



INTERNATIONAL ATOMIC ENERGY AGENCY
UNITED NATIONS EDUCATIONAL, SCIENTIFIC AND CULTURAL ORGANIZATION



INTERNATIONAL CENTRE FOR THEORETICAL PHYSICS

MIRAMARE - P.O.B. 586 - 34100 TRIESTE (ITALY) - TELEPHONES: 224281/2/3/4/5/6 - CABLE: CENTRATOM



SMR/23/54

AUTUMN COURSE ON APPLICATIONS OF ANALYSIS
TO MECHANICS

22 September - 26 November 1976 - - - -

IDENTIFICATION OF SYSTEMS:

Application to oil field problems

G. CHAVENT
I R I A
Le Chesnay
France

These are preliminary lecture notes intended for participants only.
Copies are available outside the Publications Office (T-floor) or from
room 112.

I Identification as an optimization problem.

We first want to show, on a very simple example, how an identification problem can be set as an optimization problem.

Let $\Omega \subset \mathbb{R}^n$ be an isotropic body with heat conductivity $a(x)$ at point $x \in \Omega$. Suppose that on the part Γ_1 of the boundary, the body Ω is at a (given) zero temperature, and that on the complementary part Γ_2 of Ω , the body receives at each point s of Γ_2 a known heat flow $g(s)$. The equation giving the steady-state temperature $u(x)$ at every point $x \in \Omega$ is:

$$(P) \begin{cases} -\sum_{i=1}^n \frac{\partial}{\partial x_i} (a(x) \frac{\partial u}{\partial x_i}) = 0 & \text{in } \Omega \\ u(s) = 0 & \text{for } s \in \Gamma_1 \\ a(s) \frac{\partial u}{\partial n}(s) = g(s) & \text{for } s \in \Gamma_2 \end{cases}$$

Starting from (P), it is possible to define two types of optimization problems:

a) A Control Problem: One wants to find a heating law g such that the temperature repartition $u(x)$ be as close as possible to a given repartition $z(x)$, and this without spending too much energy. Define:

Error-Cost function $J(g) = \int_{\Omega} (u(x;g) - z(x))^2 dx + \nu \int_{\Gamma_2} g^2(s) ds$

Set of admissible controls $: L^2(\Gamma_2)$

so that the control problem may be set as

$$(c) \quad \boxed{\begin{array}{l} \text{Min } J(g) \\ g \in L^2(\Gamma_2) \end{array}}$$

Remark 1: As the mapping $g \longrightarrow u$ is linear, the functional $J(g)$ is quadratic, so that it can be proved, under nice hypotheses, that the optimization problem (c) has a unique solution $\bar{g}$. ■

b) An Identification problem. The resolution of the control problem (c) requires a knowledge of all the parameters appearing in the state equations (P), that is of the heat conductivity function $a(x)$. As this function is not accessible to direct experimental measure, we shall try an indirect determination using some temperature measurements which are available. Suppose (for simplicity only) that the temperature is measured at every point $x \in \Omega$, which results in an observed temperature $z(x)$. Define

Error function $J(a) = \int (u(x;a) - z(x))^2 dx$

Set of admissible parameters: we know a priori (strictly positive) lower and upper bounds α and M for the heat conductivity, so that we set

$$\mathcal{Q}_{ad} = \{ a(x) / 0 < \alpha \leq a(x) \leq M \text{ a.e. on } \Omega \}$$

and we may define the identification problem as:

$$(I) \quad \boxed{\text{Min } J(a)}$$

Remark 2 : Even in the simple case of the linear state equation (P), the mapping $a \rightarrow u$ is non linear, so that the functional J is generally non quadratic, and it may happen that the optimization problem (I) do not have a solution ■

The aim of this course is to see how to use the formulation (I) of an identification problem to solve it actually, i.e. we shall go up to its numerical resolution. The tools we shall develop apply of course to the off-line resolution of control problems.

In the following paragraph, we shall give some notions about optimization algorithms which may be used for the resolution of (I).

② Two Unconstrained Optimization Algorithms

The numerical optimization algorithm is always applied to the discretized version (see § V) of the continuous problem (I). So, from a practical point of view, it's not a restriction to suppose that the parameter a is finite dimensional (instead of infinite dimensional as in problem (I)). This hypothesis will hold throughout all the § ② and ③.

