact subset K of R^p such that $\{\hat{q}_k\}$ is tight in K .

Comments:

- (a) Assumptions H3, H5, and H8 correspond to assumptions (A1) through (A3) of GM91; assumptions H4 and H9 supply the hypotheses stated in GM91's Theorem 2; and the definitions of a_k and q_k given in equation (2) correspond to those used in GM91. Since we will show that assumption (A4) of GM91 is satisfied by our algorithm, this allows us to use the conclusion of their Theorem 2.
- (b) The domain of g given in H4 is one commonly assumed for convergence results (e.g., Spall (1992)).

We can now state our first theorem as follows:

Theorem 1: Under hypotheses H1 through H9, $\hat{q}_k$ converges in probability to the set of global minima of $L(q)$.

Proof: See Maryak and Chin (1999), and the remark on convergence in probability in GM91, p. 1003.

3. SPSA WITHOUT INJECTED NOISE AS A GLOBAL OPTIMIZER

As indicated in the introduction above, the injection of noise into an algorithm, while providing for global optimization, introduces some difficulties such as the need for more “tuning” of the extra terms and

retarded convergence in the vicinity of the solution, due to the continued addition of noise. This effect on the rate of convergence of an algorithm using injected noise is technically subtle, but may have an important influence on the algorithm's performance. In particular, Yin (1999) shows that an algorithm of the form (2) converges at a rate proportional to $\sqrt{\log \log(k + \text{const})}$, while the nominal local convergence rate for an algorithm *without* injected noise is $k^{1/3}$, i.e., $k^{1/3}(\hat{q}_k - q^*)$ converges in distribution (Spall (1992)). These rates indicate a significant difference in performance between the two algorithms.

A certain characteristic of the SPSA gradient approximation led us to question whether SPSA needed to use injected noise for global convergence. Although this gradient approximation tends to work very well in an SA recursion, the SPSA gradient, evaluated at any single point in q -space, tends to be less accurate than the standard finite-difference gradient approximation evaluated at q . So, one is led to consider whether the *effective* noise introduced (automatically) into the recursion by this inaccuracy is sufficient to provide for global convergence *without* a further injection of additive noise. It turns out that *basic* SPSA (i.e., *without* injected noise) does indeed achieve the same type of global convergence as in Theorem 1, but under a different, and more difficult to check, set of conditions.

In this Section, we designate Kushner (1987) as K87, and Kushner and Yin (1997) as KY97. Here we are working with the basic SPSA algorithm having the same form as equation (1):

$$\hat{q}_{k+1} = \hat{q}_k - a_k \hat{g}_k(\hat{q}_k), \quad (4)$$

where $\hat{g}_k(\bullet)$ is the simultaneous-perturbation approximate gradient defined in Section 2, and now (obviously) no extra noise is injected into the algorithm. For use in the subsequent discussion, it will be convenient to define

$$b_k(\hat{q}_k) \equiv E(\hat{g}_k(\hat{q}_k) - g(\hat{q}_k) | \mathfrak{s}_k), \text{ and } e_k(\hat{q}_k) \equiv \hat{g}_k(\hat{q}_k) - E(\hat{g}_k(\hat{q}_k) | \mathfrak{s}_k),$$

where $\mathfrak{s}_k$ denotes the s -algebra generated by $\{\hat{q}_1, \hat{q}_2, \dots, \hat{q}_k\}$, which allows us to write equation (4) as

$$\hat{q}_{k+1} = \hat{q}_k - a_k [g(\hat{q}_k) + e_k(\hat{q}_k) + b_k(\hat{q}_k)]. \quad (5)$$

Another key element in the subsequent discussion is the ordinary differential equation (ODE):

$$\dot{q} = g(q), \quad (6)$$

which, in Lemma 1 of the Appendix is shown to be the "limit mean" ODE for algorithm (4).

