



Damage and fracture evolution in brittle materials by shape optimization methods

Grégoire Allaire, François Jouve, Nicolas van Goethem

► To cite this version:

Grégoire Allaire, François Jouve, Nicolas van Goethem. Damage and fracture evolution in brittle materials by shape optimization methods. *Journal of Computational Physics*, 2011, 230 (12), pp.5010–5044. hal-00784048

HAL Id: hal-00784048

<https://inria.hal.science/hal-00784048>

Submitted on 4 Feb 2020

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

Damage and fracture evolution in brittle materials by shape optimization methods

Grégoire Allaire¹ François Jouve²
Nicolas Van Goethem^{1,3}

November 9, 2010

¹ Centre de Mathématiques Appliquées, École Polytechnique, 91128 Palaiseau, France.
Email: gregoire.allaire@polytechnique.fr

² Laboratoire J.-L. Lions, Université Paris Diderot (Paris 7), 75205 Paris, France.
Email: jouve@math.jussieu.fr

³ Universidade de Lisboa, Faculdade de Ciências, Departamento de Matemática, CMAF, Av. Prof. Gama Pinto, 1649-003 Lisboa, Portugal.
Email: vangoeth@ptmat.fc.ul.pt

Abstract

This paper is devoted to a numerical implementation of the Francfort-Marigo model of damage evolution in brittle materials. This quasi-static model is based, at each time step, on the minimization of a total energy which is the sum of an elastic energy and a Griffith-type dissipated energy. Such a minimization is carried over all geometric mixtures of the two, healthy and damaged, elastic phases, respecting an irreversibility constraint. Numerically, we consider a situation where two well-separated phases coexist, and model their interface by a level set function that is transported according to the shape derivative of the minimized total energy. In the context of interface variations (Hadamard method) and using a steepest descent algorithm, we compute local minimizers of this quasi-static damage model. Initially, the damaged zone is nucleated by using the so-called topological derivative. We show that, when the damaged phase is very weak, our numerical method is able to predict crack propagation, including kinking and branching. Several numerical examples in $2d$ and $3d$ are discussed.

1 Introduction

Fracture mechanics is a field of paramount importance which is the subject of intense research efforts, see [22, 25, 46] and reference therein. While many works address the issue of microscopic modelling of fractures and the coupling of some defect atomistic models with macroscopic elasto-plastic models, we focus on purely macroscopic models in the framework of continuum mechanics. Roughly

speaking such continuum models can be classified in two main categories. On the one hand, there are models of crack growth and propagation which assume that the crack is a $(d - 1)$ -dimensional hypersurface in dimension d (a curve in the plane, and a surface in the three-dimensional space). On the other hand, one can consider models of damage where there is a competition between the initial healthy elastic phase and another damaged elastic phase. The transition from healthy to damaged can be smooth (i.e., there is a continuous damage variable which measures to what extent, or local proportion, the material is damaged) or sharp (i.e., there is an interface between a fully healthy and fully damaged zones). The Francfort-Marigo model [35] of quasi-static damage evolution for brittle materials pertains to the latter category and it is the purpose of this work to propose a numerical implementation of such a model. One of our main conclusion is that, although the Francfort-Marigo model is a damage model, it is able to describe crack propagation, when the damaged phase is very weak, and it gives quite similar results to those obtained in [21, 22]. This is not so much a surprise (although not a proof, of course) since the numerical approach in these papers is based on a Γ -convergence approximation (à la Ambrosio-Tortorelli) which amounts to replace the original fracture model by a damage model.

Section 2 gives a complete description of the Francfort-Marigo damage model that we briefly summarize now. A smooth body $\Omega \subset \mathbb{R}^d$ ($d = 2, 3$) is filled with two elastic phases: the undamaged or “healthy” phase, and the damaged one which is much weaker. The damaged zone is $\Omega^0 \subset \Omega$, with characteristic function $\chi(x)$, and the healthy zone is the remaining region $\Omega^1 = \Omega \setminus \Omega^0$. The behavior of such a mixture is assumed to be linearly elastic with a perfect interface (i.e., natural transmission conditions take place at the interface). Starting from an initial configuration of damaged and healthy phases mixture $\chi_{init}(x)$ and for a given set of loads, the new damaged configuration $\chi_{opt}(x)$ is obtained by minimizing a total energy

$$J(\chi) = J_{elast}(\chi) + \kappa \int_{\Omega} \chi \, dV, \quad (1.1)$$

which is the sum of the elastic energy and a Griffith-like bulk energy for the creation of the damaged region (where $\kappa > 0$ is a material parameter representing the energy density released at the onset of damage), under an irreversibility constraint which forbids an initially damaged zone to become healthy anew, i.e.,

$$\chi(x) \geq \chi_{init}(x).$$

A quasi-static damage evolution model is then obtained by a time discretization of the force loading and by applying the previous constrained minimization at each time step.

For numerical purposes we represent the interface Σ between the damaged and healthy regions, Ω^0 and Ω^1 respectively, by a level set function. The level set method for front propagation, as introduced by S. Osher and J. Sethian [53], is well-known to be very convenient for this purpose, including the possibility of topology changes. Here, we take advantage of another feature of the level set

method, namely the local character of front displacement. In other words, we do not seek global minimizers of (1.1) but rather local minimizers obtained from the initial configuration χ_{init} by transporting it using the level set method. Although global minimization is the ultimate goal in many optimization problems (like, for example, shape optimization [3, 5]), it turns out to be an undesirable feature in the present problem of damage evolution. Indeed, as explained in [21], global minimization is mechanically not sound for a quasi-static evolution problem where meta-stable states should be preferred to globally stable states attained by crossing a high energy barrier.

In the context of the level set method, at each time step, the new damage configuration χ_{opt} is obtained from the initialization χ_{init} by solving a transport Hamilton-Jacobi equation with a normal velocity which is minus the shape derivative of the total energy (1.1). Section 3 is devoted to the computation of such a shape derivative, following Hadamard method of geometric optimization (see e.g. [3, 41, 49, 62]). Remark that this computation is not standard (and indeed new in the elasticity context, to the best of our knowledge) since it is an interface between two materials, rather than a boundary, which is moved and since the full strain and stress tensors are not continuous through the interface. Note however that, for continuous fields, the derivation with respect to the shape of an interface is already known, see e.g. [56, 63]. The numerical algorithm for the level set method is by now standard and is briefly recalled in Section 6.

One of the inconveniences of the level set method, as well as of most numerical methods for crack propagation, is its inability to nucleate damage and start a front evolution if there is no initial interface. Therefore, we use another ingredient to initialize our computations when no initial damaged zone is prescribed. Namely, we use the notion of topological derivative as introduced in [33, 37, 61], and applied to the case of elastic inclusions in [9, 10, 18] for inverse problems, and to cracks in [64]. The topological derivative aims at determining whether it is worth or not nucleating an infinitesimal damage inclusion in the healthy zone Ω^1 . This information is complementary to that obtained by shape variation since, on the one hand, the shape derivative cannot nucleate new inclusions and, on the other hand, once an inclusion is created, only the shape derivative can expand it further on. The notion of topological derivative will be detailed in Section 4.

The resulting numerical algorithm is somehow similar to previous algorithms in structural optimization [5, 65]. When the damaged phase is much weaker than the healthy phase (say, with a 10^{-3} ratio between the Young moduli) and for a suitably chosen Griffith energy release parameter κ (which scales like the inverse of the mesh size Δx), our numerical results are very similar to those of [21] which were obtained for a fracture model. Therefore we claim that our numerical implementation of the Francfort-Marigo damage model is able to simulate crack propagation. Numerical experiments, including a study of convergence under mesh refinement, are performed in Section 6. We believe our approach is simpler and computationally less intensive than other classical methods for crack propagation [1, 14, 17, 39, 40, 50, 51]. Let us emphasize that

level-set methods have already been used in fracture mechanics [19, 39, 40], usually in conjunction with the extended finite element method [47]. However, one novelty of our work is that we use a single level-set function instead of two for parametrizing the crack and that the weak damage phase avoids the use of discontinuous finite elements. After completion of this work we learned that similar ideas were independently introduced in [15] and [45]. A different approach, called eigendeformation, was recently proposed in [59]: it uses two fields, like in [21], and relies on a scaling resembling ours (see (2.11) below). Eventually Section 7 draw some conclusions on our numerical experiments which yield comparable but different results from those obtained by the back-tracking algorithm for global minimization proposed in [20], [22]. Our results, including some computations in $2d$, were announced in [7].

2 The Francfort-Marigo model of damage

2.1 Description of the model

This section gives a comprehensive description of the Francfort-Marigo model [35] of quasi-static damage evolution for brittle materials. In a smooth domain $\Omega \subset \mathbb{R}^d$ this damage model is stated as a macroscopic phase transition problem between a damaged phase occupying a subset $\Omega^0 \subset \Omega$ and an healthy phase in the remaining region $\Omega^1 = \Omega \setminus \Omega^0$. To simplify the presentation, in a first step we consider a static problem starting from a healthy configuration (namely, without any irreversibility constraint). The characteristic function of Ω^0 is denoted by $\chi(x)$. The healthy and damaged phases are both assumed to be linear, isotropic and homogeneous, so we work in a linearized elasticity framework and the Lamé tensor of elasticity in Ω is

$$A_\chi = A^1(1 - \chi) + A^0\chi,$$

where $0 < A^0 < A^1$ are the Lamé tensors of isotropic elasticity in the damaged and healthy regions, respectively, defined by

$$A^{0,1} = 2\mu^{0,1}I_4 + \lambda^{0,1}I_2 \otimes I_2$$

where I_2 and I_4 denote the identity 2nd and 4th order tensors, respectively.

The boundary of the body is made of two parts, $\partial\Omega = \Gamma_D \cup \Gamma_N$, where a Dirichlet boundary condition u_D is imposed on Γ_D and a Neumann boundary condition g is imposed on Γ_N . We assume that $u_D \in H^1(\Omega; \mathbb{R}^d)$, $g \in L^2(\partial\Omega; \mathbb{R}^d)$ and we consider also a body force $f \in L^2(\Omega; \mathbb{R}^d)$. (Slightly stronger regularity assumptions on the data f, g, u_D will be made in the sequel.) We denote by n the unit normal vector on $\partial\Omega$. We introduce the affine space of kinematically admissible displacement fields

$$V = \{u \in H^1(\Omega; \mathbb{R}^d) \text{ such that } u = u_D \text{ on } \Gamma_D\}.$$

As usual, the strain and stress tensors associated to a displacement u write as

$$e(u) = \frac{1}{2}(\nabla u + \nabla^T u), \quad \sigma(u) = A_\chi e(u). \quad (2.1)$$

The elasticity system reads as

$$\begin{cases} -\operatorname{div}(A_\chi e(u_\chi)) = f & \text{in } \Omega, \\ u_\chi = u_D & \text{on } \Gamma_D, \\ A_\chi e(u_\chi)n = g & \text{on } \Gamma_N. \end{cases} \quad (2.2)$$

It is well-known that (2.2) can be restated as a minimum potential energy principle, that is, the displacement field $u_\chi \in V$ minimizes in V the energy functional

$$\mathcal{P}_\chi(u) = \int_\Omega \left(\frac{1}{2} A_\chi e(u) \cdot e(u) - f \cdot u \right) dV - \int_{\Gamma_N} g \cdot u dS,$$

i.e.,

$$\mathcal{P}_\chi(u_\chi) = \min_{u \in V} \mathcal{P}_\chi(u).$$

The Francfort-Marigo model amounts to minimize jointly over u and χ a total energy which is the sum of the elastic potential energy and of a Griffith-type energy (accounting for the creation of the damaged region), writing as

$$\mathcal{J}(u, \chi) = \mathcal{P}_\chi(u) + \kappa \int_\Omega \chi dV, \quad (2.3)$$

where κ is a positive material parameter which represents the release of elastic energy due to the decrease of rigidity at the onset of damage and can be interpreted as a density of dissipated energy of the damaged region. We call κ the Griffith energy release parameter. In other words, the Francfort-Marigo model is based on the minimization over $\chi \in L^\infty(\Omega; \{0, 1\})$ of

$$J(\chi) = \mathcal{J}(u_\chi, \chi) = \min_{u \in V} \mathcal{J}(u, \chi). \quad (2.4)$$

Instead of writing (2.4), we can first minimize in χ and later in u (since (2.3) is doubly minimized, the order of minimization does not matter). Since $\chi(x)$ takes only the values 0 and 1, the minimization is easy, provided that we know u_χ (which is of course never the case). Indeed, minimizing (2.4) is equivalent to the following local minimization at each point $x \in \Omega$

$$\min_{\chi \in \{0, 1\}} \left\{ \frac{1}{2} A_\chi e(u_\chi) \cdot e(u_\chi) + \kappa \chi \right\}(x),$$

providing a transition criterion from the healthy to the damaged phase as soon as the release of elastic energy is larger than the threshold κ . More precisely, a point x is damaged if and only if

$$\frac{1}{2} A^1 e(u_\chi) \cdot e(u_\chi)(x) - \frac{1}{2} A^0 e(u_\chi) \cdot e(u_\chi)(x) \geq \kappa. \quad (2.5)$$

After minimization in χ we obtain a non-linear non-convex functional to be minimized in V

$$\begin{aligned} \mathcal{E}(u) &= \frac{1}{2} \int_\Omega \min(A^1 e(u) \cdot e(u), A^0 e(u) \cdot e(u) + 2\kappa) dV \\ &\quad - \int_\Omega f \cdot u dV - \int_{\Gamma_N} g \cdot u dS. \end{aligned} \quad (2.6)$$

In truth the Francfort-Marigo model is quasi-static which means that we consider a sequence of minimization problems of the above type, with an additional thermodynamic irreversibility constraint. The time is discretized by an increasing sequence $(t_i)_{i \geq 1}$, with $t_1 = 0$ and $t_i < t_{i+1}$. At each time t_i the loads are denoted by f_i and g_i , the imposed boundary displacement is $u_{D,i}$, the affine space of kinematically admissible displacement fields is V_i , the characteristic function of the damaged phase is χ_i and the corresponding displacement is u_{χ_i} , solution of (2.2) with loads f_i and g_i and Dirichlet boundary condition $u_{D,i}$. The initial damaged zone is given and characterized by χ_0 .

The model is irreversible which means that a material point $x \in \Omega$ which is damaged at a previous time must remain damaged at a later time t_i , i.e.,

$$\chi_i(x) \geq \chi_{i-1}(x). \quad (2.7)$$

Therefore, introducing $\mathcal{J}_i$ and J_i , which are defined as (2.3) and (2.4) with the loads at time t_i , the Francfort-Marigo model is a sequence, indexed by $i \geq 1$, of minimization problems

$$\inf_{\chi \in L^\infty(\Omega; \{0,1\}), \chi \geq \chi_{i-1}} J_i(\chi) = \inf_{u \in V_i, \chi \in L^\infty(\Omega; \{0,1\}), \chi \geq \chi_{i-1}} \mathcal{J}_i(u, \chi), \quad (2.8)$$

with minimizers χ_i and u_{χ_i} (if any).

2.2 Mathematical properties of the model

The Francfort and Marigo model is ill-posed, namely, there does not exist any minimizer of (2.8) in most cases. This can easily be seen because (2.8) is equivalent to the minimization of the non-linear elastic energy (2.6) which is not convex, neither quasi-convex. Actually, one of the main purposes of the seminal paper [35] of Francfort and Marigo was to relax the minimization problem (2.8) and show the existence of suitably generalized solutions. The relaxation of (2.8) amounts to introduce composite materials, obtained by a fine mixing of the two phases, as competitors in the minimization of the total energy. Such composite materials include the limits, in the sense of homogenization, of minimizing sequences of (2.8): they are characterized by a phase volume fraction in the range $[0, 1]$ and a homogenized elasticity tensor which is the output of the microstructure at given volume fractions. It turns out that optimal microstructures are found in the class of sequential laminates. For further details we refer to [35] for the first time step and to [34] for the following time steps (where the irreversibility constraint plays a crucial role). This relaxed approach has been used for numerical computations of damage evolution in [4].

One drawback of the Francfort-Marigo approach is that it relies on global minimization, i.e., at each time step t_i the functional $\mathcal{J}_i(u, \chi)$ is globally minimized with respect to both variables u and χ . There is no true mechanical motivation for insisting on global minimization with respect to χ . Because of global minimization, damage might occur at time step t_i in a region far away from the initially damaged zone at the previous time step t_{i-1} , whereas, in most

circumstances, it seems more natural from a physical viewpoint to have expansion of the previously damaged area. Therefore, in a quasi-static regime which may favor metastability effects, it seems reasonable to prefer local minimization (with respect to χ) instead of global minimization. In the context of fracture mechanics it was proved in [51] that criticality solutions of the Griffith model are different from the energy globally minimizing solutions proposed by Francfort and Marigo.

Unfortunately, for a scalar-valued version of our damage model (antiplane elasticity), it was recently proved in [38] that local minima are actually global ones (both in the original setting of characteristic functions or in the relaxed setting of composite materials, locality being evaluated in the $L^1(\Omega)$ -norm). However this last result of [38] does not prevent the possibility of a different framework in which local minimizers would not be global ones (see, for example, the notion of ε -stable minimizer in [44]). In the present paper we propose such a framework based on the notion of front propagation in the original case of a macroscopic distribution of healthy and damaged phases (i.e., not considering composite materials). Instead of representing a damaged zone by a characteristic function $\chi \in L^\infty(\Omega; \{0, 1\})$ we rather introduce the interface Σ between the healthy and the damaged regions. Admissible variations of this interface are obtained in the framework of Hadamard method of shape variations [3], [41], [49], [55], [62] (see Section 3 below). More precisely, the minimization in (2.8) is restricted to configurations which are obtained by a Lipschitz diffeomorphism from a reference or an initial configuration. This is a severe restriction of the space of admissible designs since, for example, all configurations share the same topology as the reference one. As a consequence there cannot be nucleation of new damaged zones away from the initial one. This leaves open the possibility of the existence of local, but not global, minimizers. We shall not prove anything rigorously on this issue but our numerical simulations indicate that they do indeed exist. Let us remark that the chosen numerical approach by level sets allows for topology changes by breaking a damage region in two parts, but never by creating a new damage region.

On the other hand, working in the framework of Hadamard method of front representation does not help at all concerning the existence of (local or global) minimizers. Once again we are speechless on this issue. Of course, one simple remedy is to add a surface energy in the minimized total energy

$$\mathcal{J}_{reg}(u, \chi) = \mathcal{P}_\chi(u) + \kappa \int_\Omega \chi \, dV + \kappa' TV(\chi), \quad (2.9)$$

with the total variation norm defined by

$$TV(\chi) = \sup_{\substack{\phi \in C^1(\Omega; \mathbb{R}^d) \\ |\phi| \leq 1 \text{ in } \Omega}} \int_\Omega \chi \operatorname{div} \phi \, dV.$$

When χ is the characteristic function of a smooth subset Ω^0 , the number $TV(\chi)$ is also the perimeter of Ω^0 . A possible justification of this new term in (2.9)

is to consider a Griffith surface energy on top of the previous Griffith bulk energy. We call "regularized" the energy in (2.9) since it is well-known to admit minimizers χ in the class $L^\infty(\Omega; \{0, 1\})$ [13]. In truth, if we mention this additive surface energy, this is because the unavoidable numerical diffusion of our computational algorithm has precisely the effect of adding such a surface energy. For our numerical tests, we shall not rely on (2.9) and rather we use the standard energy (2.3).

2.3 Goal of the present study

The goal of this paper is to propose and test the following numerical method for the damage model of Francfort and Marigo. At each time step t_i the minimization (2.8) is performed by Hadamard method of shape sensitivity. In other words, we compute the shape derivative of the objective function J_i with respect to the interface between the healthy and damaged phases and, applying a steepest descent algorithm, we move this interface in (minus) the direction of the shape gradient. The minimization of J_i is stopped when the shape gradient is (approximately) zero, i.e., at a stationary point (a local minimizer in numerical practice) of the objective function. We use a level set approach to characterize the interface between the healthy and damaged phases. As is well known, it allows for large deformations of the interface with possibly topology changes. After convergence at time t_i , we pass to the next discrete time t_{i+1} by changing the loads and we start a new minimization of J_{i+1} , taking into account the irreversibility constraint (2.7). We iterate until a final time $t_{i_{final}}$ which we choose when the structure is almost entirely damaged.

We propose two possible ways of initializing our computations. Either we start from an initial damaged zone χ_0 at time $t_1 = 0$, or, in case the initial structure is not damaged at all, we nucleate a small damaged zone by using the notion of topological derivative. This nucleation step takes place before we start the first minimization of J_1 . In particular, the resulting initial damaged zone is usually not a local minimizer of the total energy (2.3). We are thus able to predict damage propagation without prescribing any initial crack as is commonly done in engineering practice.