Let:

(2.1) J be an application from $\mathbb{R}^N$ into $\mathbb{R}$

We shall denote by $\nabla J(a)$ the gradient of J (ie the column vector of partial derivatives of J with respect to $a_1, a_2, \dots, a_N$) and by $\langle \cdot, \cdot \rangle$ the scalar product in $\mathbb{R}^N$.

We are concerned with the following unconstrained optimization problem:

(2.2) find $\hat{a} \in \mathbb{R}^N$ such that $J(\hat{a}) \leq J(a) \quad \forall a \in \mathbb{R}^N$.

Thm 1: Under the hypothesis:

(2.3) J is continuous

(2.4) $J(a)$ tends to $+\infty$ when a tends to infinity

ie $\forall K > 0, \exists L > 0$ st $\|a\| > L \Rightarrow J(a) \geq K$

there exists at least a solution $\hat{a}$ to the problem (2.2)

If moreover J is strictly convex, $\hat{a}$ is unique.

Proof: choose any $\tilde{a} \in \mathbb{R}^N$ and define:

$$C = \{a \in \mathbb{R}^N \mid J(a) \leq J(\tilde{a})\}$$

5

α is clearly a closed bounded, hence compact subset of $\mathbb{R}^N$ (use (2.3) and (2.4)), so J attains its minimum on α , which is obviously its minimum on all $\mathbb{R}^N$. ■

We want now to design a method for actually computing $\hat{a}$. Suppose that

(2.5) J is a differentiable mapping

Then noticing that, for every $a \in \mathbb{R}^N$:

$$\left. \frac{d}{dp} [J(a - p \nabla J(a))] \right|_{p=0} = - \|\nabla J(a)\|^2 \leq 0$$

we see on the figure 1 that, for sufficiently small p :

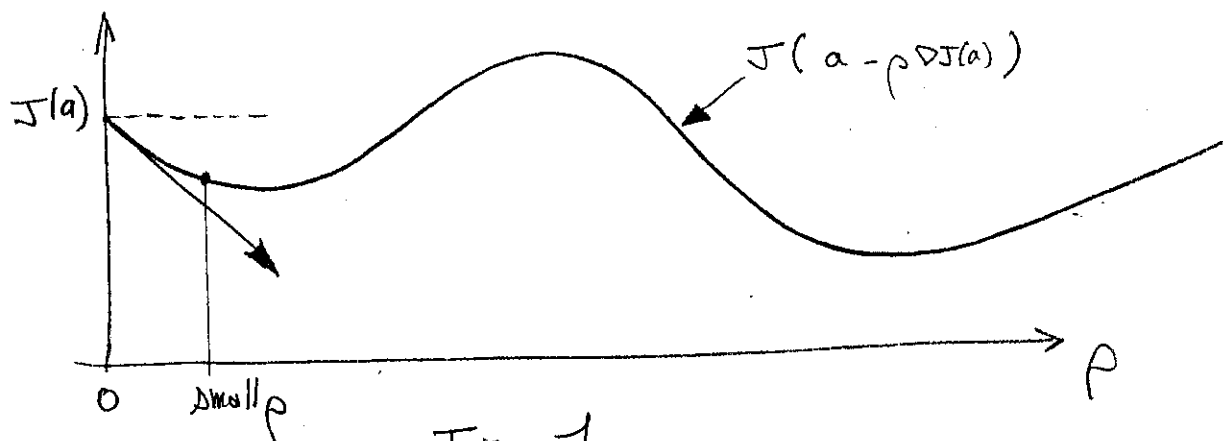
$$J(a - p \nabla J(a)) \leq J(a)$$


Fig 1

This leads to the

Simple Gradient Algorithm

- choose a gradient step $p > 0$
- choose an initial guess $a^0 \in \mathbb{R}^N$
- define a sequence $a^0, a^1, \dots, a^k, \dots$ etc... by

$$(2.6) \quad a^{k+1} = a^k - p \nabla J(a^k)$$

Does a^k converge toward $\hat{a}$? An answer is given

Thm 2 Under the hypothesis (2.4) (2.5) and:

$$(2.7) \begin{cases} \nabla J \text{ is a Lipschitz application} \\ \text{i.e. } \exists M > 0 \text{ s.t. } \forall a, b \in \mathbb{R}^N \quad \|\nabla J(a) - \nabla J(b)\| \leq M \|a - b\| \end{cases}$$

then:

i) there exists at least a solution $\hat{a}$ to the problem (2.2), such that:

$$(2.8) \quad \nabla J(\hat{a}) = 0$$

ii) $\exists \bar{\rho} > 0$ s.t. $\forall \rho \in]0, \bar{\rho}[$, $\forall a^0 \in \mathbb{R}^N$, the sequence $\{a^k\}$ defined by (2.6) contains a subsequence a^N s.t.:

$$(2.9) \quad a^N \rightarrow a^* \quad \text{with} \quad \nabla J(a^*) = 0$$

iii) If moreover J is strictly convex, all the sequence a^k converges toward the unique solution $\hat{a}$ of problem (2.2).

Proof: i) is a simple corollary of Thm 1

ii) Let us first show that one may choose ρ so that the sequence $\{J(a^k)\}$ is decreasing. Define ρ_0 by:

$$(2.10) \quad J(a^{k+1}) - J(a^k) = \langle \nabla J(a^k), a^{k+1} - a^k \rangle + \rho_0$$

Applying the finite increment theorem to the mapping $a \rightarrow J(a) - \nabla J(a^k)a$ one gets:

$$|\rho_0| \leq \|a^{k+1} - a^k\| \sup_{a \in [a^k, a^{k+1}]} \|\nabla J(a) - \nabla J(a^k)\|$$

Using (2.7) one gets:

$$(2.11) \quad |\rho_0| \leq M \|a^{k+1} - a^k\|^2$$

Using (2.6) (2.10) (2.11) we have:

$$(2.12) \quad J(a^{k+1}) - J(a^k) \leq -\rho(1 - \rho M) \|\nabla J(a^k)\|^2$$

Define:

$$(2.13) \quad \bar{\rho} = \frac{1}{M}$$

Then for every $\rho \in]0, \bar{\rho}[$ the sequence $\{J(a^k)\}$ is

decreasing (cf (2.12)) and bounded below (from part i) of Thm),⁷
and hence converging, so that $J(a^k) - J(a^{k+1}) \rightarrow 0$. As

$$(2.14) \quad 0 \leq \|\nabla J(a^k)\|^2 \leq \frac{1}{\rho(1-\rho^2)} (J(a^k) - J(a^{k+1}))$$

we see that

$$(2.15) \quad \nabla J(a^k) \rightarrow 0 \quad \text{as } k \rightarrow +\infty.$$

As the sequence $\{J(a^k)\}$ is converging, it is bounded above; From (2.4) we deduce then that the sequence $\{a^k\}$ is bounded, and that there exists a subsequence $a^{k'}$ converging to some a^* . As ∇J is Lipschitz, it's continuous, hence:

$$(2.16) \quad \nabla J(a^{k'}) \rightarrow \nabla J(a^*)$$

Comparing (2.15) and (2.16) we see that $\nabla J(a^*) = 0$, which ends the proof of ii)

iii) If J is strictly convex, $\bar{a}$ is the only solution of the equation $\nabla J(a) = 0$, so that from every subsequence $\{a^{k'}\}$ of $\{a^k\}$, we can extract a subsequence $\{a^{k''}\}$ converging toward $\bar{a}$; this proves that all the sequence a^k converges to $\bar{a}$. ■

Remark 3: For $N=2$ and for a function J whose level lines are ellipses, the simple gradient algorithm propagates as shown in fig 2:

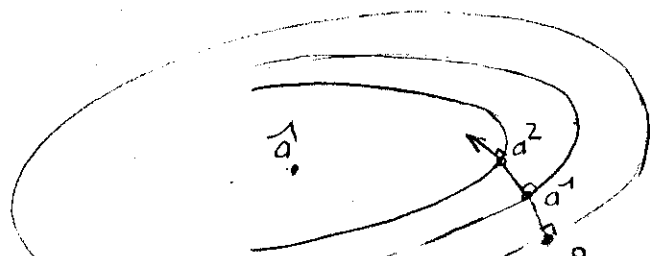


Fig 2

8

The practical choice of ρ in the simple gradient algorithm is somewhat difficult, as you don't know $\bar{\rho}$:

— if ρ is chosen very small, your algorithm shall be stable ($\rho \in]0, \bar{\rho}[$ very probably!), but the speed of convergence is slow (the amount of change of the parameters at each gradient iteration is small)

— if ρ is chosen large, the speed of convergence shall increase, but the risk of divergence of the algorithm ($\rho \gg \bar{\rho}$) increases simultaneously.