Now we can state our assumptions for Theorem 2, as follows:

J.1 . Let $\Delta_k \in R^p$ be a vector of p mutually independent mean-zero random variables

$\{\Delta_{k1}, \Delta_{k2}, \dots, \Delta_{kp}\}'$ such that $\{\Delta_k\}$ is a mutually independent sequence and Δ_k is independent of the sequence $\{\hat{q}_1, \dots, \hat{q}_{k-1}\}$, and such that Δ_{ki} is $\forall i, k$ symmetrically distributed about zero, $|\Delta_{ki}| \leq a_1 < \infty$ a.s. and $E|\Delta_{ki}^{-2}| \leq a_2 < \infty$.

J.2 Let $e_k^{(+)}$ and $e_k^{(-)}$ represent random measurement noise terms that satisfy $E((e_k^{(+)} - e_k^{(-)}) | \mathfrak{s}_k) = 0$ a.s. $\forall k$. The $\{e_k^{(\pm)}\}$ sequences need not be assumed independent. Assume that

$$E((e_k^{(\pm)})^2 | \mathfrak{s}_k) \leq a_3 < \infty \text{ a.s. } \forall k.$$

J.3 (a). $L(q)$ is thrice continuously differentiable and the individual elements of the third derivative satisfy $|L_{i_1 i_2 i_3}^{(3)}(q)| \leq a_5 < \infty$.

(b). $|L(q)| \rightarrow \infty$ as $|q| \rightarrow \infty$.

J.4 The algorithm parameters satisfy the following: the gains $a_k > 0$, $a_k \rightarrow 0$ as $k \rightarrow \infty$, and

$$\sum_{k=1}^{\infty} a_k = \infty. \text{ The sequence } \{c_k\} \text{ is of form } c_k = c/k^g, \text{ where } c > 0 \text{ and } g \in [1/6, 1/2], \text{ and}$$

$$\sum_{k=0}^{\infty} (a_k / c_k)^2 < \infty.$$

J.5 The gradient $g(q)$ is bounded and Lipschitz continuous.

J.6 The ODE (6) has a unique solution for each initial condition.

J.7 For the ODE (6), suppose that there exists a finite collection of disjoint compact stable invariant sets (see K87) $K_1, K_2, \dots, K_m$, such that $\bigcup_i K_i$ contains all the limit sets for (6). These sets are interpreted as closed sets containing all local (including global) minima of the loss function.

J.8 For any $h > 0, Z^h < \infty$; p^h has a unique weak limit p as $h \rightarrow 0$ (Z^h and p^h are defined in Section 2).

J.9 $E |\sum_{i=1}^k e_i(\hat{q}_i)| < \infty \forall k.$

J.10 For any asymptotically stable (in the sense of Liapunov) point, $\bar{q}$, of the ODE (6), there exists a neighborhood of the origin in R^p such that the closure, Q_2 , of that neighborhood satisfies

$\bar{q} + Q_2 \equiv \{\bar{q} + y : y \in Q_2\} \subset \Theta$, where $\Theta \subset R^p$ denotes the allowable q -region. There is a neighborhood, Q_1 , of the origin in R^p and a real-valued function $H_1(y_1, y_2)$, continuous in $Q_1 \times Q_2$, whose y_1 -derivative is continuous on Q_1 for each fixed $y_2 \in Q_2$, and such that the following limit holds. For any $c, \Delta > 0$, with c being an integral multiple of Δ , and any functions $y_1(\bullet), y_2(\bullet)$ taking values in $Q_1 \times Q_2$ and being constant on the intervals $[i\Delta, (i+1)\Delta]$, $i\Delta < c$, we have

$$\int_0^c H_1(y_1(s), y_2(s)) ds = \lim_{m,n} \sup \frac{\Delta}{m} \log E \exp \left[\sum_{i=0}^{(c/\Delta)-1} y_1'(i\Delta) \sum_{j=im}^{im+m-1} b_{n+j}(\hat{q}_{n+j}) \right]. \quad (7)$$

Also, there is a function $H_2(y_3)$ that is continuous and differentiable in a small neighborhood of the origin, and such that

$$\int_0^c H_2(y_1(s)) ds = \lim_{m,n} \sup \frac{\Delta}{m} \log E \exp \left[\sum_{i=0}^{(c/\Delta)-1} y_1'(i\Delta) \sum_{j=im}^{im+m-1} e_{n+j}(\hat{q}_{n+j}) \right]. \quad (8)$$