Although the considered model has been designed in the framework of damage mechanics, it turns out to be able to accurately describe crack propagation in some specific regimes. More precisely, when the damaged phase is very weak (its rigidity A_0 is negligible) and the energy release rate is large enough, the results of our numerical computations are cracks rather than damaged sub-domains. In other words, the damaged zone is a thin hypersurface with a thickness of a few mesh cells concentrating along a curve in $2d$ or a surface in $3d$. However, our model, based on the minimization of (2.3), has no intrinsic lengthscale as opposed to other fracture models where there is a competition between bulk (elastic) energy and surface (crack extension) energy [22]. Therefore we must introduce some characteristic lengthscale in our model if we want to support our claim that it is able to predict crack propagation. We do this at a numerical level by requiring that our fracture results are convergent under mesh refine-

ment, a necessary condition for any reasonable numerical algorithm. To obtain such a convergence we scale the Griffith bulk energy release parameter κ like $1/\Delta x$, where Δx is the mesh size which is refined. More precisely, we introduce a characteristic lengthscale ℓ and we define a new material parameter γ which can be thought of as a Griffith surface energy release parameter (or fracture toughness in the language of fracture mechanics)

$$\gamma = \kappa \ell. \quad (2.10)$$

Then, instead of minimizing (2.3), we minimize (assuming, for simplicity, that there are only bulk forces)

$$\mathcal{J}_{\Delta x}(u, \chi) = \int_{\Omega} \left(\frac{1}{2} A_{\chi} e(u) \cdot e(u) - f \cdot u \right) dV + \frac{\gamma}{\Delta x} \int_{\Omega} \chi dV, \quad (2.11)$$

where $\gamma/\Delta x$ has the same physical units than κ . Although (2.11) has been written in a continuous framework, we are actually interested in its discretized version for a mesh of size Δx obtained, for example, with piecewise affine continuous Lagrange finite elements for u and piecewise constant finite elements for χ . In other words, rather than (2.11) we consider

$$\mathcal{J}_{\Delta x}(u_{\Delta x}, \chi_{\Delta x}) = \int_{\Omega} \left(\frac{1}{2} A_{\chi} e(u_{\Delta x}) \cdot e(u_{\Delta x}) - f \cdot u_{\Delta x} \right) dV + \frac{\gamma}{\Delta x} \int_{\Omega} \chi_{\Delta x} dV \quad (2.12)$$

where the minimization carries over the fields $u_{\Delta x}$ and $\chi_{\Delta x}$ belonging to the above finite element spaces (of finite dimension, linked to the mesh size Δx). When Δx goes to zero, we expect that, for a minimizing sequence $\chi_{\Delta x}$, the last term of (2.11) converges to a surface energy

$$\lim_{\Delta x \rightarrow 0} \frac{\gamma}{\Delta x} \int_{\Omega} \chi_{\Delta x} dV = \gamma \int_{\Gamma} dS,$$

where Γ is the crack curve in $2d$ or surface in $3d$. The numerical examples of Section 6 show that it is indeed the case: the damage zone concentrates around a surface Γ with a thickness of a few cells Δx . We believe that the discrete scaled energies (2.12) converges, in some sense to be made precise, as Δx and A_0 go to zero, to the fracture model

$$\min_{u, \Gamma} \int_{\Omega \setminus \Gamma} \left(\frac{1}{2} A_1 e(u) \cdot e(u) - f \cdot u \right) dV + \gamma \int_{\Gamma} dS \quad (2.13)$$

where the displacement field u may be discontinuous through the crack Γ . We are not able to prove such a result which would first require to order the speed of convergence of Δx and A_0 to zero. Remark however that similar results of convergence of a sequence of discrete energies to a continuous limit energy have already been obtained, e.g., in the context of image segmentation for discrete Mumford-Shah energies [29] or for spin systems [24]. A natural candidate for the type of convergence of (2.12) to (2.13) would be of course Γ -convergence.

However, since our numerical approach relies on some type of local minimizers, whereas Γ -convergence only deals with global minimizers, one should pertain to a variant of Γ -convergence for "local minimizers" (a notion to be made precise) as in the recent works [23], [43], [58]. Although a convergence of (2.12) to (2.13) would probably be difficult and quite technical to prove, our numerical results are a clear indication that it may hold true.

This conjectured link between the damage model (2.11) and the fracture model (2.13) is, of course, reminiscent (but not equivalent) of the numerical approach in [21], [22] where a fracture model is numerically approximated by a damage model (based on the Γ -convergence result of [8]).

3 Shape derivative

3.1 On the notion of shape gradient

Shape differentiation is a classical topic [3], [41], [49], [55], [62]. We briefly recall its definition and main results in the present context. Here, the overall domain Ω is fixed and we consider a smooth open subset $\omega \subset \Omega$ which may vary. Denoting by χ the characteristic function of ω , we consider variations of the type

$$\chi_\theta = \chi \circ (Id + \theta), \text{ i.e., } \chi_\theta(x) = \chi(x + \theta(x)),$$

with $\theta \in W^{1,\infty}(\Omega; \mathbb{R}^d)$ such that θ is tangential on $\partial\Omega$ (this last condition ensures that $\Omega = (Id + \theta)\Omega$). It is well known that, for sufficiently small θ , $(Id + \theta)$ is a diffeomorphism in Ω .

Definition 3.1 *The shape derivative of a function $J(\chi)$ is defined as the Fréchet derivative in $W^{1,\infty}(\Omega; \mathbb{R}^d)$ at 0 of the application $\theta \rightarrow J(\chi \circ (Id + \theta))$, i.e.*

$$J(\chi \circ (Id + \theta)) = J(\chi) + J'(\chi)(\theta) + o(\theta) \quad \text{with} \quad \lim_{\theta \rightarrow 0} \frac{|o(\theta)|}{\|\theta\|_{W^{1,\infty}}} = 0,$$

where $J'(\chi)$ is a continuous linear form on $W^{1,\infty}(\Omega; \mathbb{R}^d)$.

Lemma 3.1 ([41], [62]) *Let ω be a smooth bounded open subset of Ω and $\theta \in W^{1,\infty}(\Omega; \mathbb{R}^d)$. Let $f \in H^1(\Omega)$ and $g \in H^2(\Omega)$ be two given functions. Assume that Σ is a smooth subset of $\partial\omega$ with boundary $\partial\Sigma$. The shape derivatives of*

$$J_1(\omega) = \int_{\omega} f dV \quad \text{and} \quad J_2(\Sigma) = \int_{\Sigma} g dS$$

are $J'_1(\omega) = \int_{\partial\omega} f \theta \cdot n dS$ and

$$J'_2(\Sigma) = \int_{\Sigma} \left(\frac{\partial g}{\partial n} + gH \right) \theta \cdot n dS + \int_{\partial\Sigma} g \theta \cdot \tau dL, \quad (3.1)$$

respectively, where n is the exterior unit vector normal to $\partial\omega$, H is the mean curvature and τ is the unit vector tangent to $\partial\omega$ such that τ is normal to both $\partial\Sigma$ and n , and dL is the $(d-2)$ -dimensional measure along $\partial\Sigma$.

3.2 Main result

To simplify the notations we forget the time index t_i in this section. Although the state equation and the cost function of the Francfort-Marigo model are (2.2) and (2.4) respectively, we consider a slightly more general setting in this section (to pave the way to more general models in the future). More precisely, we consider a state equation

$$\begin{cases} -\operatorname{div}(A_\chi e(u_\chi)) = f_\chi & \text{in } \Omega, \\ u_\chi = u_D & \text{on } \Gamma_D, \\ A_\chi e(u_\chi)n = g_\chi & \text{on } \Gamma_N, \end{cases} \quad (3.2)$$

where

$$f_\chi := (1 - \chi)f^1 + \chi f^0 \quad \text{and} \quad g_\chi := (1 - \chi)g^1 + \chi g^0$$

with $f^k \in H^1(\Omega; \mathbb{R}^d) \cap C^{0,\alpha}(\Omega; \mathbb{R}^d)$ and $g^k \in H^2(\Omega; \mathbb{R}^d) \cap C^{1,\alpha}(\Omega; \mathbb{R}^d)$, $k = 0$ or 1 ($0 < \alpha < 1$). We also assume that u_D belongs to $H^2(\Omega; \mathbb{R}^d)$ and that the subset Ω^0 (with characteristic function χ) is smooth. Under these assumptions the solution u_χ of (3.2) belongs to $H^2(\Omega^0; \mathbb{R}^d)$ and $H^2(\Omega^1; \mathbb{R}^d)$ and is of class $C^{2,\alpha}$ away from the boundary and from the interface. The cost function is taken as

$$J(\chi) = \frac{1}{2} \int_{\Omega} A_\chi e(u_\chi) \cdot e(u_\chi) dV + \int_{\Omega} j_\chi(x, u_\chi) dV + \int_{\partial\Omega} h_\chi(x, u_\chi) dS, \quad (3.3)$$

where

$$j_\chi := (1 - \chi)j^1 + \chi j^0 \quad \text{and} \quad h_\chi := (1 - \chi)h^1 + \chi h^0$$

with $j^k(x, u)$ and $h^k(x, u)$, $k = 0, 1$, twice differentiable functions with respect to u , satisfying the following growth conditions

$$\begin{aligned} |j^k(x, u)| &\leq C(|u|^2 + 1), & |(j^k)'(x, u)| &\leq C(|u| + 1), & |(j^k)''(x, u)| &\leq C, \\ |h^k(x, u)| &\leq C(|u|^2 + 1), & |(h^k)'(x, u)| &\leq C(|u| + 1), & |(h^k)''(x, u)| &\leq C. \end{aligned} \quad (3.4)$$

where $'$ denotes the partial derivative with respect to $u \in \mathbb{R}^d$. To avoid some unnecessary technicalities we also assume that $h^1(x, u_D(x)) = h^0(x, u_D(x))$ on Γ_D so that the objective function is equal to

$$J(\chi) = \frac{1}{2} \int_{\Omega} A_\chi e(u_\chi) \cdot e(u_\chi) dV + \int_{\Omega} j_\chi(x, u_\chi) dV + \int_{\Gamma_N} h_\chi(x, u_\chi) dS + C$$

where C is a constant which does not depend on χ .

Remark 3.1 *When the imposed displacement on Γ_D vanishes, $u_D = 0$, the cost function of the Francfort-Marigo model simplifies and reduces to a multiple of the compliance. Indeed, the energy equality for the state equation (2.2) (which is valid only if $u_D = 0$), namely*

$$\int_{\Omega} A_\chi e(u_\chi) \cdot e(u_\chi) dV = \int_{\Omega} f \cdot u_\chi dV + \int_{\Gamma_N} g \cdot u_\chi dS,$$

implies that the cost function (2.4) (with $j_\chi = \kappa\chi - f \cdot u_\chi$ and $g_\chi = -g \cdot u_\chi$) reduces to

$$J(\chi) = \kappa \int_{\Omega} \chi dV - \frac{1}{2} \left(\int_{\Omega} f \cdot u_\chi dV + \int_{\Gamma_N} g \cdot u_\chi dS \right). \quad (3.5)$$

Of course, the study of (3.5) is much simpler than that of the general objective function (3.3). However, since many numerical tests involve non-homogeneous boundary displacements, $u_D \neq 0$, we must study (3.3) and not merely (3.5).

We need to introduce the so-called adjoint problem

$$\begin{cases} -\operatorname{div}(A_\chi e(p_\chi)) = f_\chi + j'_\chi(x, u_\chi) & \text{in } \Omega, \\ p_\chi = 0 & \text{on } \Gamma_D, \\ A_\chi e(p_\chi)n = g_\chi + h'_\chi(x, u_\chi) & \text{on } \Gamma_N. \end{cases} \quad (3.6)$$

We denote by Σ the interface between the damaged and healthy regions Ω^0 and Ω^1 . We define $n = n^0 = -n^1$ the outward unit normal vector to Σ . We use the jump notation

$$[\alpha] = \alpha^1 - \alpha^0 \quad (3.7)$$

for a quantity α that has a jump across the interface Σ .

The shape derivative of (3.3) will be an integral on the interface Σ as is clear from Lemma 3.1. The state u_χ and adjoint p_χ are continuous on Σ but not all their derivatives. Actually the tangential components of their deformation tensors are continuous as well as the normal vector of their stress tensors. To make this result precise, at each point of the interface Σ we introduce a local basis made of the normal vector n and a collection of unit tangential vectors, collectively denoted by t , such that (t, n) is an orthonormal basis. For a symmetric $d \times d$ matrix $\mathcal{M}$, written in this basis, we introduce the following notations

$$\mathcal{M} = \begin{pmatrix} \mathcal{M}_{tt} & \mathcal{M}_{tn} \\ \mathcal{M}_{nt} & \mathcal{M}_{nn} \end{pmatrix}$$

where $\mathcal{M}_{tt}$ stands for the $(d-1) \times (d-1)$ minor of $\mathcal{M}$, $\mathcal{M}_{tn}$ is the vector of the $(d-1)$ first components of the d -th column of $\mathcal{M}$, $\mathcal{M}_{nt}$ is the row vector of the $(d-1)$ first components of the d -th row of $\mathcal{M}$, and $\mathcal{M}_{nn}$ the (d, d) entry of $\mathcal{M}$. Let us recall that dV, dS and dL indicate volume integration in $\mathbb{R}^d$, and surface (or line, according to the value of d) integration in $\mathbb{R}^{d-1}$ and $\mathbb{R}^{d-2}$, respectively.

Lemma 3.2 *Let e and σ denote the strain and stress tensors of the solution to the state equation (3.2) or adjoint state equation (3.6). All components of σ_{nt} , σ_{nn} , and e_{tt} are continuous across the interface Σ (assumed to be smooth) while all other entries have jumps through Σ , rewritten in terms of these continuous quantities as*

$$\begin{cases} [e_{nn}] &= [(2\mu + \lambda)^{-1}] \sigma_{nn} - [\lambda(2\mu + \lambda)^{-1}] \operatorname{tr} e_{tt} \\ [e_{tn}] &= [(2\mu)^{-1}] \sigma_{tn} \\ [\sigma_{tt}] &= [2\mu] e_{tt} + ([2\mu\lambda(2\mu + \lambda)^{-1}] \operatorname{tr} e_{tt} + [\lambda(2\mu + \lambda)^{-1}] \sigma_{nn}) I_2^{d-1} \end{cases} \quad (3.8)$$

where I_2^{d-1} is the identity matrix of order $d - 1$.

Proof. By standard regularity theory, on both sides of the smooth interface Σ the solution, as well as its deformation and stress tensors e and σ , are smooth. This implies that the continuity of the displacement through the interface yields the continuity of e_{tt} . The transmission condition implies that σ_{tn} and σ_{nn} are also continuous on the interface. The other quantities have jumps (3.8) which are computed through the strain-stress relation (2.1). $\square$

Theorem 3.1 *Let Ω be a smooth bounded open set, Σ be a smooth hypersurface in Ω , $\gamma = \Sigma \cap \Gamma_N$ and $\theta \in W^{1,\infty}(\Omega; \mathbb{R}^d)$. The shape derivative in the direction θ of the objective function $J(\chi)$, as given by (3.3), is*

$$\begin{aligned} J'(\chi)(\theta) &= \int_{\Sigma} \mathcal{D}(x) \theta \cdot n \, dS \\ &+ \int_{\Sigma} ((f^0 - f^1) \cdot p_{\chi} + (j^0 - j^1)(x, u_{\chi})) \theta \cdot n \, dS \\ &+ \int_{\gamma} ((g^0 - g^1) \cdot p_{\chi} + (h^0 - h^1)(x, u_{\chi})) \theta \cdot \tau \, dL \end{aligned} \quad (3.9)$$

with

$$\begin{aligned} \mathcal{D}(x) &= -\left[\frac{1}{(\lambda + 2\mu)}\right] \sigma_{nn}(u_{\chi}) \sigma_{nn}(p_{\chi}) - \left[\frac{1}{\mu}\right] \sigma_{tn}(u_{\chi}) \cdot \sigma_{tn}(p_{\chi}) \\ &+ [2\mu] e_{tt}(u_{\chi}) \cdot e_{tt}(p_{\chi}) + \left[\frac{2\lambda\mu}{(\lambda + 2\mu)}\right] \text{tre}_{tt}(u_{\chi}) \text{tre}_{tt}(p_{\chi}) \\ &+ \left[\frac{\lambda}{(\lambda + 2\mu)}\right] (\sigma_{nn}(u_{\chi}) \text{tre}_{tt}(p_{\chi}) + \sigma_{nn}(p_{\chi}) \text{tre}_{tt}(u_{\chi})) \\ &+ \left[\frac{1}{2(\lambda + 2\mu)}\right] (\sigma_{nn}(u_{\chi}))^2 + \left[\frac{1}{2\mu}\right] |\sigma_{tn}(u_{\chi})|^2 - [\mu] |e_{tt}(u_{\chi})|^2 \\ &- \left[\frac{\lambda\mu}{\lambda + 2\mu}\right] (\text{tre}_{tt}(u_{\chi}))^2 - \left[\frac{\lambda}{\lambda + 2\mu}\right] \sigma_{nn}(u_{\chi}) \text{tre}_{tt}(u_{\chi}), \end{aligned} \quad (3.10)$$

where u_{χ} and p_{χ} are the solutions of the state equation (3.2) and adjoint equation (3.6), respectively, and where the brackets denotes the jump as defined by (3.7). Equivalently, $\mathcal{D}(x)$ can be rewritten as

$$\begin{aligned} \mathcal{D}(x) &= -\sigma_{nn}(u_{\chi})[e_{nn}(p_{\chi})] + e_{tt}(u_{\chi}) \cdot [\sigma_{tt}(p_{\chi})] - 2[e_{tn}(u_{\chi})] \cdot \sigma_{tn}(p_{\chi}) \\ &+ \frac{1}{2} \left(\sigma_{nn}(u_{\chi})[e_{nn}(u_{\chi})] - e_{tt}(u_{\chi}) \cdot [\sigma_{tt}(u_{\chi})] + 2\sigma_{tn}(u_{\chi}) \cdot [e_{tn}(u_{\chi})] \right). \end{aligned} \quad (3.11)$$

Remark 3.2 *A formula, partially symmetric to (3.11), holds true*

$$\begin{aligned} \mathcal{D}(x) &= -\sigma_{nn}(p_{\chi})[e_{nn}(u_{\chi})] + e_{tt}(p_{\chi}) \cdot [\sigma_{tt}(u_{\chi})] - 2[e_{tn}(p_{\chi})] \cdot \sigma_{tn}(u_{\chi}) \\ &+ \frac{1}{2} \left(\sigma_{nn}(u_{\chi})[e_{nn}(u_{\chi})] - e_{tt}(u_{\chi}) \cdot [\sigma_{tt}(u_{\chi})] + 2\sigma_{tn}(u_{\chi}) \cdot [e_{tn}(u_{\chi})] \right). \end{aligned} \quad (3.12)$$

The main interest of (3.11), or (3.12), compared to (3.10), is that it does not involve jumps of the Lamé coefficients which blow up when the damaged phase degenerate to zero.

Indeed, it is interesting to investigate the limit of the shape derivative in Theorem 3.1 when A^0 converges to zero. In such a case, we recover previously known formulas, used in shape optimization [3], [41], [62]. As A^0 tends to zero, it is well known that, on the interface Σ , the normal stress $\sigma n = (\sigma_{tn}, \sigma_{nn})$ converges also to zero, while the deformation tensor e remains bounded. Therefore, the limit formula of (3.11), or (3.12), is

$$\mathcal{D}(x) = e_{tt}(u_\chi) \cdot \sigma_{tt}(p_\chi) - \frac{1}{2} e_{tt}(u_\chi) \cdot \sigma_{tt}(u_\chi). \quad (3.13)$$

The proof of Theorem 3.1 is given in the next subsection (except some technical computations which are postponed to Appendix A). Similar results in the conductivity setting (scalar equations) appeared in [16], [42], [54].

Let us now restate Theorem 3.1 for the Francfort-Marigo cost function, in which case we have

$$j^k(x, u) = -f^k \cdot u^k + \kappa \delta_{k0} \quad \text{and} \quad h_k(x, u) = -g^k \cdot u^k$$

where δ_{k0} is the Kronecker symbol, equal to 0 if $k = 1$ and to 1 if $k = 0$. It turns out that the problem is self-adjoint, i.e., there is no need of an adjoint state. More precisely, in this context we find that $p_\chi = 0$. We further simplify the previous Theorem 3.1 by taking forces which are the same in the damaged and healthy regions, i.e., $f^0 = f^1$ and $g^0 = g^1$. Then, we obtain

Corollary 3.1 *Let $f^0 = f^1$ and $g^0 = g^1$. The shape derivative of (2.4) in the direction θ is*

$$J'(\chi)(\theta) = \int_{\Sigma} \mathcal{D}(x) \theta \cdot n \, dS$$

with

$$\mathcal{D}(x) = \kappa + \frac{1}{2} \left(\sigma_{nn}(u_\chi) [e_{nn}(u_\chi)] - e_{tt}(u_\chi) \cdot [\sigma_{tt}(u_\chi)] + 2\sigma_{tn}(u_\chi) \cdot [e_{tn}(u_\chi)] \right). \quad (3.14)$$

Furthermore, if $A^0 \leq A^1$, then $(\mathcal{D}(x) - \kappa) \leq 0$ on Σ .

The last result of Corollary 3.1 implies that, upon neglecting the Griffith energy release rate, i.e. taking $\kappa = 0$, one should take $\theta \cdot n \geq 0$ to get a negative shape derivative. In other words, the damaged phase should fill the entire domain in order to minimize the energy functional (2.4) (which is clear from the minimization (2.5)).

3.3 The Lagrangian approach to shape differentiation

This section is devoted to the proof of Theorem 3.1 by means of a Lagrangian method which, in the context of shape optimization, is described in e.g. [3], [5], [28]. It amounts to introduce a Lagrangian which, as usual, is the sum of the objective function and of the constraints multiplied by suitable Lagrange multipliers. In shape optimization the state equation is seen as a constraint and the corresponding Lagrange multiplier is precisely the adjoint state at optimality. The shape derivative $J'(\chi)(\theta)$ is then obtained as a simple partial derivative of the Lagrangian $\mathcal{L}$. This approach is also very convenient to guess the exact form of the adjoint problem.