That is why gradient algorithms have been designed which use a different gradient step ρ^k at each gradient iteration. A popular among those algorithms is the

Steepest Descent Algorithm

- choose an initial guess $a^0 \in \mathbb{R}^n$
- define a sequence $a^0, a^1, \dots, a^k, \dots$ by:

$$(2.17) \quad a^{k+1} = a^k - \rho^k \nabla J(a^k)$$

where ρ^k is defined by

$$(2.18) \quad J(a^k - \rho^k \nabla J(a^k)) \leq J(a^k - \rho \nabla J(a^k)) \quad \forall \rho > 0.$$

Remark 4: Graphical interpretation (compare with Remark 3)

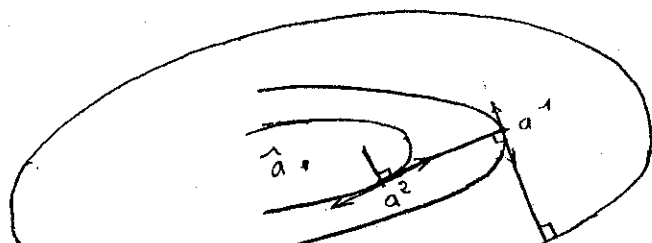


Fig 3

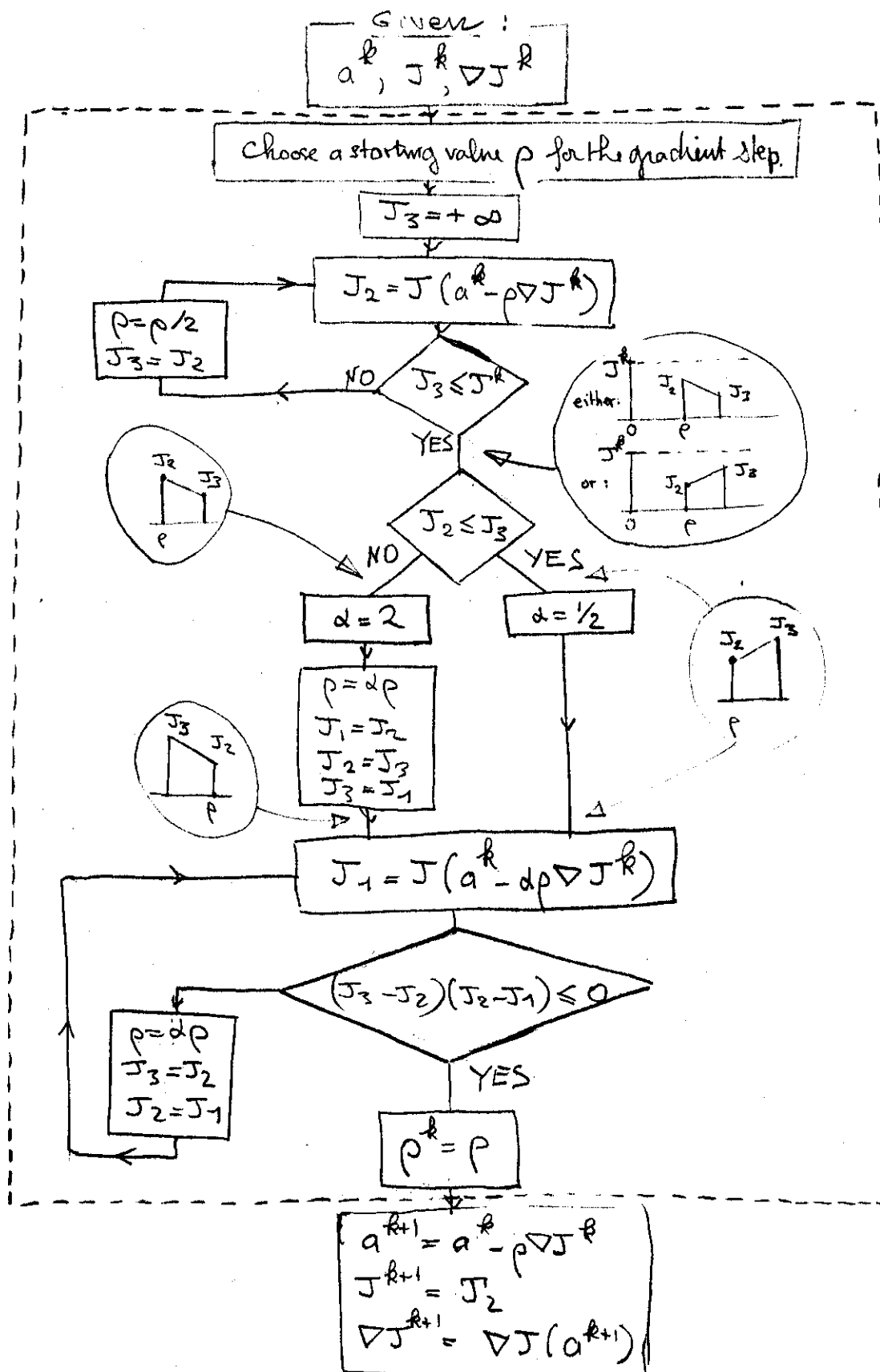
As for the simple gradient algorithm, it's possible to give a proof of convergence of the steepest descent algorithm under similar hypothesis.