A bit more notation is needed. Let $T > 0$ be interpreted such that $[0, T]$ is the total time period under consideration in ODE (6). Let

$$\begin{aligned} \bar{H}(y_1, y_2) &= 0.5[H_1(2y_1, y_2) + H_2(2y_1)] \\ \bar{L}(b, y_2) &= \sup_{y_1} [y_1'(b - g(y_2)) - \bar{H}(y_1, y_2)] \end{aligned} \quad (9)$$

and, for $f(0) = x \in R^1$, define the function

$$S(T, f) = \int_0^T \bar{L}(f(s), f(s)) ds,$$

if $f(\bullet)$ is a real-valued absolutely-continuous function on $[0, T]$ and to take the value ∞ otherwise. $S(T, f)$ is the usual action functional of the theory of large deviations (adapted to our context). Define $t_n \equiv \sum_{i=0}^{n-1} a_i$, and $t_k^n = \sum_{i=0}^{k-1} a_{n+i}$. Define $\{\hat{q}_k^n\}$ and $q^n(\bullet)$ by

$$\hat{q}_0^n = x \in \Theta, \hat{q}_{k+1}^n = \hat{q}_k^n - a_{n+k} \hat{g}_{n+k}(\hat{q}_k^n), \text{ and } q^n(t) = \hat{q}_k^n \text{ for } t \in [t_k^n, t_{k+1}^n).$$

Now we can state the last two assumptions for Theorem 2:

J.11 For each $d > 0$ and $i = 1, 2, \dots, m$, there is a r -neighborhood of K_i , denoted $N_r(K_i)$, and $d_r > 0, T_r < \infty$ such that, for each $x, y \in N_r(K_i)$, there is a path, $f(\bullet)$, with $f(0) = x, f(T_y) = y$, where $T_y \leq T_r$ and $S(T_r, f) \leq d$.

J.12 There is a sphere, D_1 , such that D_1 contains $\bigcup_i K_i$ in its interior, and the trajectories of $q^n(\bullet)$ stay in D_1 . All paths of ODE (6) starting in D_1 stay in D_1 .

Note 1. Assumptions J1, J2, and J3(a) are from Spall (1992), and are used here to characterize the noise terms $b_k(\hat{q}_k)$ and $e_k(\hat{q}_k)$. Assumption J3(b) is used on page 178 of K87. Assumption J4 expresses standard conditions on the algorithm parameters (see Spall (1992)), and implies hypothesis (A10.2) in KY97, p. 174. Assumptions J5 and J6 correspond to hypothesis (A10.1) in KY97, p. 174. Assumption J7 is from K87, p. 175. Assumption J8 concerns the limiting distribution of $\hat{q}_k$. Assumption J9 is used to establish the “mean” criterion for the martingale sequence in Lemma 2. Assumptions J11 and J12 are the “controllability” hypothesis A4.1 and the hypothesis A4.2, respectively, of K87, p. 176.

Note 2. Assumption J10 corresponds to hypotheses (A10.5) and (A10.6) in KY97, pp. 179-181. Although these hypotheses are standard forms for this type of large deviation analysis, it is important to justify their reasonableness. The first part (equation (7), involving noise terms $b_k(\hat{q}_k)$) of J10 is justified by the discussion in KY97, p. 174, which notes that the results of their subsection 6.10 are valid if the noise terms (that they denote x_n) are bounded. This discussion is applicable to our algorithm since the $b_k(\hat{q}_k)$ noise terms were shown by Spall (1992) to be $O(c_k^2)$ ($c_k \rightarrow 0$) a.s. The second part (equation (8), involving noise terms $e_k(\hat{q}_k)$) is justified by the discussion in KY97, p. 174, which notes that the results in their subsection 6.10 are valid if the noise terms they denote dM_n (corresponding to our noise terms $e_k(\hat{q}_k)$) satisfy the martingale difference property that we have established in Lemma 2 of the Appendix.

Now we can state our main theorem:

Theorem 2. Under assumptions J1 through J12, $\hat{q}_k$ converges in probability to the set of global minima of $L(q)$.

