

Selection Mechanisms for Microstructures and Reversible Martensitic Transformations



Francesco M. G. Della Porta
Keble College
University of Oxford

A thesis submitted for the degree of
Doctor of Philosophy
Trinity 2018

Acknowledgements

Thanks to John Ball, for leading me without pushing me through my research. The inestimable freedom he has granted me, was always supported by a strong quiet guidance and precious suggestions. It has been a pleasure to see how he approaches problems.

I am profoundly indebted also to Sherry Chen for inviting me to spend some time in her lab, and giving me the possibility to observe martensitic transformations with my eyes. Also, I would like to thank Sherry for the time spent in useful discussions, which have inspired Chapter 2.

I thank Richard James and Jan Kristensen for reading my thesis, assessing my work, and providing me with constructive feedback.

I would also like to acknowledge Tomonari Inamura, Richard James, Konstantinos Koumatos and Angkana Rüland for the useful discussions and suggestions, as much as Jan Kristensen for carefully supervising my second mini-project during my first year.

Thanks also to Xian Chen, Yintao Song and Richard James for kindly allowing me to include Figure 2.2 in my thesis.

I acknowledge Maurizio Grasselli for giving my ego the necessary boost to undertake this PhD.

This work would not exist without the support of the Engineering and Physical Sciences Research Council.

I am profoundly grateful to my officemates Luca and Tabea, for saving me from madness, for sharing good and bad moods, jokes and, also, food. Our office has been full of discussions, maths related and non. Thanks also to Giacomo Canevari, his brilliant mind and infinite patience have helped me many times during his stay in Oxford. Thanks also to my HK-officemates Chenbo, Zhuo Hui, Mostafa and Gang-peng for treating me as family, and to all other people who have helped me to discover Hong Kong. I would like to acknowledge my friends from Milano who tried to keep my interests wide, especially Saret, Lussy, Cochi, Tai, Lore and Snannolo who came

to visit. Thanks also to all the people I have met in Oxford, and who have made my staying in Oxford so cosy. Among all, thanks to Adam, Cecilia, Chris, Emile, Guus, Hamish, Michael, Peter and Valerio.

I would like to thank my parents and my brothers, for their unconditional support, and for pointing me always in the direction of happiness. Thanks to Ludovica because climbing any peak with her is a piece of cake.

Abstract

The work in this thesis is inspired by the fabrication of $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$. This is the first alloy undergoing ultra-reversible martensitic transformations and closely satisfying the cofactor conditions, particular conditions of geometric compatibility between phases, which were conjectured to influence reversibility. With the aim of better understanding reversibility, in this thesis we study the martensitic microstructures arising during thermal cycling in $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$, which are complex and different in every phase transformation cycle.

Our study is developed in the context of continuum mechanics and nonlinear elasticity, and we use tools from nonlinear analysis.

The first aim of this thesis is to advance our understanding of conditions of geometric compatibility between phases. To this end, first, we further investigate cofactor conditions and introduce a physically-based metric to measure how closely these are satisfied in real materials. Secondly, we introduce further conditions of compatibility and show that these are nearly satisfied by some twins in $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$. These might influence reversibility as they improve compatibility between high and low temperature phases.

Martensitic phase transitions in $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$ are a complex phenomenon, especially because the crystalline structure of the material changes from a cubic to a monoclinic symmetry, and hence the energy of the system has twelve wells. There exist infinitely many energy-minimising microstructures, limiting our understanding of the phenomenon as well as our ability to predict it. Therefore, the second aim of this thesis is to find criteria to select physically-relevant energy minimisers. We introduce two criteria or selection mechanisms. The first involves a moving mask approximation, which allows one to describe some experimental observations on the dynamics, while the second is based on using vanishing interface energy. The moving mask approximation reflects the idea of a moving curtain covering and uncovering microstructures during the phase transition, as appears to be the case for $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$, and many

other materials during thermally induced transformations. We show that the moving mask approximation can be framed in the context of a model for the dynamics of nonlinear elastic bodies. We prove that every macroscopic deformation gradient satisfying the moving mask approximation must be of the form

$$\mathbf{1} + \mathbf{a}(\mathbf{x}) \otimes \mathbf{n}(\mathbf{x}), \quad \text{a.e. } \mathbf{x}.$$

With regards to vanishing interface energy, we consider a one-dimensional energy functional with three wells, which simplifies the physically relevant model for martensitic transformations, but at the same time highlights some key issues. Our energy functional admits infinitely many minimising gradient Young measures, representing energy-minimising microstructures. In order to select the physically relevant ones, we show that minimisers of a regularised energy, where the second derivatives are penalised, generate a unique minimising gradient Young measure as the perturbation vanishes.