In the present setting it is the shape of the subdomains Ω^0 and Ω^1 which is varying, or equivalently the interface Σ . Differentiating with respect to the position of this interface is more complicated than differentiating with respect to the outer boundary as in usual shape optimization problems. The additional difficulty, which was recognized in [54] (see also [16], [42]) is that the solution u_χ of the state equation (3.2) is not shape differentiable in the sense of Definition 3.1. The reason is that some spatial derivatives of u_χ are discontinuous at the interface (because of the jump in the material properties): thus, when we additionally differentiate with respect to the position of Σ , we obtain that those spatial derivatives of $u'_\chi(\theta)$ have a part which is a measure concentrated on the interface, and consequently $u'_\chi(\theta)$ "escapes" from the functional space V in which we differentiate. The remedy is simply to rewrite the state equation (3.2) as a transmission problem. We thus introduce the restrictions u^0 to Ω^0 , and u^1 to Ω^1 , of the solution u_χ of (3.2). In other words, they satisfy $u_\chi = (1 - \chi)u^1 + \chi u^0$ and are solutions of the transmission problem

$$\begin{cases} -\operatorname{div}(A^1 e(u^1)) = f^1 & \text{in } \Omega^1 \\ u^1 = u_D & \text{on } \Gamma_D^1 = \Gamma_D \cap \partial\Omega^1 \\ A^1 e(u^1) n^1 = g^1 & \text{on } \Gamma_N^1 = \Gamma_N \cap \partial\Omega^1 \\ u^1 = u^0 & \text{on } \Sigma = \partial\Omega^0 \cap \partial\Omega^1 \\ A^1 e(u^1) n^1 + A^0 e(u^0) n^0 = 0 & \text{on } \Sigma \end{cases} \quad (3.15)$$

and

$$\begin{cases} -\operatorname{div}(A^0 e(u^0)) = f^0 & \text{in } \Omega^0 \\ u^0 = u_D & \text{on } \Gamma_D^0 = \Gamma_D \cap \partial\Omega^0 \\ A^0 e(u^0) n^0 = g^0 & \text{on } \Gamma_N^0 = \Gamma_N \cap \partial\Omega^0 \\ u^0 = u^1 & \text{on } \Sigma \\ A^0 e(u^0) n^0 + A^1 e(u^1) n^1 = 0 & \text{on } \Sigma \end{cases}, \quad (3.16)$$

which is equivalent to (3.2). Recall that $n = n^0 = -n^1$ denotes the outward unit normal vector to the interface Σ .

Introducing the notations $\sigma^i(v^i) = A^i e(v^i)$ and $\sigma^i(q^i) = A^i e(q^i)$, the general

Lagrangian is defined as

$$\begin{aligned}
\mathcal{L}(v^1, v^0, q^1, q^0, \Sigma) &= \int_{\Omega^1} \left[j^1(x, v^1) + \frac{1}{2} \sigma^1(v^1) \cdot e(v^1) - \sigma^1(v^1) \cdot e(q^1) + f^1 \cdot q^1 \right] dV \\
&+ \int_{\Omega^0} \left[j^0(x, v^0) + \frac{1}{2} \sigma^0(v^0) \cdot e(v^0) - \sigma^0(v^0) \cdot e(q^0) + f^0 \cdot q^0 \right] dV \\
&+ \int_{\Gamma_N^0} [g^0 \cdot q^0 + h^0(x, v^0)] dS + \int_{\Gamma_N^1} [g^1 \cdot q^1 + h^1(x, v^1)] dS \\
&- \frac{1}{2} \int_{\Sigma} (\sigma^1(v^1) + \sigma^0(v^0)) n \cdot (q^1 - q^0) dS \\
&- \frac{1}{2} \int_{\Sigma} (\sigma^1(q^1) + \sigma^0(q^0)) n \cdot (v^1 - v^0) dS \\
&+ \frac{1}{2} \int_{\Sigma} (\sigma^1(v^1) + \sigma^0(v^0)) n \cdot (v^1 - v^0) dS, \tag{3.17}
\end{aligned}$$

where q^0 and q^1 play the role of Lagrange multiplier or, at optimality, of the adjoint state p^0 and p^1 (on the same token, at optimality v^0, v^1 are equal to u^0, u^1). The functions v^0, v^1 satisfy non homogeneous Dirichlet boundary conditions and belong to the affine space V , while the other functions q^0, q^1 vanishes on Γ_D and thus belong to the vector space V_0 defined as

$$V_0 = \{u \in H^1(\Omega; \mathbb{R}^d) \text{ such that } u = 0 \text{ on } \Gamma_D\}.$$

Of course, differentiating the Lagrangian with respect to q^0 and q^1 , and equaling it to 0, provides the state equations (3.15) and (3.16). The next result states that differentiating the Lagrangian with respect to v^0 and v^1 , and equaling it to 0, yields the adjoint equation.

Lemma 3.3 *The optimality condition*

$$\frac{\partial \mathcal{L}}{\partial v^1}(u^1, u^0, p^1, p^0, \chi) = \frac{\partial \mathcal{L}}{\partial v^0}(u^1, u^0, p^1, p^0, \chi) = 0$$

for variations in V_0 is equivalent to the adjoint problem (3.6).

Proof. This is a classical computation [3], [28], [54] which we do not detail. Differentiating the Lagrangian with respect to v^0 and v^1 and equaling it to zero yields

$$\begin{cases} -\operatorname{div}(A^i e(p^i)) = j'_i(x, u^i) - \operatorname{div}(A^i e(u^i)) & \text{in } \Omega^i \\ p^i = 0 & \text{on } \Gamma_D^i \\ A^i e(p^i) n^i = h'_i(x, u^i) + A^i e(u^i) n^i & \text{on } \Gamma_N^i \\ p^0 = p^1 & \text{on } \Sigma \\ A^0 e(p^0) n = A^1 e(p^1) n & \text{on } \Sigma \end{cases} \tag{3.18}$$

which is equivalent to (3.6). $\square$

As we already said, the solution u_χ of (3.2) is not shape differentiable. However its Lagrangian or transported counterpart, namely $\theta \rightarrow u_{\chi \circ (Id + \theta)} \circ (Id + \theta)$, is actually differentiable by a simple application of the implicit function theorem (see chapter 5 in [41]). As a consequence, upon a suitable extension outside Ω^i , the solution u^i of (3.15-3.16) are indeed shape differentiable.

Lemma 3.4 *The solutions u^1 of (3.15) and u^0 of (3.16) are shape differentiable.*

The main interest of the Lagrangian is that its partial derivative with respect to the shape χ , evaluated at the state u_χ and adjoint p_χ , is equal to the shape derivative of the cost function.

Lemma 3.5 *The cost function $J(\chi)$ admits a shape derivative which is given by*

$$J'(\chi)(\theta) = \frac{\partial \mathcal{L}}{\partial \chi}(u^1, u^0, p^1, p^0, \chi)(\theta), \quad (3.19)$$

where (u^1, u^0, p^1, p^0) are the solutions of the state equation (3.15-3.16) and adjoint equation (3.18).

Proof. This is again a classical result [3], [28] which we briefly recall. We start from the identity

$$J(\chi)(\theta) = \mathcal{L}(u^1, u^0, q^1, q^0, \chi) \quad (3.20)$$

where q^1, q^0 are any functions in V . We differentiate (3.20) with respect to the shape. By virtue of Lemma 3.4 we obtain

$$J'(\chi)(\theta) = \frac{\partial \mathcal{L}}{\partial \chi}(u^1, u^0, q^1, q^0, \chi)(\theta) + \langle \frac{\partial \mathcal{L}}{\partial v^{0,1}}(u^1, u^0, q^1, q^0, \chi), \frac{\partial u^{0,1}}{\partial \chi}(\theta) \rangle. \quad (3.21)$$

The notation $\frac{\partial \mathcal{L}}{\partial \chi}$ means that it is a shape partial derivative, i.e., we differentiate $\mathcal{L}$ in the sense of Definition 3.1 while keeping the other arguments (u^1, u^0, q^1, q^0) fixed. Taking now $(q^1, q^0) = (p^1, p^0)$ cancels the last term in (3.21) because it is the variational formulation of the adjoint problem by virtue of Lemma 3.3. We thus obtain (3.19). $\square$

To finish the proof of Theorem 3.1 it remains to compute the partial shape derivative of the Lagrangian. It is a conceptually simple application of Lemma 3.1 which, nevertheless, is quite tedious. Therefore the proof of the following Lemma is postponed to Appendix A.

Lemma 3.6 *The partial shape derivative of the Lagrangian*

$$\frac{\partial \mathcal{L}}{\partial \chi}(u^1, u^0, p^1, p^0, \chi)(\theta),$$

is precisely equal to the right hand side of (3.9).

4 Topological derivative

The aim of this section is to evaluate the sensitivity of the cost function to the introduction of an infinitesimal damaged region ω_ρ inside the healthy region Ω^1 . In theory the shape of the smooth inclusion can be arbitrary. However, for practical and numerical purposes it will be assumed to be a ball in $\mathbb{R}^d$.

4.1 Main result

Let ω be a smooth open subset of $\mathbb{R}^d$. Let $\rho > 0$ be a small positive parameter which is intended to go to zero. For a point $x_0 \in \Omega^1$ we define a rescaled inclusion

$$\omega_\rho = \{x \in \mathbb{R}^d : \frac{x - x_0}{\rho} \in \omega\}, \quad (4.1)$$

which, for small enough ρ is strictly included in Ω^1 and disconnected from Ω^0 . The total damaged zone is thus $\Omega_\rho^0 := \Omega^0 \cup \omega_\rho$ and the healthy phase is $\Omega_\rho^1 := \Omega \setminus \Omega_\rho^0$. Let $\chi_\rho, \chi, \chi_{\omega_\rho}$ denote the characteristic functions of Ω_ρ^0, Ω^0 and ω_ρ , respectively (verifying $\chi_\rho = \chi + \chi_{\omega_\rho}$). In the sequel, in order to distinguish integration in the variables x and $y := \frac{x - x_0}{\rho}$, the symbol dV will sometimes be replaced by $dV(y)$ or $dV(\eta)$ (where η is a dummy variable similar to y).

Let us recall the notations for the non-perturbed domain $\Omega = \Omega^0 \cup \Omega^1$ (i.e., without the damage inclusion). The cost function then writes as

$$J(\chi) = \frac{1}{2} \int_{\Omega} A_\chi e(u_\chi) \cdot e(u_\chi) dV + \int_{\Omega} j_\chi(x, u_\chi) dV + \int_{\partial\Omega} h_\chi(x, u_\chi) dS, \quad (4.2)$$

where $j_\chi = j^0\chi + j^1(1-\chi)$, $h_\chi = h^0\chi + h^1(1-\chi)$, and the so-called “background” solution u_χ solves the state equation (3.2) on $\Omega = \Omega^0 \cup \Omega^1$. As in the previous section, we assume that the integrands $j^0, j^1(x, u)$ and $h^0, h^1(x, u)$ are twice differentiable functions with respect to u , satisfying the growth conditions (3.4). Moreover, let us recall that the so-called “background” dual solution p_χ solves the adjoint problem (3.6) on $\Omega = \Omega^0 \cup \Omega^1$.

On the perturbed domain $\Omega = \Omega_\rho^0 \cup \Omega_\rho^1$, the cost function is

$$J(\chi_\rho) = \frac{1}{2} \int_{\Omega} A_{\chi_\rho} e(u_{\chi_\rho}) \cdot e(u_{\chi_\rho}) dV + \int_{\Omega} j_{\chi_\rho}(x, u_{\chi_\rho}) dV + \int_{\partial\Omega} h_\chi(x, u_{\chi_\rho}) dS,$$

because $\chi_\rho \equiv \chi$, and thus $h_{\chi_\rho} \equiv h_\chi$, on $\partial\Omega$ (the inclusion ω_ρ is away from the boundary), and where u_{χ_ρ} solves

$$\begin{cases} -\operatorname{div}(A_{\chi_\rho} e(u_{\chi_\rho})) = f & \text{in } \Omega \\ u_{\chi_\rho} = u_D & \text{on } \Gamma_D \\ A_{\chi_\rho} e(u_{\chi_\rho}) n = g & \text{on } \Gamma_N \end{cases} \quad (4.3)$$

with $A_{\chi_\rho} = A^0\chi_\rho + A^1(1-\chi_\rho)$, the Lamé tensor of the material with the inclusion.

Definition 4.1 *If the objective function admits the following so-called topological asymptotic expansion for small $\rho > 0$:*

$$J(\chi_\rho) - J(\chi) - \rho^d DJ(x_0) = o(\rho^d),$$

then the number $DJ(x_0)$ is called the topological derivative of J at x_0 for the inclusion shape ω .

The main result of this section is the following theorem.

Theorem 4.1 *The topological derivative $DJ(x_0)$ of the general cost function (4.2), evaluated at x_0 for an inclusion shape ω , has the following expression:*

$$\begin{aligned} DJ(x_0) := Me(u_\chi)(x_0) \cdot e(p_\chi)(x_0) & - \frac{1}{2} Me(u_\chi)(x_0) \cdot e(u_\chi)(x_0) \\ & + |\omega|(j^0 - j^1)(x_0, u_\chi(x_0)), \end{aligned} \quad (4.4)$$

where u_χ and p_χ are the solution to the primal and dual problems (3.2) and (3.6), respectively, and where M is the so-called elastic moment tensor as defined below by (4.10). Moreover, M is positive if $[A]$ is positive, and negative if $[A]$ is negative.

In the case of our damage model, the cost function is (2.4), i.e., $j_\chi(x, u_\chi) = \kappa\chi - f \cdot u_\chi$ and $h_\chi(x, u_\chi) = -g \cdot u_\chi$. The problem is then known to be self-adjoint, i.e., the adjoint p_χ is equal to 0. In such a case Theorem 4.1 simplifies as follows.

Corollary 4.1 *The topological derivative of the cost function (2.4) at x_0 for an inclusion shape ω is*

$$DJ(x_0) := |\omega|\kappa - \frac{1}{2} Me(u_\chi)(x_0) \cdot e(u_\chi)(x_0), \quad (4.5)$$

where u_χ is the background solution of (3.2) and M is the elastic moment tensor defined below by (4.10).

In 2d, the elastic moment tensor M for a unit disk-inclusion ω has been computed in [11]. The topological derivative (4.5) for a disk-inclusion is:

$$\begin{aligned} DJ(x_0) = \pi\kappa - 2\pi \frac{\mu_1[\mu](\lambda_1 + 2\mu_1)}{\lambda_1(\mu_0 + \mu_1) + \mu_1(\mu_1 + 3\mu_0)} e(u_\chi) \cdot e(u_\chi)(x_0) \\ + \frac{\pi}{2} \left(-\frac{(\lambda_1 + 2\mu_1)[\lambda + \mu]}{\lambda_0 + \mu_0 + \mu_1} + 2 \frac{\mu_1[\mu](\lambda_1 + 2\mu_1)}{\lambda_1(\mu_0 + \mu_1) + \mu_1(\mu_1 + 3\mu_0)} \right) \text{tre}(u_\chi) \text{tre}(u_\chi)(x_0). \end{aligned}$$

In 3d, the elastic moment tensor M for a unit ball-inclusion ω has also been computed in [12]. The topological derivative (4.5) for a ball-inclusion is:

$$DJ(x_0) = \frac{4\pi}{3}\kappa - \frac{2\pi}{3b} \left(2[\mu]e(u) \cdot e(u) + \frac{[\lambda]b - 2[\mu]a}{(3a + b)} \text{tre}(u) \text{tre}(u) \right),$$

with $\nu_1 = \frac{\lambda_1}{2(\lambda_1 + \mu_1)}$ and

$$a := -\frac{5\mu_1\nu_1[\lambda] - \lambda_1[\mu]}{15\lambda_1\mu_1(1 - \nu_1)}, \quad b := \frac{15\mu_1(1 - \nu_1) - 2[\mu](4 - 5\nu_1)}{15\mu_1(1 - \nu_1)} > 0.$$

In order to prove Theorem 4.1 we need several technical tools detailed in the next subsections.

4.2 Elastic moment tensor

The goal of this subsection is to define the elastic moment tensor as a 4th order tensor expressing the leading behaviour in the far field of w_ξ , solution to the canonical problem (4.9) of a unit damage inclusion ω in a uniform healthy background.

We introduce a microscopic variable $y = \frac{x - x_0}{\rho}$ in order to rescale the problem with a unit inclusion ω . This rescaling, centered on the inclusion, in the limit as ρ goes to zero, transforms the elasticity problem posed on Ω in a problem posed on $\mathbb{R}^d$. The symbols e_y , div_y etc. are used to specify the derivation w.r.t. y .

We begin by recalling the Green tensor for linear elasticity in a uniform infinite material.

Notations 4.1 (Green tensor of elasticity) *The fundamental tensor of linear elasticity $\Gamma := (\Gamma_{ij})_{1 \leq i, j \leq d}$ reads:*

$$\Gamma_{ij}(y) := \begin{cases} -\frac{\alpha}{4\pi} \frac{\delta_{ij}}{|y|^{d-2}} - \frac{\beta}{4\pi} \frac{y_i y_j}{|y|^d} & \text{if } d \geq 3 \\ \frac{\alpha}{2\pi} \delta_{ij} \ln |y| - \frac{\beta}{2\pi} \frac{y_i y_j}{|y|^2} & \text{if } d = 2 \end{cases}, \quad (4.6)$$

where

$$\alpha = \frac{1}{2} \left(\frac{1}{\mu^1} + \frac{1}{2\mu^1 + \lambda^1} \right) \quad \text{and} \quad \beta = \frac{1}{2} \left(\frac{1}{\mu^1} - \frac{1}{2\mu^1 + \lambda^1} \right).$$

The component Γ_{ij} represents the i th Cartesian component of the fundamental solution in the free-space with a unit Dirac load δ_0 at the origin in the direction of vector $-e_j$, that is,

$$-\text{div} \left(A^1 e_y \left(\sum_{i=1}^d \Gamma_{ij} e_i \right) \right) = -e_j \delta_0, \quad (4.7)$$

where e_k denotes the k th element of the canonical basis of $\mathbb{R}^d$.

We introduce the following Hilbert space (so-called Deny-Lions or Beppo-Levi space)

$$W := \{w \in H_{loc}^1(\mathbb{R}^d; \mathbb{R}^d) \text{ such that } e(w) \in L^2(\mathbb{R}^d; \mathbb{R}^{d \times d})\}, \quad (4.8)$$

equipped with the scalar product of $L^2(\mathbb{R}^d)$ for the deformation tensor $e(w)$, which is well adapted to elasticity problems posed in the whole space $\mathbb{R}^d$. For any symmetric matrix ξ we introduce $w_\xi(y)$, solution to the canonical problem

$$\begin{cases} -\operatorname{div}_y (A_{\chi_\omega} e_y(w_\xi)) = -\operatorname{div}_y (\chi_\omega [A] \xi) & \text{in } \mathbb{R}^d, \\ w_\xi \in W, \end{cases} \quad (4.9)$$

which is easily seen to be well-posed. The fact that w_ξ belongs to W implies it has some decay properties at infinity (by embedding of W in some Lebesgue space, see [32], [2]). We shall not dwell on them since Lemma 4.1 below improve these decay properties.

Lemma 4.1 (Far field expression) *The solution w_ξ of the canonical problem (4.9) has the following pointwise behavior at infinity:*

$$w_\xi = -\partial_p \Gamma_q(y) M_{pqkl} \xi_{kl} + \mathcal{O}(|y|^{-d}) \quad \text{as } |y| \rightarrow \infty, \quad (4.10)$$

where $\Gamma_q := \Gamma_{kq} e_k$ is the fundamental Green's tensor of linear elasticity of the healthy material, and M is the 4th order elastic moment tensor with respect to inclusion ω , independent of ξ , defined by

$$M = [A] (N + |\omega| I_4) \quad (4.11)$$

with a 4th order tensor N defined by

$$N\xi := \int_{\omega} e_y(w_\xi) dV(y). \quad (4.12)$$

Remark 4.1 *Lemma 4.1 tells us that, because the right hand side in (4.9) has zero average, w_ξ behaves like $\mathcal{O}(|y|^{-d+1})$ at infinity. The interest of the canonical problem for us is that, by denoting $\xi_0 = e(u_\chi)(x_0)$, we shall prove in some sense*

$$u_{\chi_\rho}(x) \approx u_\chi(x) + \rho w_{\xi_0} \left(\frac{x - x_0}{\rho} \right).$$

Remark 4.2 *The elastic moment tensor M as defined by (4.11) is exactly the same tensor as introduced in [9] and [11] (by means of layer potential techniques) or in [27] (by means of a variational approach in the conductivity setting).*

Proof of Lemma 4.1. Let us consider an inclusion ω located in the free-space $\mathbb{R}^d$ and introduce a smooth open set U strictly containing ω and a cut-off function $\varphi \in C^\infty(\mathbb{R}^d)$ such that $\varphi \equiv 0$ on ω , $\varphi \equiv 1$ on $\mathbb{R}^d \setminus U$. We define a function $f(y)$ by

$$f := -\operatorname{div}_y (A_{\chi_\omega} e(\varphi w_\xi)), \quad (4.13)$$

which has compact support in U because of (4.9) and the fact that $\varphi \equiv 1$ on $\mathbb{R}^d \setminus U$. Since $\varphi \equiv 0$ on ω we deduce that

$$\begin{cases} -\operatorname{div}_y (A^1 e_y(\varphi w_\xi)) = f & \text{in } \mathbb{R}^d, \\ \varphi w_\xi \in W, \end{cases} \quad (4.14)$$

We can thus use the Green tensor to compute the k th component of the solution of (4.14)

$$\varphi(y)e_k \cdot w_\xi(y) = - \int_{\mathbb{R}^d} \Gamma_{kq}(y-\eta) f_q(\eta) dV(\eta). \quad (4.15)$$

It turns out that

$$\begin{aligned} \int_{\mathbb{R}^d} f(y) dV(y) &= \int_U f(y) dV(y) = - \int_{\partial U} A_{\chi_\omega} e(\varphi w_\xi) n dS(y) = \\ &= - \int_U \operatorname{div} (A_{\chi_\omega} e(w_\xi)) dV(y) \stackrel{(4.9)}{=} - \int_U \operatorname{div} (\chi_\omega [A] \xi) dV(y) \\ &= - \int_{\partial U} \chi_\omega [A] \xi n dS(y) = 0 \end{aligned}$$

with n denoting the usual normal unit vector to ∂U . By Taylor expansion of the Green function $\Gamma_{kq}(y-\eta)$ in terms of $\Gamma_{kq}(y)$ and its derivatives, taking into account that f has zero average and compact support in U , and since $\varphi \equiv 1$ away from U , (4.15) yields that

$$e_k \cdot w_\xi(y) = \partial_p \Gamma_{kq}(y) \int_{\mathbb{R}^d} \eta_p f_q(\eta) dV(\eta) + \mathcal{O}(|y|^{-d}). \quad (4.16)$$