The idea of the proof is as follows (see the Appendix for the details). This theorem follows from results (in a different context) in K87 for an algorithm $\hat{q}_{k+1} = \hat{q}_k - a_k[g(\hat{q}_k) + z_k]$, where z_k is i.i.d. Gaussian (injected) noise. In order to prove our Theorem 2, we start by writing the SPSA recursion as $\hat{q}_{k+1} = \hat{q}_k - a_k[g(\hat{q}_k) + z_k^*]$, where $z_k^* \equiv \hat{g}_k(\hat{q}_k) - g(\hat{q}_k)$ is the “effective noise” introduced by the inaccuracy of the SPSA gradient approximation. So, our algorithm has the same form as that in K87. However, since z_k^* is not i.i.d. Gaussian, we cannot use K87’s result directly. Instead, we use material in Kushner and Yin (1997) to establish a key “large deviation” result related to our algorithm (4), which allows the result in K87 to be used with z_k^* replacing the z_k in his algorithm.

4. NUMERICAL STUDIES: SPSA WITHOUT INJECTED NOISE

4.1. Two-Dimensional Problem

A study was done to compare the performance of SPSA to a recently published application of the popular genetic algorithm (GA). The loss function is the well-known Griewank function (see Haataja (1999)) defined for a two-dimensional $q = (t_1, t_2)'$, by:

$$L(q) = \cos(t_1 - 100) \cos[(t_2 - 100) / \sqrt{2}] - [(t_1 - 100)^2 + (t_2 - 100)^2] / 4000 - 1,$$

which has thousands of local minima in the vicinity of a single global minimum at $q = (100, 100)'$ at which $L(q) = 0$. Haataja (1999) describes the application of a GA to this function (actually, to find the *maximum* of $-L(q)$) based on noise-free evaluations of $L(q)$ (i.e., $e_k = 0$). This study achieved a success rate of 66% (see Haataja's Table 1.3, p.16) in 50 independent trials of the GA, using 300 generations and 9000 $L(q)$ evaluations in each run of the GA. Haataja's definition of a successful solution is a reported solution where the norm of the solution minus the correct value, q^* , is less than 0.2, and the value of the loss function at the reported solution is within 0.01 of the correct value of zero. We examined the performance of basic SPSA (without adding injected noise) on this problem, using $a_k = a / (A + k)^a$, with $A = 60$, $a = 100$ and $a = .602$, a slowly decreasing gain sequence of a form that has been used in many applications (see Spall (1998)). For the gradient approximation (equation (3)), we chose each component of Δ_k to be an independent sample from a Bernoulli (± 1) distribution, and $c_k = c / k^g$, with $c = 10$ and $g = .101$. Since we used the exact loss function, the e_k noise terms were zero. We ran SPSA, allowing 3000 function evaluations in each of 50 runs, and starting the algorithm (each time) at a point randomly chosen in the domain $[-200, 400] \times [-200, 400]$. Haataja's q -domain was also constrained to lie in a box, but the dimensions of the box were not specified. Hence we chose a domain that is a cube centered at the global minimum, in which there are many local minima of $L(q)$ (as seen in Haataja's (1999) Figure 1.1). SPSA successfully located the global minimum in all 50 runs (100% success rate).

4.2. Ten-Dimensional Problem

For a more ambitious test of the global performance of SPSA, we applied SPSA to a loss function given in Example 6 of Styblinski and Tang (1990), which we will designate for convenience as ST90. The loss function is:

$$L(q) = (2p)^{-1} \sum_{i=1}^p t_i^2 - 4p \prod_{i=1}^p \cos(t_i),$$

where $p = 10$ and $q = (t_1, \dots, t_p)'$. This function has the global minimum value of -40 at the origin, and a large number of local minima. As in the two-dimensional study above, we used the exact loss function. Our goal is to compare the performance of SPSA without injected noise with simulated annealing and with a GA.