The results developed in this thesis are motivated by the study of $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$, but their relevance is not limited to this material. The results on the cofactor conditions developed here can help for the understanding of new alloys undergoing ultra-reversible transformations, and as a guideline for the fabrication of future materials. Furthermore, the selection mechanisms studied in this work can be useful in selecting physically relevant microstructures not only in $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$, but also in other materials undergoing martensitic transformations, and other phenomena where pattern formation is observed.

Contents

1	Introduction	1
1.1	Martensitic transformations	1
1.2	Calculus of variations and quasiconvexity	6
1.3	Nonlinear elasticity model	12
1.3.1	Lack of quasiconvexity and relaxation	15
1.3.2	Hadamard jump conditions	16
1.4	Twinning theory and the cofactor conditions	18
1.5	Overview	23
2	On the cofactor conditions and further conditions of super-compatibility between phases	26
2.1	Martensitic transformations from cubic to monoclinic lattices	30
2.2	Cofactor conditions for two wells	35
2.2.1	Compound twins	36
2.2.2	Type I/II twins	39
2.3	Type I and II star twins	42
2.4	Star twins in cubic to monoclinic transformations	46
2.5	How closely does a material satisfy the cofactor conditions?	57
3	A model for the evolution of highly reversible martensitic transformations	66
3.1	Introduction	66
3.2	Dynamics of martensitic phase transitions in continuum solid mechanics	68
3.3	Hypotheses, boundary conditions and approximations	69
3.4	Existence of weak solutions	72
3.5	Convergence as $k \rightarrow +\infty$ to a limiting constrained theory	83
3.6	Connections with the moving mask assumption	87
3.7	The evolution of a simple laminate in a one-dimensional context . . .	90

3.7.1	An example: the simple laminate	90
3.7.2	Behaviour of solutions to the 1-dimensional problem	93
4	Analysis of a moving mask approximation for martensitic transformations	100
4.1	Introduction	100
4.2	Some preliminaries on k -rectifiable sets	102
4.3	The moving mask assumption	103
4.4	Generalized Hadamard conditions	106
4.5	Basic properties of microstructures	117
4.6	Macroscopic moving interfaces	130
4.7	Experimental evidence	134
5	Vanishing interfacial energy as a selection mechanism for minimising gradient Young measures	136
5.1	Proof of the first Γ -limit	143
5.2	Construction of an upper bound	144
5.3	Construction of a lower bound	152
5.4	The second Γ -limit	170
5.5	Selecting Minimizing sequences without Γ -convergence	173
5.6	Some remarks on the assumptions	176
5.6.1	Two examples	179
	Bibliography	180

List of Figures

1.1	On the left cubic austenite, stable at high temperature; on the right tetragonal martensite, stable at low temperatures.	2
1.2	On the left, cubic austenite, which under a sufficiently large vertical force can have greater free energy than particular martensitic variants, on the right.	3
1.3	Example of shape memory. A curved wire is stretched and deformed into a straight wire. The deformation is apparently plastic as the wire does not recover its curved shape when the stretching force is released. When the sample is heated, it recovers its original shape. The shape remains unchanged when the material is cooled down again.	3
1.4	Representation of shape memory at a microscopic scale.	4
1.5	Example of minimising sequence $\nabla \mathbf{y}_k$ involving both austenite and martensite. Here, $\mathbf{R}_1 \mathbf{U}_1 - \mathbf{R}_2 \mathbf{U}_2 = \mathbf{b} \otimes \mathbf{m}$, and $(1 - \lambda) \mathbf{R}_1 \mathbf{U}_1 + \lambda \mathbf{R}_2 \mathbf{U}_2 = \mathbf{1} + \mathbf{a} \otimes \mathbf{n}$. Therefore compatibility between phases is achieved on average. As k tends to infinity, the grey region, the only one where the energy is different from zero and the material is not stress-free, vanishes.	15
1.6	Example of incompatible deformation gradients (on the left) and compatible deformation gradients (on the right) at a crystalline length-scale.	17
1.7	Example of triple junctions for type I and type II twins when the cofactor conditions are satisfied (see Theorem 1.4.2 and Theorem 1.4.3).	21

