



AVERTISSEMENT

Ce document est le fruit d'un long travail approuvé par le jury de soutenance et mis à disposition de l'ensemble de la communauté universitaire élargie.

Il est soumis à la propriété intellectuelle de l'auteur. Ceci implique une obligation de citation et de référencement lors de l'utilisation de ce document.

D'autre part, toute contrefaçon, plagiat, reproduction illicite encourt une poursuite pénale.

Contact : ddoc-theses-contact@univ-lorraine.fr

LIENS

Code de la Propriété Intellectuelle. articles L 122. 4

Code de la Propriété Intellectuelle. articles L 335.2- L 335.10

http://www.cfcopies.com/V2/leg/leg_droi.php

<http://www.culture.gouv.fr/culture/infos-pratiques/droits/protection.htm>

Thèse présentée à l'UNIVERSITÉ DE METZ

pour l'obtention du

Doctorat en Mathématiques
mention Mathématiques Appliquées

par

Mr Vincent LÉCUYER
professeur Agrégé

ÉTUDE ANALYTIQUE ET NUMÉRIQUE
DE PROBLÈMES OSCILLANTS
EN CALCUL DES VARIATIONS

Thèse soutenue le 10 Novembre 2000 devant les Professeurs :

M. CHIPOT, Université de Zürich..... *directeur de thèse*

B. DACOROGNA, École Polytechnique Fédérale de Lausanne }
J.-L. MALLET, École de Géologie de Nancy..... } *rapporteurs*

B. BRIGHI, Université de Mulhouse..... }
J.-L. GREFFE, École des Mines de Nancy } *examineurs*
A. HENROT, Institut Élie Cartan de Nancy }

à Maman

BIBLIOTHEQUE UNIVERSITAIRE - METZ	
N° inv.	2000 134 S
Cote	S/M ₃ 00/44
Loc	Magasin

REMERCIEMENTS

Je tiens à adresser mes remerciements les plus chaleureux à M. Chipot pour son enthousiasme et la confiance qu'il a toujours su me témoigner. Au travers de ce travail, initié en particulier par de nombreuses et fructueuses discussions avec B. Brighi, il a su me faire découvrir une application du Calcul des Variations, et de l'Analyse Fonctionnelle plus généralement, à un domaine riche et surprenant, celui de la gémellation des phases cristallines.

Je remercie également très sincèrement Mrs B. Dacorogna et J.-L. Mallet de m'avoir fait l'honneur d'accepter d'être les rapporteurs de ce travail, ainsi que Mrs J.-L. Greffe et A. Henrot pour avoir amicalement accepté de faire partie du Jury.

Je veux encore remercier tout le groupe de développement du logiciel GOCAD® de l'École de Géologie de Nancy pour avoir permis les premières visualisations graphiques des résultats numériques.

Je remercie enfin ma famille et mes amis, et en particulier mon épouse pour son soutien et ses encouragements infaillibles.

TABLE DES MATIÈRES

Page	Titre
	Remerciements
1	Table des matières
2	I - INTRODUCTION
3	I.1. Origine Physique du Problème
7	I.2. L'exemple du Fer (Transition Cubique ↔ Quadratique)
10	II - MATHÉMATISATION DU PROBLÈME
11	II.1. Modèle Mathématique
13	II.2. Énoncé et Discrétisation du Problème
16	II.3. Mesure de Young
17	II.4. Rang-1 Convexité et Quasi-Convexité
19	II.5. L'Exemple choisi
26	III - SIMULATION ET ÉTUDE NUMÉRIQUE
27	III.1. Présentation des Méthodes de Descente
29	III.2. Description du Logiciel développé pour cette Étude
42	III.3. Mise en Œuvre du Logiciel et Résultats Numériques
57	IV - ESTIMATIONS DE L'ÉNERGIE
58	IV.1. Estimation de l'Énergie en Dimension 2
64	IV.2. Application Numérique
69	IV.3. Un Exemple en Dimension 3
80	CONCLUSION
82	Index des tableaux
83	Références bibliographiques

I - INTRODUCTION

I.1. ORIGINE PHYSIQUE DU PROBLÈME

Un cristal est un solide dont la structure atomique est ordonnée et périodique dans les trois directions de l'espace. Il peut donc être considéré comme un réseau de points (ou nœuds), qui reproduit périodiquement un motif élémentaire (ou maille).

Les cristaux présentent des éléments de symétrie (centres, plans ou axes) d'ordre 2, 3, 4 ou 6. L'étude de ces symétries a montré qu'il existe 7 systèmes cristallins caractérisés par leur maille élémentaire (cf tableaux 1 et 2).

Sous certaines conditions (changement de température, de pression, excitation électrique pour les cristaux liquides), le réseau du cristal peut subir un réarrangement : c'est la transition de phase, pendant laquelle les molécules changent de position. Si cette transition de phase s'accompagne d'un changement des symétries du cristal, cette transition est plus "traumatisante" pour le cristal. Dans ce cas, la phase la plus symétrique est appelée Austénite, et la phase la moins symétrique Martensite.

Considérons, par exemple, une transition de phase liée à la température du cristal. La position de référence des molécules du cristal est choisie en phase austénitique à la température critique de transition ($\theta = \theta_c$). On cherche les gradients de déformation minimisant l'énergie potentielle de cette déformation.

Dans la phase austénitique (généralement : $\theta > \theta_c$), la déformation naturelle du cristal (donc celle qui assure une énergie potentielle de déformation minimale) est une homothétie (dilatation thermique), dont le gradient est de la forme

$$U(\theta) = k(\theta).I, \text{ où } I = \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix}.$$

La référence (déformation nulle) étant prise à la température de transition, on a, en particulier :

$$k(\theta_c) = 1 \text{ et donc } U(\theta_c) = I.$$

Par ailleurs, tant que $\theta \geq \theta_c$, $U(\theta)$, et en particulier $U(\theta_c)$, sont invariants par le groupe des symétries de la phase austénitique.

Dans la phase martensitique (généralement : $\theta < \theta_c$), le gradient de la déformation est noté $U'(\theta)$. $U'(\theta)$, et en particulier $U'(\theta_c)$, sont invariants par le groupe des symétries de la phase martensitique.

Mais, la transition s'accompagne ici d'une diminution des symétries du cristal dans un rapport n égal au rapport entre les ordres respectifs des groupes de symétries des phases austénitique et martensitique. L'application du groupe des symétries de la phase austénitique à $U'(\theta_c)$ n'est donc pas univalente, mais conduit à un ensemble de n gradients $\{w_i, 1 \leq i \leq n\}$ possédant de nombreuses propriétés, dont celle de contenir $U(\theta_c) = I$ dans leur enveloppe convexe.

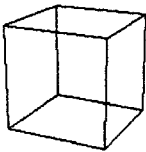
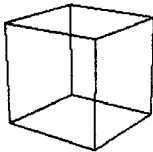
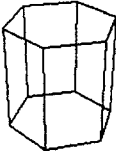
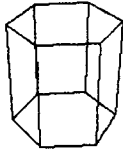
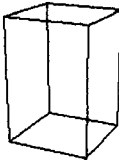
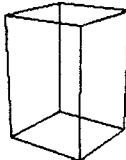
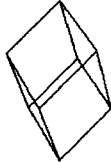
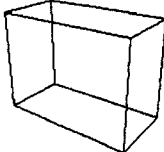
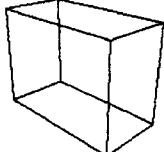
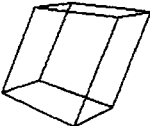
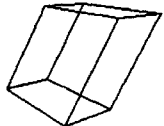
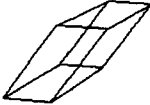
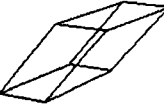
type	maille	stéréogramme	
cubique <i>cubic</i>	cube		
hexagonal <i>hexagonal</i>	prisme droit à base losange dont un angle mesure 120°		
quadratique <i>tetragonal</i>	prisme droit à base carrée		
rhomboédrique <i>trigonal</i>	rhomboèdre		
orthorhombique <i>orthorhombic</i>	parallélépipède rectangle		
monoclinique <i>monoclinic</i>	prisme oblique à base rectangle		
triclinique <i>triclinic</i>	parallélépipède scalène		

tableau 1 : les 7 structures cristallines tri-dimensionnelles

type	Éléments de G	Ordre de G
cubique <i>cubic</i>	Id, $\text{Rot}(\pm\pi/2, e_i)$, $\text{Rot}(\pi, e_i)$ pour $1 \leq i \leq 3$, $\text{Rot}(\pi, e_i \pm e_j)$ pour $i \neq j$, $\text{Rot}(\pm 2\pi/3, e_1 \pm e_2 \pm e_3)$	24
hexagonal <i>hexagonal</i>	Id, $\text{Rot}(\pi, e_i)$ pour $1 \leq i \leq 3$, $\text{Rot}(\pi, e_1 \pm \sqrt{3}.e_2)$, $\text{Rot}(\pi, \sqrt{3}.e_1 \pm e_2)$, $\text{Rot}(\pm\pi/3, e_3)$, $\text{Rot}(\pm 2\pi/3, e_3)$	12
quadratique <i>tetragonal</i>	Id, $\text{Rot}(\pi, e_i)$ pour $1 \leq i \leq 3$, $\text{Rot}(\pm\pi/2, e_3)$, $\text{Rot}(\pi, e_1 \pm e_2)$	8
rhomboédrique <i>trigonal</i>	Id, $\text{Rot}(\pi, e_1)$, $\text{Rot}(\pm 2\pi/3, e_3)$, $\text{Rot}(\pi, e_1 \pm \sqrt{3}.e_2)$	6
orthorhombique <i>orthorhombic</i>	Id, $\text{Rot}(\pi, e_i)$ pour $1 \leq i \leq 3$	4
monoclinique <i>monoclinic</i>	Id, $\text{Rot}(\pi, e_1)$	2
triclinique <i>triclinic</i>	Id	1

tableau 2 : groupe des symétries des 7 structures cristallines tri-dimensionnelles

I.2. L'EXEMPLE DU FER (TRANSITION CUBIQUE ↔ QUADRATIQUE)

Dans le fer, la phase austénitique possède la structure *cubique centrée* (cristal cubique, dont le groupe de symétries est d'ordre 24), alors que la phase martensitique est *cubique à face centrée* (cristal quadratique, dont le groupe de symétries est d'ordre 8). D'après ce qui a été dit, la transition de phase doit engendrer 3 variants "jumeaux" de la Martensite. La déformation de la maille cubique centrée en maille cubique à face centrée peut, effectivement, se faire dans les 3 directions principales du cristal, ainsi que le montrent les stéréogrammes suivants :

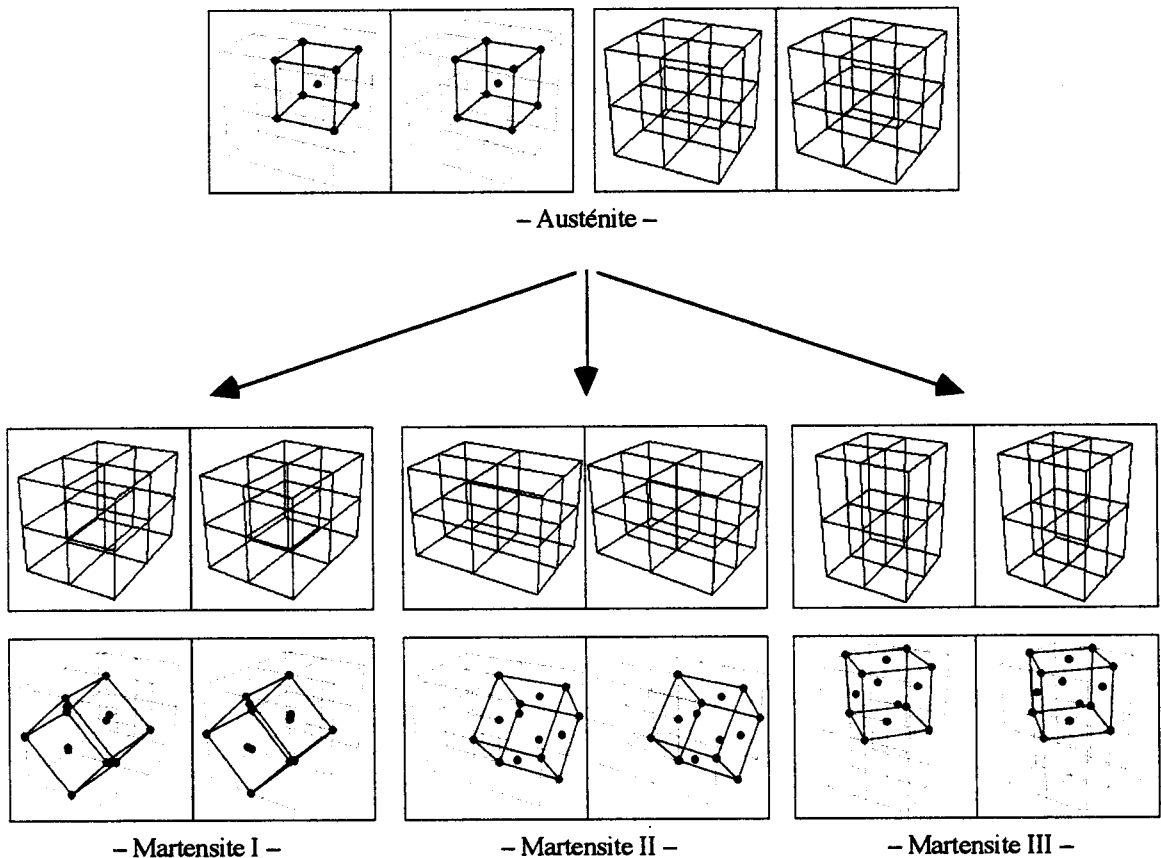


tableau 3 : transition de phase *cubique centrée* vers *cubique à faces centrées*

Si les 3 variants de Martensite n'étaient pas présents dans des proportions identiques (gémellation, ou *twinning*, de la Martensite), le cristal subirait une déformation macroscopique due au changement de proportions de la maille cubique centrée pendant la transition de phase. Les observations microscopiques font apparaître une structure rappelant le dessin d'un tissu, et baptisée "structure en *tweed*", de ces différents variants (illustrée ici dans des alliages où la Martensite possède deux variants : long et mince / petit et large) :

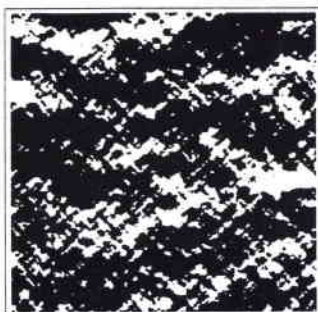


fig. 1 : observation expérimentale de *tweed* au microscope électronique dans un alliage Ni-Al
(photographie : Lee Tanner [22])

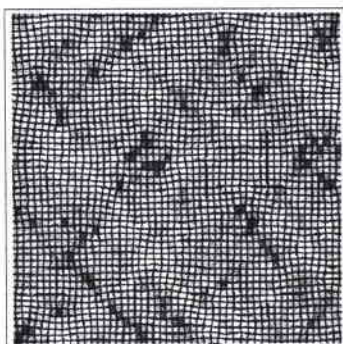


fig. 2 : modélisation d'un alliage Fe-Pd où apparaît une répartition en *tweed* des deux variants martensitiques.

(Jim Krumhansl, Sivan Kartha et James P. Sethna [18])

Dans un système d'axes principaux du cristal, le gradient de la transition de phase de l'Austénite vers l'un des variants de Martensite est donc de l'une des 3 formes suivantes :

$$(\nabla u)_I = \begin{pmatrix} \mu & 0 & 0 \\ 0 & \lambda & 0 \\ 0 & 0 & \lambda \end{pmatrix}, (\nabla u)_{II} = \begin{pmatrix} \lambda & 0 & 0 \\ 0 & \mu & 0 \\ 0 & 0 & \lambda \end{pmatrix}, (\nabla u)_{III} = \begin{pmatrix} \lambda & 0 & 0 \\ 0 & \lambda & 0 \\ 0 & 0 & \mu \end{pmatrix},$$

où les coefficients λ et μ sont tels que :

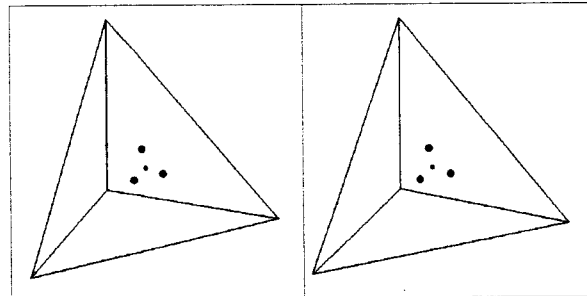
- la masse volumique du cristal soit respectée : $\det(\nabla u) = \det(I)$, soit $\lambda^2 \mu = 1$;
- la maille martensitique soit un cube : $\lambda \cdot \sqrt{2} = \mu$.

Ces deux conditions conduisent aux valeurs de λ et μ :

$$\lambda = 2^{-1/6} \approx 0,891 \text{ et } \mu = 2^{1/3} \approx 1,260.$$

On peut remarquer ici que les 3 gradients $(\nabla u)_I$, $(\nabla u)_{II}$ et $(\nabla u)_{III}$ ne sont pas rang 1-compatibles, *i.e.* la différence entre deux quelconques d'entre eux est une matrice de rang > 1 . Cette remarque est importante pour la suite de l'étude (*cf* remarque 2.3).

Le stéréogramme ci-dessous représente $(\nabla u)_I$, $(\nabla u)_{II}$ et $(\nabla u)_{III}$, ainsi que la matrice identité (puisque $\nabla u = I$ avant la transition de phase), dans les coordonnées de leurs trois coefficients diagonaux.



II - MATHÉMATISATION DU PROBLÈME

II.1. MODÈLE MATHÉMATIQUE

Prenons comme référence le cristal dans sa phase austénitique, au voisinage des conditions de transition de phase, et confiné dans $\mathbb{R}^3$ à l'intérieur d'un ouvert Ω . Soit

$$u : \Omega \rightarrow \mathbb{R}^3$$

la fonction de déformation du cristal pendant la transition de phase, c'est-à-dire que, pour un atome situé à la position x dans la phase austénitique (avant la transition), $u(x)$ représente la nouvelle position de ce même atome dans la phase martensitique (après la transition).

La fonction u est supposée continue et injective. On suppose en outre qu'elle respecte l'orientation de l'espace. Son gradient est noté

$$U = \nabla u = \left(\frac{\partial u_i}{\partial x^j} \right)$$

et le tenseur symétrique positif de la déformation est $C = U^T \cdot U$ (tenseur de Cauchy-Green). L'énergie potentielle associée à la déformation u , due à la transition de phase, est de la forme :

$$J(u) = \int_{\Omega} W(\nabla u(x)) \, dx,$$

où W est une certaine densité d'énergie, possédant, en particulier, les propriétés suivantes :

- invariance dans tout changement de repère direct :

$$\forall R \in O^+, W(R \cdot U) = W(U) ;$$

- invariance par les isométries du groupe des symétries du cristal en phase austénitique :

$$\forall H \in G, W(U \cdot H) = W(U) ;$$

- W est positive et de minimum nul, ce minimum étant atteint en les gradients w_i :

$$\forall R \in O^+, \forall H \in G, W(R \cdot w_i \cdot H) = 0.$$

Remarque 2.1 : les gradients w_i correspondant aux minima de la densité d'énergie potentielle W sont appelés les puits de W .

Exemples de densités d'énergie en dimension 2 : pour des raisons de généralisation du problème, les puits w_i ne seront pas nécessairement des matrices issues de l'étude physique d'un cristal réel, mais des matrices "bien choisies" dans le cadre d'un type de problème variationnel à résoudre. En outre, on pose :

$$W(U) = \psi(C) \text{ et } C_i = w_i^T \cdot w_i$$

- exemple 1 : $w_1 = \begin{pmatrix} 1 - \varepsilon & 0 \\ 0 & 1 + \varepsilon \end{pmatrix}$, $w_2 = \begin{pmatrix} 1 + \varepsilon & 0 \\ 0 & 1 - \varepsilon \end{pmatrix}$, et

$$W(U) = \psi(C) = \left(\left(\frac{C_{11} - C_{22}}{2} \right)^2 - \varepsilon^2 \right)^2$$

- exemple 2 : $\{w_i, 1 \leq i \leq n\}$ sont les n puits souhaités pour W , et

$$W(U) = \psi(C) = \prod_{i=1}^n |U - w_i|^2$$

- exemple 3 : $\{w_i, 1 \leq i \leq n\}$ sont les n puits de W , et

$$W(U) = \psi(C) = \min_{1 \leq i \leq n} |U - w_i|$$

II.2. ÉNONCÉ ET DISCRÉTISATION DU PROBLÈME

Rappelons les notations suivantes : si Ω est un domaine de $\mathbb{R}^n$,

- $W^{1,\infty}(\Omega;\mathbb{R}^n)$ est l'espace de Sobolev des fonctions mesurables $u : \Omega \rightarrow \mathbb{R}^n$ telles que u , ainsi que sa dérivée $\nabla u : \Omega \rightarrow \mathbb{R}^{n \times n}$ soient essentiellement bornées sur Ω ;

- $W_0^{1,\infty}(\Omega;\mathbb{R}^n) = \{ u \in W^{1,\infty}(\Omega;\mathbb{R}^n) / \forall x \in \partial\Omega, u(x) = 0 \}$;

- $W^{1,\infty}(\Omega) = W^{1,\infty}(\Omega;\mathbb{R})$ et $W_0^{1,\infty}(\Omega) = W_0^{1,\infty}(\Omega;\mathbb{R})$.

Problème continu : soit une fonction borélienne W , positive, et possédant k puits w_i , $1 \leq i \leq k$, i.e. :

$$\begin{cases} \forall i \in \{1, \dots, k\}, W(w_i) = 0 \\ \forall F \in M_n(\mathbb{R}) \setminus \{w_i, 1 \leq i \leq k\}, W(F) > 0 \end{cases} \quad (2.1)$$

Il s'agit de chercher la fonction $u \in W_0^{1,\infty}(\Omega;\mathbb{R}^n)$, minimisant l'énergie globale J sur Ω , i.e. telle que :

$$\forall v \in W_0^{1,\infty}(\Omega;\mathbb{R}^n), J(u) \leq J(v).$$

Remarque 2.2 : dans le cadre des déformations d'un cristal, il est nécessaire d'imposer de plus

$$\det(\nabla u) > 0$$

qui correspond à un respect de l'orientation (la matière ne s'"inter-pénètre" pas). Cependant, cette condition est difficile à mettre en œuvre, et sera omise en un premier temps.

Problème discret : en vue d'un traitement numérique, on utilise des éléments finis standards sur le domaine Ω , supposé polygonal. On définit sur Ω une triangulation T par :

- T réalise une partition de $\overline{\Omega}$ en polyèdres de $\mathbb{R}^n$ tous d'intérieur non vide ;
- chaque face d'un élément de T est, soit l'une des faces d'un autre élément de T , soit une partie du bord $\partial\Omega$ de Ω .

Le diamètre d'un polyèdre $K \in T$ est défini comme étant le diamètre de la plus petite boule contenant K . On note h le plus grand des diamètres des polyèdres K de T .

On définit alors l'espace

$$V_0^h = \{ v \in W_0^{1,\infty}(\Omega; \mathbb{R}^n) / \forall K \in T, v|_K \in \mathcal{P}_1(K) \},$$

où $\mathcal{P}_1(K)$ est l'ensemble des fonctions polynomiales de degré 1 sur le polyèdre K . On cherche désormais une fonction $u \in V_0^h$, minimisant l'énergie J , i.e. telle que :

$$\forall v \in V_0^h, J(u) \leq J(v).$$

Lemme 2.1 : soient K_1 et K_2 deux éléments de T ayant une face commune. Si $v \in V_0^h$, les gradients $(\nabla v)_1$ et $(\nabla v)_2$ de v respectivement sur K_1 et sur K_2 sont rang 1-compatibles, i.e. :

$$\text{rg}((\nabla v)_2 - (\nabla v)_1) \leq 1.$$

Preuve : soient K_1 et K_2 deux éléments de T ayant une face commune.