For the simulated annealing algorithm, we use the results reported in ST90. They used an advanced form of simulated annealing called fast simulated annealing (FSA). According to ST90, FSA has proven to be much more efficient than classical simulated annealing due to using Cauchy (rather than Gaussian) sampling and using a fast (inversely linear in time) cooling scheme. For more details on FSA,

see ST90. The results of their application of FSA to the above $L(q)$ are given in Table 1 below (FSA values taken from Table 10 of ST90). Table 1 shows the results of 10 independent runs of each algorithm. In each case (each run of each algorithm), the best value of $L(q)$ found by the algorithm is shown. In their study, although FSA was allowed to use 50,000 function evaluations for each of the runs, the algorithm showed very limited success in locating the global minimum. It should be noted that the main purpose of the ST90 paper was to examine a relatively new algorithm, stochastic approximation combined with convolution smoothing. This algorithm, which they call SAS, was much more effective than FSA, yielding results between those shown in Table 1 for GA and SPSA.

For the genetic algorithm (GA), we implemented a GA using the popular features of elitism (elite members of the old population pass unchanged into the new population), tournament selection (tournament size = 2), and real-number encoding (see Mitchell (1996), pp. 168, 170, and 157, respectively). After considerable experimentation, we found the following settings for the GA algorithm to provide the best performance on this problem. The population size was 100, the number of elite members (those carried forward unchanged) in each generation was 10, the crossover rate was 0.8, and mutation was accomplished by adding a Gaussian random variable with mean zero and standard deviation 0.01 to each component of the offspring. The original population of 100 (10-dimensional) q - vectors was created by uniformly randomly generating points in the 10-dimensional hypercube centered at the origin, with edges of length 6 (so, all components had absolute value less than or equal to 3 radians). We constrained all component values in subsequent generations to be less than or equal to 4.5 in absolute value. This worked a bit better than constraining them to be less than 3, since, with the tighter constraints, the GA got stuck at the constraint boundary and could not reach local minima that were just over the boundary. All runs of the GA algorithm reported here used 50,000 evaluations of the loss function. The results of the 10 independent runs of GA are shown in Table 1. Although the algorithm did reasonably well in getting close to the minimum loss value of -40 , it only found the global minimum in one of the 10 runs (run #8). In the other nine cases, a few (typically two or four) of the components were trapped in a local minimum (around $\pm \pi$ radians), while the rest of the components (approximately) achieved the correct value of zero. Note that the nature of the loss function is such that the value of $L(q)$ is very close to an integer (e.g., -39.0 or -38.0) when an even number (e.g., 2 or 4) of components of q are near $\pm \pi$ radians.

We examined the performance of basic SPSA (without adding injected noise), using the algorithm parameters defined in Subsection 4.1 with $A=60$, $a=1$, $\alpha=.602$, $c=2$, and $g=.101$. We started q at $t_i=3$ radians, $i=1,\dots,p$, resulting in an initial loss function value of -31 . This choice of starting point was at the outer boundary of the domain in which we chose initial values for the GA algorithm, and we did not constrain the search space for SPSA as we did for GA (the initialization and search space for FSA were not reported in ST90). We ran 10 Monte Carlo trials (i.e., randomly varying the choices of Δ_k). The results are tabulated in Table 1. The results of these numerical studies show a strong performance of the basic SPSA algorithm in difficult global optimization problems

Table 1. Best Loss Function Value in Each of 10 Independent Runs of Three Algorithms

Run	SPSA	GA	FSA
1	-40.0	-38.0	-24.9
2	-40.0	-39.0	-15.5
3	-40.0	-39.0	-29.0
4	-40.0	-38.0	-32.1
5	-40.0	-37.0	-30.2
6	-40.0	-39.0	-30.1
7	-40.0	-38.0	-27.9
8	-40.0	-40.0	-20.9
9	-40.0	-38.0	-28.5
10	-40.0	-39.0	-34.6
Average Value	-40.0	-38.5	-27.4
Number of Function Evaluations	2,500	50,000	50,000

5. SUMMARY

SPSA is an efficient gradient-free SA algorithm that has performed well on a variety of complex optimization problems. We showed in Section 2 that, as with some standard SA algorithms, adding injected noise to the basic SPSA algorithm can result in a global optimizer. More significantly, in Section 3, we showed that, under certain conditions, the basic SPSA recursion can achieve global convergence *without the need for injected noise*. The use of basic SPSA as a global optimizer can ease the implementation of the global optimizer (no need to tune the injected noise) and result in a significantly faster rate of convergence (no extra noise corrupting the algorithm in the vicinity of the solution). In the numerical studies, we found significantly better performance of SPSA as a global optimizer than for the popular simulated annealing and genetic algorithm methods, which are often recommended for global optimization. In particular, in the case of a 10-dimensional optimization parameter (q), the fast simulated annealing and genetic algorithms generally failed to find the global solution.