- 2.1 Example of average compatibility between three different laminates occurring in type II star twins. Here $\mathbf{R}_1, \dots, \mathbf{R}_6 \in SO(3)$ and $\mathbf{V}_1, \dots, \mathbf{V}_6 \in \mathbb{R}_{Sym+}^{3 \times 3}$ are six deformation gradients corresponding to six martensitic variants (possibly $\mathbf{V}_i = \mathbf{V}_j$ for $i \neq j$). We have $\mathbf{R}_i \mathbf{V}_i = \mathbf{1} + \mathbf{a}_i \otimes \mathbf{n}_\alpha$ for $i = 1, 2$, $\mathbf{R}_i \mathbf{V}_i = \mathbf{1} + \mathbf{a}_i \otimes \mathbf{n}_\beta$ for $i = 3, 4$, and $\mathbf{R}_i \mathbf{V}_i = \mathbf{1} + \mathbf{a}_i \otimes \mathbf{n}_\gamma$ for $i = 5, 6$, for some $\mathbf{a}_1, \dots, \mathbf{a}_6 \in \mathbb{R}^3$, $\mathbf{n}_\alpha, \mathbf{n}_\beta, \mathbf{n}_\gamma \in \mathbb{S}^2$. Furthermore, there exist $\mu \in (0, 1)$, $\mathbf{a}^* \in \mathbb{R}^3$ such that $\mu \mathbf{R}_1 \mathbf{V}_1 + (1 - \mu) \mathbf{R}_2 \mathbf{V}_2 = \mathbf{1} + \mathbf{a}^* \otimes \mathbf{n}_\alpha$, $\mu \mathbf{R}_3 \mathbf{V}_3 + (1 - \mu) \mathbf{R}_4 \mathbf{V}_4 = \mathbf{1} + \mathbf{a}^* \otimes \mathbf{n}_\beta$ and $\mu \mathbf{R}_5 \mathbf{V}_5 + (1 - \mu) \mathbf{R}_6 \mathbf{V}_6 = \mathbf{1} + \mathbf{a}^* \otimes \mathbf{n}_\gamma$, where $\text{rank}(\mathbf{R}_{2i} \mathbf{V}_{2i} - \mathbf{R}_{2i-1} \mathbf{V}_{2i-1}) = 1$ for $i = 1, 2, 3$. In the centre we have a transition layer which makes the two laminates exactly compatible in the limit $\varepsilon \rightarrow 0$. This image is a projection of a plane, but we remark that $\mathbf{n}_\alpha, \mathbf{n}_\beta, \mathbf{n}_\gamma$ are linearly independent. 29
- 2.2 Examples of junctions between three different laminates observed in $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$ (cf. Figure 2.1). Courtesy of Xian Chen, Yintao Song and Richard James. 30
- 2.3 Three different views of a pyramidal microstructure constructed with type II half-star twins (also possible with star twins). The polyhedron is divided into a blue, a yellow and a green region, in each of which we have a different exactly compatible laminate, namely $\mathbf{F}_1 = \mathbf{1} + \mathbf{a}^* \otimes \mathbf{n}_\alpha$, $\mathbf{F}_2 = \mathbf{1} + \mathbf{a}^* \otimes \mathbf{n}_\beta$ and $\mathbf{F}_3 = \mathbf{1} + \mathbf{a}^* \otimes \mathbf{n}_\gamma$, for some $\mathbf{a}^* \in \mathbb{R}^3 \setminus \{\mathbf{0}\}$. The three constant average deformation gradients $\mathbf{F}_1, \mathbf{F}_2, \mathbf{F}_3$ satisfy (2.6), that means that $\mathbf{n}_\alpha, \mathbf{n}_\beta, \mathbf{n}_\gamma$ are linearly independent. At the three faces of the pyramid with normals $\mathbf{n}_\alpha, \mathbf{n}_\beta, \mathbf{n}_\gamma$ the martensitic laminates are compatible with austenite without an interface layer. Such pyramid can be used as a building block for three-dimensional curved interfaces. 47
- 2.4 Two different views of a pyramidal microstructure constructed with type II star twins. The polyhedron is divided into a blue, a yellow, a red and a green region, in each of which we have a different exactly compatible laminate, namely $\mathbf{F}_1 = \mathbf{1} + \mathbf{a}^* \otimes \mathbf{n}_\alpha$, $\mathbf{F}_2 = \mathbf{1} + \mathbf{a}^* \otimes \mathbf{n}_\beta$, $\mathbf{F}_3 = \mathbf{1} + \mathbf{a}^* \otimes \mathbf{n}_\gamma$ and $\mathbf{F}_4 = \mathbf{1} + \mathbf{a}^* \otimes \mathbf{n}_\delta$, for some $\mathbf{a}^* \in \mathbb{R}^3 \setminus \{\mathbf{0}\}$. The four constant average deformation gradients $\mathbf{F}_1, \mathbf{F}_2, \mathbf{F}_3, \mathbf{F}_4$ satisfy (2.6), that means that $\mathbf{n}_\alpha, \mathbf{n}_\beta, \mathbf{n}_\gamma, \mathbf{n}_\delta$ are linearly independent. At the three faces of the pyramid with normals $\mathbf{n}_\alpha, \mathbf{n}_\beta, \mathbf{n}_\gamma, \mathbf{n}_\delta$ the martensitic laminates are compatible with austenite without an interface layer. Such pyramid can be used as a building block for three-dimensional curved interfaces. 47