Puisque $v|_{K_1} \in \mathcal{P}_1(K_1)$ et que $v|_{K_2} \in \mathcal{P}_1(K_2)$, il existe des réels $(a_i^j)_{\substack{0 \leq i \leq n \\ 1 \leq j \leq n}}$ et $(b_i^j)_{\substack{0 \leq i \leq n \\ 1 \leq j \leq n}}$ tels

que (avec la convention de sommation d'Einstein) :

$$(i) \forall x \in K_1, \forall j \in \{1, \dots, n\}, w(x) = a_0^j + a_i^j \cdot x^i, \text{ si bien que, sur } K_1, (\nabla v)_1 = (a_i^j)_{\substack{1 \leq i \leq n \\ 1 \leq j \leq n}};$$

$$(ii) \forall x \in K_2, \forall j \in \{1, \dots, n\}, w(x) = b_0^j + b_i^j \cdot x^i, \text{ si bien que, sur } K_2, (\nabla v)_2 = (b_i^j)_{\substack{1 \leq i \leq n \\ 1 \leq j \leq n}};$$

$$(iii) \forall x \in K_1 \cap K_2, \forall j \in \{1, \dots, n\}, w(x) = a_0^j + a_i^j \cdot x^i = b_0^j + b_i^j \cdot x^i, \text{ i.e. } (b_i^j - a_i^j) \cdot x^i = a_0^j - b_0^j.$$

La condition (iii) exprime que l'ensemble des solutions du système linéaire de n équations à n inconnues

$$(b_i^j - a_i^j).x^j = a_0^j - b_0^j, 1 \leq j \leq n,$$

contient $K_1 \cap K_2$, la face commune aux polyèdres K_1 et K_2 . Or, $K_1 \cap K_2$ est, à son tour, un polyèdre contenu dans un hyperplan de $\mathbb{R}^n$, d'intérieur non vide en tant que polyèdre de $\mathbb{R}^{n-1}$, et donc de dimension $(n - 1)$. Ce système linéaire est donc de rang au plus 1, ainsi que sa matrice

$$(b_i^j - a_i^j)_{\substack{1 \leq i \leq n \\ 1 \leq j \leq n}} = (\nabla v)_2 - (\nabla v)_1$$

d'après les conditions (i) et (ii). ■

Remarque 2.3 : la première difficulté de la simulation numérique apparaît maintenant :

- 1) comme il vient d'être dit, les gradients de toute fonction $v \in V_0^h$, et du minimiseur u en particulier, sur 2 éléments contigus de la triangulation, doivent être rang 1-compatibles ;
- 2) puisque la densité d'énergie W est minimale en ses puits w_i , les gradients du minimiseur u sur chaque élément de la triangulation devront être aussi voisins que possible de l'un des puits ;
- 3) or, les puits eux-mêmes ne sont pas, en général, rang 1-compatibles ; c'est le cas, en particulier, de la transition de phase détaillée au paragraphe "I.2. L'EXEMPLE DU FER".

II.3. MESURE DE YOUNG

Il est possible (voir la thèse de M. Collins [15]) de construire des suites minimisantes pour le problème précédent, *i.e.* une suite (u_k) de fonctions de Ω dans $\mathbb{R}^n$ telles que :

$$\lim_{k \rightarrow \infty} J(u_k) = 0.$$

Toutefois, on constate que, si l'énergie potentielle W possède plusieurs puits, on obtient, en général :

$$J(\lim_{k \rightarrow \infty} u_k) > 0.$$

Cette remarque conduit à penser que la suite minimisante fait apparaître des structures fines, laissant à la matrice $\nabla u_k(x)$ la possibilité de "s'approcher" des puits de W , mais que ces structures de plus en plus fines disparaissent à la limite.

On considère alors le gradient de la déformation limite comme une mesure de Young $v_x = v_x(x)$ (voir [1]), définie sur l'ensemble $M = M_n(\mathbb{R})$ des matrices carrées $n \times n$, de sorte que l'on puisse écrire :

$$\lim_{k \rightarrow \infty} \int_{\Omega} W(\nabla u_k(x)) \, dx = \int_{\Omega} \int_M W(A) \, dv_x(A) \, dx \quad (2.2)$$

Et, puisque l'énergie minimale est supposée nulle, il vient :

$$\int_{\Omega} \int_M W(A) \, dv_x(A) \, dx = 0 \quad (2.3)$$

i.e.

$$\text{supp}(v_x) \subset \{w_i\} \quad p.p. \text{ sur } \Omega. \quad (2.4)$$

Remarque 2.4 : il convient cependant de se souvenir que, dans le cas d'un cristal, les structures fines des suites minimisantes ne peuvent être raffinées à l'infini, en raison de la nature discrète du réseau cristallin.

II.4. RANG-1 CONVEXITÉ ET QUASI-CONVEXITÉ

Rappelons les définitions suivantes : si φ est une fonction de $\mathbb{R}^{n \times n}$ dans $\mathbb{R}$,

1) φ est rang 1-convexe si, pour tout couple (A, B) de matrices de $\mathbb{R}^{n \times n}$, rang 1-compatibles (*i.e.* $\text{rang}(A - B) \leq 1$) on a :

$$\forall \lambda \in [0, 1], \varphi(\lambda A + (1 - \lambda)B) \leq \lambda \varphi(A) + (1 - \lambda)\varphi(B);$$

2) φ est quasi-convexe si, pour tout domaine borné D de $\mathbb{R}^n$ et toute fonction $v \in W^{1, \infty}(D; \mathbb{R}^n)$ on a :

$$\forall A \in \mathbb{R}^{n \times n}, \int_D \varphi(A + \nabla v(x)) dx \geq |D| \varphi(A);$$

3) φ est poly-convexe si $\varphi(A)$ est une fonction convexe des mineurs de la matrice A ;

4) φ est convexe si, pour tout couple (A, B) de matrices de $\mathbb{R}^{n \times n}$, on a :

$$\forall \lambda \in [0, 1], \varphi(\lambda A + (1 - \lambda)B) \leq \lambda \varphi(A) + (1 - \lambda)\varphi(B).$$

Rappelons également quelques résultats, qui seront utilisés plus loin :

Proposition 2.1 : pour une fonction $\varphi : \mathbb{R}^{m \times n} \rightarrow \mathbb{R}$, on a les implications suivantes :

$$\varphi \text{ convexe} \Rightarrow \varphi \text{ poly-convexe} \Rightarrow \varphi \text{ quasi-convexe} \Rightarrow \varphi \text{ rang 1-convexe}$$

et

$$\varphi \text{ quasi-convexe} \Leftarrow \varphi \text{ rang 1-convexe} \text{ sauf si } m = n = 2.$$

Preuve : voir, en particulier, la démonstration de la réciproque par Vladimir Šverák [23], et une tentative de contre-exemple dans le cas $m = n = 2$ par Vladimir Šverák et Pablo Pedregal [21].

Corollaire 1.1 : si $n = 1$, convexité $\Leftrightarrow$ rang 1-convexité, et donc les quatre notions coïncident.

Corollaire 1.2 : si on note respectivement φ^{**} , $P\varphi$ et $Q\varphi$ et $R\varphi$ les enveloppes convexe, poly-convexe, quasi-convexe et rang 1-convexe de φ , alors

$$\varphi^{**} \leq P\varphi \leq Q\varphi \leq R\varphi \leq \varphi$$

Proposition 2.2 : soit une fonction $\varphi : \mathbb{R}^{n \times n} \rightarrow \mathbb{R}$. Alors :

$$\forall A \in \mathbb{R}^{n \times n}, Q\varphi(A) = \inf_{v \in W_0^{1,\infty}(\Omega, \mathbb{R}^n)} \frac{1}{|\Omega|} \int_{\Omega} \varphi(A + \nabla v(x)) \, dx \quad (2.5)$$

En particulier, si $A = 0$:

$$Q\varphi(0) = \inf_{v \in W_0^{1,\infty}(\Omega, \mathbb{R}^n)} \frac{1}{|\Omega|} \int_{\Omega} \varphi(\nabla v(x)) \, dx \quad (2.6)$$

Corollaire 2.1 : l'énergie potentielle minimale d'un cristal étant égale à

$$J_{\min} = \inf_{v \in W_0^{1,\infty}(\Omega, \mathbb{R}^n)} \int_{\Omega} \varphi(\nabla v(x)) \, dx,$$

on peut en fait écrire :

$$J_{\min} = |\Omega| \cdot Q\varphi(0). \quad (2.7)$$

II.5. L'EXEMPLE CHOISI

Le principal exemple étudié dans ce travail a déjà été envisagé plusieurs fois dans d'autres travaux. Il s'agit d'un problème bi-dimensionnel ($n = 2$), faisant intervenir une densité d'énergie à 4 puits non 2 à 2 rang 1-compatibles :

$$w_1 = \begin{pmatrix} 1 & 0 \\ 0 & 2 \end{pmatrix}, w_2 = \begin{pmatrix} 2 & 0 \\ 0 & -1 \end{pmatrix}, w_3 = \begin{pmatrix} -1 & 0 \\ 0 & -2 \end{pmatrix}, w_4 = \begin{pmatrix} -2 & 0 \\ 0 & 1 \end{pmatrix}. \quad (2.8)$$

Pour visualiser la rang 1-incompatibilité de ces 4 puits, on les représente dans le plan de leurs coefficients diagonaux :

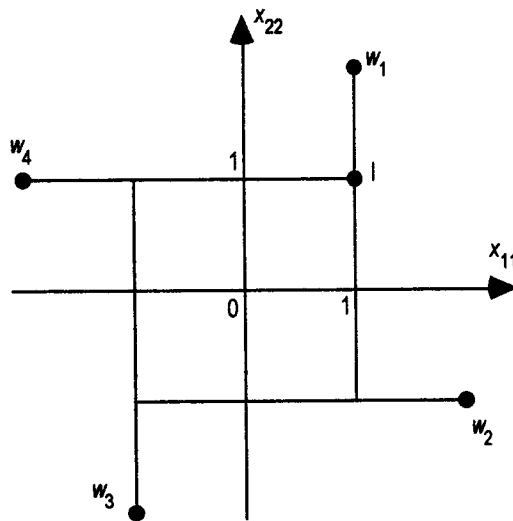


fig. 3 : représentation des 4 puits du problème

Sur une telle représentation, deux matrices diagonales rang 1-compatibles sont alignées sur une droite parallèle à l'un des deux axes de coordonnées. On y observe ainsi, par exemple, la rang 1-compatibilité de la matrice identité I avec w_1 et avec w_4 .

Les raisons qui ont guidé le choix de cet exemple sont les suivantes :

- comme on l'a vu, les puits d'énergie w_i ne sont pas rang 1-compatibles, ce qui correspond à la situation physique développée en introduction ;
- on connaît la valeur de l'infimum d'énergie (théorème 2.1), ainsi que la mesure de Young minimisante (théorème 2.3) ;
- on sait que cet infimum ne peut être atteint par une fonction de $W_0^{1,\infty}(\Omega; \mathbb{R}^2)$ (théorème 2.2), et donc *a fortiori* par une fonction discrétisée de V_0^h , ce qui ne sera pas sans poser des difficultés dans l'étude numérique.

Théorème 2.1 : dans le cas $n = 2$, si la densité d'énergie W possède les 4 puits (2.8), alors :

$$\inf_{v \in W_0^{1,\infty}(\Omega; \mathbb{R}^2)} \int_{\Omega} W(\nabla v(x)) \, dx = 0.$$

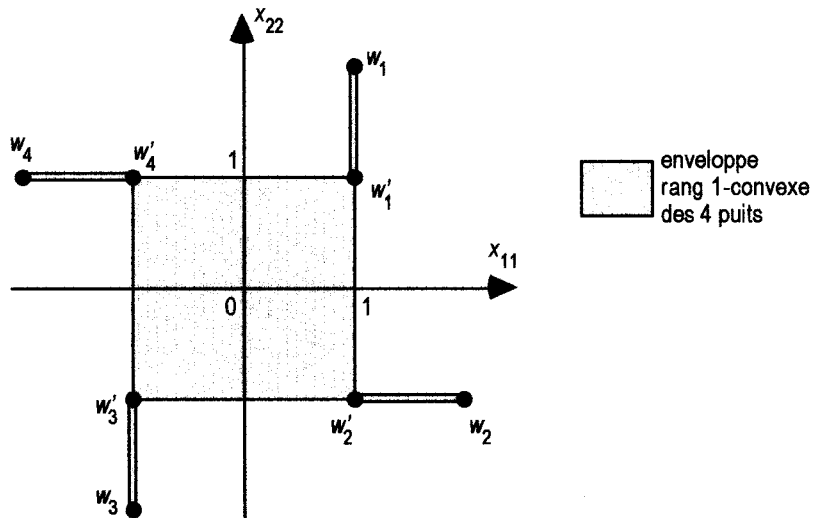


fig. 4 : représentation de l'enveloppe rang 1-convexe des 4 puits

Preuve : on définit 4 autres matrices

$$w'_1 = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}, w'_2 = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}, w'_3 = \begin{pmatrix} -1 & 0 \\ 0 & -1 \end{pmatrix}, w'_4 = \begin{pmatrix} -1 & 0 \\ 0 & 1 \end{pmatrix}. \quad (2.9)$$

D'après la proposition 2.1, on obtient :

- pour tout indice i , $1 \leq i \leq 4$: $0 \leq QW(w_i) \leq W(w_i) = 0$, soit $QW(w_i) = 0$;
- puisque QW est quasi-convexe, elle est également rang 1-convexe.

Ainsi, par exemple :

$$w'_1 = \frac{1}{3} (2w_1 + w'_2) \Rightarrow 0 \leq QW(w'_1) \leq \frac{2}{3} QW(w_1) + \frac{1}{3} QW(w'_2)$$

et, puisque $QW(w_1) = 0$:

$$0 \leq QW(w'_1) \leq \frac{1}{3} QW(w'_2).$$

Il est facile de vérifier que l'on obtient ainsi :

$$0 \leq QW(w'_1) \leq \frac{1}{3} QW(w'_2) \leq \left(\frac{1}{3}\right)^2 QW(w'_3) \leq \left(\frac{1}{3}\right)^3 QW(w'_4) \leq \left(\frac{1}{3}\right)^4 QW(w'_1),$$

soit en fait :

$$QW(w'_i) = 0, \text{ pour tout } i \in \{1, 2, 3, 4\}.$$

Soit alors un couple de réels $(\alpha, \beta) \in [-1, 1]^2$, et les matrices M_α , N_α et $P_{\alpha, \beta}$ définies respectivement par :

$$M_\alpha = \frac{1+\alpha}{2} w'_1 + \frac{1-\alpha}{2} w'_4 = \begin{pmatrix} \alpha & 0 \\ 0 & 1 \end{pmatrix}$$

$$N_\alpha = \frac{1+\alpha}{2} w'_2 + \frac{1-\alpha}{2} w'_3 = \begin{pmatrix} \alpha & 0 \\ 0 & -1 \end{pmatrix}$$

$$P_{\alpha, \beta} = \frac{1+\beta}{2} M_\alpha + \frac{1-\beta}{2} N_\alpha = \begin{pmatrix} \alpha & 0 \\ 0 & \beta \end{pmatrix}$$

Puisque QW est rang 1-convexe : $\forall (\alpha, \beta) \in [0,1]^2$,

$$\bullet 0 \leq QW(M_\alpha) \leq \frac{1+\alpha}{2} QW(w'_1) + \frac{1-\alpha}{2} QW(w'_4) = 0, \text{ soit } QW(M_\alpha) = 0,$$

$$\bullet 0 \leq QW(N_\alpha) \leq \frac{1+\alpha}{2} QW(w'_2) + \frac{1-\alpha}{2} QW(w'_3) = 0, \text{ soit } QW(N_\alpha) = 0,$$

ce qui permet d'obtenir le résultat pour les matrices $P_{\alpha,\beta}$:

$$\bullet 0 \leq QW(P_{\alpha,\beta}) \leq \frac{1+\beta}{2} QW(M_\alpha) + \frac{1-\beta}{2} QW(N_\alpha) = 0, \text{ soit } QW(P_{\alpha,\beta}) = 0.$$

En particulier : $QW(P_{0,0}) = QW(0) = 0$. La relation (2.7) permet de conclure. ■

Théorème 2.2 (voir [10]) : l'infimum n'est pas atteint dans $W_0^{1,\infty}(\Omega; \mathbb{R}^2)$.

Preuve : par l'absurde : supposons qu'il existe une fonction $u \in W_0^{1,\infty}(\Omega; \mathbb{R}^2)$ telle que

$$\int_{\Omega} W(\nabla u(x)) \, dx = \inf_{v \in W_0^{1,\infty}(\Omega; \mathbb{R}^2)} \int_{\Omega} W(\nabla v(x)) \, dx = 0.$$

Rappelons que :

$$\forall F \in M_2(\mathbb{R}) \setminus \{w_i, 1 \leq i \leq 4\}, W(F) > 0.$$

On en déduit que

$$\nabla u(x) \in \{w_i, 1 \leq i \leq 4\} \quad p.p. \, x \in \Omega.$$

En particulier, les deux coefficients non diagonaux des 4 puits sont nuls, donc :

$$\frac{\partial u_1}{\partial x^2}(x) = \frac{\partial u_2}{\partial x^1}(x) = 0 \quad p.p. \, x \in \Omega.$$

Mais alors les fonctions composantes u_1 et u_2 sont indépendantes respectivement de x^2 et de x^1 : pour un certain $x_0 = (x_0^1, x_0^2)$ quelconque dans $\Omega \cup \partial\Omega$,

$$u_1(x) = u_1(x^1, x_0^2) \quad \text{et} \quad u_2(x) = u_2(x_0^1, x^2).$$

Or, la fonction u s'annule sur $\partial\Omega \neq \emptyset$. Donc :

• pour tout $x \in \Omega$, on choisit un $x_0 \in \partial\Omega$ tel que $x^1 = x_0^1$; pour un tel choix :

$$u_1(x) = u_1(x^1, x_0^2) = u_1(x_0^1, x_0^2) = u(x_0) = 0 ;$$

• on prouve de la même façon que, pour tout $x \in \Omega$:

$$u_2(x) = 0.$$

Donc, en fait, u_1 et u_2 sont identiquement nulles sur Ω , soit finalement :

$$u = 0,$$

puis

$$\nabla u(x) = \begin{pmatrix} 0 & 0 \\ 0 & 0 \end{pmatrix} \text{ p.p. } x \in \Omega,$$

ce qui est absurde puisque $\nabla u(x) \in \{w_i, 1 \leq i \leq 4\}$ p.p. $x \in \Omega$. ■

Remarque 2.5 : on verra plus loin que la déformation nulle $u = 0$ est effectivement un attracteur puissant des suites minimisantes construites numériquement.

Théorème 2.3 (voir [10]) : soit une suite minimisante (u_k) , i.e. telle que

$$\lim_{k \rightarrow \infty} J(u_k) = 0.$$

Alors,

$$u_k \text{ tend vers } 0 \text{ uniformément sur } \Omega. \quad (2.10)$$

De plus, la mesure de Young associée à une telle suite est donnée par :

$$\nu_x = \frac{1}{4} (\delta_{w_1} + \delta_{w_2} + \delta_{w_3} + \delta_{w_4}), \quad (2.11)$$

où δ_F représente la masse de Dirac en $F \in M_2(\mathbb{R})$.

Preuve : puisque l'injection canonique de $W_0^{1,\infty}(\Omega;\mathbb{R}^2)$ dans $C^0(\Omega;\mathbb{R}^2)$ est compacte, il suffit, pour justifier la première affirmation, de prouver que 0 est la seule limite possible pour une suite minimisante.

Soit alors une suite minimisante (u_k) . Quitte à extraire une sous-suite, il existe une fonction u telle que :

$$\begin{cases} \lim_{k \rightarrow \infty} (u_k) = u \in W_0^{1,\infty}(\Omega;\mathbb{R}^2) \text{ uniformément sur } \Omega \\ \lim_{k \rightarrow \infty} (\nabla u_k) = \nabla u \text{ dans } (L^\infty(\Omega))^4 \text{ — faible } *$$

Si ν_x est la mesure de Young associée à (u_k) , alors on rappelle que, pour toute fonction continue φ de $\mathbb{R}^{2 \times 2}$ ou $M_2(\mathbb{R})$ dans $\mathbb{R}$:

$$\lim_{k \rightarrow \infty} \varphi(\nabla u_k(x)) = \varphi(\nabla u(x)) = \int_{M_2(\mathbb{R})} \varphi(A) d\nu_x(A) \text{ p.p. } x \in \Omega \text{ dans } L^\infty(\Omega) \text{ — faible } *.$$

En particulier, pour $\varphi = W$, on sait que

$$\int_{\Omega} \int_{M_2(\mathbb{R})} W(A) d\nu_x(A) dx = 0.$$

Donc $\text{supp}(\nu_x) \subset \{w_1, w_2, w_3, w_4\}$, i.e. il existe 4 fonctions mesurables positives α_i telles que

$$\nu_x = \sum_{i=1}^4 \alpha_i(x) \cdot \delta_{w_i} \text{ p.p. } x \in \Omega,$$

et donc, en fait, pour toute fonction continue φ :

$$\varphi(\nabla u(x)) = \sum_{i=1}^4 \alpha_i(x) \cdot \varphi(w_i) \text{ p.p. } x \in \Omega. \quad (2.12)$$

En particulier, pour $(k,l) \in \{1, 2\}^2$, les fonctions $\varphi_{kl} : A \mapsto A_{kl}$, sont continues. On peut leur appliquer la relation (2.12) :

$$(\nabla u(x))_{kl} = \sum_{i=1}^4 \alpha_i(x) \cdot (w_i)_{kl} \text{ p.p. } x \in \Omega,$$

soit en fait

$$\nabla u(x) = \sum_{i=1}^4 \alpha_i(x) \cdot w_i \text{ p.p. } x \in \Omega, \quad (2.13)$$

et en particulier pour les coefficients non diagonaux :

$$(\nabla u(x))_{12} = (\nabla u(x))_{21} = 0 \text{ p.p. } x \in \Omega.$$

L'argument déjà utilisé pour prouver le théorème 2.2 conduit finalement à

$$u = 0 \text{ sur } \Omega,$$

ce qui finit de prouver l'affirmation (2.10). Mais alors, la relation (2.13) s'écrit aussi :

$$\sum_{i=1}^4 \alpha_i(x) \cdot w_i = 0 \quad p.p. \ x \in \Omega,$$

ou, de manière équivalente sur les coefficients diagonaux des matrices :

$$\begin{cases} \alpha_1(x) + 2\alpha_2(x) - \alpha_3(x) - 2\alpha_4(x) = 0 \\ 2\alpha_1(x) - \alpha_2(x) - 2\alpha_3(x) + \alpha_4(x) = 0 \end{cases} \quad p.p. \ x \in \Omega. \quad (2.14)$$

La fonction constante : $A \mapsto 1$, et le déterminant étant continues, on peut leur appliquer la relation (2.12) :

$$\begin{cases} 1 = \sum_{i=1}^4 \alpha_i(x) \\ \det(\nabla u(x)) = 0 = \sum_{i=1}^4 \alpha_i(x) \cdot \det(w_i) \end{cases} \quad p.p. \ x \in \Omega,$$

soit en fait

$$\begin{cases} \alpha_1(x) + \alpha_2(x) + \alpha_3(x) + \alpha_4(x) = 1 \\ 2\alpha_1(x) - 2\alpha_2(x) + 2\alpha_3(x) - 2\alpha_4(x) = 0 \end{cases} \quad p.p. \ x \in \Omega. \quad (2.15)$$

On vérifie facilement que les systèmes (2.14) et (2.15) conduisent à

$$\alpha_1(x) = \alpha_2(x) = \alpha_3(x) = \alpha_4(x) = \frac{1}{4} \quad p.p. \ x \in \Omega.$$

■

III - SIMULATION ET ÉTUDE NUMÉRIQUE

III.1. PRÉSENTATION DES MÉTHODES DE DESCENTE

S'agissant d'un problème de recherche d'enveloppe quasi-convexe, on applique des méthodes issues de l'optimisation convexe : les méthodes de type "gradient" ou "gradient conjugué" (cf [13]).

Principe : ce sont des méthodes itératives.

On considère une fonctionnelle dérivable $J : \Omega \subset \mathbb{R}^n \rightarrow \mathbb{R}$, et un élément $u_0 \in \Omega$, appelé vecteur initial. On construit, par récurrence, une suite (u_n) de vecteurs de $\mathbb{R}^n$ telle que la suite $(J(u_k))$ soit décroissante :

- connaissant u_k ($k \geq 0$), on détermine une direction de descente $d_k \in \mathbb{R}^n$, et on détermine le minimum de la fonction

$$\varphi : \mathbb{R}^+ \rightarrow \mathbb{R}, \lambda \mapsto J(u_k - \lambda \cdot d_k); \quad (3.1)$$

- si φ atteint son minimum en λ_k , on pose $u_{k+1} = u_k - \lambda_k \cdot d_k$;
- on répète cette récurrence jusqu'à une certaine condition d'arrêt.

Différentes méthodes du gradient : il existe plusieurs façons de choisir la direction de descente d_k .

1) *Gradient à pas optimal* : la direction de descente est donnée par l'opposé du gradient de J en u_k , soit

$$d_k = -\nabla J(u_k).$$

Il s'agit, en effet, de la direction de plus forte descente locale en u_k , mais qui, d'un point de vue plus global, n'est pas optimale.

2) *Gradient conjugué* : pour améliorer cette descente, on combine le gradient de J en u_k avec les directions de descente antérieures en une direction optimisée :

$$d_k = - \left\{ \nabla J(u_k) + \frac{\|\nabla J(u_k)\|^2}{\|\nabla J(u_{k-1})\|^2} \cdot d_{k-1} \right\} \quad (3.2)$$

(méthode de Polak-Ribière)

ou

$$d_k = - \left\{ \nabla J(u_k) + \frac{\langle \nabla J(u_k) - \nabla J(u_{k-1}), \nabla J(u_k) \rangle}{\|\nabla J(u_{k-1})\|^2} \cdot d_{k-1} \right\} \quad (3.3)$$

(méthode de Fletcher-Reeves)

Remarques.