APPENDIX (LEMMAS RELATED TO THEOREM 2 AND PROOF OF THEOREM 2)

In this Appendix, we designate Kushner (1987) as K87, and Kushner and Yin (1997) as KY97. Here we are working with the basic SPSA algorithm as defined in equation (4):

$$\hat{q}_{k+1} = \hat{q}_k - a_k \hat{g}_k(\hat{q}_k).$$

We first establish an important preliminary result that is needed in order to apply the results from K87 and KY97 in the proof of Theorem 2.

Lemma 1. The ordinary differential equation (eq. (6) above),

$$\dot{q} = g(q),$$

is the “limit mean ODE” for algorithm (4).

Proof: Examining the definition of limit mean ODE given in KY97, pp. 174 & 138, it is clear that we need to prove that $\frac{1}{m} \sum_{k=n}^{m+n-1} [g(q) + e_k(\hat{q}_k) + b_k(\hat{q}_k)] \rightarrow g(q)$ w.p. 1 as $m, n \rightarrow \infty$. Since Spall

(1992) has shown that $b_k(\hat{q}_k) \rightarrow 0$ w.p. 1, we can conclude using Cesaro summability that the contribution of the $b_k(\hat{q}_k)$ terms to the limit is zero w.p. 1. For the $e_k(\hat{q}_k)$ terms, we have by definition that $E[e_k(\hat{q}_k)] = 0$; hence, by the law of large numbers, the contribution of the $e_k(\hat{q}_k)$ terms to the limit is also zero. Q.E.D.

Our next Lemma relates to Note 2 in Section 3.

Lemma 2. Under assumptions J1, J3(a), and J9, the sequence $\{e_k(\hat{q}_k)\}$ is a $\mathfrak{N}_k$ -martingale difference.

Proof: It is sufficient to show that $M_k \equiv \sum_{i=1}^k e_k(\hat{q}_k)$ is a $\mathfrak{N}_k$ -martingale. Assumption J9 satisfies the first requirement (see KY97, p.68) of the martingale definition, that $E|M_k| < \infty$. For the main

requirement, we have for any k :

$$\begin{aligned} E(M_{k+1} | M_k, \dots, M_1) &= E[e_{k+1}(\hat{q}_{k+1}) + M_k | M_k, \dots, M_1] \\ &= M_k + E[e_{k+1}(\hat{q}_{k+1}) | M_k, \dots, M_1] \\ &= M_k + E\{\hat{g}_{k+1}(\hat{q}_{k+1}) - E(\hat{g}_{k+1}(\hat{q}_{k+1}) | \hat{q}_{k+1}) | M_k\} \\ &= M_k + E_{\hat{q}_{k+1}}\{[\hat{g}_{k+1}(\hat{q}_{k+1}) - E(\hat{g}_{k+1}(\hat{q}_{k+1}) | \hat{q}_{k+1})] | M_k, \hat{q}_{k+1}\} \\ &= M_k + E_{\hat{q}_{k+1}}\{E(\hat{g}_{k+1}(\hat{q}_{k+1}) | M_k, \hat{q}_{k+1}) - E_{\hat{q}_{k+1}}[E(\hat{g}_{k+1}(\hat{q}_{k+1}) | M_k, \hat{q}_{k+1})]\} = M_k, \end{aligned}$$

where $E_{\hat{q}_{k+1}}$ denotes expectation conditional on $\hat{q}_{k+1}$, and all equalities concerning conditional expectations are w.p. 1. Q.E.D.