2.5	Given matrices $\mathbf{U}_i \in \mathbb{R}_{Sym+}^{3 \times 3}$ as in (2.8) describing cubic to monoclinic phase transitions, with eigenvalues $1, \lambda, d$, we plot in cyan and in blue the curves of lattice parameters such that respectively type I star and type II star twins exist. In green and in black the curves of lattice parameters such that respectively type I half-star and type II half-star twins exist. In red the curve of lattice parameters such that both, type I and type II twins exist (see Remark 2.1.1). In magenta and in yellow we plot the curves where $a = c$ (cubic to orthorhombic transformation), and the cofactor conditions are satisfied for type II and type I twin respectively. Finally, the dashed line is the curve where $\det \mathbf{U}_i = 1$.	56
2.6	Example of triple junctions for type I and type II twins with perturbed austenite.	59
3.1	Picture of a section of Ω_C , on the left $\Omega_M(t)$ where $\nabla \mathbf{y} = \mathbf{1} + \mathbf{a} \otimes \mathbf{n}$, on the right Ω_A , where $\nabla \mathbf{y} = \mathbf{1}$	91
4.1	Points which are not regular: $\mathbf{x}_a$ is in a smooth k -graph contained in $\Gamma(t)$, but is not separating Ω_A from Ω_M . $\Gamma(t)$ does not coincide with a Lipschitz function differentiable at $\mathbf{x}_b$.	108
4.2	Picture of two parallel interfaces moving in opposite directions and meeting at a certain interface where $\mathbf{a}, \mathbf{n}$ are discontinuous. In this case, $\nabla \cdot (\mathbf{n} \otimes \mathbf{a}) = \mathbf{0}$, but $\nabla \cdot \mathbf{a} \neq 0$.	120
5.1	Minimizing sequences for $\mathcal{E}$ when $0 \notin \mathcal{Z}$.	138
5.2	Piecewise approximation of the constructed function: in Figure 5.2a we show the function constructed in $(0, l_0)$, whose gradient oscillates between z_1 and z_2 . In Figure 5.2b we show the function constructed in $(l_0, 1)$, whose gradient oscillates between z_1 and z_3 .	148
5.3	Example of L -interval and D -interval defined in Definition 5.3.2. In this picture, the red, the blue and the green intervals are respectively the sets of points where $ v_x - z_1 \leq \eta$, $ v_x - z_2 \leq \eta$, and $ v_x - z_3 \leq \eta$. The $B_{\pm}^{\eta}$ -transition layers are coloured in yellow.	154