1) Les méthodes de gradient conjugué sont plus efficaces que la méthode de gradient à pas optimal. Par exemple, lorsque la fonctionnelle J est elliptique*, on montre que la méthode de gradient à pas optimal converge, alors que les méthodes de gradient conjugué convergent *en au plus n itérations*.

2) Il a été plusieurs fois constaté empiriquement que la méthode de Polak-Ribière l'emporte sur la méthode de Fletcher-Reeves.

* J est dite *elliptique* si elle est continûment dérivable et si, pour un certain $\alpha > 0$:

$$\langle \nabla J(v) - \nabla J(u), v - u \rangle \geq \alpha \|v - u\|^2, \forall u, v.$$

III.2. DESCRIPTION DU LOGICIEL DÉVELOPPÉ POUR CETTE ÉTUDE

En vue d'une étude numérique du problème discrétisé, un logiciel spécifique a été développé en langage Pascal sur Macintosh™.

Dans un premier menu figurent les commandes de pilotage du logiciel, accompagnées chacune de leur "raccourci-clavier" :

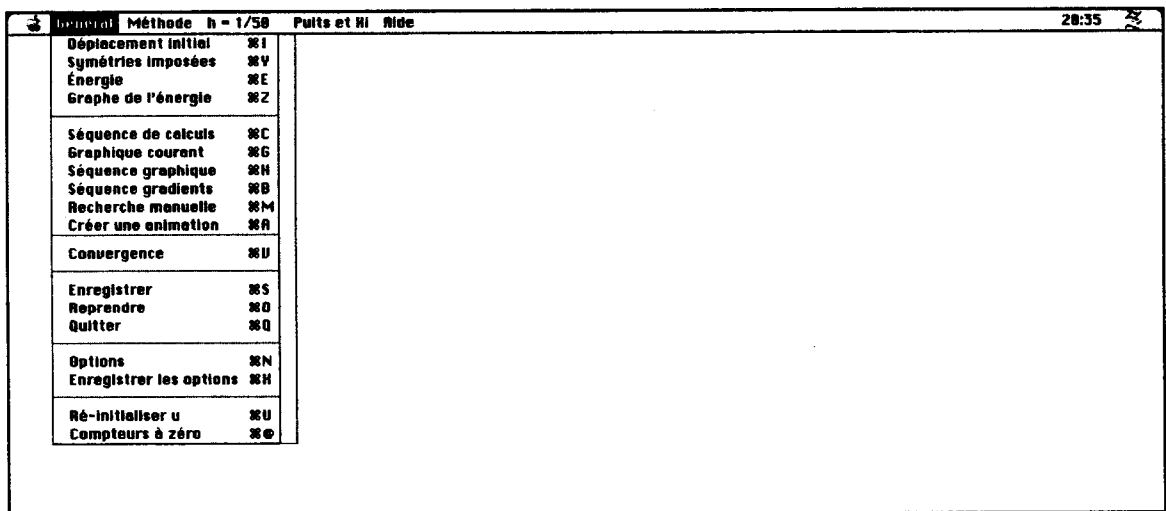


fig. 5 : menu principal

Ce menu est composé de plusieurs blocs.

Le premier comprend 4 commandes permettant d'agir sur tous les paramètres qui vont définir la suite minimisante : déformation initiale, contrainte à certaines symétries, et choix du modèle d'énergie, que complète une commande donnant accès à une visualisation interactive de l'énergie choisie.

Le second bloc permet de contrôler la séquence de calcul : séquence automatique ou manuelle, séquence de calculs, des déformations du cristal ou des représentations des gradients, avec la possibilité de créer une animation Quicktime® de la séquence :

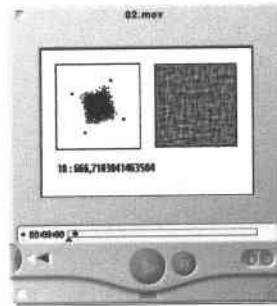


fig. 6 : panneau de contrôle d'une animation Quicktime®

Le bloc suivant ne comprend qu'une seule commande, celle qui permet de visionner la courbe de descente en énergie du cristal.

Enfin, les trois derniers blocs gèrent différentes options de l'interface avec l'utilisateur, et la ré-initialisation de la séquence.

III.2.1. INITIALISATIONS DE LA SÉQUENCE

Dans une première boîte de dialogue, on peut de définir la déformation initiale :

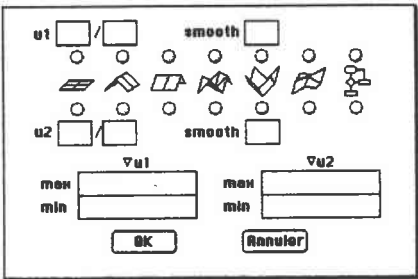


fig. 7 : dialogue d'initialisation

On dispose de 7 possibilités pour chaque composante, éventuellement réglables par l'intermédiaire de 3 paramètres, et dont l'échelle est déterminée en fixant les valeurs extrêmes du gradient correspondant :

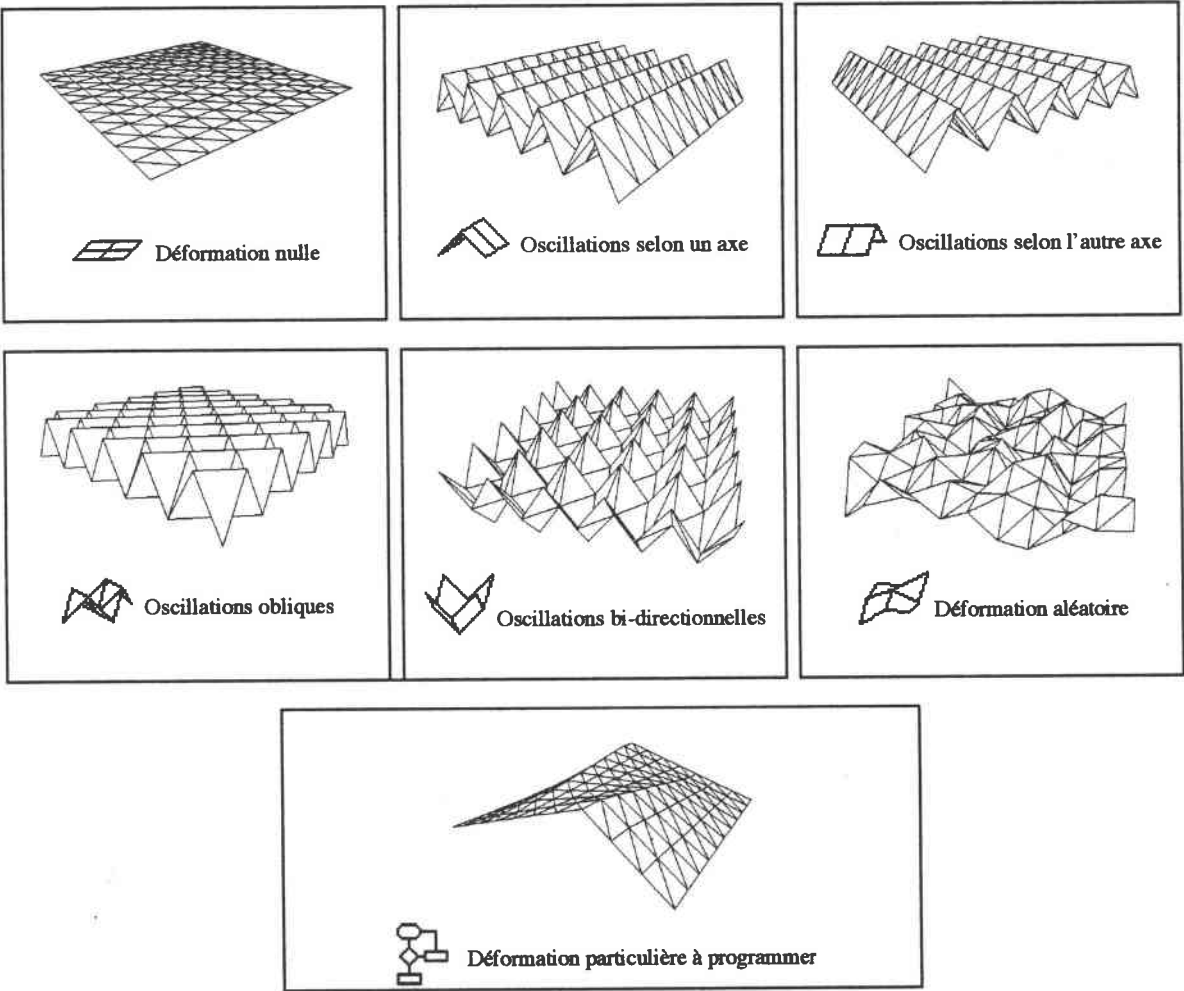


fig. 8 : les différents choix de la déformation initiale

III.2.2. SYMÉTRIES IMPOSÉES

Par ailleurs, il est possible de contraindre la déformation à respecter certaines symétries pendant toute la séquence :

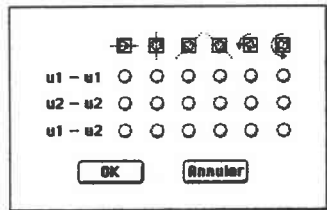


fig. 9 : conditions de symétrie

Il est proposé 6 contraintes, applicables à u_1 , u_2 ou entre u_1 et u_2 , et cumulables entre elles : 4 symétries axiales et 2 rotations.

III.2.3. CHOIX ET VISUALISATION DU MODÈLE D'ÉNERGIE

Enfin, il est possible de choisir l'un des 5 modèles d'énergie suivants. Pour chacun d'eux, il est en outre possible d'utiliser leur logarithme, et de fixer par la suite leurs éventuels paramètres (α , γ , λ et μ) :

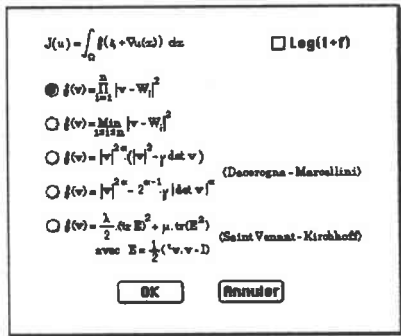


fig. 10 : choix de la fonctionnelle d'énergie

Lorsque le modèle est choisi, il est possible de le visualiser au-dessus du plan des 2 coordonnées diagonales des gradients 2×2 . Cette visualisation est interactive grâce à quelques touches virtuelles, permettant d'agir sur la position du graphique dans 2 directions de l'espace (la direction verticale de l'écran  et une direction en profondeur , mais aussi sur l'échelle de ce graphique , et sur le domaine utilisé ( ). On dispose enfin de la possibilité de voir  ou de masquer  les lignes cachées.

La figure suivante montre le graphe obtenu avec le premier modèle d'énergie pour 4 puits non rang-1 compatibles, auquel on applique la fonction logarithme pour en faire ressortir les extrema. Les lignes cachées sont masquées et le domaine utilisé est $[-3,3]^2$:

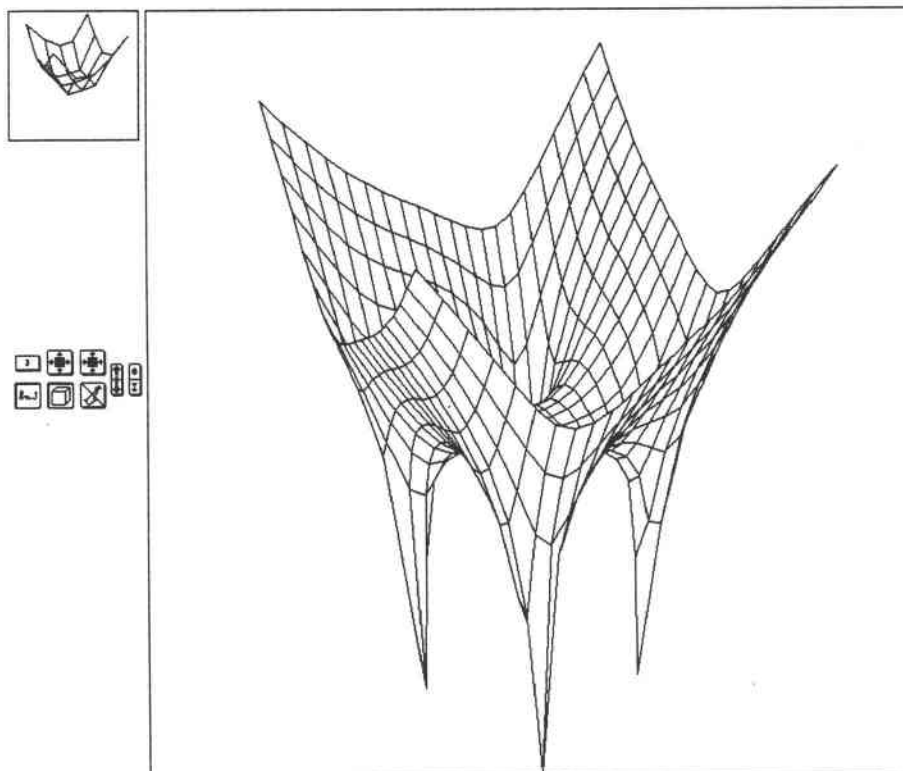


fig. 11 : visualisation de la fonctionnelle d'énergie

III.2.4. AUTRES PARAMÈTRES

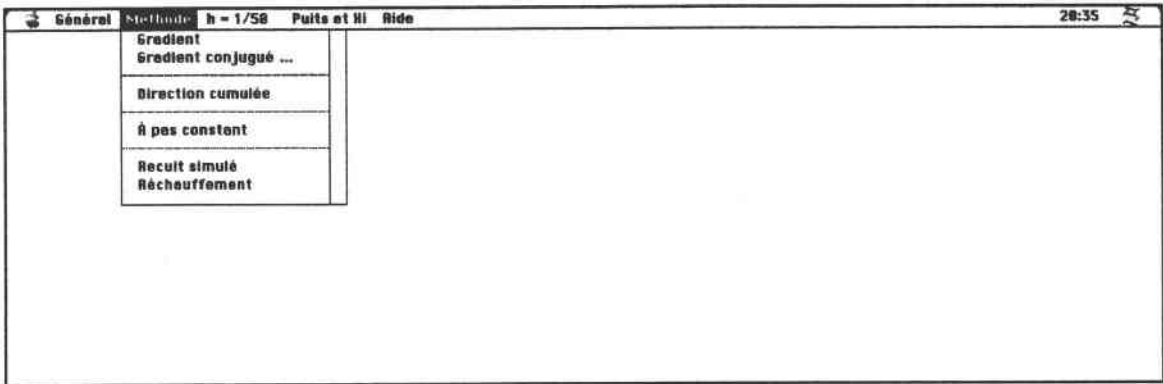


fig. 12 : choix de la méthode

Dans le second menu sont proposées plusieurs options concernant la méthode de calcul : méthodes de gradient à pas optimal ou gradient conjugué, méthode de la direction cumulée (extension du gradient conjugué qui définit la direction de descente à partir d'un certain nombre des directions précédentes, et non plus seulement à partir de la seule direction précédente), ou à pas constant. À toute étape de la séquence, il est par ailleurs possible d'appliquer un recuit simulé (*annealing*) ou un simple "réchauffement" du cristal (perturbation aléatoire et réglable en amplitude de la déformation).

Un troisième menu propose les options de définition du maillage :

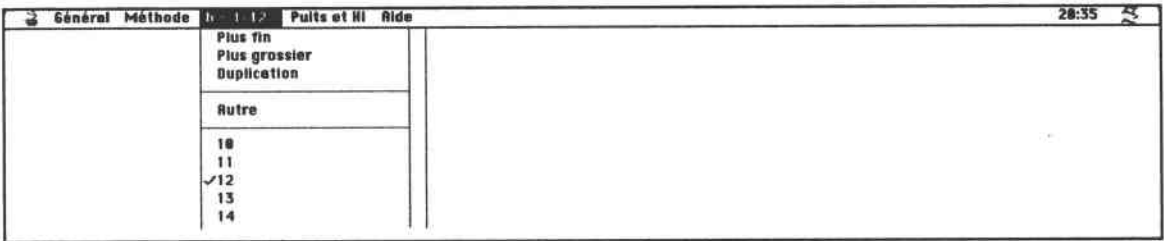


fig. 13 : définition du maillage

Ce menu permet de choisir le nombre de subdivisions sur chaque axe (n , de 10 à 100), mais également trois opérations automatiques applicables à tout moment de la séquence :

- pour affiner le maillage en doublant le nombre de subdivisions par interpolation linéaire ;
- pour alléger le maillage en divisant le nombre de subdivisions par 2 ;
- de dupliquer le maillage par deux translations parallèles aux axes, le maillage final ayant deux fois plus de subdivisions que le maillage initial, et présentant 4 fois le motif de celui-ci.

Enfin, un dernier menu permet de définir le nombre (de 2 à 6) et la position des éventuels puits de la fonctionnelle d'énergie, ainsi que la condition au bord ξ :

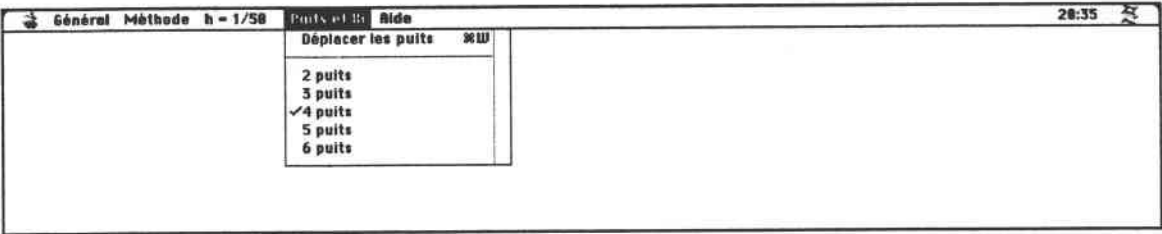


fig. 14 : définition du nombre de puits

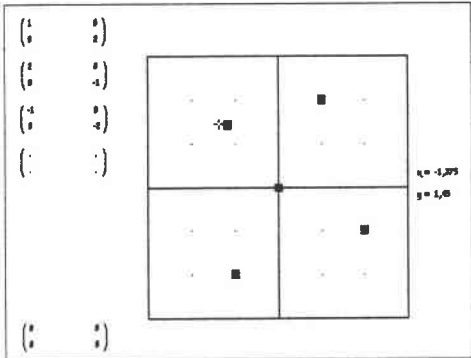


fig. 15 : interface de positionnement des puits et de la condition au bord ξ ; ici, dans le cas de 4 puits, ξ et les trois premiers puits sont déjà fixés, et l'utilisateur déplace le 4^{ème}.

III.2.5. EXEMPLE D'UTILISATION

Voici un exemple montrant les différentes visualisations des résultats qu'offre le logiciel.

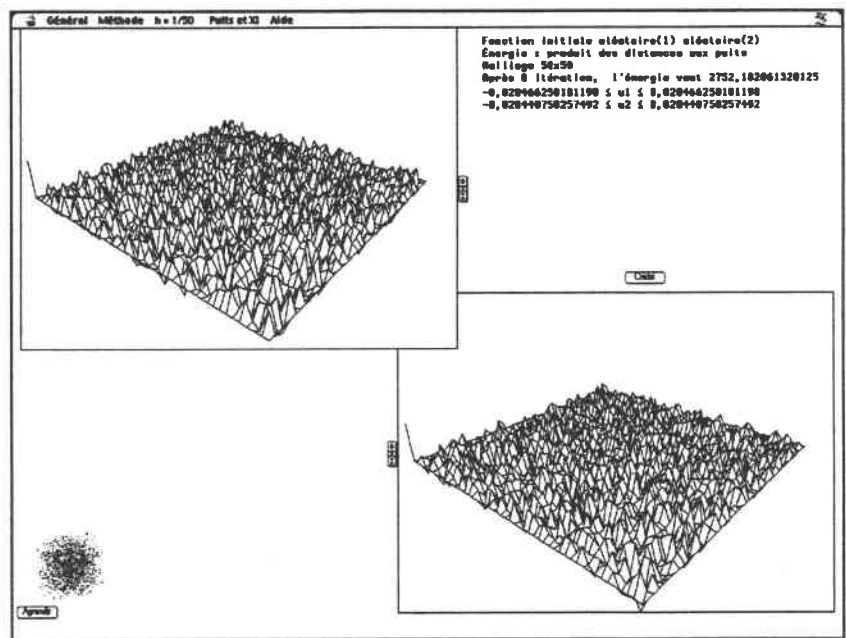


fig. 16 : graphes (manœuvrables en 3D) des deux composantes de la déformation initiale.
Les gradients sont représentés en bas et à gauche dans le plan de leurs 2 coordonnées diagonales.

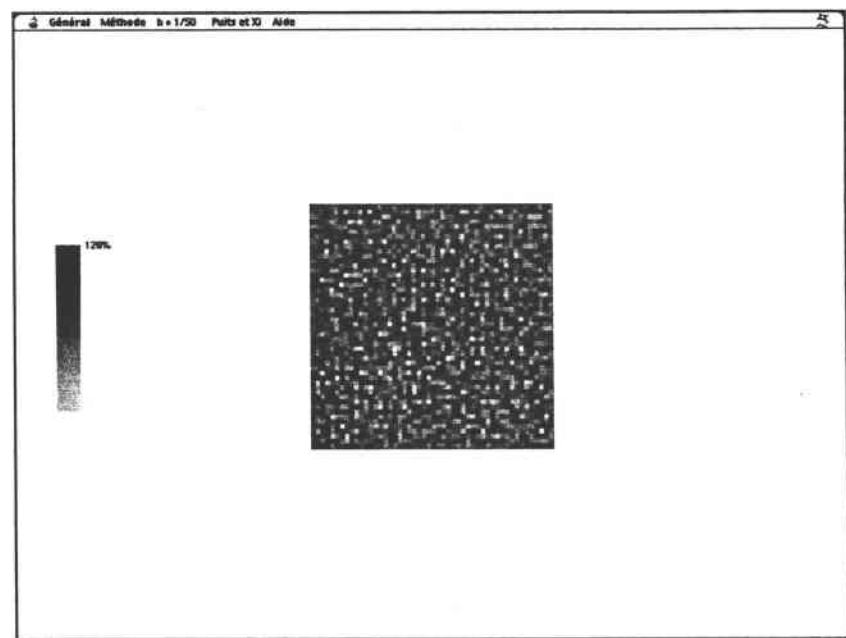


fig. 17 : une représentation en fausses couleurs du cristal correspondant à cette déformation initiale.

Les "fausses couleurs" ou "niveaux de gris" employés rendent compte de la superficie de chacune des mailles du cristal : foncée si elle est dilatée, claire si elle est comprimée :

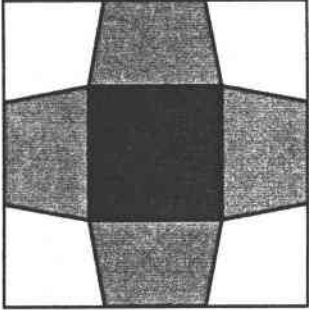
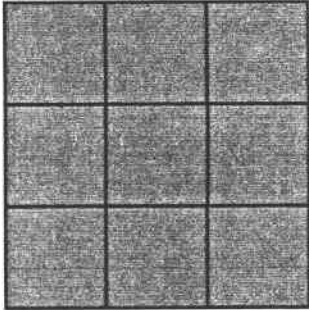
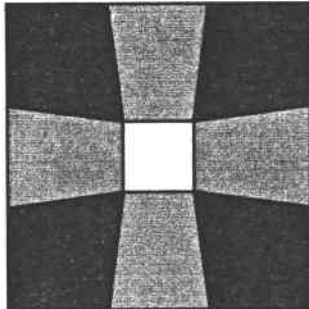
gris foncé	maille dilatée, où le gradient de la déformation est voisin du puits $\begin{pmatrix} 1 & 0 \\ 0 & 2 \end{pmatrix}$	
gris moyen	maille peu déformée, où le gradient de la déformation est voisin de l'un des puits $\begin{pmatrix} 2 & 0 \\ 0 & -1 \end{pmatrix} \text{ ou } \begin{pmatrix} -2 & 0 \\ 0 & 1 \end{pmatrix}$	
gris clair	maille comprimée, où le gradient de la déformation est voisin du puits $\begin{pmatrix} -1 & 0 \\ 0 & -2 \end{pmatrix}$	

fig. 18

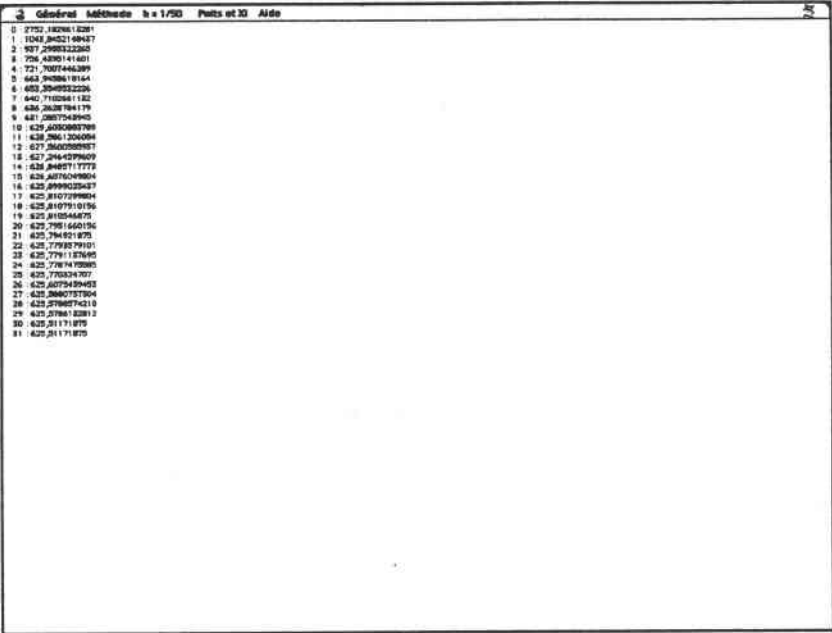


fig. 19 : séquence numérique des énergies successives.