Let us now evaluate $\int_{\mathbb{R}^d} \eta_p f_q(\eta) dV(\eta)$ that for the sake of calculus is rewritten as $\int_{\mathbb{R}^d} B^{pq} \eta \cdot f(\eta) dV(\eta)$, where $B^{pq} := e_q \otimes e_p$ is a second order tensor. By (4.13) and since $A^1 = A_{\chi_\omega}$ on ∂U ,

$$\begin{aligned} \int_U B^{pq} \eta \cdot f(\eta) dV(\eta) &= \int_U A_{\chi_\omega} e(\varphi w_\xi) \cdot e(B^{pq} \eta) dV(\eta) - \int_{\partial U} A_{\chi_\omega} e(\varphi w_\xi) n \cdot B^{pq} \eta dS(\eta) \\ &= \int_U A^1 e(\varphi w_\xi) \cdot e(B^{pq} \eta) dV(\eta) - \int_{\partial U} A^1 e(\varphi w_\xi) n \cdot B^{pq} \eta dS(\eta) \\ &= - \int_U \underbrace{\operatorname{div} (A^1 e(B^{pq} \eta))}_{=0} \varphi w_\xi dV(\eta) + \int_{\partial U} A^1 e(B^{pq} \eta) n \cdot \varphi w_\xi dS(\eta) \\ &\quad - \int_{\partial U} A^1 e(\varphi w_\xi) n \cdot B^{pq} \eta dS(\eta) \\ &= \int_{\partial U} (A_{\chi_\omega} e(B^{pq} \eta) n \cdot w_\xi - A_{\chi_\omega} e(w_\xi) n \cdot B^{pq} \eta) dS(\eta) \\ &= \int_U \underbrace{\operatorname{div} (A_{\chi_\omega} e(B^{pq} \eta))}_{=\operatorname{div}(-\chi_\omega [A] B^{pq})} \cdot w_\xi dV(\eta) - \int_U \underbrace{\operatorname{div} (A_{\chi_\omega} e(w_\xi))}_{=\operatorname{div}(\chi_\omega [A] \xi)} \cdot B^{pq} \eta dV(\eta) \\ &= \int_\omega [A] B^{pq} \cdot e(w_\xi) dV(\eta) + \int_\omega [A] \xi \cdot B^{pq} dV(\eta) \\ &= [A] B^{pq} \cdot \int_\omega (e(w_\xi) + \xi) dV(\eta). \end{aligned} \quad (4.17)$$

Introducing M_{ijkl} defined as

$$M_{ijkl} := [A]_{ijmn} (N + |\omega| I_4)_{mnkl} \quad (4.18)$$

we obtain (4.10). $\square$

Lemma 4.2 (Symmetry and signature of M) *The elastic moment tensor M , defined by (4.11), is symmetric and positive if $A^0 < A^1$ while negative if $A^0 > A^1$.*

Proof of Lemma 4.2. Let us multiply (4.9) by the solution $w_{\xi'}$ for another symmetric tensor ξ' , integrate by parts and observe that, by the symmetry property of the left hand side, we have

$$\begin{aligned} \int_{\mathbb{R}^d} A_{\chi_\omega} e(w_\xi) \cdot e(w_{\xi'}) dV &= [A] \xi \cdot N \xi' = [A] \xi' \cdot N \xi \\ &= [A] N \cdot \xi \otimes \xi' = [A] N \cdot \xi' \otimes \xi, \end{aligned} \quad (4.19)$$

the symmetry of $[A]N$ and hence of M immediately follows. Take $\xi = \xi'$ in (4.19), then $[A]N$ is clearly positive. Therefore if $[A] > 0$, then M is obviously positive. Assume now that $[A] < 0$. The solution w_ξ of (4.9) is the minimizer of the following energy

$$I(w) = \frac{1}{2} \int_{\mathbb{R}^d} A_{\chi_\omega} e(w) \cdot e(w) dV - \int_{\mathbb{R}^d} \chi_\omega [A] \xi \cdot e(w) dV$$

and its minimal value is, by (4.12), $I(w_\xi) = -\frac{1}{2}[A]N\xi \cdot \xi$. On the other hand as soon as we rewrite

$$A_{\chi_\omega} e(w) \cdot e(w) = -\chi_\omega [A] e(w) \cdot e(w) + A^1 e(w) \cdot e(w)$$

we obtain the lower bound $I^-(w)$:

$$I(w) \geq I^-(w) := -\frac{1}{2} \int_{\mathbb{R}^d} \chi_\omega [A] e(w) \cdot e(w) dV - \int_{\mathbb{R}^d} \chi_\omega [A] \xi \cdot e(w) dV.$$

It is easily seen that $e(w) = -\xi$ is a critical point in ω of the above lower bound, which, by the negative character of $[A]$, turns out to be the unique minimizer, thereby providing the minimal value $\frac{1}{2}|\omega|[A]\xi \cdot \xi$. Thus we deduce that

$$|\omega|[A]\xi \cdot \xi \leq -[A]N\xi \cdot \xi$$

which implies the desired result $M < 0$. $\square$

4.3 Asymptotic analysis in the perturbed domain

This subsection is aimed at comparing the solutions of elasticity problems in the perturbed and non-perturbed domains. We define the difference $v := u_{\chi_\rho} - u_\chi$ between the perturbed (u_{χ_ρ}) and the background (u_χ) displacement fields. The equation satisfied by v is

$$\begin{cases} -\operatorname{div}(A_{\chi_\rho} e(v)) = -\operatorname{div}(\chi_{\omega_\rho} [A] e(u_\chi)) & \text{in } \Omega \\ v = 0 & \text{on } \Gamma_D \\ A_{\chi_\rho} e(v) n = 0 & \text{on } \Gamma_N \end{cases} \quad (4.20)$$

Let us introduce a tensor $\xi_0 := e(u_\chi)(x_0)$ and let $w_{\xi_0}(y)$ be the solution of (4.9) for $\xi = \xi_0$. We define a rescaled function $w_{\xi_0}^\rho(x) := \rho w_{\xi_0}(\frac{x-x_0}{\rho})$ which is a solution to

$$-\operatorname{div}\left(A_{\chi_\rho}e(w_{\xi_0}^\rho)\right) = -\operatorname{div}\left([A]\chi_{\omega_\rho}e(u_\chi)(x_0)\right) \quad \text{in } \Omega,$$

satisfying non-homogeneous, but small, boundary conditions. This function $w_{\xi_0}^\rho$ is the leading term of a so-called inner asymptotic expansion for v as stated by the following Lemma.

Lemma 4.3 *For any cut-off function $\theta \in C_c^\infty(\Omega)$ such that $\theta \equiv 1$ in a neighborhood U of x_0 , there exists a constant $C > 0$ independent of ρ such that we have*

$$v = \theta w_{\xi_0}^\rho + \delta, \quad (4.21)$$

with

$$\|\delta\|_{H^1(\Omega)} \leq C\rho^{d/2+1}. \quad (4.22)$$

Moreover

$$\|w_{\xi_0}^\rho\|_{L^2(\Omega)} \leq C \begin{cases} \rho^2 \sqrt{\log \rho} & \text{if } d = 2 \\ \rho^{d/2+1} & \text{if } d \geq 3 \end{cases} \quad \text{and} \quad \|e(w_{\xi_0}^\rho)\|_{L^2(\Omega)} \leq C\rho^{d/2} \quad (4.23)$$

Remark 4.3 *In the vicinity of the inclusion ω_ρ , we have $\theta \equiv 1$ for sufficiently small ρ , and (4.21) can be restated as*

$$v(x) = \rho w_{\xi_0}\left(\frac{x-x_0}{\rho}\right) + o_{H^1}(\rho),$$

which is an inner asymptotic expansion for v , solution of (4.20). The L^2 -norms of δ and $w_{\xi_0}^\rho$ are of the same order (at least for $d \geq 3$) but the L^2 -norm of $\nabla \delta$ is smaller by a factor ρ than that of $\nabla w_{\xi_0}^\rho$, which explains the $o(\rho)$ remainder in the above approximation of v .

Proof of Lemma 4.3. The estimates (4.23) on $w_{\xi_0}^\rho$ are simply obtained by rescaling and by the decay properties of w_{ξ_0} . We obtain

$$\|e(w_{\xi_0}^\rho)\|_{L^2(\Omega)}^2 = \int_{\Omega} |e_y(w_{\xi_0})(\frac{x}{\rho})|^2 dV = \rho^d \int_{\Omega/\rho} |e_y(w_{\xi_0})(y)|^2 dV(y) \leq C\rho^d.$$

Similarly

$$\|w_{\xi_0}^\rho\|_{L^2(\Omega)}^2 = \rho^2 \int_{\Omega} |w_{\xi_0}(\frac{x}{\rho})|^2 dV = \rho^{d+2} \int_{\Omega/\rho} |w_{\xi_0}(y)|^2 dV(y).$$

However, Lemma 4.1 tells us that the behaviour at infinity of w_{ξ_0} is such that it does not belong to $L^2(\mathbb{R}^d)$ but is of the order of $\mathcal{O}(|y|^{-d+1})$. Therefore, using

the radial coordinate $r = |y|$ yields

$$\begin{aligned} \|w_{\xi_0}^\rho\|_{L^2(\Omega)}^2 &\leq C\rho^{d+2} \int_{\Omega/\rho} \frac{1}{1+|y|^{2(d-1)}} dV(y) \leq C\rho^{d+2} \int_1^{1/\rho} \frac{dr}{r^{d-1}} \\ &\leq C \begin{cases} \rho^4 |\log \rho| & \text{if } d=2 \\ \rho^{d+2} & \text{if } d \geq 3 \end{cases} \end{aligned} \quad (4.24)$$

which is the desired result. Furthermore, since $w_{\xi_0} = \mathcal{O}(|y|^{-d+1})$ and $\nabla w_{\xi_0} = \mathcal{O}(|y|^{-d})$ at infinity, we also deduce by rescaling that

$$\|w_{\xi_0}^\rho\|_{L^\infty(\Omega \setminus U)} \leq C\rho^d \quad \text{and} \quad \|\nabla w_{\xi_0}^\rho\|_{L^\infty(\Omega \setminus U)} \leq C\rho^d. \quad (4.25)$$

We now write the equation satisfied by δ :

$$\begin{cases} -\operatorname{div}(A_{\chi_\rho} e(\delta)) = -\operatorname{div}([A] \chi_{\omega_\rho} (e(u_\chi)(x) - e(u_\chi)(x_0))) + g & \text{in } \Omega \\ \delta = 0 & \text{on } \Gamma_D \\ A_{\chi_\rho} e(\delta) n = 0 & \text{on } \Gamma_N \end{cases} \quad (4.26)$$

where

$$g = \operatorname{div}[A_{\chi_\rho} e(\theta w_{\xi_0}^\rho)] - \theta \operatorname{div}[\chi_{\omega_\rho} [A] e(u_\chi)(x_0)]. \quad (4.27)$$

Let us multiply (4.26) and (4.27) by δ and integrate by parts, in such a way that

$$\begin{aligned} C\|e(\delta)\|_{L^2(\Omega)}^2 &\leq \left| \int_{\Omega} A_{\chi_\rho} e(\delta) \cdot e(\delta) dV \right| \leq \int_{\omega_\rho} \left| [A] (e(u_\chi)(x) - e(u_\chi)(x_0)) \cdot e(\delta) \right| dV \\ &\quad + \left| \int_{\Omega} g \cdot \delta dV \right|. \end{aligned} \quad (4.28)$$

for $C > 0$. Let us remark that, away from the interface between the two phases, u_χ is of class $\mathcal{C}^{2,\alpha}$ for some $\alpha > 0$ (since we assume the forces to be of class $\mathcal{C}^{0,\alpha}$). Furthermore, the inclusion ω_ρ is smooth, so the $\mathcal{C}^{2,\alpha}$ regularity of u_χ holds up to the interface in the inclusion, and hence

$$|e(u_\chi)(x) - e(u_\chi)(x_0)| \leq C\rho \quad \text{in } \omega_\rho, \quad (4.29)$$

which implies

$$\int_{\omega_\rho} \left| [A] (e(u_\chi)(x) - e(u_\chi)(x_0)) \cdot e(\delta) \right| dV \leq C\rho^{d/2+1} \|e(\delta)\|_{L^2(\Omega)}. \quad (4.30)$$

Moreover by (4.27), it results that

$$\begin{aligned} \int_{\Omega} g \cdot \delta dV &= - \int_{\Omega} A_{\chi_\rho} e(\theta w_{\xi_0}^\rho) \cdot e(\delta) dV + \int_{\Omega} \chi_{\omega_\rho} [A] e(u_\chi)(x_0) \cdot e(\theta \delta) dV \\ &= - \int_{\Omega} A_{\chi_\rho} e(\theta w_{\xi_0}^\rho) \cdot e(\delta) dV + \int_{\Omega} A_{\chi_\rho} e(w_{\xi_0}^\rho) \cdot e(\theta \delta) dV \\ &= \int_{\Omega} A_{\chi_\rho} \left(e(w_{\xi_0}^\rho) \cdot (\delta \otimes \nabla \theta)^s - e(\delta) \cdot (w_{\xi_0}^\rho \otimes \nabla \theta)^s \right) dV, \end{aligned} \quad (4.31)$$

where the superscript “s” stands for the symmetric part. Hence, since $\nabla\theta$ vanishes on a neighborhood U of ω_ρ , by Korn inequality and by estimates (4.25), it follows that

$$\begin{aligned} \left| \int_{\Omega} g \cdot \delta dV \right| &\leq C \left(\|w_{\xi_0}^\rho\|_{L^\infty(\Omega \setminus U)} + \|e(w_{\xi_0}^\rho)\|_{L^\infty(\Omega \setminus U)} \right) \|\nabla\theta\|_{L^2(\Omega \setminus U)} \|e(\delta)\|_{L^2(\Omega)} \\ &\leq C \rho^d \|e(\delta)\|_{L^2(\Omega)}. \end{aligned} \quad (4.32)$$

Therefore, by (4.28)-(4.32) and since $d/2 + 1 \leq d$ for $d \geq 2$, the following global estimate holds

$$\|e(\delta)\|_{L^2(\Omega)} \leq C \left(\rho^d + \rho^{d/2+1} \right) \leq C \rho^{d/2+1}, \quad (4.33)$$

completing the proof by Korn and Poincaré inequalities. $\square$

Similarly, we shall need a comparison between the perturbed and background adjoints. However, the adjoint in the perturbed domain (with an inclusion) is not the standard one. Rather, we introduce a slightly different adjoint problem

$$\begin{cases} -\operatorname{div}(A_{\chi_\rho} e(\tilde{p}_{\chi_\rho})) = f_\chi + j'_\chi(x, u_\chi) & \text{in } \Omega \\ \tilde{p}_{\chi_\rho} = 0 & \text{on } \Gamma_D \\ A_{\chi_\rho} e(\tilde{p}_{\chi_\rho}) n = g_\chi + h'_\chi(x, u_\chi) & \text{on } \Gamma_N \end{cases} \quad (4.34)$$

whose solution $\tilde{p}_{\chi_\rho}$ depends on the inclusion since the Lamé tensor A_{χ_ρ} corresponds to the perturbed domain $\Omega = \Omega_\rho^0 \cup \Omega_\rho^1$. Nevertheless, $\tilde{p}_{\chi_\rho}$ is different from p_{χ_ρ} , defined by (3.6) with χ_ρ instead of χ , because the right hand side of (4.34) depends only on χ and not on χ_ρ .

We define the difference between the above perturbed adjoint and the “true” background adjoint, $q := \tilde{p}_{\chi_\rho} - p_\chi$, which is the solution of

$$\begin{cases} -\operatorname{div}(A_{\chi_\rho} e(q)) = -\operatorname{div}(\chi_{\omega_\rho} [A] e(p_\chi)) & \text{in } \Omega \\ q = 0 & \text{on } \Gamma_D \\ A_{\chi_\rho} e(q) n = 0 & \text{on } \Gamma_N \end{cases} \quad (4.35)$$

We introduce the tensor $\xi'_0 := e(p_\chi)(x_0)$ and the rescaled function $w_{\xi'_0}^\rho(x) := \rho w_{\xi'_0}^\rho(\frac{x-x_0}{\rho})$ which is the leading term of an inner asymptotic expansion for q . Lemma 4.3 can then be generalized as follows.

Lemma 4.4 *For any cut-off function $\theta \in C_c^\infty(\Omega)$ such that $\theta \equiv 1$ in a neighborhood U of x_0 , there exists a constant $C > 0$ independent of ρ such that we have*

$$q = \theta w_{\xi'_0}^\rho + \delta,$$

with

$$\|\delta\|_{H^1(\Omega)} \leq C \rho^{1+d/2}. \quad (4.36)$$

Moreover

$$\|w_{\xi'_0}^\rho\|_{L^2(\Omega)} \leq C \begin{cases} \rho^2 \sqrt{\log \rho} & \text{if } d = 2 \\ \rho^{d/2+1} & \text{if } d \geq 3 \end{cases} \quad \text{and} \quad \|e(w_{\xi'_0}^\rho)\|_{L^2(\Omega)} \leq C \rho^{d/2}.$$

4.4 Proof of Theorem 4.1

We combine the ingredients of the two previous subsections to prove Theorem 4.1 on the topological derivative. Let us recall that we assume the integrands of the objective function, $j^0, j^1(x, u)$ and $h^0, h^1(x, u)$, to be C^2 functions with respect to u with adequate growth conditions.

Recalling that $h_{\chi_\rho} \equiv h_\chi$ on $\partial\Omega$ because the inclusion does not touch the boundary, we write a second-order Taylor expansion of the objective function

$$\begin{aligned}
J(\chi_\rho) &= \frac{1}{2} \int_{\Omega} A_\chi e(u_\chi + v) \cdot e(u_\chi + v) dV - \frac{1}{2} \int_{\omega_\rho} [A] e(u_\chi + v) \cdot e(u_\chi + v) dV \\
&+ \int_{\Omega} j_\chi(u_\chi + v) dV + \int_{\partial\Omega} h_\chi(u_\chi + v) dS + \int_{\omega_\rho} (j^0 - j^1)(u_\chi + v) dV \\
&= J(\chi) + \int_{\Omega} A_\chi e(u_\chi) \cdot e(v) dV + \frac{1}{2} \int_{\Omega} A_\chi e(v) \cdot e(v) dV \\
&- \frac{1}{2} \int_{\omega_\rho} [A] (e(u_\chi) \cdot e(u_\chi) + 2e(u_\chi) \cdot e(v)) dV - \frac{1}{2} \int_{\omega_\rho} [A] e(v) \cdot e(v) dV \\
&+ \int_{\Omega} j'_\chi(x, u_\chi) \cdot v dV + \int_{\partial\Omega} h'_\chi(x, u_\chi) \cdot v dS + \int_{\omega_\rho} (j^0 - j^1)(u_\chi) dV \\
&+ \frac{1}{2} \int_{\Omega} j''_\chi(\bar{u}_\chi) v \cdot v dV + \frac{1}{2} \int_{\partial\Omega} h''_\chi(\bar{u}_\chi) v \cdot v dS \\
&+ \int_{\omega_\rho} (j^0 - j^1)'(u_\chi) \cdot v dV + \frac{1}{2} \int_{\omega_\rho} (j^0 - j^1)''(\bar{u}_\chi) v \cdot v dV, \tag{4.37}
\end{aligned}$$

where $\bar{u}_\chi = u_\chi + \zeta v$ with $0 < \zeta < 1$. From assumption (3.4) we know that j''_χ and h''_χ are bounded on $\bar{\Omega}$ and thus

$$\left| \int_{\Omega} j''_\chi(\bar{u}_\chi) v \cdot v dV \right| \leq C \|v\|_{L^2(\Omega)}^2 \leq o(\rho^d)$$

and, since $v = \delta$ on $\partial\Omega$,

$$\left| \int_{\partial\Omega} h''_\chi(\bar{u}_\chi) v \cdot v dS \right| = \left| \int_{\partial\Omega} h''_\chi(\bar{u}_\chi) \delta \cdot \delta dS \right| \leq C \|\delta\|_{H^1(\Omega)}^2 \leq o(\rho^d)$$

by Lemma 4.3. A similar estimate holds for the last term of (4.37). The penultimate term is bounded by

$$\left| \int_{\omega_\rho} (j^0 - j^1)'(u_\chi) \cdot v dV \right| \leq C \rho^{d/2} (\|u_\chi\|_{L^\infty(\omega_\rho)} + 1) \|v\|_{L^2(\Omega)} \leq o(\rho^d)$$

because the background solution u_χ is smooth on ω_ρ (it does not "see" the inclusion). Thus, the two last lines of (4.37) are small of the order of $o(\rho^d)$. All other terms in (4.37) contribute to the final result, formula (4.4). First, by rescaling and continuity of u_χ on ω_ρ , we have

$$\int_{\omega_\rho} (j^0 - j^1)(u_\chi) dV = \rho^d |\omega| (j^0 - j^1)(u_\chi(x_0)) + o(\rho^d).$$

Similarly, by continuity of $e(u_\chi)$, using the notation $\xi_0 = e(u_\chi)(x_0)$, and since $v = \theta w_{\xi_0}^\rho + \delta$ with $\theta \equiv 1$ in ω_ρ , we deduce

$$\begin{aligned} \frac{1}{2} \int_{\omega_\rho} [A] \left(e(u_\chi) \cdot e(u_\chi) + 2e(u_\chi) \cdot e(v) \right) dV &= \frac{\rho^d}{2} \int_{\omega} [A] \left(\xi_0 \cdot \xi_0 + 2\xi_0 \cdot e_y(w_{\xi_0}) \right) dV(y) \\ &\quad + \int_{\omega_\rho} [A] e(u_\chi) \cdot e(\delta) dV + o(\rho^d). \end{aligned}$$