5.4	The aim of Lemma 5.3.2 is to show that, on the interval (a, b) , the L^2 -norm of the red function in this picture is bigger or equal than the L^2 -norm of the blue one. Here both, the red and the blue functions, are piecewise linear, with derivative in $\{0, \tau_1, \tau_2\}$ for almost every $x \in (a, b)$. Furthermore, $\mathcal{C}_1, \mathcal{C}_2$ are the sets where the derivative of the red function is respectively τ_1 and τ_2 . We have $l_1 = \mathcal{L}(\mathcal{C}_1)$, $l_2 = \mathcal{L}(\mathcal{C}_2)$, that is, the measure of the set where the blue curve has gradient τ_1 (resp. τ_2) is equal to the measure of the set where the blue curve has gradient τ_1 (resp. τ_2).	155
5.5	Classification of D -intervals. In this picture, the red, the blue and the green intervals are respectively the sets of points where $ v_x - z_1 \leq \eta$, $ v_x - z_2 \leq \eta$, and $ v_x - z_3 \leq \eta$. The $B_{\pm}^{\eta}$ -transition layers are coloured in yellow.	159
5.6	The set Σ_i constructed in the proof of Proposition 5.3.1: $\Sigma_{i+1} = \emptyset$, while $\Sigma_i = \Sigma_i^d \cup \Sigma_i^m$, where $\Sigma_i^d \subset G_i^d$, $\Sigma_i^m \subset G_i^m$. In this picture, the red, the blue and the green intervals are respectively the sets of points where $ v_x - z_1 \leq \eta$, $ v_x - z_2 \leq \eta$, and $ v_x - z_3 \leq \eta$. The $B_{\pm}^{\eta}$ -transition layers are coloured in yellow.	163
5.7	5.7a. Microstructure of lower energy in case (H7) does not hold. Microstructures of low energy in case (H8) does not hold are represented in Figures 5.7b and 5.7c, respectively in the case where $\hat{y} \leq \sqrt{\frac{z_2 z_{21}}{z_3 z_{31}}}$ and $\hat{y} > \sqrt{\frac{z_2 z_{21}}{z_3 z_{31}}}$.	177
5.8	Verification of the hypotheses (H6)–(H8) for the examples of Section 5.6.1. Figure 5.8a is the graphical verification of (H6) for the two examples: in blue the case $z_2 = \frac{1}{3}$, in red the one with $z_2 = \frac{1}{2}$. Figure 5.8b and 5.8c verify (H7) and (H8) in the example with $z_2 = \frac{1}{3}$.	179

Chapter 1

Introduction

1.1 Martensitic transformations

Crystals are solid materials where atoms are highly ordered in three-dimensional periodic structures called lattices. In nature crystalline solids appear as single crystals or as polycrystals. While in the former the periodic structure is repeated without changes in the whole sample, the latter can be divided in regions, called grains, where the same lattice structure has different orientations.

In certain alloys or ceramics, the crystalline structure is different at different temperatures. The high-temperature crystalline structure is called austenite, or parent phase, while the low-temperature one is usually called martensite. Austenite has often a crystalline structure with cubic symmetry, and is more ordered than martensite, in the sense that its crystalline structure enjoys more symmetries. When the temperature is moved across a certain critical temperature θ_T , a phase transition takes place from austenite to martensite, or vice-versa. These solid to solid phase transitions are called martensitic transformations, and are diffusionless and of first order. Diffusionless means that locally it is possible to deduce one crystalline structure as a linear deformation of the other. Of first order means that the lattice change is abrupt. Examples of alloys undergoing martensitic transformations are $\text{In}_{77}\text{Tl}_{23}$, $\text{Ni}_{50}\text{Ti}_{50}$, $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$, as well as many steels, while among ceramics we mention BaTiO_3 . As austenite has usually higher symmetry than martensite, there are different ways to deform the lattice of the former and obtain the latter (cf. Figure 1.1). We call the different symmetry-related types of martensitic lattices, obtained by deforming austenite in different ways, *variants* of martensite. Often, martensitic variants appear finely mixed, creating beautiful and interesting microstructures. However, in general, predicting the microstructures appearing during the phase transformation is

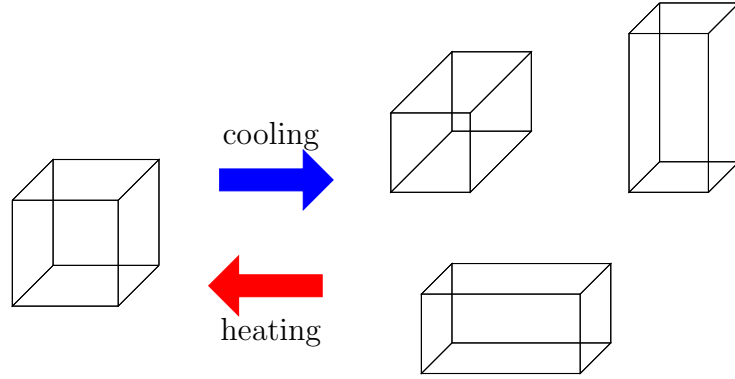


Figure 1.1: On the left cubic austenite, stable at high temperature; on the right tetragonal martensite, stable at low temperatures.

still an open problem.