A key step in the proof of our main result (Theorem 2 below) is establishing the following “large deviation” result (Lemma 3). Let B_x be a set of continuous functions on $[0, T]$ taking values in Θ and with initial value x . Let B_x^0 denote the interior of B_x , and $\bar{B}_x$ denote the closure.

Lemma 3. Under assumptions J4, J5, J6, and J10, we have

$$\begin{aligned} - \inf_{f \in B_x^0} S(T, f) &\leq \liminf_n \log P_X^n \{q^n(\bullet) \in B_x\} \\ &\leq \limsup_n \log P_X^n \{q^n(\bullet) \in B_x\} \leq - \inf_{f \in \bar{B}_x} S(T, f), \quad (9) \end{aligned}$$

where P_X^n denotes the probability under the condition that $q^n(0) = x$.

Proof: This result is adapted from Theorem 10.4 in KY97, p. 181. Note that our assumption J10 is a modified form of their assumptions (A10.5) and (A10.6), using “equals” signs rather than inequalities. The two-sided inequality in (9) follows from J10 by an argument analogous to the proof of KY87’s Theorem 10.1 (p. 178), which uses an “equality” assumption ((A10.4), p. 174) to arrive at a two-sided large deviation result analogous to (9) above. Q.E.D.

We restate our main theorem:

Theorem 2: Under hypotheses J1 through J12, $\hat{q}_k$ converges in probability to the set of global minima of $L(q)$.

Proof: This result follows from a discussion in K87. Theorem 2 of K87, (p.177) describes probabilities involving expected times for the SA algorithm (system (1.1) of K87) to transition

from one K_i to another. The SA algorithm he uses can be written in our notation as $\hat{q}_{k+1} = \hat{q}_k - a_k[g(\hat{q}_k) + z_k]$, where z_k is i.i.d. Gaussian (injected) noise. The K87 Theorem 2 uses the i.i.d. Gaussian assumption only to arrive at a large deviation result exactly analogous to our Lemma 3. The subsequent results in K87 are based on this large deviation result. Recall that the SPSA algorithm without injected noise can be written in the form $\hat{q}_{k+1} = \hat{q}_k - a_k[g(\hat{q}_k) + z_k^*]$. Since we have established Lemma 3 for SPSA, the results of K87 hold for the SPSA algorithm with its “effective” noise $\{z_k^*\}$ replacing the $\{z_k\}$ sequence used in K87. In particular, K87’s discussion (pp. 178, 179) of his Theorem 2 is applicable to our Theorem 2 context (SPSA without injected noise), which corresponds to K87’s “potential case.” Note that our formulation corresponds to the K87 setup where $b(x, x) = \bar{b}(x)$ in his notation, which, by the comment in K87, p. 179, means that his discussion is applicable to his system (1.1) and hence to our setup. In his discussion on p. 179, K87 indicates that the difference between the measure of x_n (which corresponds to our $\hat{q}_k$) and the invariant measure (which we have denoted p^h) converges asymptotically ($n, k \rightarrow \infty, h \rightarrow 0$) to the zero measure weakly. This means that, in the limit as $k \rightarrow \infty$, $\hat{q}_k$ is equivalent to p in the same sense as in Theorem 2 of Gelfand and Mitter (1991), and the desired convergence in probability follows as in Theorem 1 above. Q.E.D.

ACKNOWLEDGEMENT

This work was partially supported by the JHU/APL Independent Research and Development Program and U.S. Navy contract N00024-98-D-8124.