Using again the continuity of $e(u_\chi)$ in ω_ρ and (4.36), we bound the last term

$$\left| \int_{\omega_\rho} [A] e(u_\chi) \cdot e(\delta) dV \right| \leq C \rho^{d+1}.$$

Second, from the variational formulation of (4.20) we get

$$\begin{aligned} \frac{1}{2} \int_{\Omega} A_\chi e(v) \cdot e(v) dV - \frac{1}{2} \int_{\omega_\rho} [A] e(u_\chi(v)) \cdot e(v) dV &= \frac{1}{2} \int_{\Omega} A_{\chi_\rho} e(v) \cdot e(v) dV \\ &= \frac{1}{2} \int_{\omega_\rho} [A] e(u_\chi) \cdot e(v) dV = \frac{\rho^d}{2} \int_{\omega} [A] \xi_0 \cdot e_y(w_{\xi_0}) dV(y) + o(\rho^d), \end{aligned}$$

where we have again replaced v by $w_{\xi_0}^\rho + \delta$ in ω_ρ and neglected the δ term. Third, from (3.2) we have

$$\int_{\Omega} A_\chi e(u_\chi) \cdot e(v) dV = \int_{\Omega} f_\chi \cdot v dV + \int_{\partial\Omega} g_\chi \cdot v dS.$$

Thus, the Taylor expansion (4.37) of the objective function is rewritten

$$\begin{aligned} J(\chi_\rho) &= J(\chi) + \int_{\Omega} \left(f_\chi + j'_\chi(x, u_\chi) \right) \cdot v dV + \int_{\partial\Omega} \left(g_\chi + h'_\chi(x, u_\chi) \right) \cdot v dS \\ &\quad - \frac{\rho^d}{2} \int_{\omega} [A] \left(\xi_0 \cdot \xi_0 + \xi_0 \cdot e_y(w_{\xi_0}) \right) dV(y) \\ &\quad + \rho^d |\omega| (j^0 - j^1)(u_\chi(x_0)) + o(\rho^d). \end{aligned} \tag{4.38}$$

By Lemma 4.1 we know that

$$-\frac{\rho^d}{2} \int_{\omega} [A] \left(\xi_0 \cdot \xi_0 + \xi_0 \cdot e_y(w_{\xi_0}) \right) dV(y) = -\frac{\rho^d}{2} M \xi_0 \cdot \xi_0.$$

It remains to show that the two first integrals in the right hand side of (4.38) are of order $\mathcal{O}(\rho^d)$ and find formula (4.4) for the topological derivative. To do so, we use the adjoint problems (4.34) and (4.35) as follows. Multiplying (4.34)

by v and (4.20) by $\tilde{p}_{\chi_\rho}$ we obtain

$$\begin{aligned}
& \int_{\Omega} \left(f_{\chi} + j'_{\chi}(x, u_{\chi}) \right) \cdot v dV + \int_{\partial\Omega} \left(g_{\chi} + h'_{\chi}(x, u_{\chi}) \right) \cdot v dS = \int_{\Omega} A_{\chi_\rho} e(\tilde{p}_{\chi_\rho}) \cdot e(v) dV \\
&= \int_{\omega_\rho} [A] e(u_{\chi}) \cdot e(\tilde{p}_{\chi_\rho}) dV = \int_{\omega_\rho} [A] e(u_{\chi}) \cdot e(p_{\chi} + q) dV \\
&= \int_{\omega_\rho} [A] e(u_{\chi}) \cdot e(p_{\chi} + \theta w_{\xi'_0}^\rho + \delta) dV \\
&= \rho^d \int_{\omega} [A] \xi_0 \cdot (\xi'_0 + e_y(w_{\xi'_0})) dV(y) + o(\rho^d) \\
&= \rho^d M \xi_0 \cdot \xi'_0 + o(\rho^d),
\end{aligned}$$

by application of Lemma 4.4, rescaling, using the continuity of $e(u_{\chi})$ and $e(p_{\chi})$ in ω_ρ and thanks to the formula for M in Lemma 4.1 (recall that $\xi'_0 = e(p_{\chi})(x_0)$). Eventually we have proved

$$J(\chi_\rho) = J(\chi) - \frac{\rho^d}{2} M \xi_0 \cdot \xi_0 + \rho^d M \xi_0 \cdot \xi'_0 + \rho^d |\omega| (j^0 - j^1)(u_{\chi}(x_0)) + o(\rho^d),$$

which is precisely formula (4.4). This achieves the proof of Theorem 4.1 since the properties of M have been proved in Lemma 4.2.

5 Computational algorithm

The main task is to compute, for each discrete time t_i , $i \geq 0$, a minimizer χ_i of the Francfort-Marigo model (2.8). As we already said, we are interested in local minima. Our notion of local minima is numerical in essence, that is, we minimize (2.8) with a gradient descent algorithm in the level set framework. A minimum is thus local in the sense of perturbations of the location of the interface Σ . Our algorithm is made of two nested loops:

- (i) an outer loop corresponding to the increasing sequence of discrete times t_i , $i \geq 0$,
- (ii) an inner loop of gradient iterations for the minimization of the functional (2.8) at each fixed time step t_i .

The irreversibility constraint (2.7) on the damaged zone is taken into account in the outer loop (i), whereas the inner loop (ii) is purely numerical and is not subject to this irreversibility constraint between two successive iterates of (ii). The inner loop is performed with the level set method of Osher and Sethian [53] that we now briefly describe (it is very similar with its application in the context of shape optimization [5], [65]).

In the fixed bounded domain Ω , uniformly meshed once and for all, we parametrize the damaged zone Ω^0 by means of a level set function ψ such that

$$\begin{cases} \psi(x) = 0 & \Leftrightarrow x \in \Sigma, \\ \psi(x) < 0 & \Leftrightarrow x \in \Omega^0, \\ \psi(x) > 0 & \Leftrightarrow x \in \Omega^1. \end{cases}$$

The normal n to the damaged region Ω^0 is recovered as $\nabla\psi/|\nabla\psi|$ and the mean curvature H is given by the divergence of the normal $\operatorname{div} n$. These quantities are evaluated by finite differences since our mesh is uniformly rectangular. Although n and H are theoretically defined only on Σ , the level-set method allows to define easily their extension in the whole domain Ω .

Following the minimization process, the damaged zone is going to evolve according to a fictitious time s which corresponds to descent stepping and has nothing to do with the "real" time t_i in the outer loop (i). As is well-known, if the shape is moving with a normal velocity $\mathcal{V}$, then the evolution of the level-set function is governed by a simple Hamilton-Jacobi equation [52], [60],

$$\frac{\partial\psi}{\partial s} + \mathcal{V}|\nabla\psi| = 0, \quad (5.1)$$

which is posed in the whole body Ω , and not only on the interface Σ , when the velocity $\mathcal{V}$ is known everywhere. We now explain how we derive $\mathcal{V}$ for our specific problem.

For the minimization of (2.8) we use the shape derivative given by (3.9),

$$J'(\chi)(\theta) = \int_{\Sigma} \mathcal{D}\theta \cdot n \, dS, \quad (5.2)$$

where the integrand $\mathcal{D}(x) \in L^2(\Sigma)$ is given by Theorem 3.1 and $\theta \in W^{1,\infty}(\Omega; \mathbb{R}^d)$ in any admissible direction of derivation. Since only the normal component of θ plays a role in (5.2), we always look for a normal vector field, i.e., we restrict our attention to

$$\theta = v n \quad \text{with a scalar field } v \in W^{1,\infty}(\Omega). \quad (5.3)$$

The velocity $\mathcal{V}$ is going to be chosen as an "optimal" direction of derivation, v , such that

$$J'(\chi)(\mathcal{V}n) = \int_{\Sigma} \mathcal{D}\mathcal{V} \, dS \leq 0. \quad (5.4)$$

The simplest choice $\mathcal{V} = -\mathcal{D}$, which enforces (5.4) and is commonly used in structural optimization [5], is not satisfactory in the present situation, since $\mathcal{D}$ is defined as a jump on Σ only, without natural extension over Ω . We therefore suggest another choice based on the identification of the duality product between $J'(\chi)$ and v (recalling that $\theta = v n$) with the usual scalar product in $H^1(\Omega)$. In other words we represent $J'(\chi)$ by a scalar field $(-\mathcal{V}) \in H^1(\Omega)$ such that, for any test function v ,

$$J'(\chi)(v n) = - \int_{\Omega} (\nabla\mathcal{V} \cdot \nabla v + \mathcal{V}v) \, dV. \quad (5.5)$$

Combining (5.2) and (5.5), and requiring the descent condition (5.4), we choose the velocity $\mathcal{V}$ in (5.1) as the unique solution in $H^1(\Omega)$ of the variational formulation

$$\int_{\Omega} (\nabla\mathcal{V} \cdot \nabla v + \mathcal{V}v) \, dV = - \int_{\Sigma} \mathcal{D}v \, dS \quad \forall v \in H^1(\Omega). \quad (5.6)$$

Solving (5.6) to compute a shape derivative is a usual trick in shape optimization for regularizing derivatives [3], [55]. However, (5.6) is used here mostly for extending the “natural” velocity $\mathcal{D}$ away from the interface Σ . In practice we add a small positive coefficient (linked to the mesh size) in front of the gradient term in (5.6) in order to limit the regularization and the spreading of the velocity around the interface.

For numerical purpose, as explained in [5], the surface integral in the right hand side of (5.6) is written as a volume integral

$$\int_{\Sigma} \mathcal{D} v dS = \int_{\Omega} \delta_{\Sigma} \mathcal{D} v dV, \quad (5.7)$$

where the Dirac mass function δ_{Σ} is approximated by

$$\delta_{\Sigma}^{\epsilon} = \frac{1}{2} |\nabla(s_{\epsilon}(\psi))|$$

with the following approximation of the sign function

$$s_{\epsilon}(x) = \frac{\psi(x)}{\sqrt{\psi(x)^2 + \epsilon^2}},$$

where $\epsilon > 0$ is a small parameter chosen in order to spread the integration over a few mesh cells around the interface. The integrand $\mathcal{D}$, being actually a jump $[\mathcal{E}]$ of a discontinuous quantity $\mathcal{E}$ (see formulae (3.11) or (3.12)), requires also some special care. In (5.7) we replace $\mathcal{D} = [\mathcal{E}]$ by

$$\mathcal{D}_{\text{approx}} = [\mathcal{E}]_{\text{approx}} = 2((1 - \chi)\mathcal{E} - \chi\mathcal{E})$$

where χ is the characteristic function of the damaged phase (numerically it is always equal to 0 or 1 except in those cells cut by the interface where it is interpolated by the local proportion of damaged phase in the cell). The factor 2 in the above formula takes into account the fact that

$$\int_{\Omega} \delta_{\Sigma}^{\epsilon} \chi dV \approx \frac{1}{2} \int_{\Sigma} dS.$$

In our numerical experiments we use formula (3.11) and not (3.10) because the latter one exhibits singular jumps when the damaged phase is very weak (which is the case for our simulations of crack propagation). Of course, in the case of a degenerate (zero) damaged phase we can use the limit formula given by Remark 3.2 which are of course much simpler (we did so in our previous publication [7]).

Remark 5.1 *Note that the same problem of computing a jump of a discontinuous quantity at an interface was independently addressed in [45]. This work is also relying on the level set method and is applied to the Mumford-Shah functional in image segmentation. It can also be applied to fracture mechanics and, as our proposed approach, it relies on a fattening of the fracture path.*

Our proposed algorithm for the inner loop (ii) is an iterative method, structured as follows:

1. Initialization of the level set function ψ^0 as the signed distance to the previous damaged interface Σ_i corresponding to the characteristic function $\chi^0 \equiv \chi_i$.
2. Iteration until convergence, for $k \geq 0$:
 - (a) Computation of the state u_k by solving a problem of linear elasticity with coefficients $A_{\chi^k} = (1 - \chi^k)A^1 + \chi^k A^0$.
 - (b) Deformation of the interface by solving the transport Hamilton-Jacobi equation (5.1). The new interface Σ^{k+1} is characterized by the characteristic function χ^{k+1} or the level-set function ψ^{k+1} solution of (5.1) after a pseudo-time step Δs_k starting from the initial condition $\psi^k(x)$ with velocity $\mathcal{V}_k$ computed through (5.6) in terms of u_k . The pseudo-time step Δs_k is chosen such that $J(\chi^{k+1}) \leq J(\chi^k)$.
 - (c) Irreversibility constraint: we replace χ^{k+1} by $\max(\chi^{k+1}, \chi^0)$ where $\chi^0 \equiv \chi_i$ corresponds to the damaged zone at the previous iteration of the outer loop (i).

At each iteration of above the inner loop, for stability reasons, we also reinitialize the level-set function ψ [52], [60]. This is crucial because the integrand $\mathcal{D}$ of the shape derivative involves normal and tangential components of stress or strain tensors, which requires a precise evaluation of the normal n by formula $\nabla\psi/|\nabla\psi|$. Actually it turns out that this reinitialization step must be much more precise in the present context than for shape optimization [5]. Indeed, a poor reinitialization can influence the propagation of the damage zone. We therefore use a trick suggested in [57] for an increased accuracy of the second-order reinitialization process. The Hamilton-Jacobi equation (5.1) is solved by an explicit second order upwind scheme on a Cartesian grid. The boundary conditions for ψ are of Neumann type. Since this scheme is explicit in time, its time step is given by a CFL condition. In numerical practice we often take the descent step Δs_k of the order of the Hamilton-Jacobi time step which stabilizes the damage evolution.

6 Simulation results

Otherwise explicitly mentioned, all our numerical experiments are performed with a healthy material having Young modulus $E^1 = 1000$ and Poisson ratio $\nu^1 = 0.3$ (white material on the pictures). The damaged phase (black material on the pictures) has always Poisson ratio $\nu^0 = 0.3$ (the fact that $\nu^0 = \nu^1$ does not matter) but has different Young modulus in different places. More precisely, in Subsection 6.1 we consider a moderately weak damaged phase with $E^0 = 500$, while in the next subsections the damage phase is assumed to be almost degenerate, i.e. $E^0 = 10^{-3}$: this last case corresponds to a limit where

our model behaves almost like a brittle fracture model. Actually some models of fracture mechanics [36] are approximated by Γ -convergence techniques [21], [22], which is similar in spirit to a damage model. Therefore it is not surprising that our damage model can predict crack propagation.

In the sequel we call *critical load* the value of the applied displacement for which the damage region has completely crossed the computational domain (meaning failure of the structure), and *initiation load* the first value for which the damage zone departs from its initialization. All other intermediate load values are called *subcritical*, while values above the critical one are called *supercritical*.

In order to validate our method, two types of numerical experiments are done. On the one hand, for simple problems we check convergence under various refinements of the mesh size, of the time step, etc. On the other hand, we compare our results with a variety of existing benchmarks tested by laboratory experiments or other numerical methods.

6.1 2d damage simulation

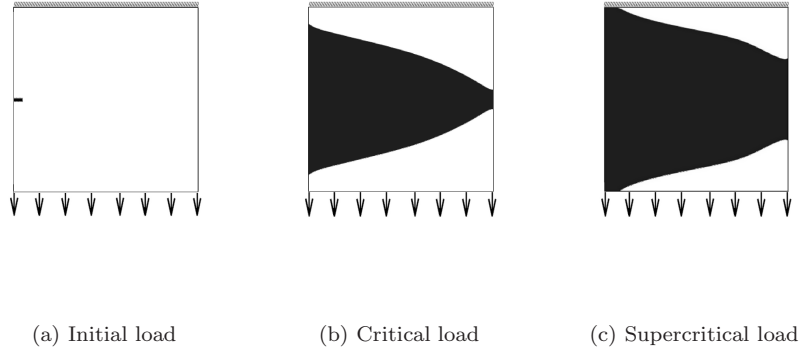


Figure 1: Mode I damage for the 320×320 mesh with 100 time steps. Initial configuration with an imposed displacement of 0.02 (a) critical load at an imposed displacement of 0.06 (b) and supercritical load with an imposed displacement of 0.072 (c).

The numerical experiments with a moderately weak damaged phase, $E^0 = 500$, are easier to perform than the ones with a degenerate phase but their results are mechanically less interesting. Therefore we content ourselves with a single experiment, namely a mode I traction (Fig. 1) in a square box of size 1 with a Griffith energy release parameter $\kappa = 1$. The imposed vertical displacement at the bottom is increased from 0.02 to 0.08 on a given time interval and shown as

abscissa in the figures. In order to study convergence under mesh refinement, four different meshes are used: 280×280 (coarse), 320×320 (intermediate 1), 400×400 (intermediate 2), 452×452 (fine). Similarly, for convergence under time step refinement, we divide the time interval successively in 100, 200 and 400 time steps. Fig. 1 displays the result for the 320×320 mesh with 100 time steps. There are no subcritical loads: the initiation load coincides with the critical load which means that, not only the appearance of damage is sudden, but the structure completely fails in just one load displacement increment. Fig. 2 shows that the results are convergent under mesh refinements. The curves are almost identical and the position of the critical load is clearly converging as the mesh is refined. Fig. 3 is concerned with convergence under time-step refinement. In particular, the critical loads for the three time refinements show very good agreement, meaning that our quasi-static numerical model seems to converge to a time-continuous model as the time step tends to zero.

The cost function (2.4), which is minimized at each time step, is the sum of the Griffith or damage energy and of the elastic energy. The damage energy, displayed on Fig. 2a, is discontinuous and increases abruptly at the critical load. Similarly, the elastic energy, displayed on Fig. 2b, is discontinuous decreasing at the critical load, which corresponds to the release of energy produced by damage. However, by comparison, the cost function, displayed on Fig. 2c, seems to be roughly continuous with respect to time (there is a small bump at the critical load).

Eventually, we have checked the following formula for the dissipation of energy (see Theorem 4.1 in [34])

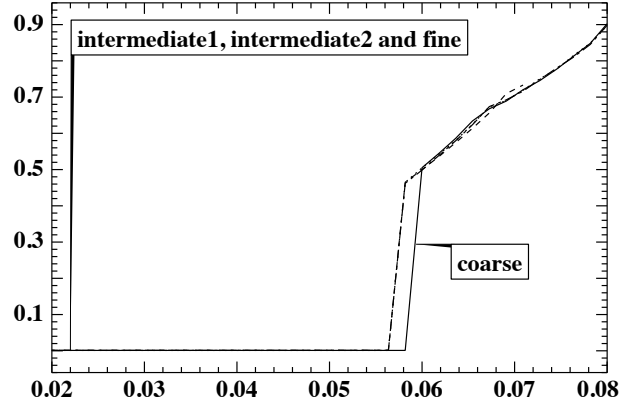
$$\min_{u, \chi} \mathcal{J}(u, \chi)(T) - \min_{u, \chi} \mathcal{J}(u, \chi)(0) = \int_0^T \int_{\Gamma_D} (\sigma n) \cdot \frac{du_D}{dt}(t) dS dt \quad (6.1)$$

where $\mathcal{J}$ is defined by (2.3) and u_D is the applied displacement. In the absence of any other applied load, formula (6.1) expresses the conservation of total energy. If we plot the right hand side of (6.1), we obtain exactly the cost function on the left hand side with a numerical precision of the order of 10^{-6} .

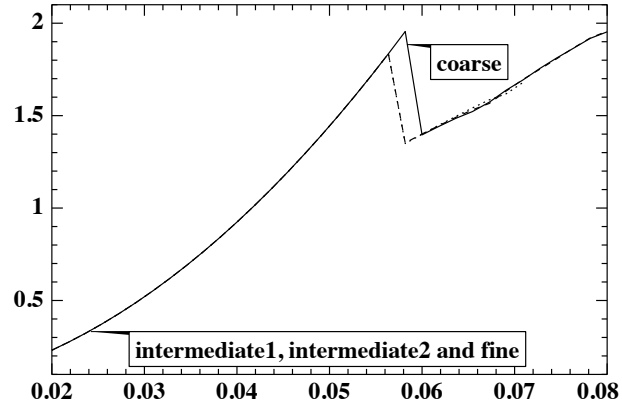
6.2 2d fracture with mode I loading

We now switch to a very weak damage phase, $E^0 = 10^{-3}$, in order to mimic crack propagation. Here the Griffith energy release parameter is $\kappa = 3.5$. We perform the same mode I traction experiment as in Subsection 6.1 with the same parameter values otherwise explicitly specified. For a mesh of size 320×320 , with an initial crack having a width of 8 mesh cells, a height of 16, and for 100 time steps, when the imposed vertical displacement at the bottom is increased from 0.005 to 0.05, we obtain a crack which breaks the structure in just one time increment (see Fig. 4). For all other values of the parameters, the same qualitative behavior is observed: the initial and critical loads are the same for a mode I crack.

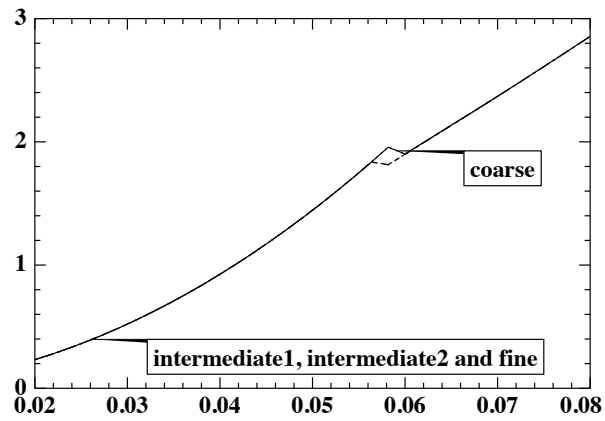
We then investigate the convergence under time-step refinement (Fig. 5). The imposed vertical displacement at the bottom is increased from 0.005 to



(a) Damage energy.



(b) Elastic energy.