From a numerical point of view, the difficulty is the determination of p^k at each gradient-iteration. Happily, a very precise determination of the p^k which minimizes J along the direction opposite to the gradient do not increase the performances of the steepest descent algorithm. So a very crude determination of p^k by dichotomy is enough. However, the determination of the global minimum of J on the direction opposite to the gradient (as required in (2.18)) is not numerically possible; and we have to be content with a p^k which gives approximately a local minimum.

The block-diagram of fig 4 gives a way to compute such a p^k , knowing a^k , J^k and ∇J^k . It requires the choice of a starting value for p to initiate the dichotomy; a simple choice for this is to take $p = p^{k-1}$, so that p has to be chosen only for the first iteration.

The upper loop is designed to eventually reduce sufficiently the starting value of p so that the corresponding criterion J be at least smaller or equal to the value J^m of the criterion at the last iteration. As is clear from fig 1, one should go out of this loop after a finite number of steps. If not, this means that the computed vector ∇J^k is not the gradient of J , which may be due to

to approximations and/or rounding errors when a^k is very near to $\hat{a}$; we shall say in that case that the 1-D approximate minimization is "unsuccessful" (see fig.)



III) A constrained optimization algorithm

As we have seen in §E), the identification problem as a generally constrained optimization problem, i.e. the parameter a has to stay in a given (generally convex) closed set $\mathcal{A}_{ad}$ of $\mathbb{R}^N$.

We shall give a slight modification of the gradient algorithm of §II) which makes it possible, in a simple but common case, to take in account those constraints.

Suppose that $a^k \in \mathcal{A}_{ad}$ (see fig 5). Then it may happen that the point $\tilde{a}^{k+1} = a^k - \rho \nabla J^k$ do not belong to $\mathcal{A}_{ad}$.

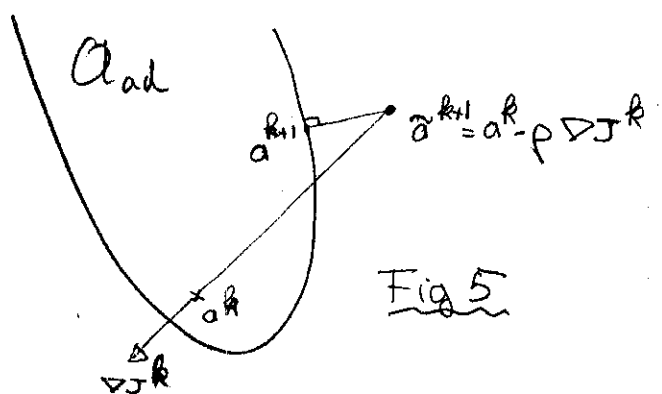


Fig 5

So, denoting by P the projection operator on $\mathcal{A}_{ad}$, we shall replace $\tilde{a}^{k+1}$ by its projection on $\mathcal{A}_{ad}$:

$$(3.1) \quad a^{k+1} = P(a^k - \rho \nabla J^k)$$

With this simple modification, both algorithms of §II) can be viewed as constrained optimization algorithm.

In practice, those constrained optimization algorithms are effective only if the computation of the projection on $\mathcal{A}_{ad}$ is numerically possible. As this computation consists in the minimization on $\mathcal{A}_{ad}$ of the distance between $\tilde{a}^{k+1}$ and a , it is generally of the same order of difficulty than the initial optimization problem.