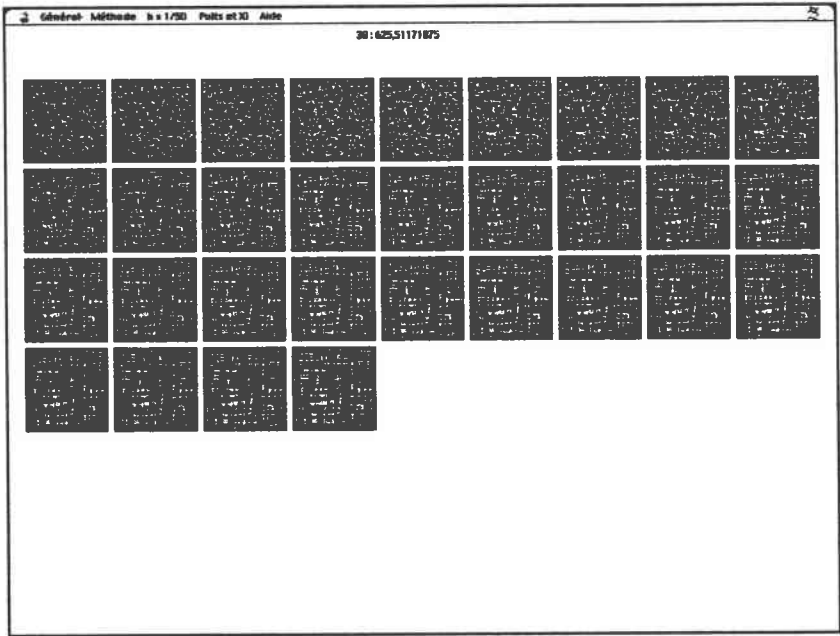
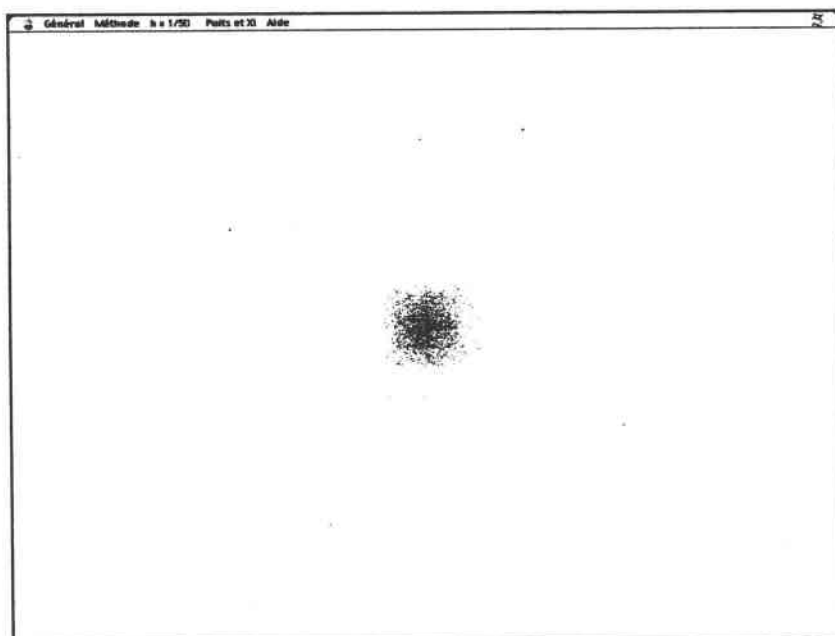
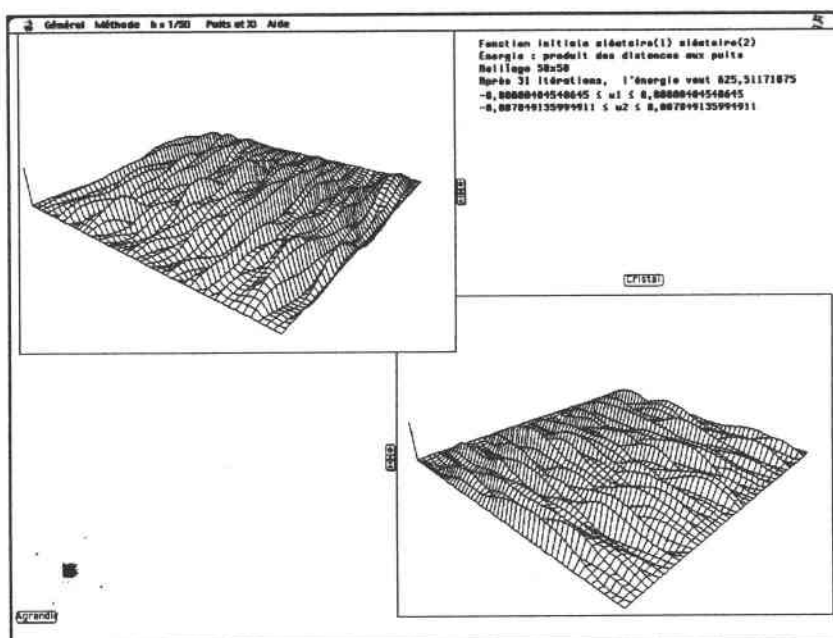


fig. 20 : évolution de la déformation du cristal.



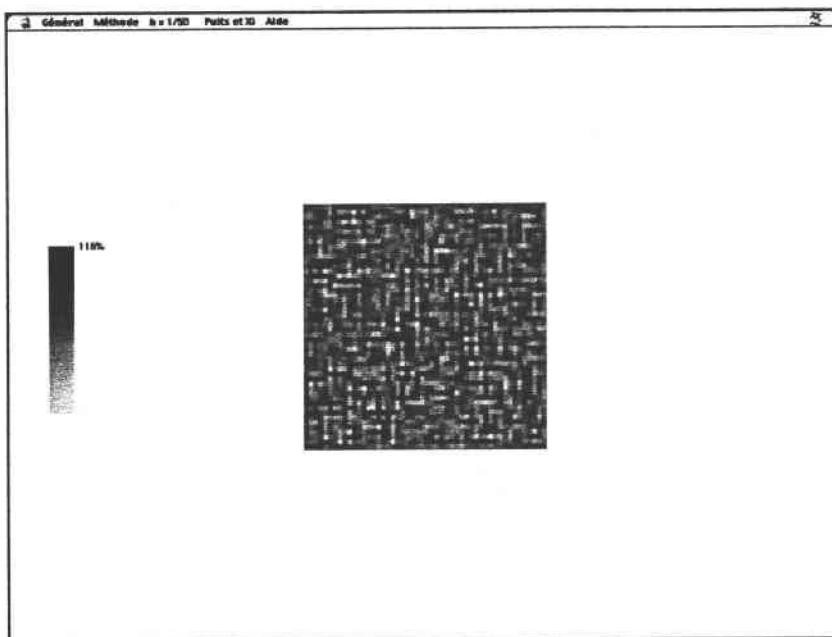


fig. 23 : déformation du cristal après 5 itérations.

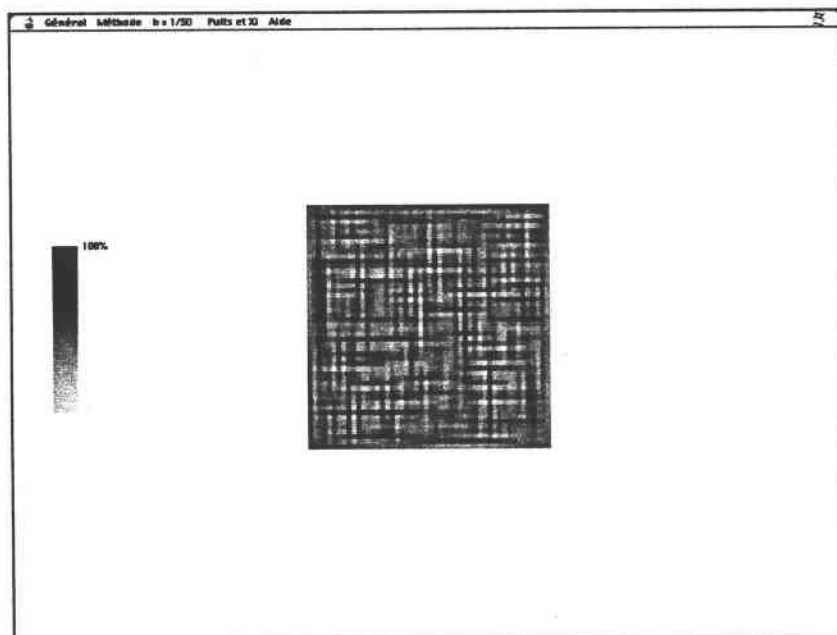


fig. 24 : déformation du cristal après 31 itérations.

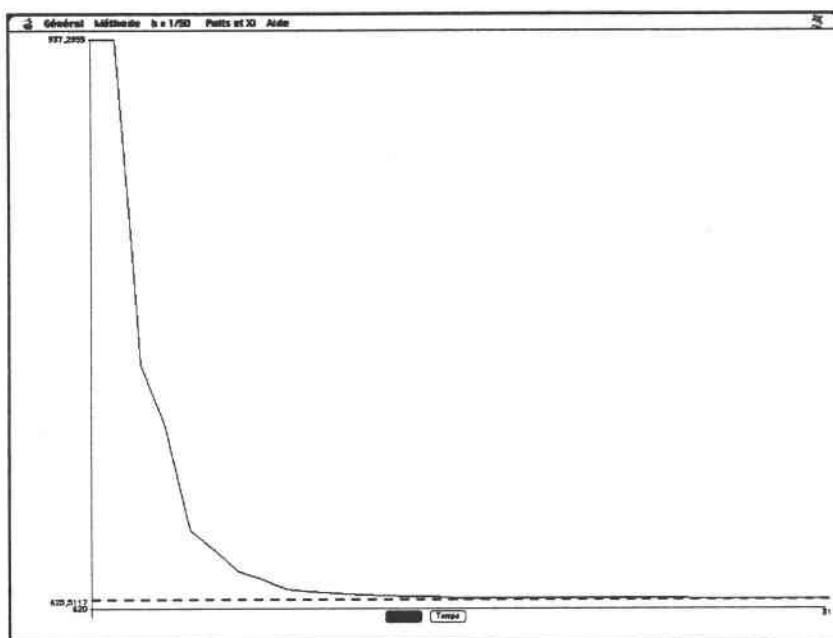


fig. 25: graphe de l'énergie du cristal en fonction du n° de l'itération.

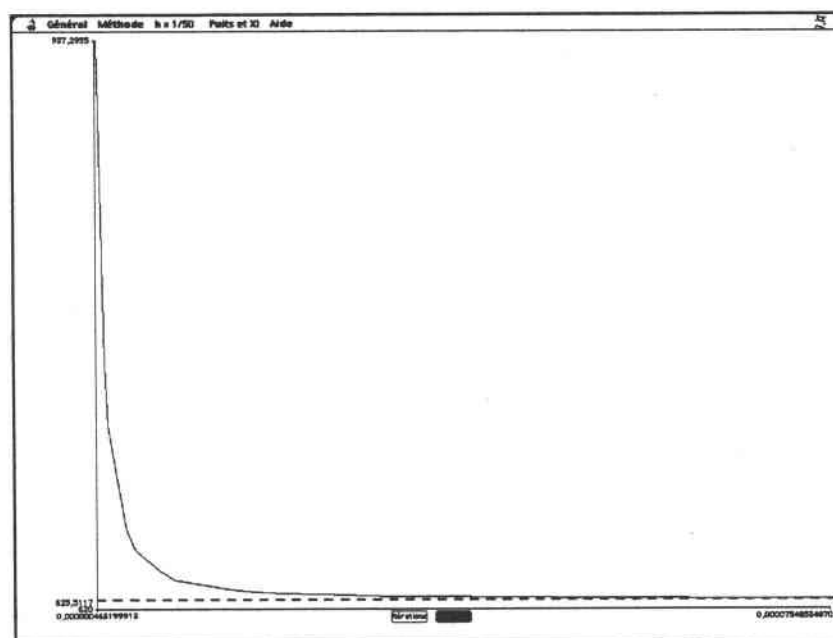


fig. 26 : graphe de l'énergie du cristal en fonction du cumul des pas de descente.

III.3. MISE EN ŒUVRE DU LOGICIEL ET RÉSULTATS NUMÉRIQUES

III.3.1. QUELQUES GÉNÉRALITÉS

Tous les calculs menés dans la suite de cette étude portent sur le cas détaillé au § II.5, à savoir une énergie de la forme

$$W(U) = \psi(C) = \prod_{i=1}^n |U - w_i|^2$$

où les w_i sont les $n = 4$ puits de cette fonctionnelle, choisis comme en (2.8) :

$$w_1 = \begin{pmatrix} 1 & 0 \\ 0 & 2 \end{pmatrix}, w_2 = \begin{pmatrix} 2 & 0 \\ 0 & -1 \end{pmatrix}, w_3 = \begin{pmatrix} -1 & 0 \\ 0 & -2 \end{pmatrix}, w_4 = \begin{pmatrix} -2 & 0 \\ 0 & 1 \end{pmatrix}.$$

Sur les différents essais, la méthode du gradient conjugué de Polak-Ribière a donné systématiquement une descente plus rapide. Elle sera donc toujours utilisée dans les résultats énoncés ici.

III.3.2. PREMIER ÉCUEIL

Pour initialiser la descente, il semblait naturel de choisir une déformation initiale nulle. Le gradient de l'énergie pour une déformation nulle étant malencontreusement nul également, les méthodes de descente conduisent alors à une séquence stationnaire et le cristal conserve sa déformation initiale nulle, dont l'énergie est obtenue facilement : si on note $\mathbf{0} = \begin{pmatrix} 0 & 0 \\ 0 & 0 \end{pmatrix}$, alors

$$\forall i \in \{1, 2, 3, 4\}, \|\mathbf{0} - w_i\|^2 = (1-0)^2 + (0-0)^2 + (0-0)^2 + (2-0)^2 = 5,$$

si bien que :

$$J(\mathbf{0}) = \prod_{i=1}^4 \|\mathbf{0} - w_i\|^2 = 5^4 = \boxed{625} \quad (3.4)$$

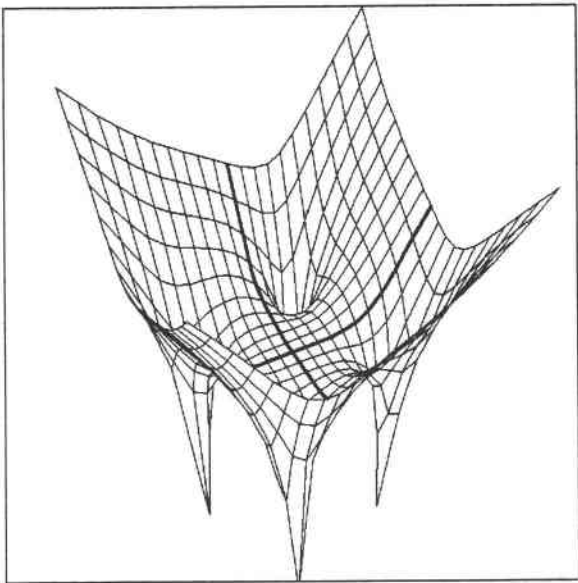


fig. 27 : point-selle de la fonctionnelle
minimisante.

Comme on l’observe sur le graphe de la fonctionnelle d’énergie(cf § 2.3), la déformation nulle (située à l’intersection des lignes en gras) correspond en fait à un point singulier (“point-selle”) de la fonctionnelle. La valeur 625 de l’énergie sera, par conséquent, considérée comme le premier écueil à franchir lors des descentes pour s’approcher de la déformation

Ce résultat est confirmé par le choix d’une déformation initiale aléatoire d’amplitude modérée, qui conduit inévitablement à l’apparition de petites oscillations, qui s’estompent après une centaine d’itérations sur un maillage 50x50 :

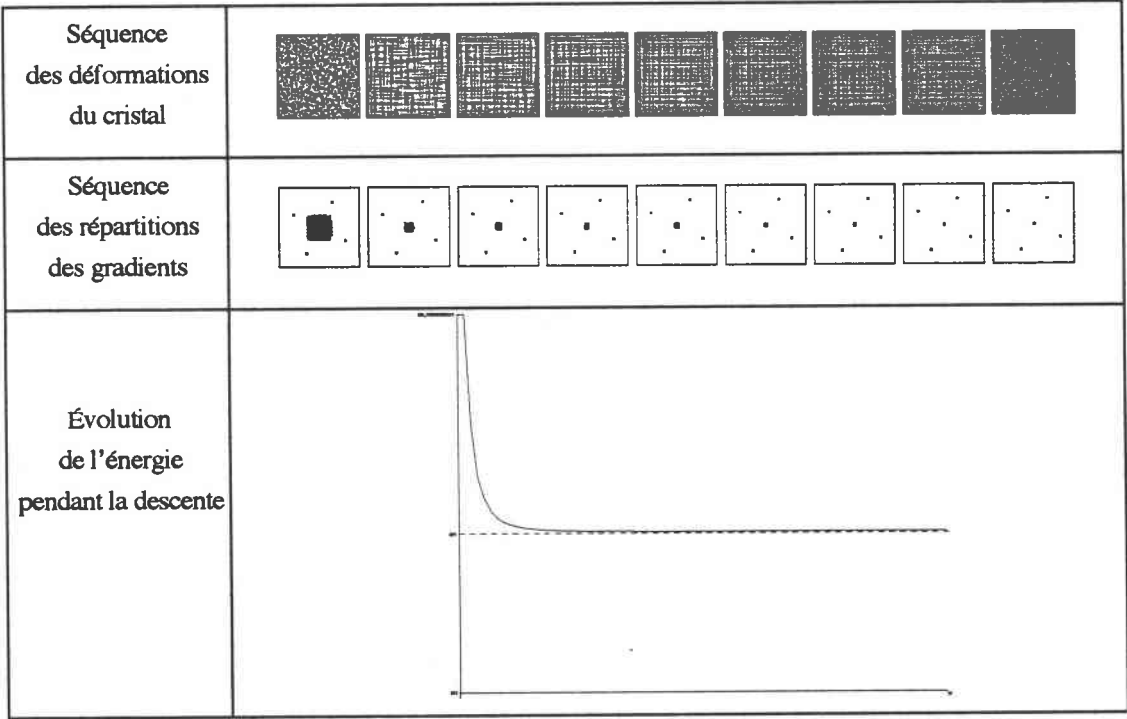


fig. 28 : descente vers la déformation nulle

III.3.3. NÉCESSITÉ DES OSCILLATIONS

Lors des initialisations aléatoires, des oscillations apparaissent, systématiquement orientées perpendiculairement sur les deux composantes de la déformation :

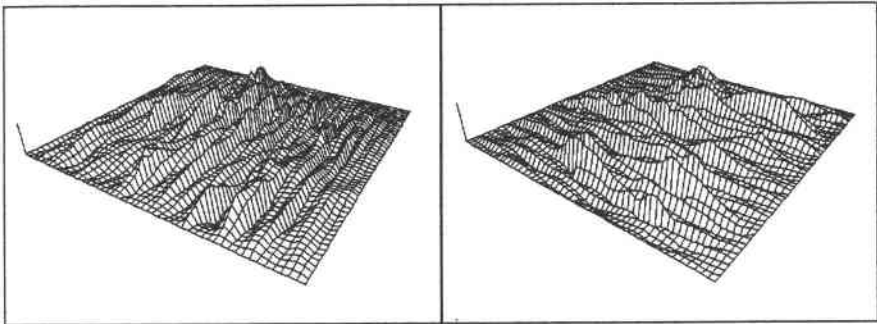


fig. 29 : oscillations issues d’une initialisation aléatoire

Les essais qui suivent confirment cette observation : en induisant des oscillations sur la déformation initiale, sur une ou deux composantes, la séquence conduit invariablement à éliminer les oscillations qui ne respectent pas les orientations précédentes, pour une énergie qui ne descend jamais sous la valeur 625 :

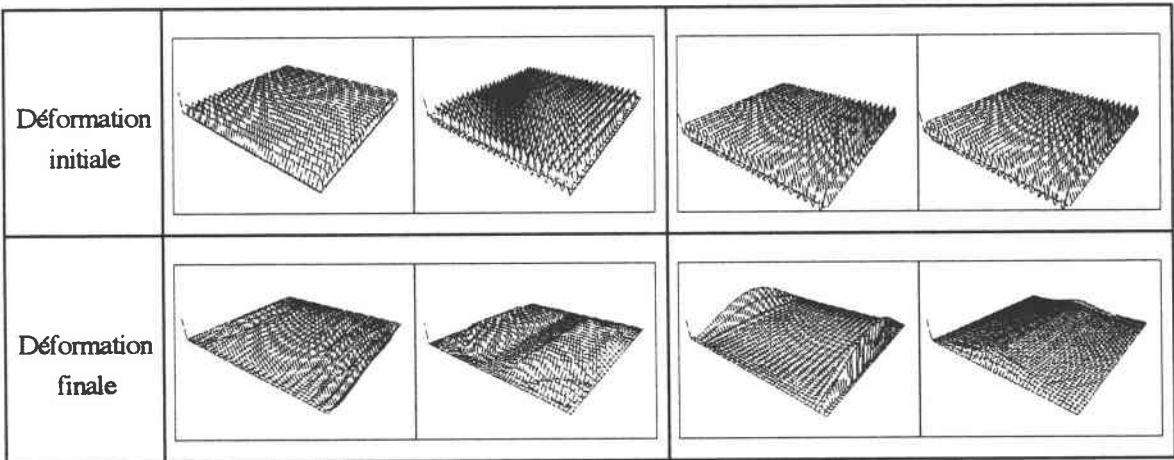


fig. 30 : disparition des oscillations lorsqu’elles sont forcées dans une “mauvaise” direction

Désormais, la déformation initiale comportera des oscillations dans les directions observées sur la figure 29.