(c) Cost function.

Figure 2: Mode I damage experiment: variations of various energies as functions of the imposed displacement. Griffith energy (a), elastic energy (b) and cost function (c) for four different mesh refinements with 100 time steps.

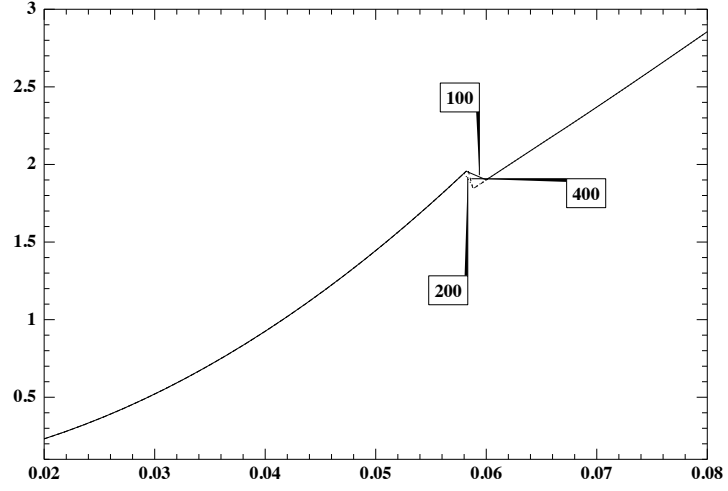


Figure 3: Mode *I* damage experiment: cost function with respect to the imposed displacement for the 320×320 mesh and for 100, 200 and 400 time steps.

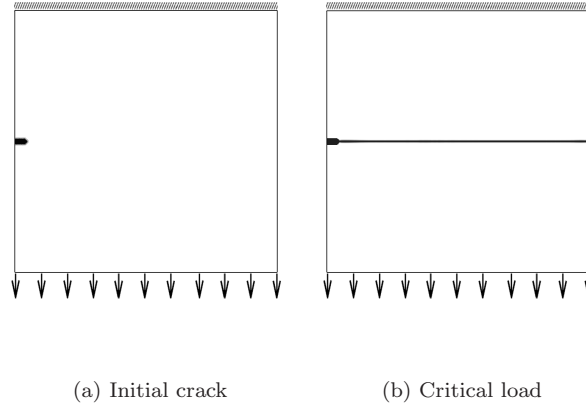


Figure 4: Mode *I* crack: initial configuration (a) and critical load at 0.0028 (b).

0.05. The time interval is successively divided in 100, 200 and 400 time steps. The mesh is of size 320×320 with an initial crack having a width of 8 mesh cells. On Fig. 5 the values of the critical loads are obviously converging as Δt goes to zero. Therefore we believe that our quasi-static numerical model, as applied to “crack-like” damage, also converges to a time-continuous model as the time step tends to zero.

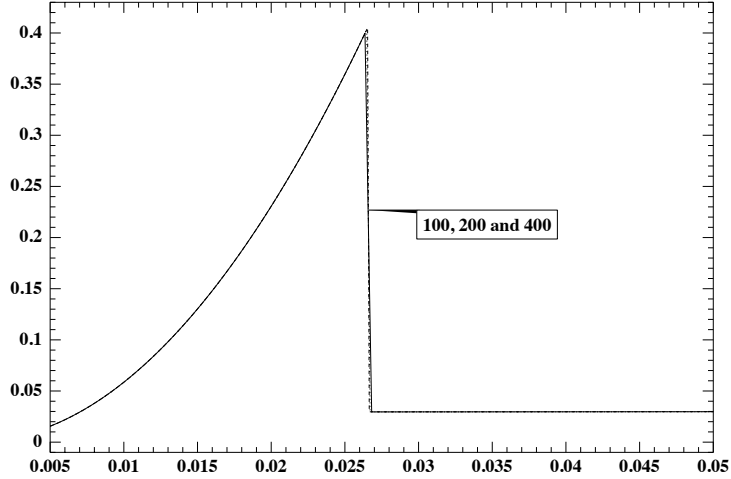


Figure 5: Time-step refinement for the mode *I* crack: cost function with respect to the imposed displacement for three time refinements.

We now perform three different test cases of convergence under mesh refinement with four meshes for each test (see Figs. 6, 7, 8). The four different meshes are: 280×280 (coarse), 320×320 (intermediate 1), 400×400 (intermediate 2), 452×452 (fine). From these three refinement processes, only the last one is fully satisfactory but the two previous ones are illuminating so we keep them in our exposition.

In the first case (Fig. 6), the given initial crack has a constant width. In other words, the number of cells in a cross-section of the initial crack is 6, 8, 10 and 12 respectively for the four different meshes. The initial crack tip is slightly rounded for the finer meshes in order to avoid the appearance of sharp corners. The imposed vertical displacement at the bottom is increased from 0.005 to 0.05. On Fig. 6 we observe that the value of the critical load is decreasing as the mesh is refined and does not seem to converge (especially when compared to the damage case in Fig. 2). Similarly, the value of the cost function at the critical load is decreasing with finer meshes because a thinner crack (on a finer mesh) costs less Griffith energy. Therefore, contrary to the damage experience of Subsection 6.1, no mesh convergence can be claimed in this first experiment.

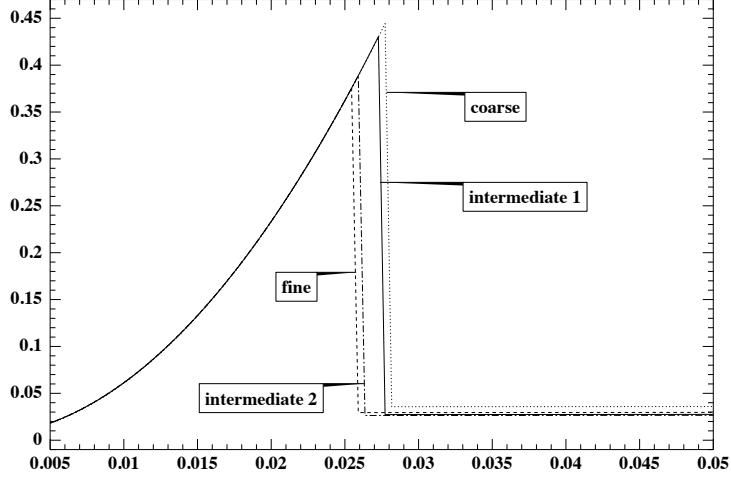


Figure 6: First mesh convergence test for the mode I crack: cost function with respect to the imposed displacement for various meshes and with 100 time steps.

In the second case (Fig. 7), we anticipate that a crack should have a thickness of the order of a few cells Δx when Δx goes to zero. Therefore, whatever the value of Δx , the initial crack is chosen with a width of two cells only, which means that the initial crack is thinner and thinner as the mesh is refined. The imposed vertical displacement at the bottom is now increased from 0.005 to 0.05. Refinements with respect to the mesh size are shown on Fig. 7. The critical load again occurs sooner for finer meshes, thereby indicating that there is no convergence under mesh refinement.

In the third case (Fig. 8), we replace the minimization of (2.4) by that of the scaled cost function (2.11) in an attempt to show that there is indeed convergence under mesh refinement. In other words we replace the Griffith energy release parameter κ by its scaled version $\frac{\gamma}{\Delta x} = \frac{\kappa \ell}{\Delta x}$ where Δx is the mesh size. We again choose $E^0/E^1 = 10^{-6}$ and take $\ell = 1/320$ (so that $\ell/\Delta x = 1$ for the “intermediate 1” mesh). In practice, this scaling implies that it is more difficult to create damage for finer meshes, a phenomenon that should balance the opposite effect displayed in the two previous cases. As explained in Section 2.3 this scaling is precisely designed so the scaled Griffith energy converges to a surface energy when Δx goes to zero. On Fig. 8 we check that the critical loads are converging, so we claim that mesh convergence is observed with this particular scaling of κ .

Eventually, we have again checked the balance of energy expressed in (6.1): the cost function perfectly matches the time integral of the dissipated energy (i.e. the the right hand side of (6.1)), up to a numerical precision of the order

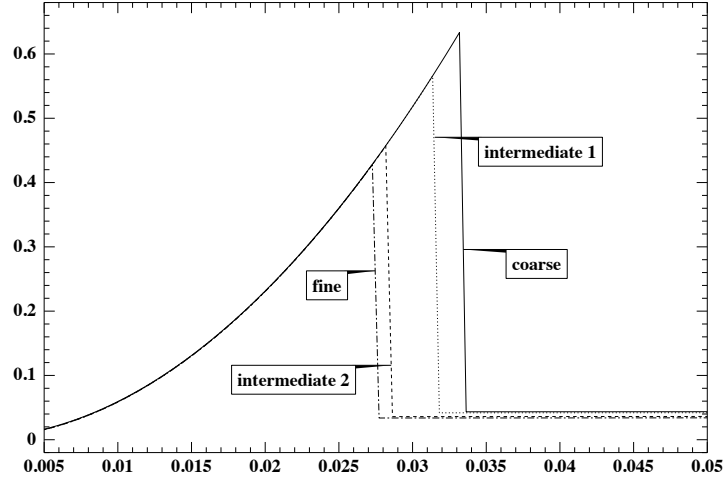


Figure 7: Second mesh convergence test for the mode I crack: cost function with respect to the imposed displacement for various meshes and with 100 time steps.

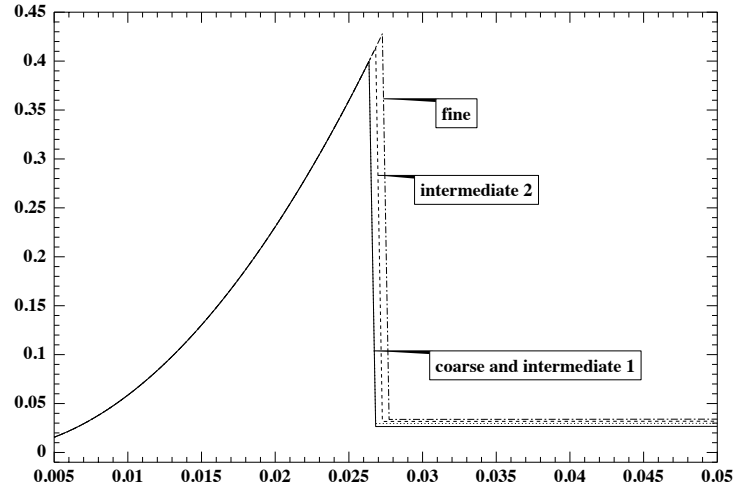


Figure 8: Third mesh convergence test for the mode I crack: cost function (2.11) with respect to simultaneous mesh, crack and κ refinements for 100 time steps.

of 10^{-6} .

6.3 2d fracture with mode II loading

We now turn to another crack experiment with a mode II loading. The dimensions of the computational domain are the same as in the above mode I experiments. The weak damage phase is again $E^0 = 10^{-6}$ while the imposed horizontal displacement at the bottom is here increased from 0.1 to 1.0 on a given time interval. On Figs. 9a and 9b the initial and critical cracks are shown for the 320×320 mesh. We emphasize that “critical” has not exactly the same meaning here as for the mode I crack: the mode II crack does not actually break the structure. The crack stops just a few cells before reaching the opposite boundary and does not move anymore as the load increases. This longest crack configuration is called critical. However, fracture is here again brutal in the sense that the initiation load coincides with the critical load. On Fig. 9 we can see that the mode II loading yields a branching of the crack. By symmetry and since the model is linearized elasticity the two crack branches are symmetric, one in compression and the other in traction. It means that another mechanical model taking into account the non-interpenetration of material would produce only the branch under traction, i.e., the lips of which are opening under the load, as it can be observed in physical experiments.

Four different meshes are used: 280×280 (coarse), 320×320 (intermediate 1), 400×400 (intermediate 2), 452×452 (fine). In our experiment (Fig. 10), we minimize the scaled version (2.11) of the cost function, i.e., κ is replaced by $\frac{\kappa \ell}{\Delta x}$, and the crack width is always exactly two mesh cells. Convergence under mesh refinement is clearly obtained. Even more, the two finest mesh curves almost coincide.

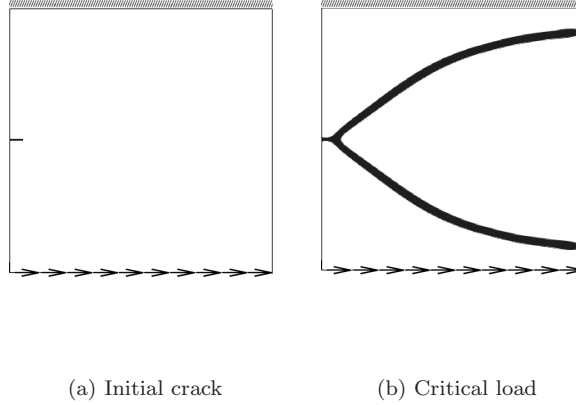


Figure 9: Mode II crack experiment: initial configuration (a) and critical load at 0.49 (b).

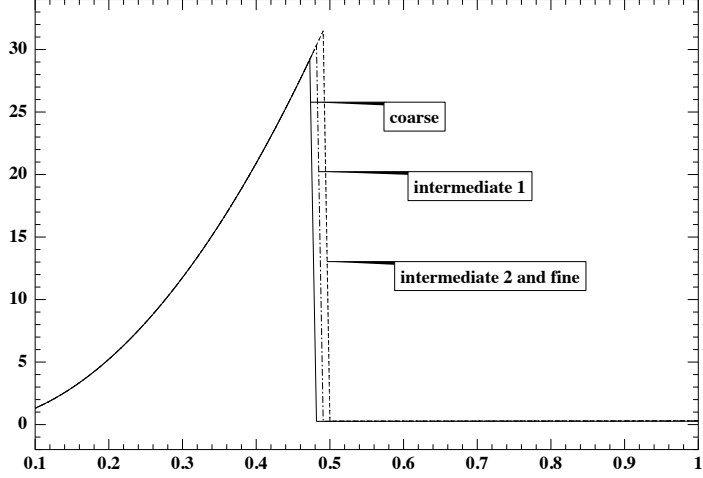


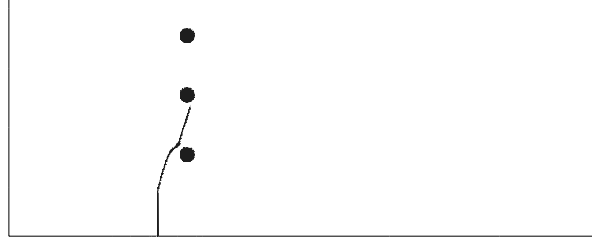
Figure 10: Mode *II* crack test: cost function (2.11) with respect to simultaneous mesh, crack and κ refinements for 100 time steps.

6.4 Bittencourt's drilled plate

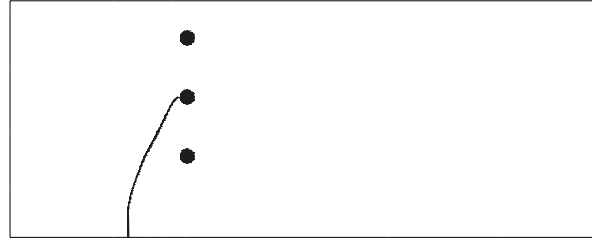
This test case has been proposed in [17] where we found all the required numerical values of the parameters. It has been reproduced in many other works, including [19, 14]. The Young modulus of the healthy phase is 3000 and its Poisson ratio 0.35. The damaged phase has a Young modulus $3 \cdot 10^{-3}$ and the same Poisson ratio. The value of the Griffith energy release parameter is $\kappa = 0.0014$. Contrary to all other numerical simulations in this paper, the present experiment has been performed with a given fixed applied force instead of a sequence of increasing displacements. The vertical unit load is applied on a single concentrated point of the upper body boundary. The value of the Griffith energy release parameter κ is such that this applied unit load is critical, i.e. a single time step produces the cracks displayed on Figs. 11a and 11b for two different crack initializations. The distance from the left face to the initial crack is denoted by a , while b is the initial crack length. The 3 holes carry a Neumann boundary condition. We use a non-uniform rectangular mesh of size 470×800 which is more refined in the vicinity of the holes. These two results are in good agreement with laboratory experiment of [17], although that of Fig. 11a shows a slightly different crack path near the second hole.

6.5 Coalescence of multiple cracks

This experiment is made on a pre-cracked sample (of size 1.6×2.2) with a vertical imposed displacement along the vertical sides (corresponding to a mode



(a) first case $a = 6, b = 1.5$



(b) second case $a = 5, b = 1$

Figure 11: The two Bittencourt's experiments:(a) first case $a = 6, b = 1.5$, (b) second case $a = 5, b = 1$,

II type loading). The healthy material has Young modulus $E^1 = 1$ and Poisson ratio $\nu^1 = 0.3$, the damaged phase has the same Poisson ratio but a smaller Young modulus $E^0 = 10^{-3}$. The value of the Griffith energy release parameter is $\kappa = 10^{-7}$. The imposed vertical displacement is increased from 0.001 to 0.005 with 100 time steps. The critical load is attained at 0.0014. Two different meshes are shown on Fig. 12.

6.6 Traction experiment on a fiber reinforced matrix

We perform a test case proposed in [21] where all precise values of the parameters can be found. The setting of Fig. 14a is the following. A unit vertical displacement is exerted on the upper layer of the solid which is also clamped at its midpoint to avoid translations and rotations. The fiber (grey inclusion on Fig. 14a) is also clamped. The healthy material has Young modulus $E^1 = 1000$ and Poisson ratio $\nu^1 = 0.3$, the damaged phase has the same Poisson ratio but a much smaller Young modulus $E^0 = 10^{-3}$. The value of the Griffith energy release parameter is $\kappa = 8000$. Excellent agreement with the numerical results of [21] are observed. Let us emphasize that this experiment is the only one using the topological derivative to initiate the damaged zone: the map of the topological gradient at the initialization is displayed on Fig. 13. More precisely, the initial body is completely healthy without any damage: the applied load is gradually increased, until damage appears because the topological derivative

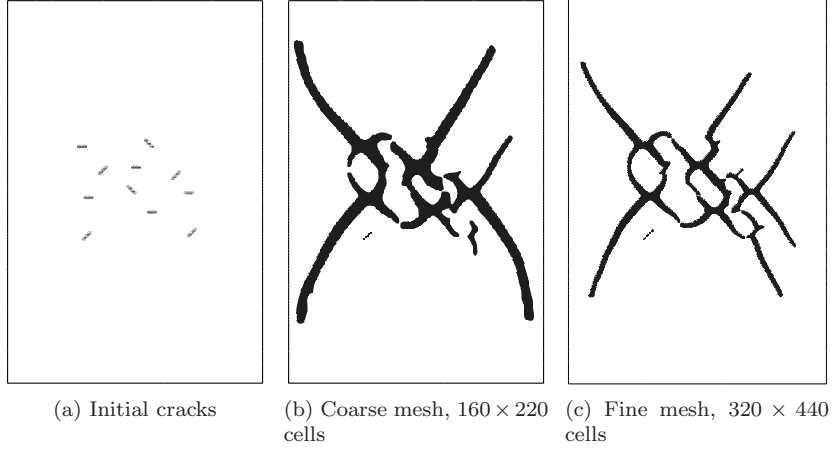


Figure 12: Multiple crack experiment with a mode II loading.

becomes negative. Once the damaged zone has been initialized we use our shape gradient method to propagate the crack without further use of the topological gradient. The final crack on Fig. 14 is very similar to that computed in [21].

There is a subtle point here in the use of the topological derivative. Such a notion is well defined for the cost function (2.3) which features a bulk Griffith-type energy. However, it is not possible to define a topological derivative for the cost function (2.13) which has a surface Griffith energy. Indeed, surface energy asymptotically dominates bulk energy for small inclusions and no balance can be established. In other words the notion of topological derivative makes sense in our damage model but is irrelevant for fracture models.

Note also that the initialization pattern suggested by the topological derivative (which is necessarily small by definition) is not a local minimizer of the cost function. Thus, such an initialization is difficult to compare with other ones in the literature and its small extension is not a contradiction with theoretical results [30] stating that an initial crack (minimizing the cost function) cannot be too small.

Therefore we now compare different (ten) initializations for the same problem. The first initialization is that given by the topological gradient and already displayed on Figure 14b. The nine other ones, displayed on Figure 15, are increasingly larger cracks obtained as intermediate inner iterates in the previous computation (they are therefore not local minimizers of the cost function). The larger cracks of Figure 15 are very similar to the initial crack pattern in [22]. For each of these initial cracks we restarted the crack evolution from a zero imposed displacement which is then gradually increased. The evolutions of the elastic energy, the Griffith energy and the cost function, when the imposed displacement is increased, are shown on Figure 16. One can clearly see the importance of the initialization, a well-known fact in the minimization of non-convex ener-

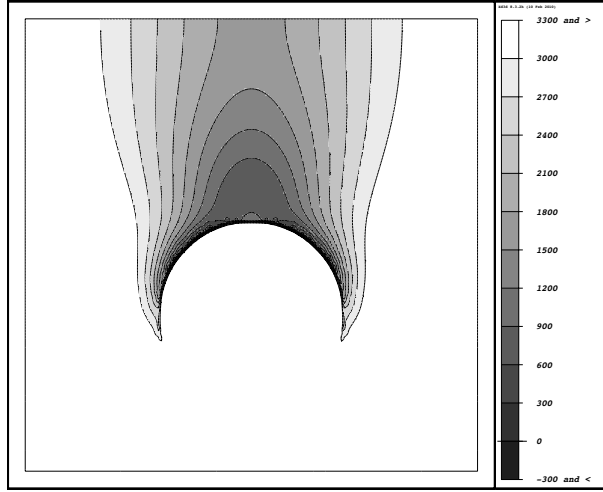


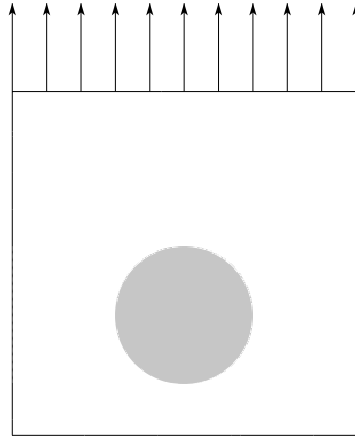
Figure 13: Fiber-reinforced matrix: isocontours of the topological derivative at the initialization.

gies. The idea of restarting the minimization at smaller loading parameters from an intermediate crack solution at a larger loading parameter has already been exploited in the so-called back-tracking algorithm of [20], [22] for global minimization. Of course, the larger the initial crack is, the smaller is the critical load. Interestingly enough we found that, for the 4th up to the 10th initializations, part of the crack evolution is smooth with respect to the loading parameter. This is actually the only occurrence in our numerical tests of a continuous crack evolution: all other examples feature a brutal fracture process. This smooth behavior is very similar to that obtained in [22] (see Figures 31 and 32 on page 92). On Figure 17, for the 5th initialization, we plot the two crack patterns obtained for the values 0.28 and 0.44 of the imposed displacement: in between the crack evolution is continuous.