REFERENCES

- [1] Alrefaei, M.H. and Andradottir, S. (1999), “A Simulated Annealing Algorithm with Constant Temperature for Discrete Stochastic Optimization,” *Management Science*, 45, pp. 748-764.
- [2] Chaing T-S., Hwang, C-R., and Sheu, S-J. (1987), “Diffusion for Global Optimization in R^n ,” *SIAM J. Control Optim.*, 25, pp. 737-753.
- [3] Chin, D.C. (1997), “Comparative Study of Stochastic Algorithms for System Optimization Based on Gradient Approximations,” *IEEE Trans. Systems, Man, and Cybernetics – Part B: Cybernetics*, 27, pp. 244-249.
- [4] Chin, D.C. (1994), “A More Efficient Global Optimization Algorithm Based on Styblinski and Tang,” *Neural Networks*, 7, pp. 573-574.
- [5] Dippon, J. and Fabian, V. (1994), “Stochastic Approximation of Global Minimum Points,” *J. Statistical Planning and Inference*, 41, pp. 327-347.
- [6] Dippon, J. and Renz, J. (1997), “Weighted Means in Stochastic Approximation of Minima,” *SIAM J. Control Optim.*, 35, pp. 1811-1827.
- [7] Fang, H., Gong, G., and Qian, M. (1997), “Annealing of Iterative Stochastic Schemes,” *SIAM J. Control Optim.*, 35, pp.1886-1907.
- [8] Gelfand, S.B. and Mitter, S.K. (1991), “Recursive Stochastic Algorithms for Global Optimization in R^d ,” *SIAM J. Control Optim.*, 29, pp. 999-1018.

- [9] Gelfand, S.B. and Mitter, S.K. (1993), "Metropolis-Type Annealing Algorithms for Global Optimization in R^d ," *SIAM J. Control Optim.*, 31, pp. 110-131.
- [10] Geman S. and Geman, D. (1984), "Stochastic Relaxation, Gibbs Distributions, and the Bayesian Restoration of Images, *IEEE Trans. Pattern Analysis and Machine Intelligence*, PAMI-6, pp. 721-741.
- [11] Haataja, J. (1999), "Using Genetic Algorithms for Optimization: Technology Transfer in Action," Chapter 1, pp. 3 – 22, in *Evolutionary Algorithms in Engineering and Computer Science*, Edited by K. Miettinen, M.M. Makela, P. Neittaanmaki, and J. Periaux, Wiley, Chichester.
- [12] Hajek, (1988), "Cooling Schedules for Optimal Annealing," *Mathematics of Operations Research*, 13, pp. 311-329.
- [13] Kushner, H.J. and Yin, G.G. (1997), *Stochastic Approximation Algorithms and Applications*, Springer, New York.
- [14] Kushner, H.J. (1987), "Asymptotic Global Behavior for Stochastic Approximation and Diffusions with Slowly Decreasing Noise Effects: Global Minimization Via Monte Carlo," *SIAM J. Appl. Math.*, 47, pp. 169-185.
- [15] Maryak, J.L. and Chin, D.C. (1999), "Efficient Global Optimization Using SPSA," *Proc. Amer. Control Conf.*, San Diego, June 2-4, pp. 890-894.
- [16] Mitchell, M. (1996), *An Introduction to Genetic Algorithms*, MIT Press, Cambridge, Mass.
- [17] Spall, J.C. (2000), "Adaptive Stochastic Approximation by the Simultaneous Perturbation Method," *IEEE Trans. Automat. Control*, 45, pp. 1839-1853.
- [18] Spall, J.C. (1998), "Implementation of the Simultaneous Perturbation Algorithm for Stochastic Optimization," *IEEE Trans. Aerospace and Electronic Systems*, 34, pp. 817-823.
- [19] Spall, J.C. (1992), "Multivariate Stochastic Approximation Using a Simultaneous Perturbation Gradient Approximation," *IEEE Trans. Automat. Control*, 37, pp. 332-341.
- [20] Styblinski, M.A. and Tang, T-S. (1990), "Experiments in Nonconvex Optimization: Stochastic Approximation with Function Smoothing and Simulated Annealing," *Neural Networks*, 3, pp. 467-483.
- [21] Yakowitz, S. (1993), "A Globally Convergent Stochastic Approximation," *SIAM J. Control Optim.* 31 pp. 30-40.
- [22] Yakowitz, S., L'Ecuyer, P., and Vazquez-Abad, F. (2000), "Global Stochastic Optimization with Low-Dispersion Point Sets," *Operations Research*, 48, pp. 939-950.
- [23] Yin, G. (1999), "Rates of Convergence for a Class of Global Stochastic Optimization Algorithms," *SIAM J. Optim.*, 10, pp. 99-120.