III.3.4. ESSAIS DE DIFFÉRENTES OSCILLATIONS INITIALES

Le problème semble avoir une extrême sensibilité vis-à-vis du choix de la déformation initiale. Suivent quelques exemples correspondant à des déformations initiales présentant différents types d'oscillations sur une composante, l'autre étant nulle. L'énergie-limite est systématiquement 625 :

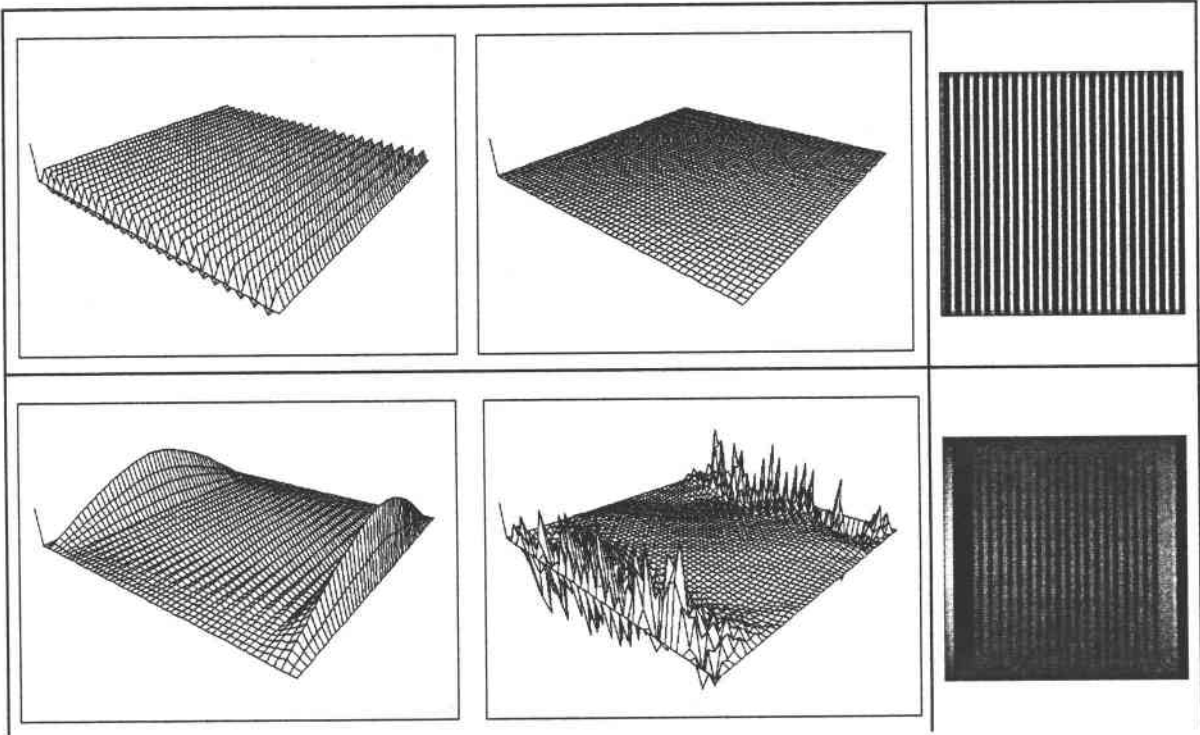


fig. 31 : la déformation finale (50 itérations) est fortement agrandie ($E \approx 625$)

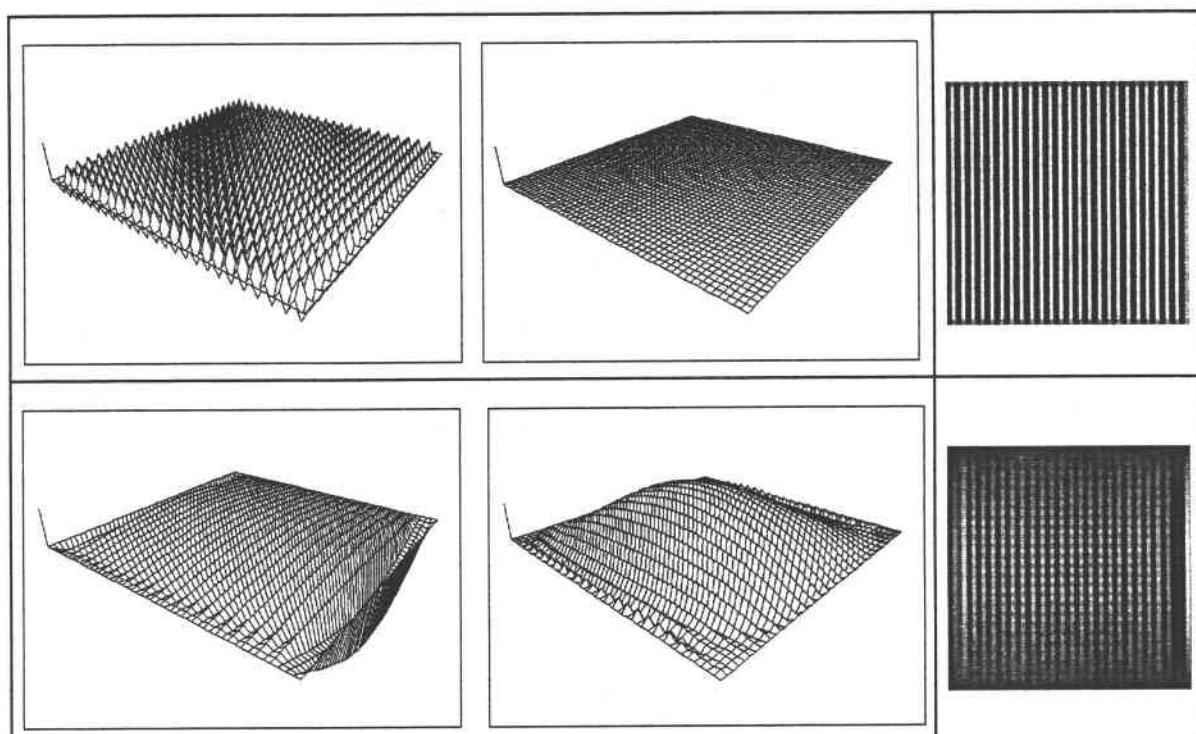


fig. 32 : la déformation finale (70 itérations) est fortement agrandie ($E \approx 625$)

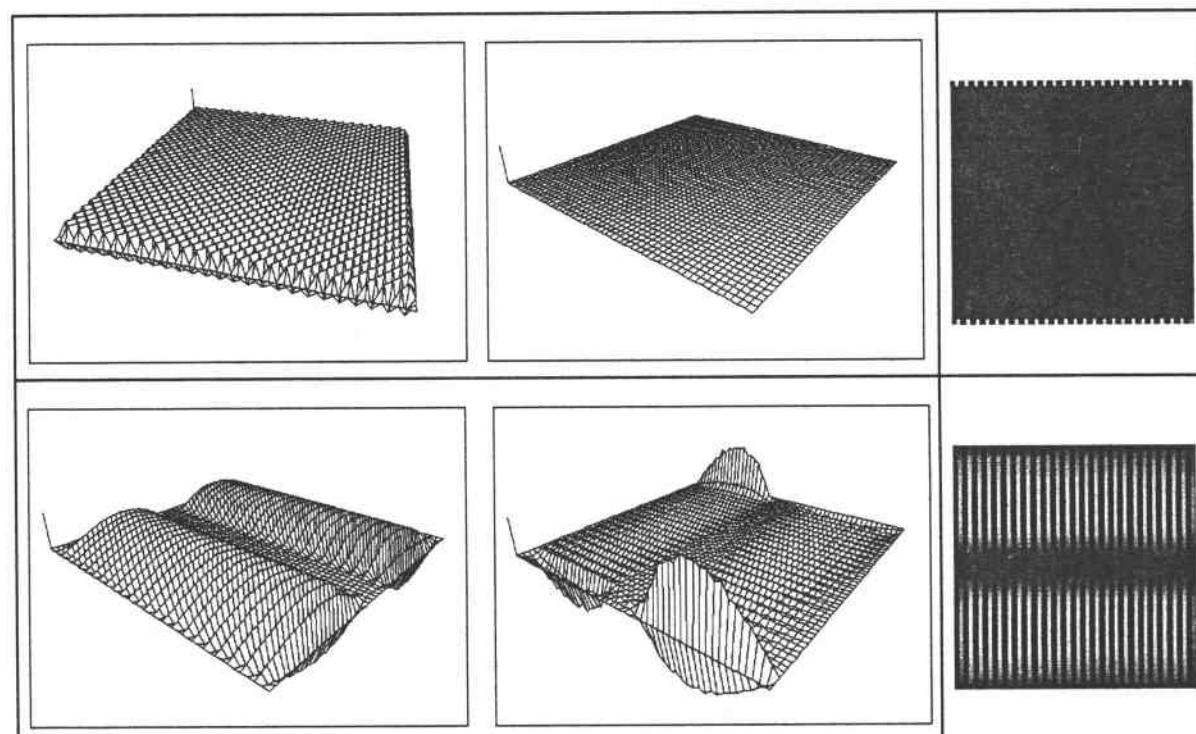


fig. 33 : la déformation finale (50 itérations) est fortement agrandie ($E \approx 625$)

III.3.5. SECOND ÉCUEIL

Ne pouvant utiliser une déformation initiale nulle, les essais ont ensuite porté sur une déformation initiale utilisant (nécessairement) des gradients rang 1-compatibles, respectant la symétrie de la mesure de Young, et aussi voisins que possible des 4 puits. En se plaçant loin des "lissages" du bord du maillage, les gradients de la déformation doivent alors être de la forme :

$$D(\varepsilon_1, \varepsilon_2) = \begin{pmatrix} \varepsilon_1 \cdot \lambda & 0 \\ 0 & \varepsilon_2 \cdot \lambda \end{pmatrix} \text{ où } \{\varepsilon_1, \varepsilon_2\} \subset \{-1, +1\}.$$

On vérifie facilement que, quels que soient les signes ε_1 et ε_2 , grâce à la symétrie des 4 puits, l'énergie d'un tel gradient prend l'expression suivante :

$$\begin{aligned} J(D) &= [(1-\lambda)^2 + (0-0)^2 + (0-0)^2 + (2-\lambda)^2] \times [(1+\lambda)^2 + (0-0)^2 + (0-0)^2 + (2-\lambda)^2] \times \dots \\ &\dots \times [(1-\lambda)^2 + (0-0)^2 + (0-0)^2 + (2+\lambda)^2] \times [(1+\lambda)^2 + (0-0)^2 + (0-0)^2 + (2+\lambda)^2] \\ &= (2\lambda^2 - 6\lambda + 5) \cdot (2\lambda^2 - 2\lambda + 5) \cdot (2\lambda^2 + 2\lambda + 5) \cdot (2\lambda^2 + 6\lambda + 5) \\ &= ((2\lambda^2 + 5)^2 - 36\lambda^2) \cdot ((2\lambda^2 + 5)^2 - 4\lambda^2) = (4\lambda^4 - 16\lambda^2 + 25) \cdot (4\lambda^4 + 16\lambda^2 + 25) \\ &= (4\lambda^4 + 25)^2 - 256\lambda^4 = 16\lambda^8 - 56\lambda^4 + 625 = 16\Lambda^2 - 56\Lambda + 625 \text{ où } \Lambda = \lambda^4 \end{aligned}$$

Ce dernier trinôme en Λ est convexe, et minimal lorsque

$$32\Lambda = 56 \Leftrightarrow \lambda = \sqrt[4]{\frac{7}{4}} \approx 1,1501633, \quad (3.5)$$

ce qui conduit aux 4 matrices :

$$D(\varepsilon_1, \varepsilon_2) = \begin{pmatrix} \varepsilon_1 \cdot \sqrt[4]{\frac{7}{4}} & 0 \\ 0 & \varepsilon_2 \cdot \sqrt[4]{\frac{7}{4}} \end{pmatrix}, \{\varepsilon_1, \varepsilon_2\} \subset \{-1, +1\} \quad (3.6)$$

pour lesquelles :

$$J(D(\varepsilon_1, \varepsilon_2)) = \boxed{576}.$$

Cette valeur est donc aussi l'énergie d'une déformation u dont les gradients sont tous égaux à l'une de ces quatre matrices dans d'égales proportions.

Suite à l'observation des oscillations apparues dans les essais précédents, la déformation initiale fut choisie de façon à comporter elle-même des oscillations régulières et orientées perpendiculairement entre les deux composantes. Comme on le voit sur les figures suivantes, la séquence ainsi obtenue franchit alors très rapidement la valeur 625, pour s'approcher de son asymptote à 576, alors que les gradients migrent massivement vers les 4 matrices rang 1-compatibles (3.6) :

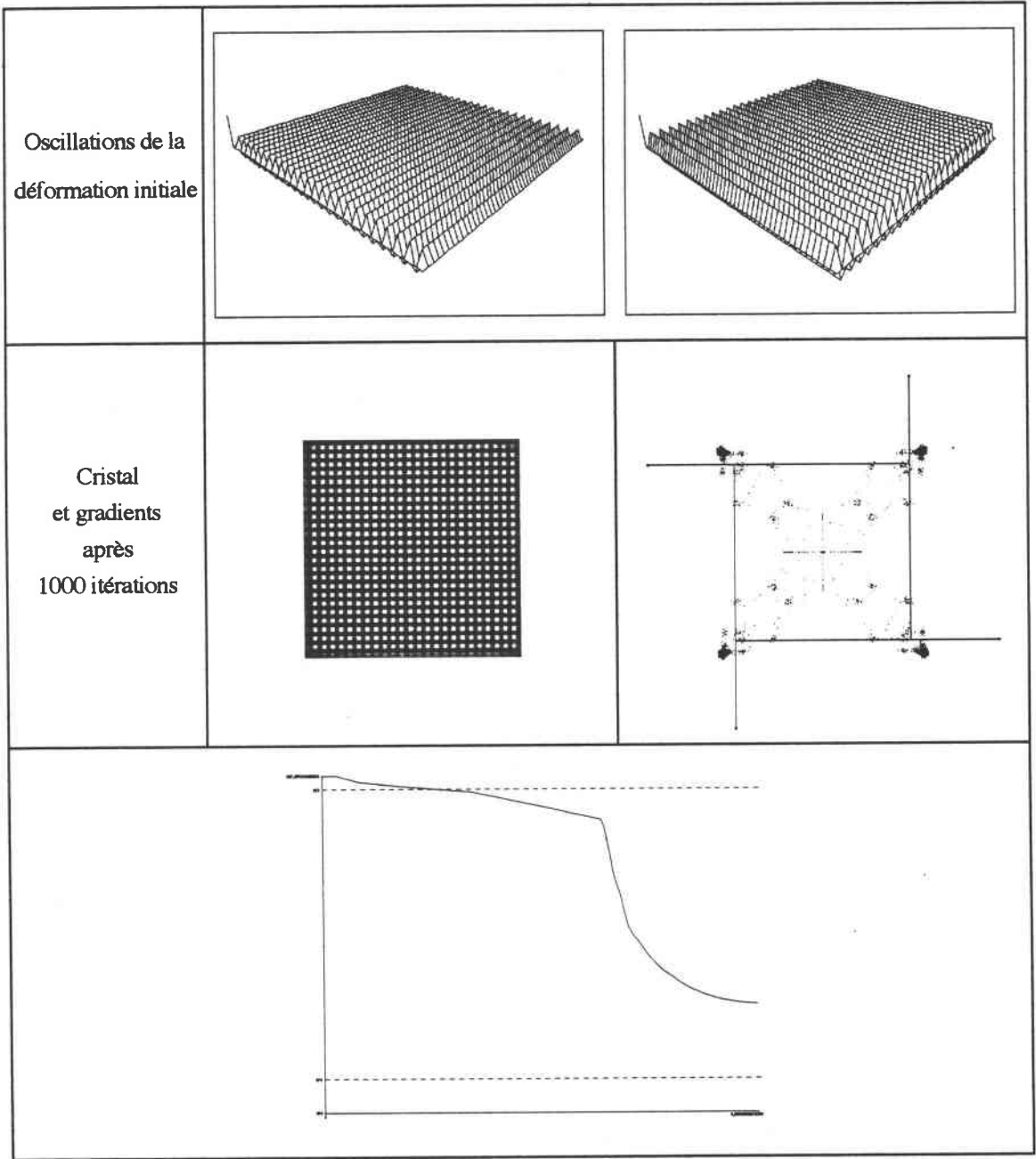


fig. 34 : descente vers le niveau d'énergie 576

III.3.6. SIMULATION DE RECUIT (ANNEALING)

Les difficultés à passer sous le niveau d'énergie 576 avec les déformations initiales envisagées ont fait envisager le recours à une technique issue de la physique : le recuit simulé (en anglais : *annealing*), ainsi nommée par sa similitude avec le réchauffement (agitation aléatoire) du cristal. Il s'agit de perturber aléatoirement, et modérément, la déformation du cristal, de façon à tenter de lui faire poursuivre une descente ayant conduit à un minimum local indésirable.

Cette technique donna quelques résultats, sans pour autant permettre de franchir la valeur 576. L'exemple suivant montre bien l'efficacité du procédé : après une descente passant tout juste sous la valeur 625, un premier recuit permet de relancer la descente de façon significative ; un second recuit permet encore d'accélérer la descente, mais ne modifie pas la limite de l'énergie de la séquence.

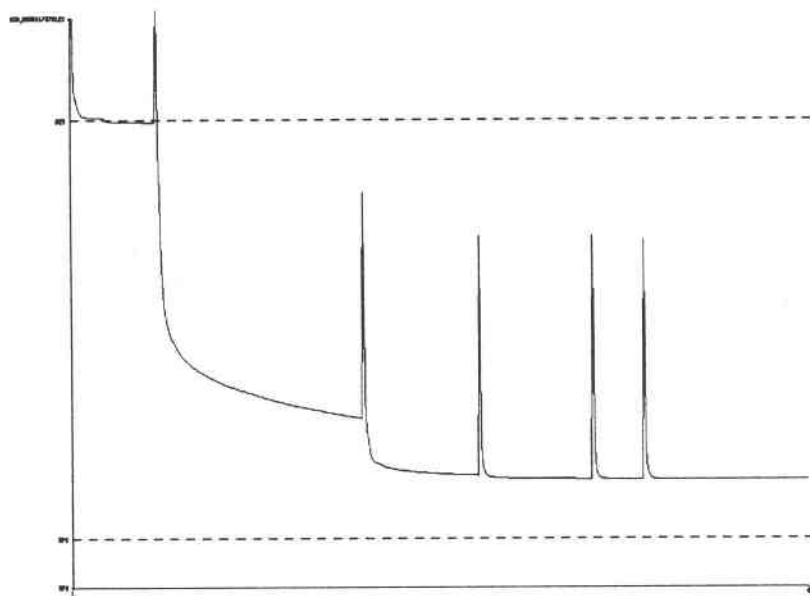


fig. 35 : amélioration de la descente par la simulation de recuit

En conclusion, la simulation de recuit permet d'atteindre plus vite l'énergie limite de la séquence, mais ne modifie pas cette limite.

III.3.7. CONDITION AU BORD ?

On constate, par ailleurs, que la valeur limite dépend de la finesse du maillage, comme l'indique le tableau suivant, où seul le nombre de mailles varie :

n	30	40	50	60	70	80	90	100
E_{∞}	595,159	590,654	587,860	585,959	584,583	583,540	582,724	582,067

tableau 4 : descente vers le niveau d'énergie 576

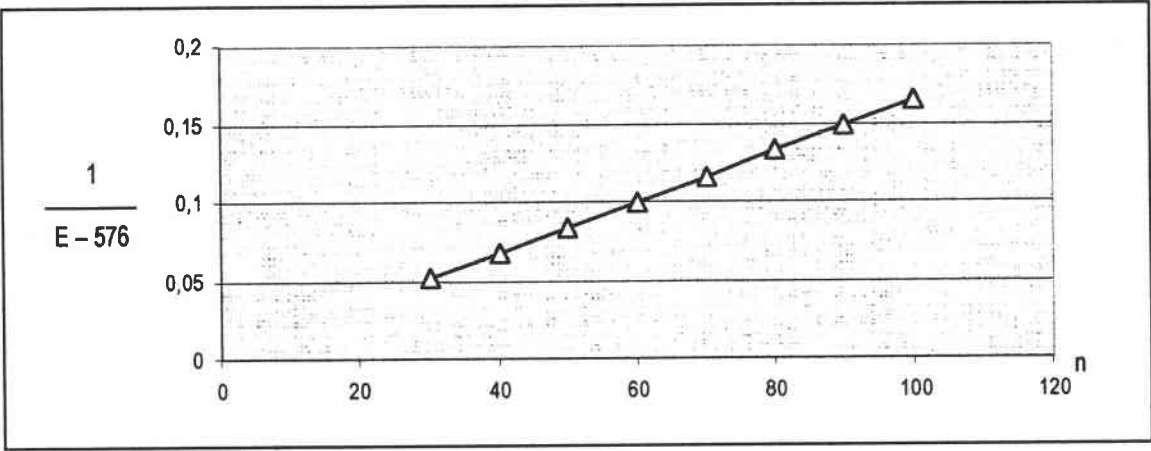


fig. 36 : corrélation linéaire entre n et $\frac{1}{E_{\infty} - 576}$

Une régression linéaire entre les variables n et $\frac{1}{E_{\infty} - 576}$ donne les résultats suivants :

• coefficient de corrélation : 0,9999998, qui assure l'excellente qualité de la relation qui suit ;

• relation : $\frac{1}{E_{\infty} - 576} = 0,0038738 + 0,00160934.n$, ou bien :

$$E_{\infty}(n) = 576 + \frac{1}{0,0038738 + 0,00160934.n}$$

(3.7)

Lorsque n tend vers l'infini, la relation (3.7) montre que la limite de l'énergie finale est bien égale à 576, dans les mêmes conditions que les calculs précédents (déformation initiale, méthode de descente, etc...).

Dans tous ces calculs, la finesse du maillage n'affecte que les effets dus à la contrainte de la condition au bord : plus le maillage est fin, mieux l'énergie parvient à s'accomoder de cette contrainte.

Il est alors apparu intéressant d'envisager des essais de descente où la condition au bord serait supprimée, au moins en un premier temps, puis ajoutée par la suite, de façon à permettre à l'énergie d'atteindre la limite de 576, voire de faire apparaître de nouvelles structures dans le cristal pouvant conduire à des énergies encore plus faibles.

Par ailleurs, la position symétrique des puits a conduit, comme il a déjà été observé plus haut, à constater que les deux composantes (horizontale et verticale) de la déformation pouvaient conserver une symétrie entre elles, à condition que la déformation initiale la possède elle-même, à savoir :

$$\forall (x_1, x_2) \in \Omega, u_2(x_1, x_2) = u_1(x_2, x_1).$$

Cette condition sera désormais imposée à toutes les déformations initiales utilisées dans cette étude, ce qui permet également de ne présenter le graphe que de l'une des deux composantes, le graphe de la seconde composante s'en déduisant par la symétrie d'axe $x_1 = x_2$.

III.3.8. APPARITION D'UNE NOUVELLE STRUCTURE

On envisage ici des déformations libres au bord du maillage, selon la remarque du paragraphe précédent. L'exemple suivant est un calcul obtenu sur un petit maillage ($n = 20$), mais déjà riche de renseignements :

1) la descente : on observe directement l'efficacité de la simulation de recuit qui, à plusieurs reprises, permet à la séquence de se dégager de minimums locaux de la fonctionnelle (discontinuités de la pente de la courbe de descente), pour atteindre un niveau d'énergie inférieur à 360 après 2000 itérations.

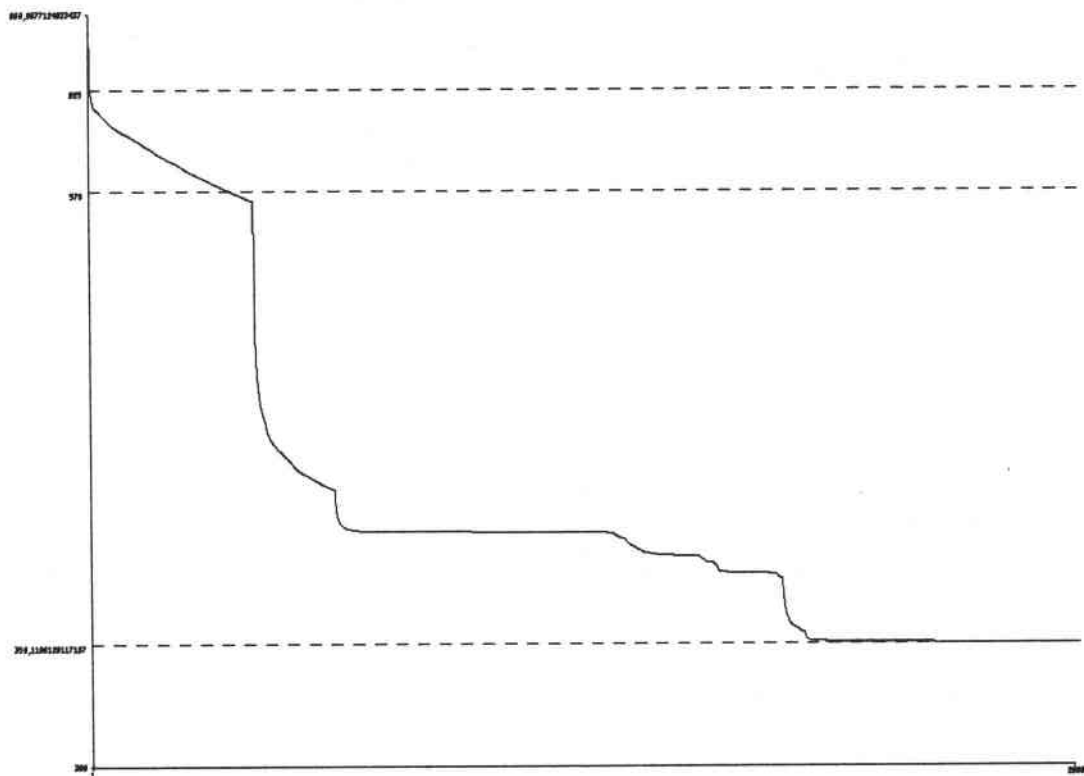


fig. 37 : descente avec simulation de recuit et sans condition au bord

2) la séquence des gradients permet de distinguer très nettement plusieurs phases :

- vues 1 à 19 : les gradients s'organisent de façon symétrique le long d'une structure en carré de plus en plus marquée utilisant les 4 matrices rang 1-compatibles (3.6) ; l'énergie s'approche de 576 ;
- vues 20 à 28 : la structure précédente est disloquée par les recuits simulés successifs, et les gradients commencent à migrer vers les 4 puits ;
- vues 29 à 54 : les gradients se concentrent très majoritairement au voisinage des 4 puits, ce qui a pour effet de diminuer l'énergie du cristal jusqu'à environ 360 , soit très au-dessous de la valeur 576 :