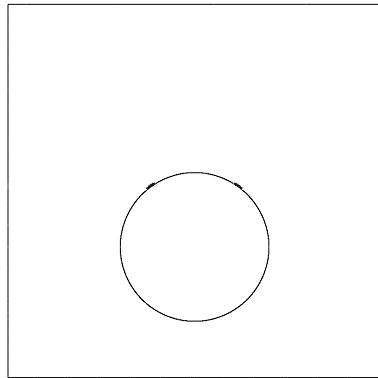
6.7 3d mode 1, mode 2 and mode 3 cracks

We eventually conclude our numerical tests by performing the 3 different mode loadings in 3d with boundary conditions (imposed displacements) shown on Fig. 18. We work with a cubic domain of size $1 \times 1 \times 1$ meshed with $80 \times 80 \times 80$ cubic cells: its left back face (circles) is fixed while a uniform displacement of modulus 0.04 is applied on its right front face. The healthy material has Young modulus $E^1 = 10^4$ and Poisson ratio $\nu^1 = 0.3$, the damaged phase has the same Poisson ratio but a smaller Young modulus $E^0 = 1$. The value of the Griffith energy release parameter is $\kappa = 1$. The initial and final cracks are shown on Fig. 19.

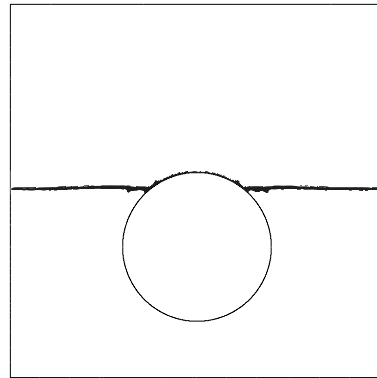
For this large test case (involving around 1.56×10^6 degrees of freedom)



(a) Initial healthy fiber-reinforced matrix.



(b) Initial damage nucleated by the topological derivative.



(c) Final crack.

Figure 14: Fiber-reinforced matrix: an example of crack initiation by the topological derivative.

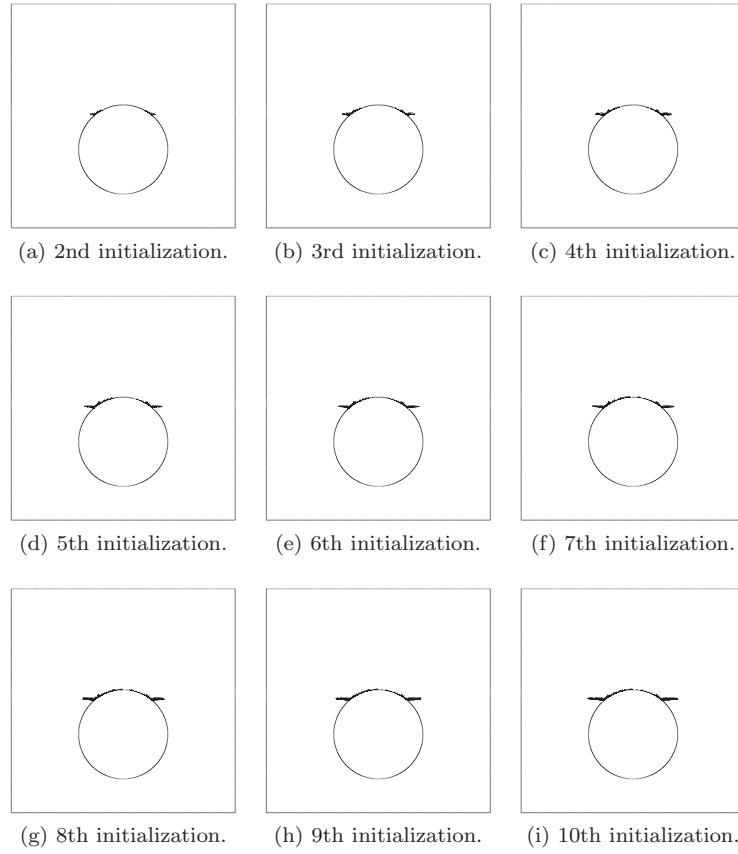
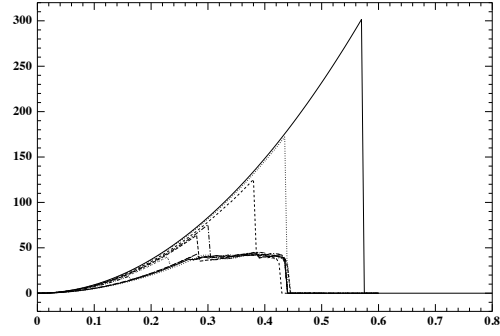
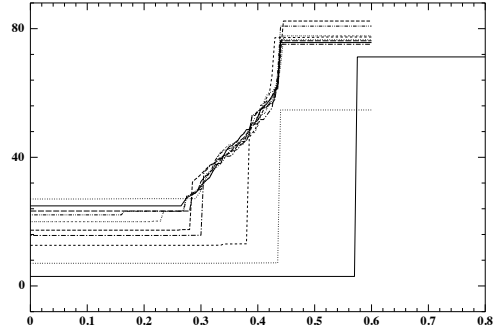


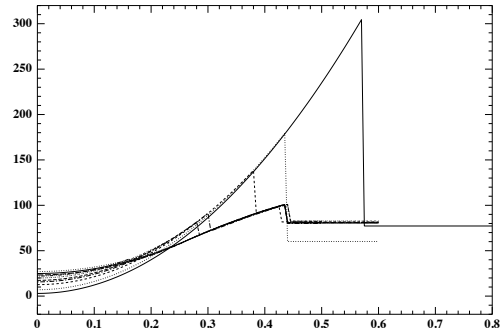
Figure 15: Fiber-reinforced matrix: different crack initializations.



(a) Elastic energy.



(b) Griffith energy.



(c) Cost function.

Figure 16: Elastic energy, Griffith energy and cost function, as a function of the imposed displacement, for the ten initializations of Figure 15. The first initialization corresponds to the largest critical load (or discontinuity) and the critical load is decreasing with the label of the initialization.

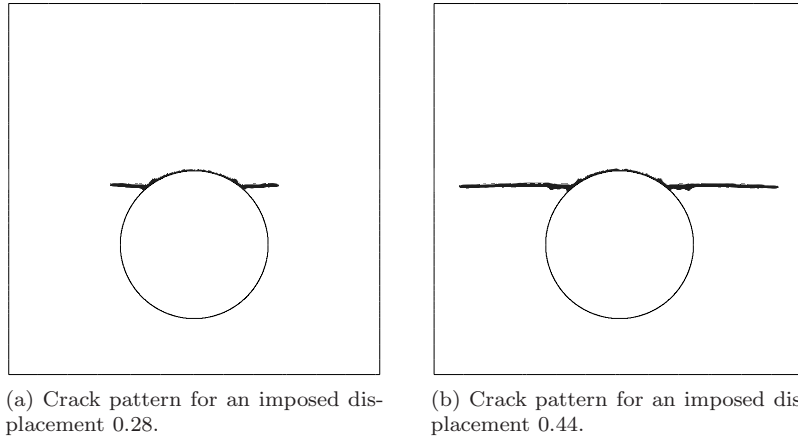


Figure 17: Fiber-reinforced matrix: two crack patterns for the 5th initialization at the beginning and at the end of a smooth evolution.

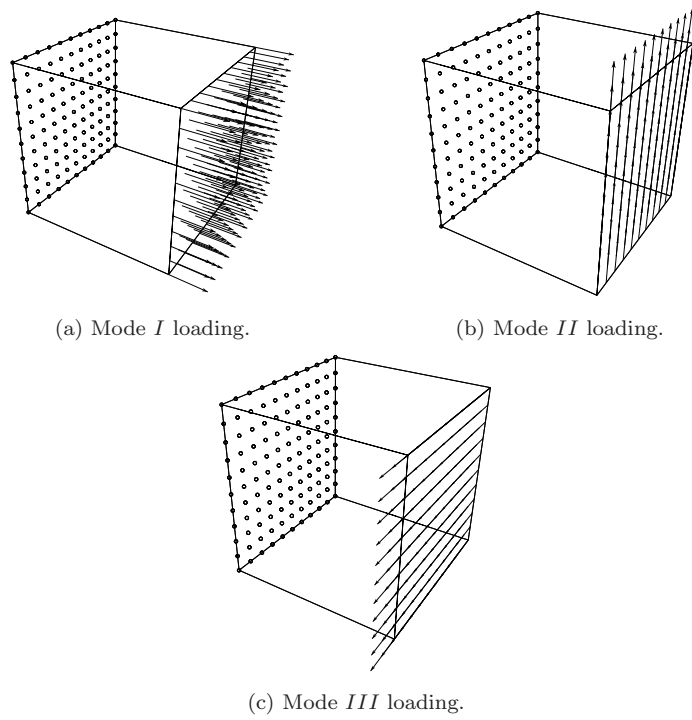


Figure 18: Boundary conditions for the modes *I*, *II* and *III* in $3d$.

we use a sparse parallel direct linear solver for solving the elasticity system, requiring $40GB$ of memory and 15 minutes on 8 Intel Xeon processors. Each of these $3d$ computations requires of the order of 150 and 200 iterations, i.e. solutions of the elasticity system, so the overall CPU time is about two days.

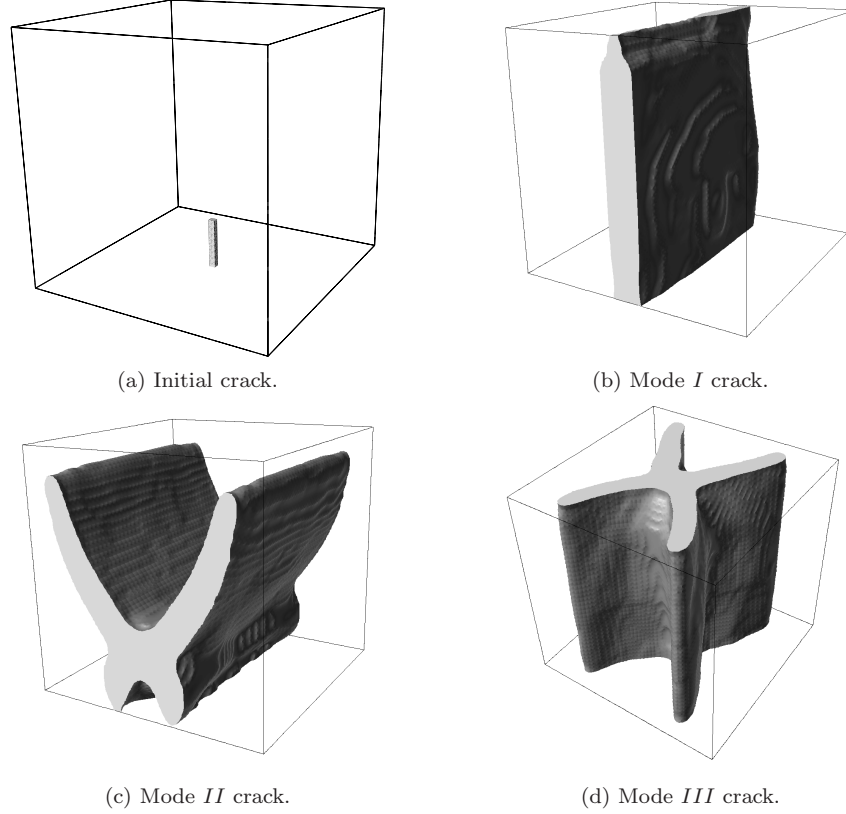


Figure 19: Initial and final cracks for the modes I , II and III in $3d$.

7 Concluding remarks

We have proposed a numerical method, based on the Francfort-Marigo damage model and using a single level set function with standard finite elements, for the simulation of damage evolution and crack propagation. Our method is computing a class of local minimizers of the Francfort-Marigo energy. The crack paths predicted by our numerical experiments are in excellent qualitative agreement with previous results in the literature. However, the quasi-static evolutions of the elastic energy, Griffith-type damage energy and total cost function are different from those we found in [22] (using global minimization), and in [51] (using

criticality). First, cracks often break the structure in a single time increment: fracture is thus a brutal process in most of our computations, which is not the case in [22], [51]. Second, as opposed to the global minimization approach of [22], the cost function (which is the sum of the elastic and damage energies) is not continuous in time: at the critical load (when the structure breaks down) some energy is thus dissipated. This phenomenon seems to be featured by all local minimization approaches and is usually interpreted as the necessity of including kinetic effects in the model.

The ill-posed character of the minimization problem (2.4) or (2.11) (which do not admit minimizers, in general) manifests itself in various aspects. First, as already said, we rarely found a local minimizer in our numerical tests unless the structure was broken (fracture is a brutal process). The only exception is the fiber reinforced matrix test of Section 6.6 where a larger than usual initial crack allows us to obtain a smooth evolution of the crack, similar to that obtained in [22]. Second, our numerical results are quite sensitive to some implementation issues. For example, it is necessary to use the complete shape derivative formulas (3.9) (which features the two phases) and not their simpler limit (3.13), obtained when the damaged phase is assumed to have zero rigidity, otherwise the minimization of the cost function is less complete and the values of the initiation or critical loads may be wrong. Another important issue is the reinitialization process which must be precise enough so that the normal and tangent vectors to the interface between the two phases are always accurately computed while the interface itself does not move at all during reinitialization (otherwise it would contradict the irreversibility constraint).

An interesting open problem is to prove the conjectured convergence of the discrete scaled energy (2.12) towards the fracture model (2.13). A natural extension of our work is to handle a non-interpenetration condition so that cracks under compression do not propagate. We have used standard $Q1$ finite elements for solving the linear elasticity system which features a large variation of the Young modulus between the two phases. It would be interesting to study if extended finite element methods (XFEM, see e.g. [39], [40]) would improve the numerical precision at a not too large expense in CPU cost. Of course, we should also perform more realistic test cases and make precise comparisons with both physical experiments and other codes, including a study of CPU cost. Eventually let us mention that shape optimization for minimizing the risk of crack propagation is also a promising field to investigate, following [48].

A Computation of the shape derivative

This appendix is devoted to the proofs of Lemma 3.6 and Corollary 3.1. We begin with Lemma 3.6 which furnishes the partial shape derivative of the Lagrangian. To prove it we use Lemma 3.1. On the one hand, the derivatives of integrals on $\Omega^{0,1}$ are simple. On the other hand, the interface Σ is either a closed surface without boundary or a surface which meets the outer boundary $\partial\Omega$: in both cases the derivative of an integral on Σ has no contribution on

its boundary $\partial\Sigma$ as in (3.1). Eventually, even if the surfaces $\Gamma_N^{0,1}$ are subset of the fixed boundary Γ_N , they can vary tangentially to Γ_N , so the derivatives of integrals on $\Gamma_N^{0,1}$ are made of the sole boundary term on $\gamma = \partial\Gamma_N^{0,1}$ in (3.1). Therefore, the partial shape derivative of $\mathcal{L}$ is

$$\begin{aligned}
\langle \frac{\partial \mathcal{L}}{\partial \chi}, \phi \rangle &= \int_{\Sigma} (\sigma^1(u^1) \cdot e(p^1) - \sigma^0(u^0) \cdot e(p^0)) \theta \cdot ndS \\
&+ \int_{\Sigma} (-f^1 \cdot p^1 + f^0 \cdot p^0 - j^1(u^1) + j^0(u^0)) \theta \cdot ndS \\
&+ \frac{1}{2} \int_{\Sigma} (\sigma^0(u^0) \cdot e(u^0) - \sigma^1(u^1) \cdot e(u^1)) \theta \cdot ndS \\
&- \frac{1}{2} \int_{\Sigma} \left(\frac{\partial}{\partial n} + H \right) \left((\sigma^1(u^1)n + \sigma^0(u^0)n) \cdot (p^1 - p^0) \right) \theta \cdot ndS \\
&- \frac{1}{2} \int_{\Sigma} \left(\frac{\partial}{\partial n} + H \right) \left((\sigma^1(p^1)n + \sigma^0(p^0)n) \cdot (u^1 - u^0) \right) \theta \cdot ndS \\
&+ \frac{1}{2} \int_{\Sigma} \left(\frac{\partial}{\partial n} + H \right) \left((\sigma^1(u^1)n + \sigma^0(u^0)n) \cdot (u^1 - u^0) \right) \theta \cdot ndS \\
&+ \int_{\gamma} (g^0 \cdot p^0 - g^1 \cdot p^1 + h^0(u^0) - h^1(u^1)) \theta \cdot \tau dL, \tag{A.1}
\end{aligned}$$

where H denotes the mean curvature and τ is the external unit vector normal to $\gamma = \partial\Gamma_N$ and n . Since $u^0 = u^1$ and $p^0 = p^1$ on Σ , the terms involving the curvature vanish on Σ . Similarly the normal component of the stress tensors are continuous through Σ . Thus, (A.1) simplifies in

$$\begin{aligned}
\langle \frac{\partial \mathcal{L}}{\partial \chi}, \phi \rangle &= \int_{\Sigma} (\sigma^1(u^1) \cdot e(p^1) - \sigma^0(u^0) \cdot e(p^0)) \theta \cdot ndS \\
&+ \int_{\Sigma} \frac{1}{2} (\sigma^0(u^0) \cdot e(u^0) - \sigma^1(u^1) \cdot e(u^1)) \theta \cdot ndS \\
&- \int_{\Sigma} \sigma(u)n \cdot \frac{\partial(p^1 - p^0)}{\partial n} \theta \cdot ndS \\
&- \int_{\Sigma} \sigma(p)n \cdot \frac{\partial(u^1 - u^0)}{\partial n} \theta \cdot ndS \\
&+ \int_{\Sigma} \sigma(u)n \cdot \frac{\partial(u^1 - u^0)}{\partial n} \theta \cdot ndS \\
&- \int_{\Sigma} ((f^1 - f^0) \cdot p + j^1(u) - j^0(u)) \theta \cdot ndS \\
&- \int_{\gamma} ((g^1 - g^0) \cdot p + h^1(u) - h^0(u)) \theta \cdot \tau dL, \tag{A.2}
\end{aligned}$$

where $u, p, \sigma(u)n, \sigma(p)n$ denotes the continuous quantities at the interface. The two last lines of (A.2) are expressed only in terms of continuous quantities through the interface, but not the five first lines that we must rewrite, using Lemma 3.2, as an explicit expression in terms of continuous functions at the

interface and jumps of the Lamé coefficients. (In the following computations, the symbol $\cdot$ will denote, according to the context, either a scalar product between two vectors, or between two matrices.)