However, in the very simple case of side constraints, i.e.

$$(3.2) \quad \mathcal{A}_{ad} = \{ a \in \mathbb{R}^N \mid \underline{a}_i \leq a_i \leq \bar{a}_i \quad \forall i=1,2,\dots,N \}$$

the projection operator is straightforward:

$$(3.3) \quad (Pa)_i = \begin{cases} \underline{a}_i & \text{if } a_i < \underline{a}_i \\ a_i & \text{if } \underline{a}_i \leq a_i \leq \bar{a}_i \\ \bar{a}_i & \text{if } a_i > \bar{a}_i \end{cases}$$

This type of constraint corresponds to the data of lower and upper bounds on the unknown coefficients, which is a very common case; the steepest descent algorithm with projection works for all of those cases; its block-diagram is given in the fig 6. The block "1-D approximate minimization of..." can be replaced by the dashed block of fig 4, provided that in this latter $a^k - p \nabla J^k$ be replaced everywhere by $P(a^k - p \nabla J^k)$

Remark 5: Change of unknown parameters. Suppose that the parameter a is an affine function of p parameters $b_1, \dots, b_p$ (generally $p \leq N$, but $p > N$ works too):

$$(3.4) \quad a = \Theta b + \tilde{a} \quad \Theta = N \times p \text{ matrix, } \tilde{a} \in \mathbb{R}^N$$

and that the optimization algorithm is written for the minimization of J with respect to a (as in fig 6)

If now we want to minimize J with respect to b , the only constraints on b being the side constraints on a which translate to side constraints on b , we have to compute:

$$(3.5) \quad b^{k+1} = b^k - p^k \frac{dJ}{db}(b^k)$$

which, by multiplication by the matrix Θ and adding $\tilde{a}$ gives:

$$(3.6) \quad a^{k+1} = a^k - p^k \Theta \frac{dJ}{db}(b^k)$$

As:

$$(3.7) \quad \frac{dJ}{db}(b^k) = {}^t \Theta \nabla J(a^k)$$

we see that the only modifications to do in the program of fig 6 are:

- i) relax the side constraints on a which are not side constraints on b by taking, for the corresponding indices i , $\underline{a}_i = -\infty$ and $\overline{a}_i = +\infty$
- ii) Insert at intersection point 1 the following block:

$$\begin{aligned} \frac{dJ}{db}(b^k) &= \Theta \nabla J(a^k) \\ \nabla J(a^k) &= \Theta \frac{dJ}{db}(b^k) \end{aligned}$$

(ii) Insert at insertion point 2 the block

$$b^{k+1} = b^k - \rho^k \frac{dJ}{db}(b^k)$$

This remark has two practical consequences:

- When writing an optimization program, take as unknown all the parameters with respect to which an optimization may be performed. The selection of the parameters with respect to which the optimization is effectively performed is then easily done with an adequate matrix Θ : if for instance only the p first parameters $a_1, \dots, a_p$ are unknown ($p \leq N$) and the $N-p$ parameters $a_{p+1}, \dots, a_N$ are known, we take

$$\Theta = \begin{bmatrix} 1 & & & 0 \\ 0 & 1 & & 0 \\ & & 1 & \\ 0 & & & 0 \\ \vdots & & & \vdots \\ 0 & & & 0 \end{bmatrix} \begin{matrix} \left. \vphantom{\begin{bmatrix} 1 & & & 0 \\ 0 & 1 & & 0 \\ & & 1 & \\ 0 & & & 0 \\ \vdots & & & \vdots \\ 0 & & & 0 \end{bmatrix}} \right\} p \\ \left. \vphantom{\begin{bmatrix} 1 & & & 0 \\ 0 & 1 & & 0 \\ & & 1 & \\ 0 & & & 0 \\ \vdots & & & \vdots \\ 0 & & & 0 \end{bmatrix}} \right\} N-p \end{matrix} \quad \tilde{a} = \begin{bmatrix} 0 \\ \vdots \\ 0 \\ a_{p+1} \\ \vdots \\ a_N \end{bmatrix}$$

so that the modifications mentioned in the points (i) and (ii) of remark 5 resume in the insertion at point 1 of:

$$\frac{\partial J(a^k)}{\partial a_i} = 0 \quad i = p+1, \dots, N$$

- When the unknown parameter is a function, take as numerical unknown vector $a = (a_1, \dots, a_N)$ when writing the identification program some discrete approximation of the function. It will make easier any change in the choice of a closed-form formula for representing that