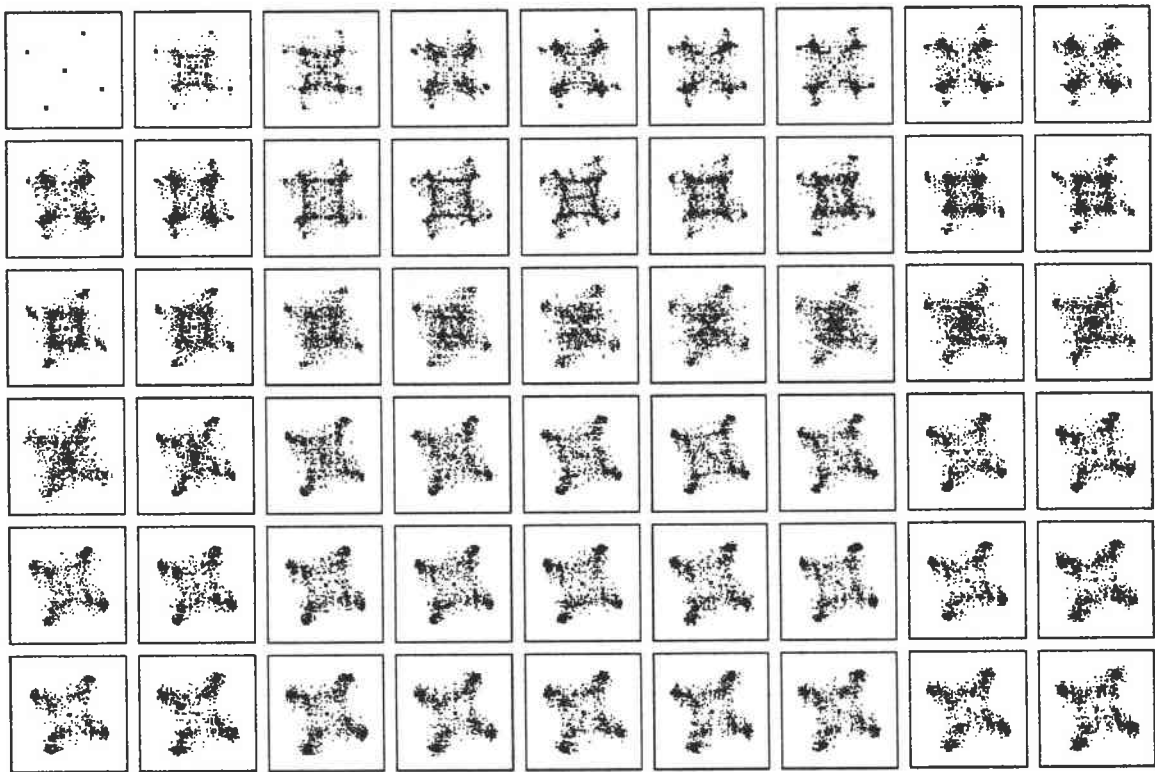


fig. 38 : séquence des gradients et numérotation des vues

1	2	3	4	5	6	7	8	9
10	11	12	13	14	15	16	17	18
19	20	21	22	23	24	25	26	27
28	29	30	31	32	33	34	35	36
37	38	39	40	41	42	43	44	45
46	47	48	49	50	51	52	53	54

3) la déformation finale laisse bien apparaître des oscillations, dans les directions désormais habituelles, mais à plusieurs échelles : une oscillation “primaire” à longue période (environ 10 mailles, soit $n/2$), à laquelle se superpose une oscillation “secondaire” à plus courte période (1 maille) :

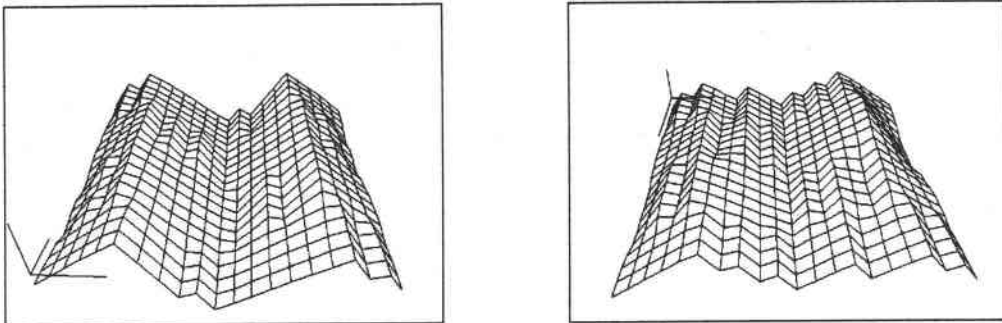


fig. 39 : composantes de la déformation finale

4) enfin, le cristal correspondant à cette déformation finale montre une structure qui n’est pas sans rappeler l’aspect d’un tissu, le *tweed* des figures 1 et 2 :

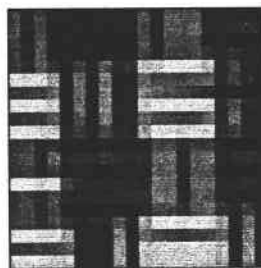


fig. 40 : déformation finale du cristal

La séquence des déformations du cristal montre comment cette structure se forme peu-à-peu à partir du bord du cristal, ce qui explique pourquoi elle n’apparaissait pas sur un cristal soumis à une contrainte à son bord :

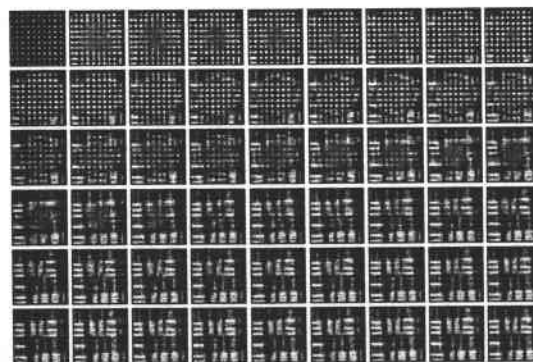


fig. 41 : séquence des déformations successives du cristal

III.3.9. OSCILLATIONS À ÉCHELLES MULTIPLES

Pour faciliter l'apparition d'une structure montrant plusieurs oscillations superposées, il fallait adapter la déformation initiale : les directions initiales des oscillations se feront toujours comme précédemment, mais sur une période de plusieurs mailles.

Les différents essais donnèrent des résultats tellement bons qu'il fut même possible d'atteindre des valeurs d'énergie assez basses tout en maintenant la condition au bord :

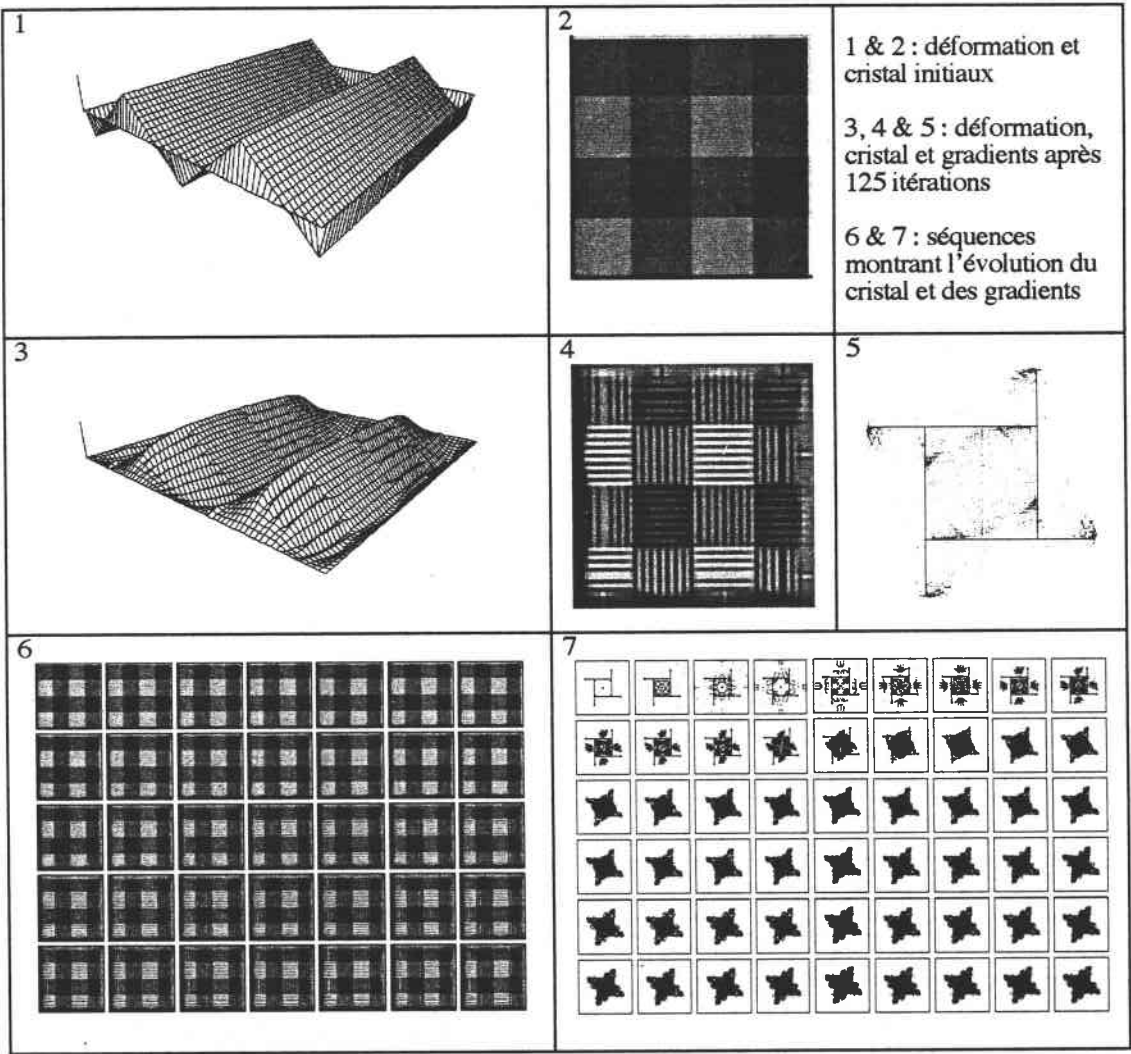


fig. 42 : 1^{er} exemple d'apparition d'oscillations courtes sur des oscillations longues.
L'énergie atteinte ici après 125 itérations vaut environ 478.

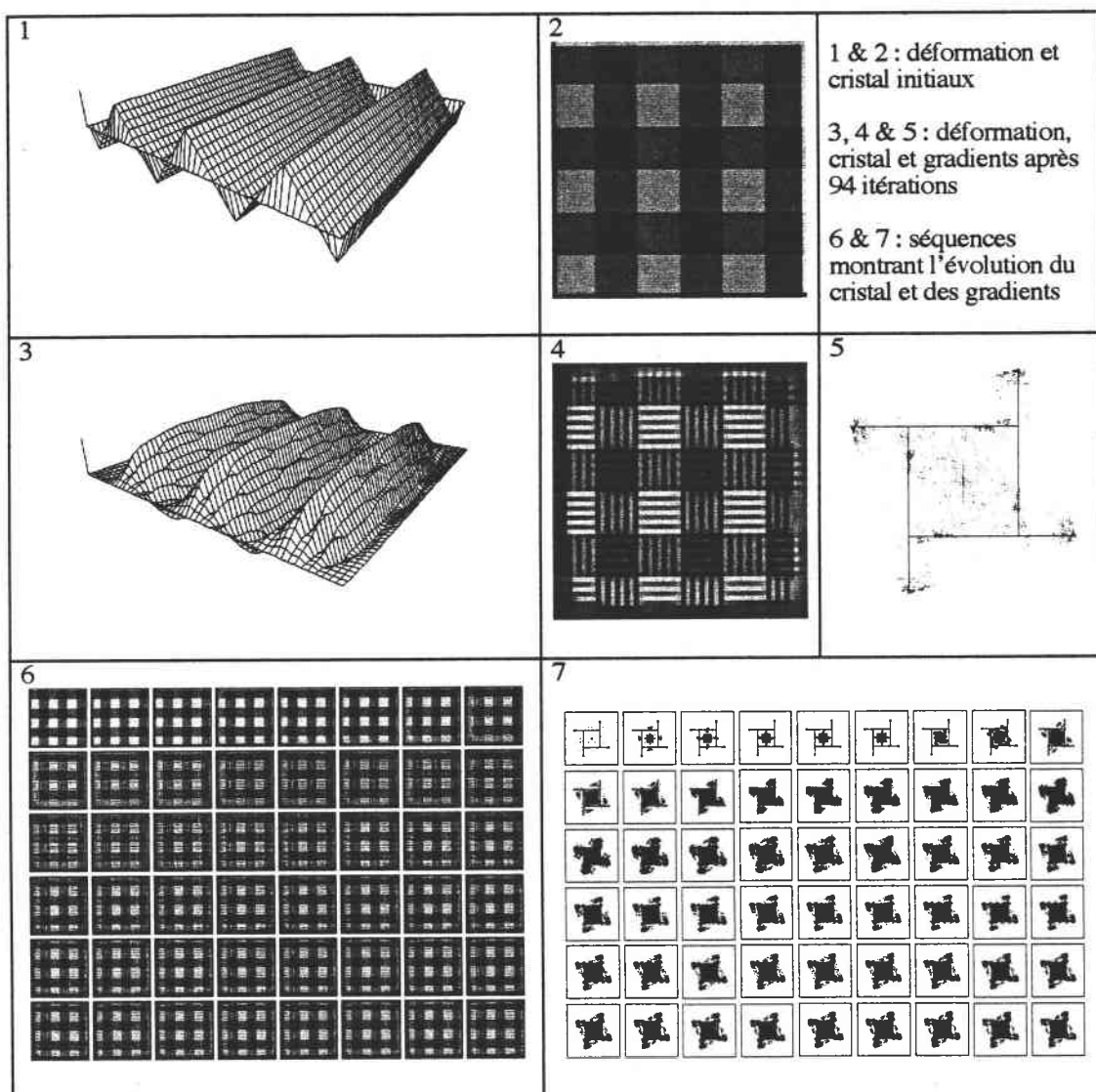


fig. 43 : 2nd exemple d'apparition d'oscillations courtes sur des oscillations longues.
L'énergie atteinte ici après 94 itérations vaut environ 461.

En fait, ces oscillations multi-échelle constituent la clé des estimations théoriques de l'énergie par rapport au pas du maillage.

IV - ESTIMATIONS DE L'ÉNERGIE

IV.1. ESTIMATION DE L'ÉNERGIE EN DIMENSION 2

On rappelle que :

- Ω est un domaine de $\mathbb{R}^2$,
- T est une triangulation de Ω de pas h (h est le plus grand des diamètres des polyèdres K de T),
- $V^h(\Omega) = \{ v \in W^{1,\infty}(\Omega; \mathbb{R}^2) / \forall K \in T, v|_K \in \mathcal{P}_1(K) \}$,
- $V_0^h(\Omega) = \{ v \in V^h(\Omega) / v = 0 \text{ sur } \partial\Omega \}$,
- φ est une fonction continue de $\mathbb{R}^{2 \times 2}$ dans $\mathbb{R}^+$ ayant les 4 puits (2.8).

Théorème 4.1 (cf [8]) : avec les définitions précédentes, il existe deux constantes C_1 et C_2 , indépendantes du pas h de la triangulation T , telles que :

$$\inf_{v \in V_0^h(\Omega, \mathbb{R}^n)} \int_{\Omega} \varphi(\nabla v(x)) \, dx \leq C_1 \cdot e^{-C_2 \cdot \|nh\|^{1/2}}$$

Preuve : pour un certain réel $\alpha \in]0,1[$, on considère le carré $[0, 2h^\alpha]^2$, sur lequel on définit une fonction u^1 par :

$$\begin{cases} u_1^1(x_1, x_2) = \min(x_1, 2h^\alpha - x_1) \\ u_2^1(x_1, x_2) = \min(x_2, 2h^\alpha - x_2) \end{cases}$$

si bien que, sur ce carré (fig. 44), ∇u coïncide avec l'une des 4 matrices rang 1-compatibles (2.9) :

$$\left\{ \begin{array}{l} \text{sur } [0, h^\alpha]^2 : \nabla u^1 = w'_1 \\ \text{sur } [0, h^\alpha] \times [h^\alpha, 2h^\alpha] : \nabla u^1 = w'_2 \\ \text{sur } [h^\alpha, 2h^\alpha] \times [0, h^\alpha] : \nabla u^1 = w'_4 \\ \text{sur } [h^\alpha, 2h^\alpha]^2 : \nabla u^1 = w'_3 \end{array} \right.$$

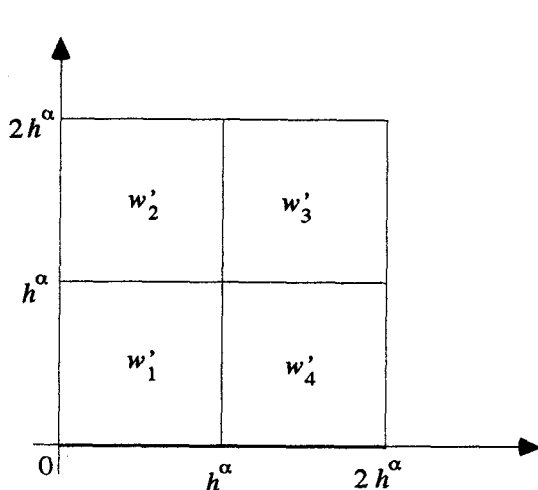


fig. 44

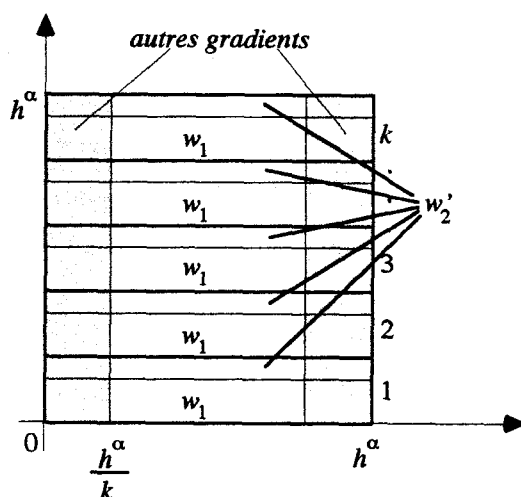


fig. 45

On choisit ensuite un certain entier $k \geq 3$, et on subdivise le carré $[0, h^\alpha]^2$ en k bandes horizontales. En réservant deux bandes verticales (en gris sur la figure 45) de largeur $\frac{h^\alpha}{k}$ pour permettre le raccord à gauche et à droite du carré avec ses voisins, on obtient ainsi k rectangles mesurant $\frac{h^\alpha}{k} \times (k-2) \cdot \frac{h^\alpha}{k}$.

Puisque $w'_1 = \frac{2}{3} w_1 + \frac{1}{3} w'_2$, on modifie la fonction u^1 en une fonction u^2 en altérant sa "pente" dans la direction de x_2 : sans modifier $\frac{\partial u^1}{\partial x_1} = \frac{\partial u^1}{\partial x_1} = 1$, on impose $\frac{\partial u^2}{\partial x_2} = 2$, soit $\nabla u^2 = w_1$ sur les deux tiers inférieurs de chacun de ces k rectangles, et $\frac{\partial u^2}{\partial x_2} = -1$, soit $\nabla u^2 = w'_2$ sur le tiers supérieur.

À leur tour, chaque bande où $\nabla u^2 = w'_2$ est un rectangle mesurant $\frac{1}{3} \frac{h^\alpha}{k} \times (k-2) \cdot \frac{h^\alpha}{k}$, pouvant donc être subdivisée en $3(k-2)$ carrés de mêmes dimensions que le carré $[0, h^\alpha/3k]$, sur lesquels on répète le processus précédent : sur chacun de ces carrés, u^2 est modifiée en une fonction u^3 , dont les gradients sont w_2 et w'_3 , sauf sur deux bandes latérales où u^3 devra simplement se raccorder aux autres parties du carré initial, y engendrant des gradients n'étant ni l'un des w_i , ni l'un des w'_i , ces gradients étant toutefois uniformément bornés indépendamment de h .

Ce qui vient d'être décrit sur le carré $[0, h^\alpha]^2$ peut être adapté aux trois autres quarts du carré $[0, 2h^\alpha]^2$ en procédant à un simple décalage sur les indices des matrices w_i et w'_i .

Cette construction engendre ainsi une suite de fonctions $u^1, u^2, \dots, u^p$ sur $[0, 2h^\alpha]^2$. Ces fonctions sont ensuite rendues périodiques de période $2h^\alpha$ dans les deux directions (et encore nommées $u^1, u^2, \dots, u^p$), et sont ainsi définies sur tout $\mathbb{R}^2 \supset \Omega$.

Si $\text{sgn}(\xi)$ représente le signe du réel ξ (-1 si $\xi < 0$, 0 si $\xi = 0$, 1 si $\xi > 0$), on pose enfin :

$$\bullet \hat{u}_i^p(x) = \text{sgn}(u_i^p(x)) \cdot \min(|u_i^p(x)|, \text{dist}(x, \partial\Omega)) \text{ pour } i = 1, 2, \quad (4.1)$$

$$\bullet u_h^p \text{ est l'interpolée de } \hat{u}^p \text{ sur la triangulation } T. \quad (4.2)$$

Notons $S = \{\nabla u_h^p \notin \{w_1, w_2, w_3, w_4\}\}$ la partie de Ω sur laquelle le gradient de u_h^p ne coïncide avec aucun des 4 puits w_i . Alors, par construction, $\hat{u}^p \in W_0^{1,\infty}(\Omega)$ et $u_h^p \in V_0^h(\Omega)$, et les gradients de u_h^p sont uniformément bornés indépendamment de h , si bien qu'il existe une constante $C > 0$ telle que :

$$\int_{\Omega} \varphi(\nabla u_h^p(x)) \, dx \leq C \cdot |S| \quad (4.3)$$

où $|S|$ est la mesure de Lebesgue de S .

S est constituée de 4 sous-ensembles S_1, S_2, S_3 et S_4 :

• S_1 est un voisinage de $\partial\Omega$ issu de la troncature au bord (4.1). Par construction de u^p :

$$|u_i^p| \leq h^\alpha \cdot \left(1 + \frac{1}{3k} + \dots + \frac{1}{(3k)^p}\right) \leq h^\alpha \cdot \frac{1}{1 - \frac{1}{3k}}, \text{ pour } i = 1, 2,$$

cette dernière expression étant majorée pour $k \geq 3$ par sa valeur en $k = 3$, soit $|u^p| \leq \frac{10}{9} h^\alpha$.

Donc, S_1 est inclus dans un domaine longeant $\partial\Omega$ et de largeur $c.h^\alpha$ pour une certaine constante $c > 0$. Il existe donc une autre constante $C > 0$ (indépendante de h) telle que :

$$|S_1| \leq C.h^\alpha \quad (4.4)$$

• S_2 est dû à l'interpolation (4.2) le long des droites $x_i = m.h^\alpha$, $m \in \mathbb{N}$. Soit respectivement N , L et l le nombre, la longueur et la largeur de telles droites. Pour une certaine constante $c > 0$:

$$N \leq c \cdot \frac{\text{diam}(\Omega)}{h^\alpha}$$

$$L \leq \text{diam}(\Omega)$$

$$l \leq h$$

puis

$$|S_2| \leq N.L.l \leq c.\text{diam}(\Omega)^2.h^{1-\alpha}$$

Donc, quitte à modifier la constante C de l'inégalité (4.4) :

$$|S_2| \leq C.h^{1-\alpha} \quad (4.5)$$

• S_3 est la partie de Ω où $\nabla u_h^p \in \{w'_1, w'_2, w'_3, w'_4\}$. Il est clair que $|S_3|$ est divisée par 3 à chaque itération de la construction de la suite, jusqu'à la p -ième fonction u^p :

$$|S_3| \leq \left(\frac{1}{3}\right)^p . |\Omega| \quad (4.6)$$

• enfin, S_4 est dû à l'interpolation (4.2) le long des droites-frontières entre les w_i et les w'_i et avec les bandes latérales contenant d'autres gradients. À la i -ème itération, l'opération s'applique à un carré Q_i de côté $\frac{h^\alpha}{(3k)^{i-1}}$. À chaque itération, un carré Q_i contient $3k(k-2)$ carrés Q_{i+1} . Donc, le carré initial de côté $2.h^\alpha$ contient 4 carrés Q_1 , soit aussi $4 \times [3k(k-2)]^{i-1}$ carrés Q_i . À l'intérieur de Ω , le nombre N_i de carrés Q_i est donc, pour une certaine constante $C > 0$:

$$N_i \leq \frac{C}{(h^\alpha)^2} [3k(k-2)]^{i-1} \leq \frac{C}{h^{2\alpha}} [3k^2]^{i-1} \quad (4.7)$$

À l'intérieur d'un carré Q , les deux bandes latérales contenant des gradients autres que les w_i et w'_i ont une mesure totale s_1 représentant la fraction $\frac{2}{k}$ de la mesure du carré lui-même :

$$s_1^i = \frac{2}{k} \cdot \left(\frac{h^\alpha}{(3k)^{i-1}} \right)^2 \quad (4.8)$$

Enfin, l'interpolation (4.2) intervient aussi le long des segments de droite séparant les rectangles contenant les gradients w_i et les gradients w'_i . Les carrés Q contiennent au plus $2k$ tels segments, autour desquels une bande de largeur au plus $2h$ peut contenir de mauvais gradients. La longueur de ces segments est inférieure au côté du carré Q , si bien que la mesure s_2 de toutes ces bandes est donc ainsi majorée :

$$s_2^i \leq 2k \cdot 2h \cdot \frac{h^\alpha}{(3k)^{i-1}} \quad (4.9)$$

Or, $|S_4| \leq \sum_{i=1}^P N_i (s_1^i + s_2^i)$. En tenant compte des résultats (4.7) à (4.9), il vient :

$$N_i \cdot s_1^i \leq \frac{C}{3^{i-1} \cdot k} \quad \text{puis} \quad \sum_{i=1}^P N_i \cdot s_1^i \leq \frac{C}{k} \quad (4.10)$$

$$N_i \cdot s_2^i \leq C \cdot k^i \cdot h^{1-\alpha} \quad \text{puis} \quad \sum_{i=1}^P N_i \cdot s_2^i \leq C \cdot k^{p+1} \cdot h^{1-\alpha} \quad (4.11)$$

Finalement :

$$|S| \leq C \cdot \left\{ h^\alpha + h^{1-\alpha} + \left(\frac{1}{3} \right)^P + \frac{1}{k} + k^{p+1} \cdot h^{1-\alpha} \right\} \quad (4.12)$$

Il reste désormais à choisir de "bonnes" valeurs pour α , k et p .