Let us compute the integrand in the first term of the right hand side of (A.2). In the local orthonormal basis (t, n) (adapted to the interface Σ and introduced in Lemma 3.2) the following decomposition holds

$$\sigma^1(u^1) \cdot e(p^1) = \sigma_{nn}^1(u^1) e_{nn}(p^1) + 2\sigma_{tn}^1(u^1) \cdot e_{tn}(p^1) + \sigma_{tt}^1(u^1) \cdot e_{tt}(p^1).$$

From Lemma 3.2 it rewrites as

$$\begin{aligned} \sigma^1(u^1) \cdot e(p^1) &= \frac{1}{\lambda^1 + 2\mu^1} \sigma_{nn}^1(u^1) (\sigma_{nn}^1(p^1) - \lambda^1 \operatorname{tre}_{tt}(p^1)) + \frac{1}{\mu^1} \sigma_{tn}^1(u^1) \cdot \sigma_{tn}^1(p^1) \\ &\quad + \left[2\mu^1 e_{tt}(u^1) + \frac{\lambda^1}{\lambda^1 + 2\mu^1} (2\mu^1 \operatorname{tre}_{tt}(u^1) + \sigma_{nn}^1(u^1)) I_2^{d-1} \right] \cdot e_{tt}(p^1) \\ &= \frac{1}{\lambda^1 + 2\mu^1} \sigma_{nn}^1(u^1) \sigma_{nn}^1(p^1) + \frac{1}{\mu^1} \sigma_{tn}^1(u^1) \cdot \sigma_{tn}^1(p^1) \\ &\quad + 2\mu^1 e_{tt}(u^1) \cdot e_{tt}(p^1) + \frac{2\lambda^1 \mu^1}{\lambda^1 + 2\mu^1} \operatorname{tre}_{tt}(u^1) \operatorname{tre}_{tt}(p^1). \end{aligned}$$

A similar computation with the index 0 instead of 1 yields the difference

$$\begin{aligned} \sigma^1(u^1) \cdot e(p^1) - \sigma^0(u^0) \cdot e(p^0) &= \left[\frac{1}{\lambda + 2\mu} \right] \sigma_{nn}(u) \sigma_{nn}(p) \\ &\quad + \left[\frac{1}{\mu} \right] \sigma_{tn}(u) \cdot \sigma_{tn}(p) + [2\mu] e_{tt}(u) \cdot e_{tt}(p) + \left[\frac{2\lambda\mu}{\lambda + 2\mu} \right] \operatorname{tre}_{tt}(u) \operatorname{tre}_{tt}(p) \end{aligned}$$

which is expressed, as desired, only in terms of continuous functions at the interface. On the same token we deduce

$$\begin{aligned} \sigma^0(u^0) \cdot e(u^0) - \sigma^1(u^1) \cdot e(u^1) &= \left[\frac{-1}{\lambda + 2\mu} \right] (\sigma_{nn}(u))^2 \\ &\quad - \left[\frac{1}{\mu} \right] |\sigma_{tn}(u)|^2 - [2\mu] |e_{tt}(u)|^2 - \left[\frac{2\lambda\mu}{\lambda + 2\mu} \right] (\operatorname{tre}_{tt}(u))^2. \end{aligned}$$

We now consider the integrand of the third, fourth or fifth line of (A.2). We use the following identity for two displacements v and q

$$\text{if } q = 0 \text{ on } \Sigma, \text{ then } \sigma(v)n \cdot \frac{\partial q}{\partial n} = 2(\sigma(v)n) \cdot (e(q)n) - \sigma_{nn}(v) e_{nn}(q) \text{ on } \Sigma.$$

We obtain

$$\begin{aligned}
\sigma(u)n \cdot \frac{\partial(p^1 - p^0)}{\partial n} &= 2(\sigma(u)n) \cdot (e(p^1)n - e(p^0)n) - \sigma_{nn}(u) (e_{nn}(p^1) - e_{nn}(p^0)) \\
&= [\frac{2}{\lambda + 2\mu}] \sigma_{nn}(u) \sigma_{nn}(p) \\
&\quad - [\frac{2\lambda}{\lambda + 2\mu}] \sigma_{nn}(u) \operatorname{tre}_{tt}(p) + [\frac{1}{\mu}] \sigma_{tn}(u) \cdot \sigma_{tn}(p) \\
&\quad - [\frac{1}{\lambda + 2\mu}] \sigma_{nn}(u) \sigma_{nn}(p) + [\frac{\lambda}{\lambda + 2\mu}] \sigma_{nn}(u) \operatorname{tre}_{tt}(p) \\
&= [\frac{1}{\lambda + 2\mu}] \sigma_{nn}(u) \sigma_{nn}(p) + [\frac{1}{\mu}] \sigma_{tn}(u) \cdot \sigma_{tn}(p) \\
&\quad - [\frac{\lambda}{\lambda + 2\mu}] \sigma_{nn}(u) \operatorname{tre}_{tt}(p). \tag{A.3}
\end{aligned}$$

We get a similar expression for $\sigma(p)n \cdot \frac{\partial(u^1 - u^0)}{\partial n}$ and $\sigma(u)n \cdot \frac{\partial(u^1 - u^0)}{\partial n}$. Summing up these contributions we deduce that the integrand of the five first lines of the right hand side of (A.2) is

$$\begin{aligned}
\mathcal{D}(x) = & - [\frac{1}{\lambda + 2\mu}] \sigma_{nn}(u) \sigma_{nn}(p) - [\frac{1}{\mu}] \sigma_{tn}(u) \cdot \sigma_{tn}(p) + [2\mu] e_{tt}(u) \cdot e_{tt}(p) \\
& + [\frac{2\lambda\mu}{\lambda + 2\mu}] \operatorname{tre}_{tt}(u) \operatorname{tre}_{tt}(p) + [\frac{\lambda}{\lambda + 2\mu}] (\sigma_{nn}(u) \operatorname{tre}_{tt}(p) + \sigma_{nn}(p) \operatorname{tre}_{tt}(u)) \\
& + [\frac{1}{2(\lambda + 2\mu)}] (\sigma_{nn}(u))^2 + [\frac{1}{2\mu}] |\sigma_{tn}(u)|^2 - [\mu] |e_{tt}(u)|^2 \\
& - [\frac{\lambda\mu}{\lambda + 2\mu}] (\operatorname{tre}_{tt}(u))^2 - [\frac{\lambda}{\lambda + 2\mu}] \sigma_{nn}(u) \operatorname{tre}_{tt}(u)
\end{aligned}$$

which is precisely formula (3.10). Using the relations (3.8) between e and σ (see Lemma 3.2) we easily deduce (3.11) from (3.10), which finishes the proof of Lemma 3.6. $\square$

Proof of Corollary 3.1. The Francfort-Marigo objective function is obtained for $j^k(u) = -f^k \cdot u + \kappa \delta_{k0}$ and $h^k(u) = -g^k \cdot u$. We deduce that $j'_k(u) = -f^k$ and $h'_k(u) = -g^k$, and thus that the adjoint state vanishes $p_\chi = 0$. If we further assume that $f^0 = f^1$ and $g^0 = g^1$, the shape derivative reduces to

$$J'(\chi)(\theta) = \int_{\Sigma} \mathcal{D}(x) \theta \cdot n \, dS$$

with

$$\begin{aligned}
\mathcal{D}(x) &= [\frac{1}{2(\lambda + 2\mu)}] (\sigma_{nn}(u_\chi))^2 + [\frac{1}{2\mu}] |\sigma_{tn}(u_\chi)|^2 - [\mu] |e_{tt}(u_\chi)|^2 \\
&\quad - [\frac{\lambda\mu}{(\lambda + 2\mu)}] (\operatorname{tre}_{tt}(u_\chi))^2 - [\frac{\lambda}{(\lambda + 2\mu)}] \sigma_{nn}(u_\chi) \operatorname{tre}_{tt}(u_\chi). \tag{A.4}
\end{aligned}$$

Assumption $A^0 \leq A^1$ is equivalent to $[\mu] \geq 0$ and $[\varkappa] \geq 0$ with $\varkappa := \lambda + \frac{2}{d}\mu$, the bulk modulus. Since $\sigma_{tn}(u)$ only appears as a square product in (A.4), it suffices to check that the combination of all other terms is indeed negative, that is,

$$\begin{aligned} \left[\frac{1}{\lambda + 2\mu}\right](\sigma_{nn}(u))^2 - [2\mu]|e_{tt}(u)|^2 &= \left[\frac{2\lambda\mu}{\lambda + 2\mu}\right](\text{tre}_{tt}(u))^2 \\ &- \left[\frac{2\lambda}{\lambda + 2\mu}\right]\sigma_{nn}(u) \text{tre}_{tt}(u) \leq 0. \end{aligned} \quad (\text{A.5})$$

Since $(\text{tre}_{tt}(u))^2 \leq (d-1)|e_{tt}(u)|^2$ (where $d = 2, 3$ is the space dimension) the left hand side of (A.5) is bounded from above by

$$\left[\frac{1}{\lambda + 2\mu}\right]\sigma_{nn}(u)^2 - \left[\frac{2\mu}{d-1} + \frac{2\lambda\mu}{\lambda + 2\mu}\right](\text{tre}_{tt}(u))^2 - \left[\frac{2\lambda}{\lambda + 2\mu}\right]\sigma_{nn}(u) \text{tre}_{tt}(u), \quad (\text{A.6})$$

which writes in term of $\varkappa$ and μ as

$$\begin{aligned} &\left[\frac{d}{d\varkappa + 2(d-1)\mu}\right]\sigma_{nn}(u)^2 - \frac{d}{d-1}\left[\frac{2d\varkappa\mu}{d\varkappa + 2(d-1)\mu}\right]\text{tr}^2 e_{tt}(u) \\ &- 2\left[\frac{d\varkappa - 2\mu}{d\varkappa + 2(d-1)\mu}\right]\sigma_{nn}(u) \text{tre}_{tt}(u). \end{aligned}$$

Since, by assumption $A^0 \leq A^1$, we have

$$\left[\frac{d}{d\varkappa + 2(d-1)\mu}\right] \leq 0 \quad \text{and} \quad \left[\frac{2d\varkappa\mu}{d\varkappa + 2(d-1)\mu}\right] \geq 0,$$

the quadratic form (A.6) is negative if and only if

$$\left[\frac{d\varkappa - 2\mu}{d\varkappa + 2(d-1)\mu}\right]^2 \leq -\frac{d}{d-1}\left[\frac{d}{d\varkappa + 2(d-1)\mu}\right]\left[\frac{2d\varkappa\mu}{d\varkappa + 2(d-1)\mu}\right]. \quad (\text{A.7})$$

Introducing the new variables $\varkappa' = \frac{d\varkappa}{2}$ and $\mu' = (d-1)\mu$, (A.7) is equivalent to

$$\left[\frac{1}{\varkappa' + \mu'}\right]\left[\frac{\varkappa'\mu'}{\varkappa' + \mu'}\right] \leq -\frac{1}{d^2}\left[\frac{(d-1)\varkappa' - \mu'}{\varkappa' + \mu'}\right]^2 = -\frac{1}{4}\left[\frac{\varkappa' - \mu'}{\varkappa' + \mu'}\right]^2. \quad (\text{A.8})$$

A brute force computation shows that (A.8) holds true. $\square$

Acknowledgments. This work has been supported by the MULTIMAT european network MRTN-CT-2004-505226 funded by the EEC. G. Allaire is a member of the DEFI project at INRIA Futurs Saclay and is partially supported by the Chair "Mathematical modelling and numerical simulation, F-EADS - Ecole Polytechnique - INRIA" and by the GNR MoMaS sponsored by ANDRA, BRGM, CEA, CNRS, EDF, and IRSN. G. Allaire and F. Jouve were partly supported by the ANR project MICA, ANR-06-BLAN-0082. N. Van Goethem was partly supported by the Fundação para a Ciência e a Tecnologia (FCT Project: PTDC/EME-PME/108751/2008). The authors would like to thank the reviewers for their constructive comments.

References

- [1] Alfaiate, J., Wells, G.N., Sluys, L.J., On the use of embedded discontinuity elements with crack path continuity for mode-I and mixed-mode fracture, *Engineering Fracture Mechanics* **69**, pp.661-686, (2002).
- [2] Allaire, G., Continuity of the Darcy's law in the low-volume fraction limit, *Ann. Scuola Norm. Sup. Pisa*, **18**, pp.475-499 (1991).
- [3] Allaire, G., *Conception optimale de structures*, Collection: Mathématiques et Applications, Vol. **58**, Springer (2007).
- [4] Allaire, G., Aubry, S., Jouve, F., Simulation numérique de l'endommagement à l'aide du modèle Francfort-Marigo, in Actes du 29ème congrès d'analyse numérique, *ESAIM Proceedings*, **3**, pp.1-9 (1998).
- [5] Allaire, G., Jouve, F., Toader, A.-M., Structural optimization using sensitivity analysis and a level-set method, *J. Comput. Phys.* **194** (1), pp.363-393, (2004).
- [6] Allaire, G., Jouve, F., de Gournay, F., Toader, A.-M., Structural optimization using topological and shape sensitivity via a level set method, *Control and Cybernetics* **34**, pp.59-80, (2005).
- [7] Allaire, G., Jouve, F., Van Goethem, N.: *A level set method for the numerical simulation of damage evolution*, Proceedings of ICIAM 2007 Zürich, R. Jeltsch and G. Wanner eds., pp.3-22, EMS, Zürich (2009).
- [8] Ambrosio, L., Tortorelli, V., Approximation of functionals depending on jumps by elliptic functionals via Γ -convergence, *Comm. Pure Appl. Math.* **43**, no. 8, pp.999-1036 (1990).
- [9] Ammari, H., Kang, H., *Reconstruction of small inhomogeneities from boundary measurements*, Lecture notes in mathematics, 1846. Springer, (2000).
- [10] Ammari, H., Kang, H., Nakamura, G., Tanuma, K., Complete asymptotic expansions of solutions of the system of elastostatics in the presence of an inclusion of small diameter and detection of an inclusion, *J. Elasticity* **67**, no. 2, pp.97-129 (2002).
- [11] Ammari, H., *An introduction to mathematics of emerging biomedical imaging*, Collection: Mathématiques et Applications, Vol. **62**, Springer (2008).
- [12] Ammari, H., Calmon, P., Iakovleva, E., Direct elastic imaging of a small inclusion, *SIAM Journal on Imaging Sciences* **1**, pp.169-187, (2008).
- [13] Ambrosio, L., Buttazzo, G., An optimal design problem with perimeter penalization, *Calc. Var.* **1**, pp.55-69 (1993).

- [14] Areias, P., Alfaiate, J., Dias-da-Costa, D., Julio, E.", Arbitrary bi-dimensional finite strain crack propagation, *Computational Mechanics* **45**, no. 1, pp.61-75 (2009).
- [15] Bernard, P.-E., Moës, N., Stolz, C., Chevaugeon, N., A thick level set approach to model evolving damage and transition to fracture, *Proceedings of ECCM 2010*, Paris, May 16-21 2010.
- [16] Bernardi, Ch., Pironneau, O., Sensitivity of Darcy's law to discontinuities. *Chinese Ann. Math. Ser. B* **24**, no. 2, pp.205-214 (2003).
- [17] Bittencourt, T.N., Ingraffea, A.R., Wawrzynek, P.A., Sousa, J.L., Quasi-automatic simulation of crack propagation for 2D LEFM problems, *Engineering Fracture Mechanics* **55**, no. 2, pp.321-324 (1996).
- [18] Bonnet, M., Constantinescu, A., Inverse problems in elasticity, *Inverse Problems* **21**, pp.1-50 (2005).
- [19] Bordas, S., Moran, B., Enriched finite elements and level sets for damage tolerance assessment of complex structures, *Engineering Fracture Mechanics* **73** (9), pp.1176-1201 (2006).
- [20] Bourdin, B., Numerical implementation of the variational formulation for quasi-static brittle fracture, *Interfaces and free boundaries*, **9**, pp.411-430 (2007).
- [21] Bourdin, B., Francfort, G. A., Marigo, J.-J., Numerical experiments in revisited brittle fracture, *J. Mech. Phys. Solids* **48**, no. 4, pp.797-826 (2000).
- [22] Bourdin, B., Francfort, G. A., Marigo, J.-J., The variational approach to fracture, *Journal of Elasticity*, **91**, pp.5148 (2008).
- [23] Braides, A., Larsen, C., Γ -convergence for stable states and local minimizers, preprint.
- [24] Braides, A., Piatnitski, A., Homogenization of surface and length energies for spin systems, preprint, available at <http://cvgmt.sns.it/papers/brapia10/>.
- [25] Bui, H.D., *Mécanique de la rupture fragile*, Masson, Paris (1983).
- [26] Burger, M., Hackl, B., Ring, W., Incorporating topological derivatives into level set methods, *J. Comp. Phys.*, **194** 1, pp.344-362 (2004).
- [27] Capdeboscq, Y., Vogelius, Michael S., Pointwise polarization tensor bounds, and applications to voltage perturbations caused by thin inhomogeneities. *Asymptot. Anal.* **50**, no. 3-4, pp.175-204 (2006).
- [28] Cea, J., Conception optimale ou identification de formes, calcul rapide de la dérivée de la fonction coût, *Math. Model. Numer. Anal.* **20** (3), pp.371-402, (1986).

- [29] Chambolle, A., Image segmentation by variational methods: Mumford and Shah functional and the discrete approximations, *SIAM J. Appl. Math.* **55**, pp.827-863 (1995).
- [30] Chambolle, A., Giacomini, A., Ponsiglione, M., Crack initiation in brittle materials, *Arch. Rat. Mech. Anal.*, **188**, pp.309-349 (2008).
- [31] Dautray, R., Lions, J.-L., *Mathematical analysis and Numerical methods for science and technology*, Vol. 2, Springer Verlag, (1985).
- [32] Deny, J., Lions, J.-L. Les espaces du type de Beppo Levi. *Ann. Inst. Fourier*, Grenoble 5 (1953–54), pp.305-370 (1955).
- [33] Eschenauer, H., Kobelev, V., Schumacher, A., Bubble method for topology and shape optimization of structures, *Structural Optimization*, **8**, pp.42-51 (1994).
- [34] Francfort, G., Garroni, A., A variational view of brittle damage evolution, *Arch. Ration. Mech. Anal.* **182**, no. 1, pp.125-152 (2006).
- [35] Francfort, G., Marigo, J.-J., Stable damage evolution in a brittle continuous medium, *Eur. J. Mech., A/Solids*, **12** (2), pp.149-189 (1993).
- [36] Francfort, G., Marigo, J.-J., Revisiting brittle fracture as an energy minimization problem, *J. Mech. Phys. Solids*, **46** (8), pp.1319-1342 (1998).
- [37] Garreau, S., Guillaume, P., Masmoudi, M., The topological asymptotic for PDE systems: the elasticity case. *SIAM J. Control Optim.* **39**, no. 6, pp.1756-1778 (2001).
- [38] Garroni, A., Larsen, C., Threshold-based Quasi-static Brittle Damage Evolution, *Arch. Ration. Mech. Anal.* **194**, no. 2, pp.585-609 (2009).
- [39] Gravouil, A., Moës, N., Belytschko, T., Non-planar 3D crack growth by the extended finite element and level sets - Part I: Mechanical model. *Int. J. Numer. Meth. Engng.*, **53**, no. 11, pp.2549-2568 (2002).
- [40] Gravouil, A., Moës, N., Belytschko, T., Non-planar 3D crack growth by the extended finite element and level sets - Part II: Level set update. *Int. J. Numer. Meth. Engng.*, **53**, no. 11, pp.2569-2586 (2002).
- [41] Henrot, A., Pierre, M., *Variation et optimisation de formes*, Mathématiques et Applications **48**, Springer Verlag, (2005)
- [42] Hettlich, F., Rundell, W., The determination of a discontinuity in a conductivity from a single boundary measurement. *Inverse Problems*, **14**, no. 1, pp.67-82 (1998).
- [43] Jerrard, R., Sternberg, P., Critical points via Gamma-convergence: general theory and applications, *Jour. Eur. Math. Soc.* **11**, pp.705-753 (2009).

- [44] Larsen, C., Epsilon-stable quasi-static brittle fracture evolution, *Comm. Pure Appl. Math.*, to appear.
- [45] Larsen, C., Richardson, C., Sarkis, M., A level set method for the Mumford-Shah functional and fracture, *SIAM J. Imag. Sci.*, to appear.
- [46] Leblond, J.-B., *Mécanique de la rupture fragile et ductile*, Hermes Science Publications, Paris (2003).
- [47] Moës, N., Dolbow, J., Belytschko, T., A finite element method for crack growth without remeshing, *Int. J. Numer. Meth. Engng.*, **46**, no. 1, pp.131-150 (1999).
- [48] Munch, A., Pedregal, P., Relaxation of an optimal design problem in fracture mechanics, to appear in ESAIM:COCV (2010).
- [49] Murat F., Simon S., Etudes de problèmes d'optimal design. *Lecture Notes in Computer Science* **41**, pp.54-62, Springer Verlag, Berlin (1976).
- [50] Negri, M., A finite element approximation of the Griffith's model in fracture mechanics, *Numer. Math.* **95**, pp.653-687 (2003).
- [51] Negri, M., Ortner, Ch., Quasi-static crack propagation by Griffith's criterion, *M3AS*, **18**, pp.1895-1925 (2008).
- [52] Osher S., Fedkiw R., *Level set methods and dynamic implicit surfaces*, Applied Mathematical Sciences, **153**, Springer-Verlag, New York (2003).
- [53] Osher, S., Sethian, J.A., Front propagating with curvature dependent speed: algorithms based on Hamilton-Jacobi formulations. *J. Comp. Phys.* **78**, pp.12-49 (1988).
- [54] Pantz, O., Sensibilité de l'équation de la chaleur aux sauts de conductivité, *C. R. Acad. Sci. Paris, Ser. I* **341**, pp.333-337 (2005).
- [55] Pironneau, O., *Optimal shape design for elliptic systems*, Springer-Verlag, New York (1984).
- [56] Pradeilles-Duval, R.-M., Stolz, C., Mechanical transformations and discontinuities along a moving surface, *J. Mech. Phys. Solids* **43**, no. 1, pp.91-121 (1995).
- [57] Russo, G., Smereka, P., A remark on computing distance functions, *J. Comp. Phys.* **163**, pp.51-67 (2000).
- [58] Sandier, E., Serfaty, S., Gamma-Convergence of Gradient Flows and Application to Ginzburg-Landau, *Comm. Pure Appl. Math.* **57**, pp.1627-1672 (2004).
- [59] Schmidt, B., Fraternali, F., Ortiz, M., Eigenfracture: an eigendeformation approach to variational fracture, *Multiscale Model. Simul.*, **7**, no. 1, pp.1237-1266 (2009).

- [60] Sethian J.A., *Level set Methods and fast marching methods: evolving interfaces in computational geometry, fluid mechanics, computer vision and materials science*, Cambridge University Press (1999).
- [61] Sokołowski, J., Żochowski, A., On the topological derivative in shape optimization, *SIAM J. Control Optim.*, **37**, pp.1251-1272 (1999).
- [62] Sokołowski J., Zolesio J.P., *Introduction to shape optimization: shape sensitivity analysis*, *Springer Series in Computational Mathematics, Vol. 10*, Springer, Berlin (1992).
- [63] Stolz, C., *Energy Methods in Non-Linear Mechanics*, Lecture Notes 11, IPPT PAN and CoE AMAS, Warsaw, 108 pages, ISSN 1642-0578, (2004).
- [64] Van Goethem, N., Novotny, A., Crack nucleation sensitivity analysis *Math. Meth. Appl. Sc.*, **33**, no. 16, (2010).
- [65] Wang, M.Y., Wang, X., Guo, D., A level set method for structural topology optimization, *Comput. Methods Appl. Mech. Engrg.*, **192**, pp.227-246 (2003).
- [66] Wang, X., Yulin, M., Wang, M.Y., Incorporating topological derivatives into level set methods for structural topology optimization, *in Optimal shape design and modeling*, T. Lewinski et al. eds., pp.145-157, Polish Academy of Sciences, Warsaw (2004).