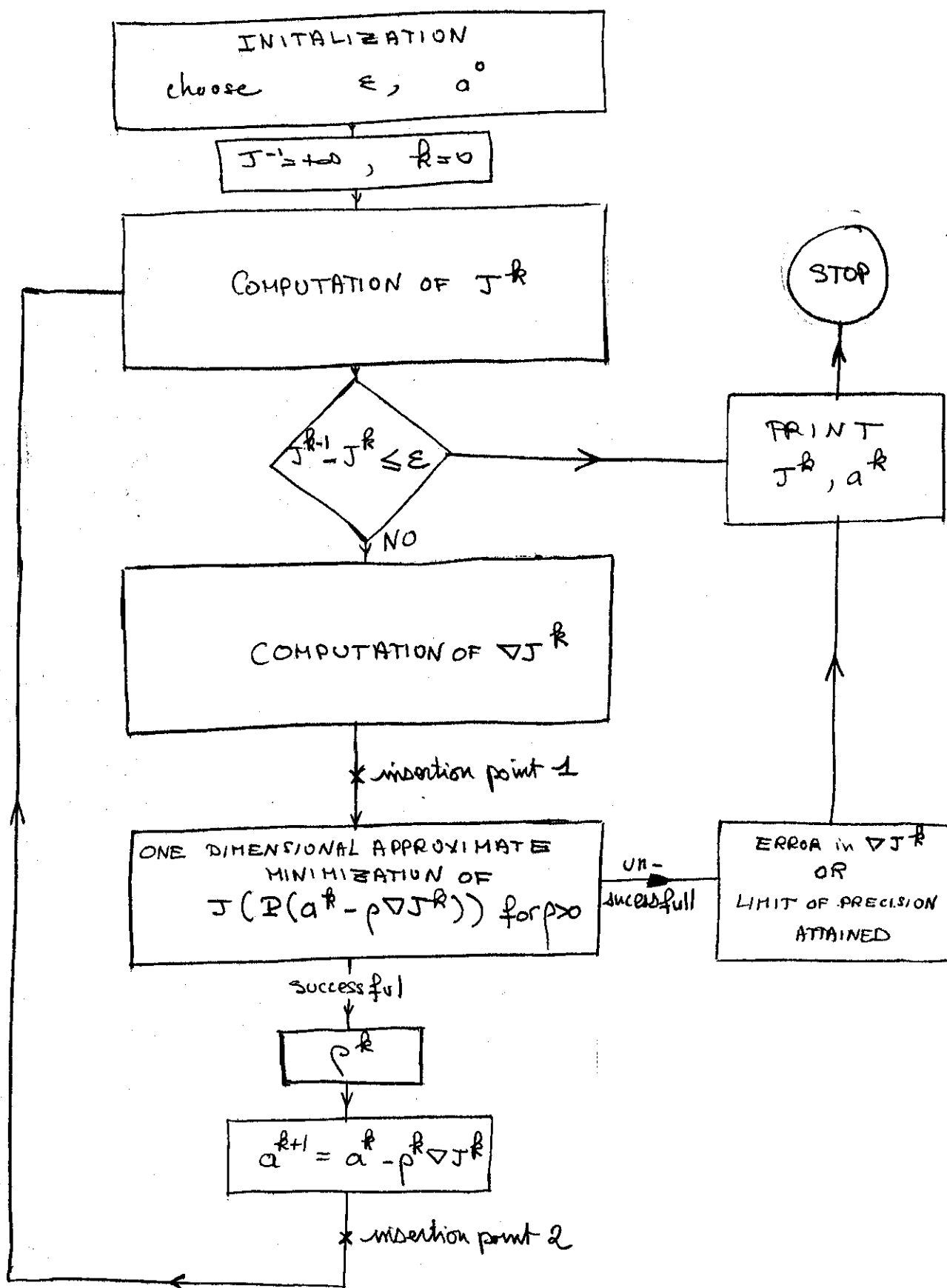


Fig 6 The Steepest Descent Algorithm with projection

In the following paragraphs, we shall study the blocks
"Computation of J^k " and "computation of ∇J^k " when J
is a function defined through the resolution of a static
equation, as in the problem (I) of §(I).