On prend $\alpha = \frac{1}{2}$, ce qui conduit, pour une autre constante $C > 0$ à :

$$|S| \leq C \cdot \left\{ \left(\frac{1}{3} \right)^P + \frac{1}{k} + k^{p+1} \cdot h^{1/2} \right\} \quad (4.13)$$

On cherche ensuite un réel λ tel que h^λ soit un majorant commun à $\frac{1}{k}$ et à $k^{p+1} \cdot h^{1/2}$:

$$\frac{1}{k} \leq h^\lambda \iff k \geq h^{-\lambda} \quad \text{et} \quad k^{p+1} \cdot h^{1/2} \leq h^\lambda \iff k \leq h^{(2\lambda-1)/2(p+1)} \quad (4.14)$$

Les deux exposants coïncident si $-\lambda = \frac{2\lambda-1}{2(p+1)}$, soit $\lambda = \frac{1}{2(p+2)}$.

On veut alors pouvoir choisir

$$k = E(h^{-1/2(p+2)}), \quad (4.15)$$

ce qui ne conduit à $k \geq 3$ que si :

$$h \leq \left(\frac{1}{3}\right)^{2(p+2)} \quad (4.16)$$

Alors, sous la condition (4.16), il vient :

$$\frac{1}{2} h^{-1/2(p+2)} \leq h^{-1/2(p+2)} - 1 < k \leq h^{-1/2(p+2)} \quad (4.17)$$

soit, simultanément :

$$\frac{1}{k} \leq 2.h^{1/2(p+2)} \quad \text{et} \quad k^{p+1}.h^{1/2} \leq h^{1/2(p+2)}$$

et donc, pour une certaine constante $C > 0$:

$$|S| \leq C. \left\{ \left(\frac{1}{3}\right)^p + h^{1/2(p+2)} \right\} \quad (4.18)$$

Or, pour

$$h = \left(\frac{1}{3}\right)^{2p(p+2)}, \quad (4.19)$$

la condition (4.16) est toujours vérifiée, et de plus $h^{1/2(p+2)} = \left(\frac{1}{3}\right)^p$, et pour une certaine constante $C > 0$:

$$|S| \leq C. \left(\frac{1}{3}\right)^p \quad (4.20)$$

Mais $h = \left(\frac{1}{3}\right)^{2p(p+2)}$ donne également $-\frac{\ln h}{2 \cdot \ln 3} = p(p+2) \leq 3p^2$ dès que $p \geq 1$, et donc, pour une

certaine constante $c > 0$:

$$p \geq c. \sqrt{-\ln h}$$

soit finalement, pour deux constantes $C_1, C_2 > 0$:

$$\inf_{v \in V_0^h(\Omega, \mathbb{R}^n)} \int_{\Omega} \varphi(\nabla v(x)) \, dx \leq C. |S| \leq C_1. e^{-C_2. \|\ln h\|^{1/2}}$$

■

IV.2. APPLICATION NUMÉRIQUE

La construction de la démonstration du théorème 4.1 a été utilisée pour définir un nouveau type de déformation initiale : pour un certain entier $q \geq 1$,

- (1) subdivision d'un carré initial $[0,1/q]^2$ en 4 quarts où ne sont utilisés que les gradients w'_i ;
- (2) subdivision de chaque carré en k bandes, avec deux bandes latérales pour les raccordements ;
- (3) modification des gradients w'_i en w_i et w'_{i+1} ;
- (4) subdivision des rectangles contenant les w'_{i+1} en carrés, et retour à l'étape (2).

L'ensemble des étapes (2) à (4) est répété p fois, puis le résultat est reproduit q fois dans chaque direction par q -périodicité pour recouvrir le domaine $\Omega = [0,1]^2$. Dans le maillage final, on a alors :

$$n = \frac{1}{h} = 2q \cdot (3k)^p \quad (4.21)$$

Remarque 4.1 : p étant le nombre d'itérations dans le processus de construction de la suite minimisante, les relations (4.15) et (4.19) permettent de déterminer le pas h du maillage et le nombre k de subdivisions des carrés nécessaires à la construction et aux majorations successives de la démonstration :

$$k = 3^p \text{ et } n = \frac{1}{h} = 3^{2p(p+2)} = (9k)^{2p}, \quad (4.22)$$

les maillages utilisés numériquement étant de dimensions très inférieures :

p	$k = 3^p$	n théorique = $(9k)^{2p}$	$h = \frac{1}{n}$	$h^\alpha = h^{1/2}$	$\frac{h^\alpha}{h} = n^{1/2}$	n empirique = $(3k)^p$
1	3	729	$\approx 0,00137$	$\approx 0,037$	27	9
2	9	43 046 721	$\approx 2,32 \cdot 10^{-8}$	$\approx 1,52 \cdot 10^{-4}$	6561	729

tableau 5 : taille des maillages en dimension 2

Voici, par exemple, la déformation obtenue avec $k = 5$, $p = 2$ et $q = 1$, soit $n = \frac{1}{h} = 450$:

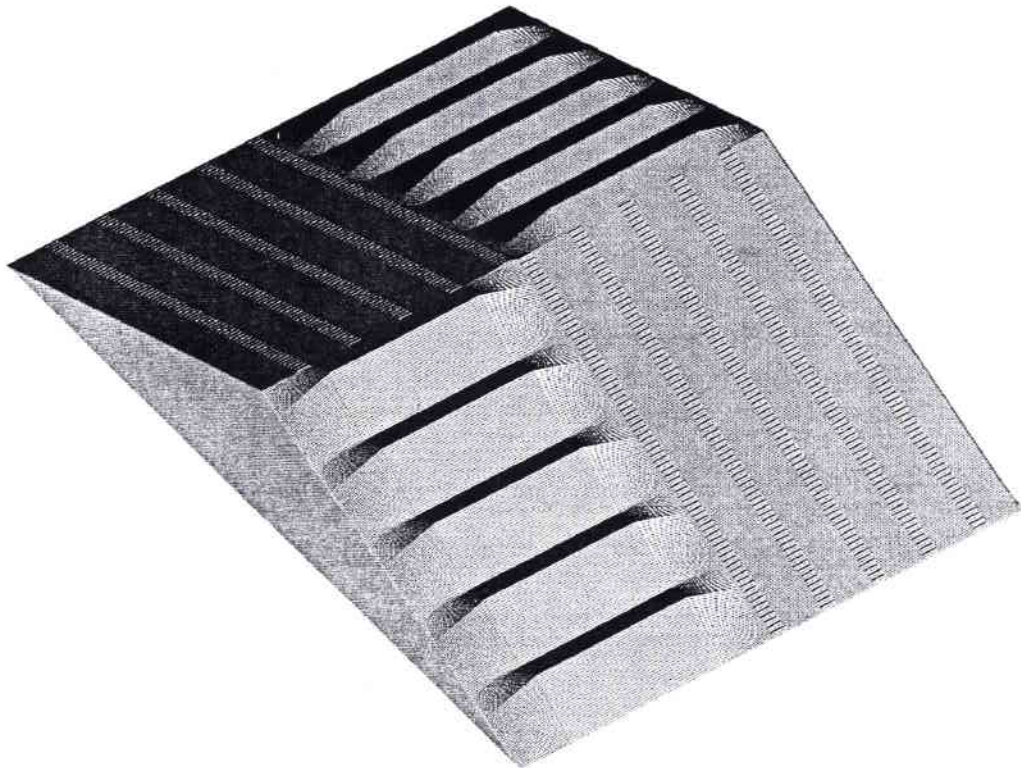


fig. 46

et le cristal correspondant :

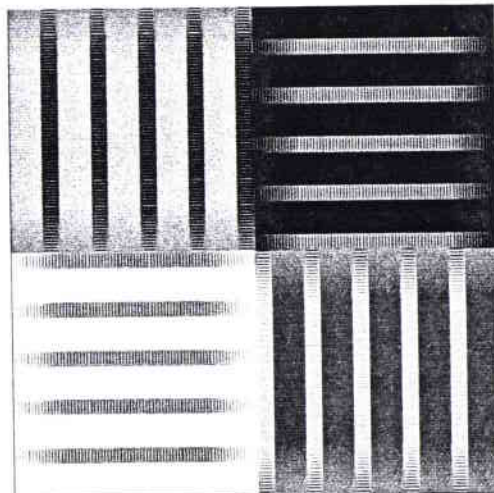


fig. 47

dont l'aspect n'est pas sans rappeler les résultats des figures 42 et 43, avec cependant une oscillation supplémentaire à l'échelle la plus fine. L'énergie de ce cristal vaut environ 572,3 .

Suit un second exemple, obtenu avec $k = 3$, $p = 2$ et $q = 3$, soit $n = \frac{1}{h} = 486$:



fig. 48



fig. 49

L'application de la méthode de descente avec la fonction précédente comme déformation initiale conduit, après 130 itérations, à une énergie d'environ 195, donc une valeur très faible par comparaison aux valeurs précédentes. Signalons toutefois que cette valeur est obtenue sans application de la condition au bord sur un seul motif de la fonction périodique, c'est-à-dire sur un carré $[m.h^\alpha, (m+1).h^\alpha]^2$ situé loin de $\partial\Omega$.

Le cristal prend alors l'aspect suivant, où l'on observe des perturbations dans les bandes latérales du découpage où se trouvent les "mauvais" gradients :

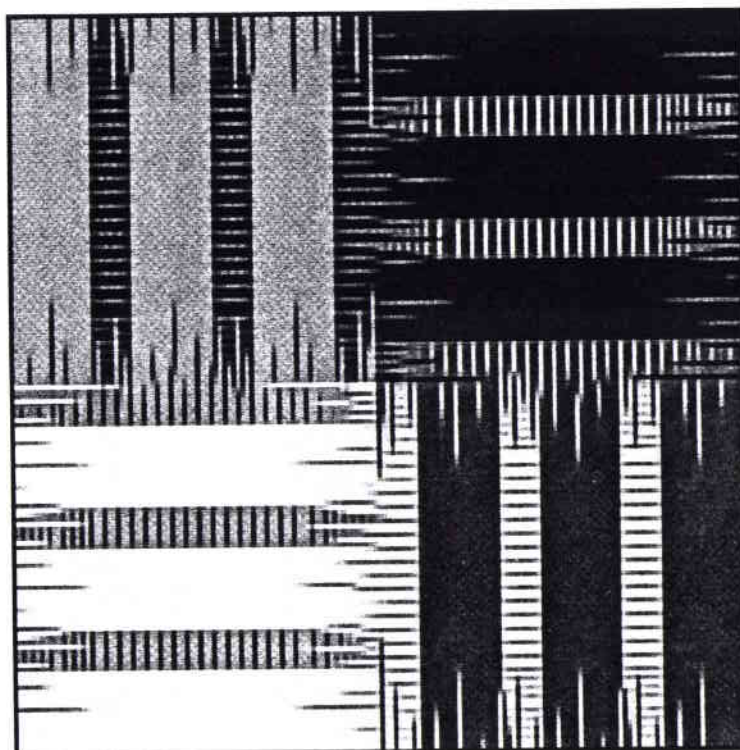


fig. 50

Signalons, pour finir, qu'il fut observé numériquement une amélioration très sensible du niveau de l'énergie en prenant, pour les bandes latérales du découpage des carrés, une largeur inférieure à celle des rectangles, de la forme $\delta \cdot \frac{h^\alpha}{k}$, où δ est un certain réel > 0 ($\delta = 1$ dans ce qui précède) :

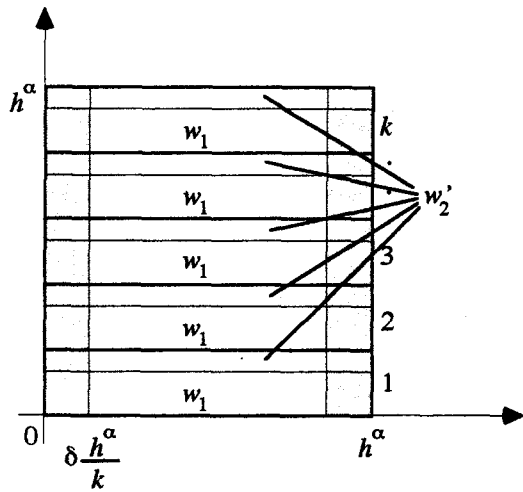


fig. 51

Voici les énergies de quelques déformations initiales (sans condition au bord) pour $p = 2$ et $q = 1$:

δ	$k = 3$	$k = 4$	$k = 5$
1/3	1589,7	1194,2	960,6
1/2	1185,4	873,1	721,3
2/3	816,5	632,1	520,0
3/4	835,9	647,0	530,5
1	877,0	687,6	572,3

tableau 6 : énergie limite selon les valeurs de δ

où il apparaît que la valeur $\delta = \frac{2}{3}$ donne les meilleurs résultats.

IV.3. UN EXEMPLE EN DIMENSION 3

Dans ce paragraphe :

- Ω est un domaine de $\mathbb{R}^3$,
- $W^{1,\infty}(\Omega; \mathbb{R}^3)$ est l'espace de Sobolev des fonctions mesurables $u : \Omega \rightarrow \mathbb{R}^3$ telles que u , ainsi que sa dérivée $\nabla u : \Omega \rightarrow \mathbb{R}^{n \times n}$ soient essentiellement bornées sur Ω ,
- $W_0^{1,\infty}(\Omega; \mathbb{R}^3) = \{ u \in W^{1,\infty}(\Omega; \mathbb{R}^3) / \forall x \in \partial\Omega, u(x) = 0 \}$,
- T est une triangulation de Ω de pas h (h est le plus grand des diamètres des polyèdres K de T),
- $V^h(\Omega) = \{ v \in W^{1,\infty}(\Omega; \mathbb{R}^3) / \forall K \in T, v|_K \in \mathcal{P}_1(K) \}$,
- $V_0^h(\Omega) = \{ v \in V^h(\Omega) / v = 0 \text{ sur } \partial\Omega \}$,
- φ est une fonction continue de $\mathbb{R}^{3 \times 3}$ dans $\mathbb{R}^+$ ayant les 5 puits non rang 1-compatibles suivants, inspirés des puits (2.8) :

$$\begin{aligned}
 w_1 &= \begin{pmatrix} 1 & 0 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 1 \end{pmatrix}, w_2 = \begin{pmatrix} 2 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 1 \end{pmatrix}, w_3 = \begin{pmatrix} -1 & 0 & 0 \\ 0 & -2 & 0 \\ 0 & 0 & 1 \end{pmatrix}, \\
 w_4 &= \begin{pmatrix} -2 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix} \text{ et } w_5 = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & -1 \end{pmatrix},
 \end{aligned} \tag{4.22}$$

c'est-à-dire que, si $w \in M_3(\mathbb{R})$, $\varphi(w) \geq 0$, et $\varphi(w) = 0$ si, et seulement si $w \in \{w_i, 1 \leq i \leq 5\}$.

On définit en outre les 4 matrices :

$$\begin{aligned}
 w'_1 &= \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix}, w'_2 = \begin{pmatrix} 1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 1 \end{pmatrix}, \\
 w'_3 &= \begin{pmatrix} -1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 1 \end{pmatrix} \text{ et } w'_4 = \begin{pmatrix} -1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix}.
 \end{aligned} \tag{4.23}$$

Les 9 matrices w_i , $1 \leq i \leq 5$, et w'_j , $1 \leq j \leq 4$, sont représentées dans les coordonnées de leurs 3 coefficients diagonaux sur le stéréogramme suivant :

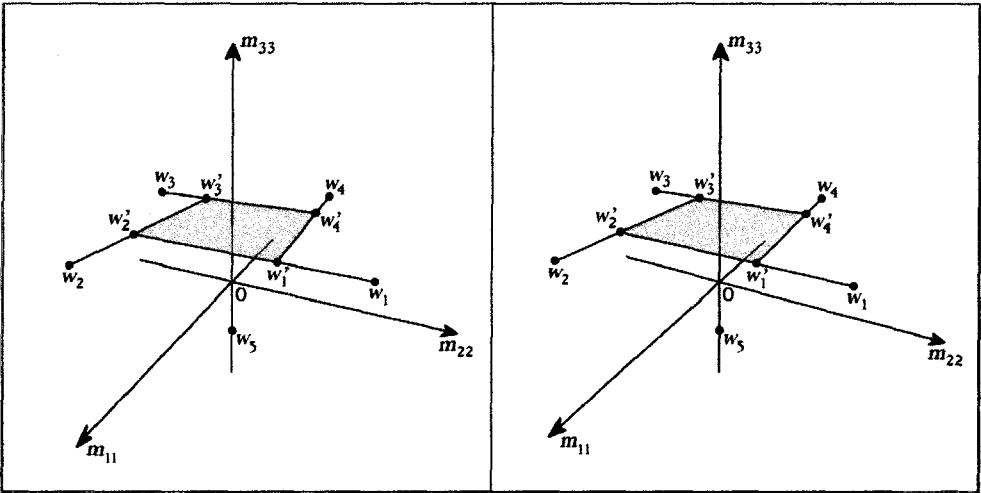


fig. 52

Théorème 4.2 : avec les définitions précédentes :

$$\inf_{v \in W_0^{1,\infty}(\Omega, \mathbb{R}^3)} \int_{\Omega} \varphi(\nabla v(x)) \, dx = 0 \quad (4.24)$$

ces infima n'étant pas atteints dans $W_0^{1,\infty}(\Omega; \mathbb{R}^3)$.

Preuve : la démonstration est identique à celle du théorème 2.1. Rappelons, en particulier, la relation (2.7), selon laquelle :

$$\forall w \in M_3(\mathbb{R}), \quad |\Omega| \cdot Q\varphi(w) = \inf_{v \in W_0^{1,\infty}(\Omega, \mathbb{R}^3)} \int_{\Omega} \varphi(w + \nabla v(x)) \, dx \quad (4.25)$$

Comme dans la démonstration du théorème 2.1, on obtient successivement :

- d'après la proposition 2.1 : $0 \leq Q\varphi(w_i) \leq \varphi(w_i) = 0$, soit

$$Q\varphi(w_i) = 0, \text{ pour tout } i \in \{1..5\};$$

- par exemple, $w'_1 = \frac{1}{3} (2w_1 + w'_2)$, donc $0 \leq Q\varphi(w'_1) \leq \frac{2}{3} Q\varphi(w_1) + \frac{1}{3} Q\varphi(w'_2) = \frac{1}{3} Q\varphi(w'_2)$

puisque $Q\varphi(w_1) = 0$, soit en itérant :

$$0 \leq Q\varphi(w'_1) \leq \frac{1}{3} Q\varphi(w'_2) \leq \left(\frac{1}{3}\right)^2 Q\varphi(w'_3) \leq \left(\frac{1}{3}\right)^3 Q\varphi(w'_4) \leq \left(\frac{1}{3}\right)^4 Q\varphi(w'_1), \text{ soit}$$

$$Q\varphi(w'_i) = 0, \text{ pour tout } i \in \{1..4\};$$

• par rang 1-convexité, pour toute matrice w appartenant à l'enveloppe rang 1-convexe des w'_i (carré grisé sur la figure 52), de la forme $w = T_{\alpha,\beta,\gamma} = \text{diag}(\alpha, \beta, \gamma)$ pour $-1 \leq \alpha, \beta \leq 1$ et $\gamma = 1$: $Q\varphi(T_{\alpha,\beta,1}) = 0$, en particulier :

$$Q\varphi(T_{0,0,1}) = Q\varphi(-w_5) = 0;$$

- par rang 1-convexité encore, puisque $Q\varphi(-w_5) = Q\varphi(w_5) = 0$:

$$0 \leq Q\varphi(0) = Q\varphi\left(\frac{(w_5) + (-w_5)}{2}\right) \leq \frac{1}{2} Q\varphi(w_5) + \frac{1}{2} Q\varphi(-w_5) = 0,$$

ce qui, compte-tenu de (4.25), prouve l'affirmation (4.24).

Supposons alors qu'il existe une fonction $u \in W_0^{1,\infty}(\Omega; \mathbb{R}^3)$ telle que

$$\int_{\Omega} \varphi(\nabla u(x)) \, dx = \inf_{v \in W_0^{1,\infty}(\Omega; \mathbb{R}^3)} \int_{\Omega} \varphi(\nabla v(x)) \, dx = 0.$$

On en déduit que

$$\nabla u(x) \in \{w_i, 1 \leq i \leq 5\} \text{ p.p. } x \in \Omega.$$

En particulier, tous les coefficients non diagonaux des 5 puits étant nuls, pour tout $i \in \{1,2,3\}$, la composante u_i n'est en fait fonction que de x_i . Mais $u = 0$ sur $\partial\Omega$, donc en fait $u = 0$ sur Ω , ce qui est absurde puisqu'alors $\nabla u = 0$, ce qui est absurde si $\nabla u(x) \in \{w_i, 1 \leq i \leq 5\}$ p.p. $x \in \Omega$. ■

Théorème 4.3 : soit une suite minimisante (u_k) de (4.24), i.e. telle que

$$u_k \in W_0^{1,\infty}(\Omega; \mathbb{R}^3) \quad \text{et} \quad \lim_{k \rightarrow \infty} \int_{\Omega} \varphi(\nabla u_k(x)) \, dx = 0.$$

Alors,

$$u_k \rightarrow 0 \text{ uniformément sur } \Omega. \quad (4.26)$$

$$\nabla u_k \rightarrow 0 \text{ dans } (L^\infty(\Omega))^9 \text{ — faible } * \quad (4.27)$$

De plus, la suite des gradients (∇u_k) engendre une mesure de Young unique donnée par :

$$\nu_x = \frac{1}{8} (\delta_{w_1} + \delta_{w_2} + \delta_{w_3} + \delta_{w_4}) + \frac{1}{2} \delta_{w_5}, \quad (4.28)$$

où δ_F représente la masse de Dirac en $F \in M_3(\mathbb{R})$.

Preuve : puisque $u_k \in W_0^{1,\infty}(\Omega; \mathbb{R}^3)$ et grâce à la compacité de l'injection canonique de $W_0^{1,\infty}(\Omega; \mathbb{R}^3)$ dans $C^0(\Omega; \mathbb{R}^3)$, quitte à extraire une sous-suite de (u_k) , il existe une fonction $u \in W_0^{1,\infty}(\Omega; \mathbb{R}^3)$ telle que :

$$\begin{cases} \lim_{k \rightarrow \infty} (u_k) = u \text{ uniformément sur } \Omega & (4.29) \\ \lim_{k \rightarrow \infty} (\nabla u_k) = \nabla u \text{ dans } (L^\infty(\Omega))^9 \text{ — faible } * & (4.30) \end{cases}$$

Si ν_x est la mesure de Young associée à (u_k) , alors on rappelle que, pour toute fonction continue f de $\mathbb{R}^{3 \times 3}$ ou $M_3(\mathbb{R})$ dans $\mathbb{R}$:

$$\lim_{k \rightarrow \infty} f(\nabla u_k(x)) = f(\nabla u(x)) = \int_{M_3(\mathbb{R})} f(A) \, d\nu_x(A) \quad p.p. \, x \in \Omega \text{ dans } L^\infty(\Omega) \text{ faible } - *.$$

En particulier, pour $f = \varphi$, on sait que, par définition de la suite (u_k) :

$$\int_{\Omega} \int_{M_3(\mathbb{R})} \varphi(A) \, d\nu_x(A) \, dx = 0.$$

Donc $\text{supp}(\nu_x) \subset \{w_1, w_2, w_3, w_4, w_5\}$, i.e. il existe 5 fonctions mesurables positives α_i telles que :

$$\nu_x = \sum_{i=1}^5 \alpha_i(x) \cdot \delta_{w_i} \quad p.p. \, x \in \Omega,$$

et donc pour toute fonction continue f de $M_3(\mathbb{R})$ dans $\mathbb{R}$, on a en fait :

$$f(\nabla u(x)) = \sum_{i=1}^5 \alpha_i(x) \cdot f(w_i) \quad p.p. x \in \Omega \quad (4.31)$$

En particulier, pour $(k, l) \in \{1, 2, 3\}^2$, les fonctions $f_{kl} : A \mapsto A_{kl}$, sont continues. On peut leur appliquer la relation (4.31) :

$$(\nabla u(x))_{kl} = \sum_{i=1}^5 \alpha_i(x) \cdot (w_i)_{kl} \quad p.p. x \in \Omega,$$

soit en fait

$$\nabla u(x) = \sum_{i=1}^5 \alpha_i(x) \cdot w_i \quad p.p. x \in \Omega, \quad (4.32)$$

et en particulier pour les coefficients non diagonaux :

$$(\nabla u(x))_{kl} = 0 \quad p.p. x \in \Omega \quad \text{et pour } k \neq l.$$

L'argument déjà utilisé pour prouver le théorème 4.2 conduit finalement à

$$u = 0 \text{ sur } \Omega,$$

ce qui finit de prouver les affirmations (4.29) et (4.30). Mais alors, la relation (4.32) s'écrit désormais :

$$\sum_{i=1}^5 \alpha_i(x) \cdot w_i = 0 \quad p.p. x \in \Omega,$$

ou, de manière équivalente sur les coefficients diagonaux des matrices :

$$\begin{cases} \alpha_1(x) + 2\alpha_2(x) - \alpha_3(x) - 2\alpha_4(x) = 0 \\ 2\alpha_1(x) - \alpha_2(x) - 2\alpha_3(x) + \alpha_4(x) = 0 \\ \alpha_1(x) + \alpha_2(x) + \alpha_3(x) + \alpha_4(x) - \alpha_5(x) = 0 \end{cases} \quad p.p. x \in \Omega. \quad (4.33)$$

La fonction constante : $A \mapsto 1$, et le déterminant étant continues, on peut leur appliquer la relation (4.31) :

$$\begin{cases} 1 = \sum_{i=1}^5 \alpha_i(x) \\ \det(\nabla u(x)) = 0 = \sum_{i=1}^5 \alpha_i(x) \cdot \det(w_i) \end{cases} \quad p.p. x \in \Omega,$$

soit en fait

$$\begin{cases} \alpha_1(x) + \alpha_2(x) + \alpha_3(x) + \alpha_4(x) + \alpha_5(x) = 1 \\ 2\alpha_1(x) - 2\alpha_2(x) + 2\alpha_3(x) - 2\alpha_4(x) = 0 \end{cases} \quad p.p. x \in \Omega. \quad (4.34)$$

On vérifie facilement que les systèmes (4.33) et (4.34) forment un système de Cramer qui conduit à :

$$\alpha_1(x) = \alpha_2(x) = \alpha_3(x) = \alpha_4(x) = \frac{1}{8} \quad \text{et} \quad \alpha_5(x) = \frac{1}{2} \quad p.p. x \in \Omega.$$

■

Théorème 4.4 : avec les définitions précédentes, il existe deux constantes C_1 et C_2 , indépendantes du pas h de la triangulation de Ω , telles que :

$$\inf_{v \in V_0^h(\Omega, \mathbb{R}^3)} \int_{\Omega} \varphi(\nabla v(x)) \, dx \leq C_1 \cdot e^{-C_2 \cdot |\ln h|^{1/2}} \quad (4.35)$$

De plus, le problème (4.35) admet un minimiseur dans $V_0^h(\Omega; \mathbb{R}^3)$.

Notation : dans ce qui suit, il a été commode de poser :

$$\text{si } A = (a_{ij})_{\substack{1 \leq i \leq 3 \\ 1 \leq j \leq 3}} \in M_3(\mathbb{R}), \text{ alors } \check{A} = (a_{ij})_{\substack{1 \leq i \leq 2 \\ 1 \leq j \leq 2}} \in M_2(\mathbb{R}). \quad (4.36)$$

Preuve du théorème 4.4 : on reprend la fonction, notée ici $\check{u}^p \in W^{1,\infty}(\mathbb{R}^2; \mathbb{R}^2)$ par cohérence avec la notation (4.36), construite dans la démonstration du théorème 4.1 : pour un certain réel $\alpha \in]0,1[$, $\check{u}^p$ est $2h^\alpha$ -périodique et, sur chaque carré du type $[2m \cdot h^\alpha, 2(m+1)h^\alpha]^2$, $m \in \mathbb{Z}$, lui-même subdivisé en 4 carrés se ramenant à $[0, h^\alpha]^2$, $\check{u}^p$ est construite en itérant p fois un procédé de lamination du carré $[0, h^\alpha]^2$ en k rectangles. Ce procédé donne à $\check{u}^p$ des gradients égaux à l'une des matrices $\check{w}_i$, $1 \leq i \leq 4$, sur une partie de $[0, h^\alpha]^2$ dont la mesure augmente à chaque itération. La première de ces itérations est résumée sur la figure 53.

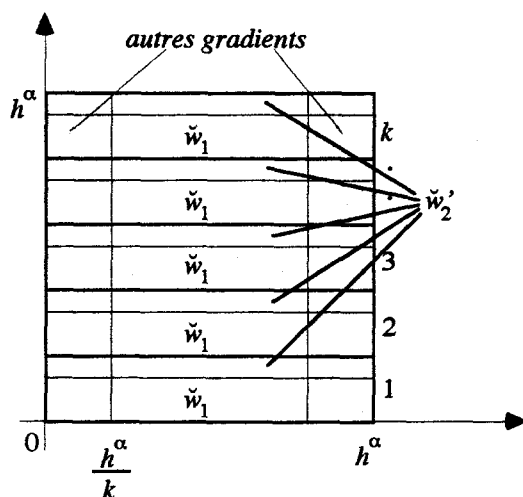


fig. 53

Pour un certain réel $\beta \in]0, \alpha[$, on définit alors la fonction u^p par :

$$\bullet i = 1, 2 : u_i^p(x_1, x_2, x_3) = \begin{cases} \text{sgn}(\check{u}_i^p(x_1, x_2)) \cdot \min(|\check{u}_i^p(x_1, x_2)|, x_3, h^\beta - x_3) & \text{si } 0 \leq x_3 \leq h^\beta \\ 0 & \text{si } h^\beta \leq x_3 \leq 2h^\beta \end{cases}$$

$$\bullet u_3^p(x_1, x_2, x_3) = \min(x_3, 2h^\beta - x_3)$$

et $2h^\beta$ -périodicité par rapport à la variable x_3 .

Notons que, pour h suffisamment petit, $h^\beta \gg h^\alpha$ puisque $\beta < \alpha$. Graphiquement, pour x_1 et x_2 fixés :

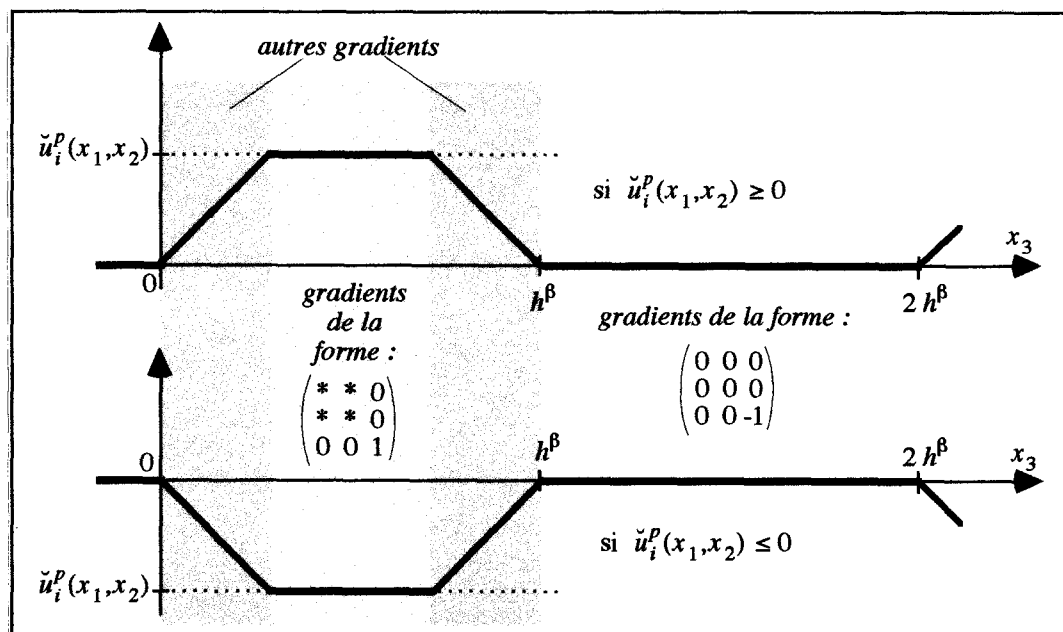


fig. 54 : graphe de $u_i^p(x_1, x_2, \cdot)$ pour $i = 1, 2$

On définit enfin :

$$\bullet \hat{u}_i^p(x) = \text{sgn}(u_i^p(x)) \cdot \min(|u_i^p(x)|, \text{dist}(x, \partial\Omega)) \text{ pour } i = 1, 2, 3, \quad (4.37)$$

$$\bullet u_h^p \text{ est l'interpolée de } \hat{u}^p \text{ sur la triangulation } T. \quad (4.38)$$

Par construction de $\hat{u}^p$, puis de u^p , tous les gradients utilisés sont bornés indépendamment de h , donc, globalement, ∇u^p est bornée indépendamment de h . Et pour une triangulation T régulière, il vient :

$$|\nabla u_h^p| \leq C \quad (4.39)$$

pour une certaine constante $C > 0$ indépendante de h .

Notons $S = \{x \in \Omega / \nabla u_h^p(x) \notin \{w_1, w_2, w_3, w_4, w_5\}\}$ la partie de Ω sur laquelle le gradient de u_h^p ne coïncide avec aucun des 5 puits w_i . Alors, par construction, $\hat{u}^p \in W_0^{1,\infty}(\Omega; \mathbb{R}^3)$ et $u_h^p \in V_0^h(\Omega; \mathbb{R}^3)$, et les gradients de u_h^p sont uniformément bornés indépendamment de h (4.39), si bien qu'il existe une constante $C > 0$ telle que :

$$\int_{\Omega} \varphi(\nabla u_h^p(x)) \, dx \leq C \cdot |S| \quad (4.40)$$

où $|S|$ est la mesure de Lebesgue de S .

S est constituée de 3 sous-sensembles S_1, S_2 et S_3 :

• S_1 est un voisinage de $\partial\Omega$, issu de la troncature au bord (4.37). Son épaisseur est inférieure à $c \cdot h^\beta$ pour un certain réel $c > 0$ puisque, par construction de u^p :

$$|u_i^p| \leq 2 \cdot h^\beta \quad (4.41)$$

si bien que, pour une certaine constante $C > 0$ tenant compte de la "superficie" de $\partial\Omega$ dans $\mathbb{R}^3$:

$$|S_1| \leq C \cdot h^\beta \quad (4.42)$$

• S_2 est la réunion de voisinages situés d'un côté de chaque plan d'équation $x_3 = k.h^\beta$ pour $k \in \mathbb{Z}$ (zones grisées sur la figure 54). Par construction de $\tilde{u}^p$:

$$|\tilde{u}_i^p| \leq h^\alpha \left(1 + \frac{1}{3k} + \dots + \frac{1}{(3k)^p} \right) \leq h^\alpha \cdot \frac{1}{1 - \frac{1}{3k}}, \text{ pour } i = 1, 2,$$

cette dernière expression étant majorée pour $k \geq 3$ par sa valeur en $k = 3$, soit $|\tilde{u}^p| \leq \frac{10}{9} h^\alpha$. L'épaisseur des voisinages précédents est donc elle aussi majorée par $\frac{10}{9} h^\alpha$. Donc, pour une certaine constante $C > 0$ indépendante de h :

$$|S_2| \leq \frac{10}{9} h^\alpha \cdot \frac{\text{diam}(\Omega)}{h^\beta} = C.h^{\alpha-\beta} \quad (4.43)$$

• Il reste à examiner le domaine $S_3 \subset \{x \in \Omega / 2k.h^\beta + 2.h^\alpha \leq x_3 \leq (2k+1).h^\beta - 2.h^\alpha, k \in \mathbb{Z}\}$, correspondant à la zone hachurée sur la figure 54, où l'on a :

$$u^p(x_1, x_2, x_3) = \left(\tilde{u}_1^p(x_1, x_2), \tilde{u}_2^p(x_1, x_2), x_3 \right) \quad (4.44)$$

soit :

$$\nabla u^p = \begin{pmatrix} \nabla \tilde{u}^p & 0 \\ 0 & 0 & 1 \end{pmatrix} \quad (4.45)$$

S_3 est inclus dans un certain nombre $N_3 \leq \frac{\text{diam}(\Omega)}{2.h^\beta}$ de "tranches" de Ω , dont l'épaisseur dans la direction x_3 vaut $h^\beta - 4.h^\alpha$, et dont les coupes pour x_3 fixé ont une mesure dans $\mathbb{R}^2$ que l'on peut évaluer avec les résultats (4.6) à (4.12) : pour une certaine constante $C > 0$,

$$N_3.(h^\beta - 4.h^\alpha) \leq \frac{\text{diam}(\Omega)}{2} \cdot \frac{h^\beta - 4.h^\alpha}{h^\beta} \leq C \quad (4.46)$$

soit, compte-tenu de (4.6) à (4.12) et pour $k \geq 3$:

$$|S_3| \leq C \cdot \left\{ h^{1-\alpha} + \left(\frac{1}{3} \right)^p + \frac{1}{k} + k^{p+1}.h^{1-\alpha} \right\} \quad (4.47)$$

• Finalement, en regroupant (4.42), (4.43) et (4.47) :

$$|S| \leq C \cdot \left\{ h^\beta + h^{\alpha-\beta} + h^{1-\alpha} + \left(\frac{1}{3} \right)^p + \frac{1}{k} + k^{p+1}.h^{1-\alpha} \right\} \quad (4.48)$$

Il reste à choisir au mieux α , β , k et p .

On choisit α et β de sorte que $\beta = \alpha - \beta = 1 - \alpha$, soit $\alpha = \frac{2}{3}$ et $\beta = \frac{1}{3}$, ce qui conduit à :

$$h^\beta + h^{\alpha-\beta} + h^{1-\alpha} + k^{p+1}.h^{1-\alpha} = h^{1/3} + h^{1/3} + h^{1/3} + k^{p+1}.h^{1/3} \leq 4.k^{p+1}.h^{1/3},$$

puis, pour une autre constante $C > 0$ à :

$$|S| \leq C. \left\{ \left(\frac{1}{3}\right)^p + \frac{1}{k} + k^{p+1}.h^{1/3} \right\} \quad (4.49)$$

On cherche ensuite un réel λ tel que h^λ soit un majorant commun à $\frac{1}{k}$ et à $k^{p+1}.h^{1/3}$:

$$\frac{1}{k} \leq h^\lambda \iff k \geq h^{-\lambda} \text{ et } k^{p+1}.h^{1/3} \leq h^\lambda \iff k \leq h^{(3\lambda-1)/3(p+1)}$$

Les deux exposants coïncident si $-\lambda = \frac{3\lambda-1}{3(p+1)}$, soit $\lambda = \frac{1}{3(p+2)}$.

On veut alors pouvoir choisir

$$k = E(h^{-1/3(p+2)}), \quad (4.50)$$

ce qui ne conduit à $k \geq 3$ que si :

$$h \leq \left(\frac{1}{3}\right)^{3(p+2)} \quad (4.51)$$

Alors, sous la condition (4.51), il vient :

$$\frac{1}{2} h^{-1/3(p+2)} \leq h^{-1/3(p+2)} - 1 < k \leq h^{-1/3(p+2)} \quad (4.52)$$

soit, simultanément :

$$\frac{1}{k} \leq 2.h^{1/3(p+2)} \text{ et } k^{p+1}.h^{1/3} \leq h^{1/3(p+2)}$$

et donc, pour une nouvelle constante $C > 0$:

$$|S| \leq C. \left\{ \left(\frac{1}{3}\right)^p + h^{1/3(p+2)} \right\} \quad (4.53)$$

Or, pour

$$h = \left(\frac{1}{3}\right)^{3p(p+2)}, \tag{4.54}$$

la condition (4.51) est toujours vérifiée, et de plus $h^{1/3(p+2)} = \left(\frac{1}{3}\right)^p$, et pour une certaine constante $C > 0$:

$$|S| \leq C.\left(\frac{1}{3}\right)^p$$

Mais $h = \left(\frac{1}{3}\right)^{3p(p+2)}$ donne également $-\frac{\ln h}{3.\ln 3} = p(p+2) \leq 3p^2$ dès que $p \geq 1$, et donc, pour une certaine constante $c > 0$:

$$p \geq c.\sqrt{-\ln h}$$

soit finalement, pour deux constantes $C_1, C_2 > 0$:

$$\inf_{v \in V_0^h(\Omega, \mathbb{R}^n)} \int_{\Omega} \varphi(\nabla v(x)) \, dx \leq C.|S| \leq C_1.e^{-C_2.\ln h^{1/2}}$$

■

Remarque 4.2 : p étant le nombre d'itérations dans le processus de construction de la suite minimisante, les relations (4.50) et (4.54) permettent de déterminer le pas h du maillage et le nombre k de subdivisions des carrés nécessaires à la construction et aux majorations successives de la démonstration :

$$k = 3^p \text{ et } n = \frac{1}{h} = 3^{3p(p+2)} = (9k)^{3p}.$$

Les valeurs de n obtenues ainsi confirment des hypothèses telles que $h^\alpha \ll h^\beta$ nécessaires à la démonstration, mais laissent aussi présager d'énormes difficultés d'exploration numérique :

p	$k = 3^p$	n théorique = $(9k)^{3p}$	$h = \frac{1}{n}$	$h^\alpha = h^{2/3}$	$h^\beta = h^{1/3}$	$\frac{h^\beta}{h^\alpha} = \frac{h^\alpha}{h}$ = $n^{1/3}$
1	3	19 683	$\approx 5.10^{-5}$	$\approx 0,00137$	$\approx 0,037$	27
2	9	$\approx 2,82.10^{11}$	$\approx 3,5.10^{-12}$	$\approx 2,3.10^{-8}$	$\approx 1,5.10^{-4}$	6561

tableau 7 : taille des maillages en dimension 3

On observera, en particulier, le rapport $\frac{h^\alpha}{h}$, identique en dimensions 2 et 3 (cf tableau 5).

CONCLUSION

Ce travail a eu, pour point de départ, une tentative d'estimation d'énergie de suites minimisantes d'un problème de Calcul des Variations résistant à l'analyse. Ce type de problème lui-même est issu de l'étude de la gémellation lors des transitions de phase dans certains cristaux et s'applique, en particulier, à des matériaux très récents : alliages à mémoire de forme, cristaux liquides, etc...

L'approche numérique fut longue et souvent décevante. Ce qui n'apparaît pas dans ce rapport est le très grand nombre de directions infructueuses qui ont été explorées numériquement. Il n'en présente que la succession apparemment assez logique de quelques essais constructifs.

Les suites minimisantes obtenues par les méthodes classiques d'optimisation convexe ont rapidement fait apparaître la nature oscillante des solutions, puis le caractère multi-échelle de ces oscillations, sans pour autant conduire à des énergies significativement proches de l'énergie théorique minimale. L'explication est venue récemment grâce aux techniques utilisées dans les démonstrations des derniers théorèmes d'estimation, aussi bien en dimension 2 qu'en dimension 3.

Le théorème 4.1, et surtout sa démonstration, permet d'obtenir enfin une estimation analytique, et fournit simultanément une méthode de construction explicite d'une suite minimisante. Cette méthode confirme certains résultats obtenus précédemment, notamment les oscillations multi-échelle, et surtout ouvre de nouvelles voies pour l'étude numérique de ces problèmes. Malheureusement, cette méthode fait intervenir des triangulations extrêmement fines (voir les tableaux 5 et 6), ce qui n'est pas sans soulever de nouvelles difficultés de programmation, dues principalement à la taille des maillages à mettre en œuvre, et qui, a posteriori, explique les difficultés de descente sur des maillages de taille "raisonnable".

Ce travail se poursuit actuellement par le développement de nouveaux outils numériques exploitant la nature fractale des suites minimisantes théoriques, notamment à l'aide de maillages dynamiques.



INDEX DES TABLEAUX

<i>Page</i>	<i>Titre</i>
5	<i>tableau 1 : les 7 structures cristallines tri-dimensionnelles</i>
6	<i>tableau 2 : groupe des symétries des 7 structures cristallines tri-dimensionnelles</i>
7	<i>tableau 3 : transition de phase cubique centrée vers cubique à faces centrées</i>
50	<i>tableau 4 : descente vers le niveau d'énergie 576</i>
64	<i>tableau 5 : taille des maillages en dimension 2</i>
68	<i>tableau 6 : énergie limite selon les valeurs de δ</i>
79	<i>tableau 7 : taille des maillages en dimension 3</i>

RÉFÉRENCES
BIBLIOGRAPHIQUES

- [1] J.M. BALL, A version of the fundamental theorem for Young Measures, In : *Partial Differential Equations and Continuum Models of Phase Transitions*, Lecture Notes in Physics, vol. 344, M. Rascle, D. Serre, M. Slemrod Eds, 1989, pp 207-215.
- [2] J.M. BALL, R.D. JAMES, *Fine phase mixtures as minimizers of energy*, Arch. Rat. Mech. Anal. **100**, 1987, pp. 13-52.
- [3] H. BREZIS, *Analyse fonctionnelle, théorie et applications*, 1983, Masson.
- [4] B. BRIGHI, *Thèse*, 1991, Université de Metz.
- [5] B. BRIGHI, M. CHIPOT, *Approximated convex envelope of a function*, SIAM J. of Numerical Analysis **31**, 1, 1994, pp. 128-148.
- [6] M. CHIPOT, *Numerical analysis of oscillations in nonconvex problems*, Numer. Math. **59**, 1991, pp. 747-767.
- [7] M. CHIPOT, *Energy estimates for Lagrange finite elements*, Atti del Seminario Matematico e Fisico dell' Università di Modena, XLIII, 1995, pp. 305-316.
- [8] M. CHIPOT, *The appearance of microstructures in problems with incompatible wells and their numerical approach*, Numer. Math. **83**, 1999, pp. 325-352.
- [9] M. CHIPOT, C. COLLINS, D. KINDERLEHRER, *Numerical analysis of oscillations in multiple wells problems*, Numer. Math. **70**, 1995, pp 259-282.
- [10] M. CHIPOT, V. LÉCUYER, *Analysis and computations in the four-well problem*, Gakuto International Series, Mathematical Sciences and Applications, vol. 7, 1995, pp 67-78.
- [11] M. CHIPOT, V. LÉCUYER, *A new scheme for quasiconvex optimization in multiple-well problems*, 8th European Conference on Numerical Mathematics and Advanced Applications, (ENU-MATH 1997) - Bock & all. Eds, World Scientific, 1998, pp. 238-249.
- [12] P.G. CIARLET, *The Finite Elements Methode for Elliptic Problems*, 1978, North Holland, Amsterdam.
- [13] P.G. CIARLET, *Introduction à l'Analyse numérique matricielle et à l'Optimisation*, 1982, Masson.
- [14] C. COLLINS, D. KINDERLEHRER, M. LUSKIN, *Numerical approximation of the solution of a variational problem with a double well potential*, SIAM J. of Numerical Analysis **28**, 1991, pp. 321-332.
- [15] C. COLLINS, *Thesis : Computation and Analysis of Twinning in crystalline Solids*, 1990, University of Minnesota.

- [16] B. DACOROGNA, *Direct Methods in the Calculus of Variations*, 1989, Springer Verlag.
- [17] B. DACOROGNA, *Introduction au calcul des variations*, 1992, Presses Polytechniques et Universitaires Romandes, Lausanne.
- [18] S. KARTHA, T. CASTÀN, J.A. KRUMHANSL, J.P. SETHNA, *The Spin-Glass Nature of Tweed Precursors in Martensitic Transformations*, Phys. Rev. Lett. **67**, 3630 (1991).
- [19] M. LUSKIN, *On the computation of crystalline microstructure*, Acta Numerica, 1996, pp. 191-257.
- [20] P. PEDREGAL, *Parametrized measures and variational principles*, Birkhäuser, 1997.
- [21] P. PEDREGAL, V. ŠVERÁK, *A Note on Quasiconvexity and Rank-one Convexity for 2x2 Matrices*, Journal of Convex Analysis, **5**, 1998, No. 1, pp. 107-117.
- [22] D. SCHRYVERS, L.E. TANNER, *An atomic model for premartensitic tweed in Ni-Al as derived from HREM images*, Proceedings of the XIIth International Congress for Electron Microscopy.
- [23] V. ŠVERÁK, *Rank-one convexity does not imply quasiconvexity*, Proc. Roy. Soc. Edin. Sect. A, **120**, 1992, pp. 185-189.
- [24] L. TARTAR, *Some remarks on separately convex functions*, in : *Microstructure and phase transitions*, D. Kinderlehrer & all. Eds, 1992, Springer Verlag, pp 191-204.
- [25] V.A. TRÉNOGUINE, *Analyse Fonctionnelle*, 1980, Éditions MIR.
B.A. ТРЕНОГИН, *ФУНКЦИОНАЛЬНЫЙ АНАЛИЗ*, 1980, Издательство «Наука», Москва.
- [26] M. WINTER, *An example of microstructure with multiple scales*, EJAM **8,2**, 1997, pp 185-208.