In reality, nucleation and propagation of martensite into austenite (or vice versa of austenite into martensite) is obstructed by energy barriers, and the phase transition does not occur at an exact temperature value. Experimentally, we usually observe that when we cool a material in its austenite phase, martensite starts to appear at a temperature M_s , while the sample is fully transformed just at $M_f < M_s$. Vice versa, when we heat up a sample in its martensitic phase, austenite starts to appear at a temperature A_s , and the transformation ends at $A_f > A_s$. In general $A_s > M_f$, and we define the thermal hysteresis of the transformation as

$$\text{thermal hysteresis} := \frac{M_s + M_f - A_s - A_f}{2}.$$

Given A_s, A_f, M_s, M_f measured experimentally, an estimate of the theoretical critical temperature is given by

$$\tilde{\theta}_T = \frac{M_s + M_f + A_s + A_f}{4}.$$

Another way to induce martensitic transformations is via external loading. Indeed, external forces may change the energy landscape, and make some variants of martensite stable and the austenite unstable, above the critical temperature (cf. Figure 1.2). In this thesis, however, we focus on thermally induced transformations.

It is important to remark that, in reality, the sample might not fully transform. This may occur for many reasons, for example due to incompatibility of microstructures in neighbouring grains [14], or due to a very small size of some grains [5]. Furthermore, forces exerted on the boundary of a grain by its neighbours may affect the phase transformation. For these reasons, when studying martensitic transformation from a mathematical perspective, authors in the literature often restrict to single crystals. The same simplification is adopted in this thesis.

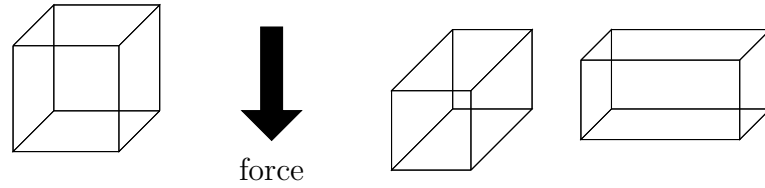


Figure 1.2: On the left, cubic austenite, which under a sufficiently large vertical force can have greater free energy than particular martensitic variants, on the right.

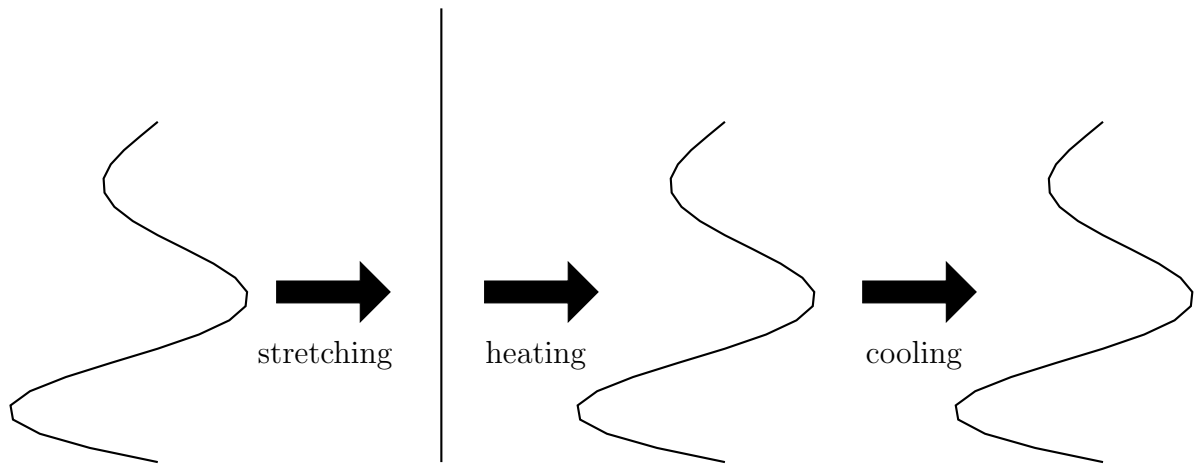


Figure 1.3: Example of shape memory. A curved wire is stretched and deformed into a straight wire. The deformation is apparently plastic as the wire does not recover its curved shape when the stretching force is released. When the sample is heated, it recovers its original shape. The shape remains unchanged when the material is cooled down again.

Shape memory

Shape memory is the ability of certain materials to recover, upon heating, deformations which are apparently plastic (cf. Figure 1.3). From a physical point of view, shape memory is based on martensitic phase transitions. Indeed, there is a one to one correspondence between the macroscopic effects presented in Figure 1.3, and the microscopic changes in the lattice depicted in Figure 1.4. The phenomenon can be described in four steps:

- we take a sample in its martensite phase, at temperature $\theta < \theta_T$;
- we stretch the sample in some direction. Locally, the martensitic microstructure changes in order to reduce the stress generated by the external force. When the external forces are released, the shape of the sample remains distorted, as the new microstructures are still energy minimising;

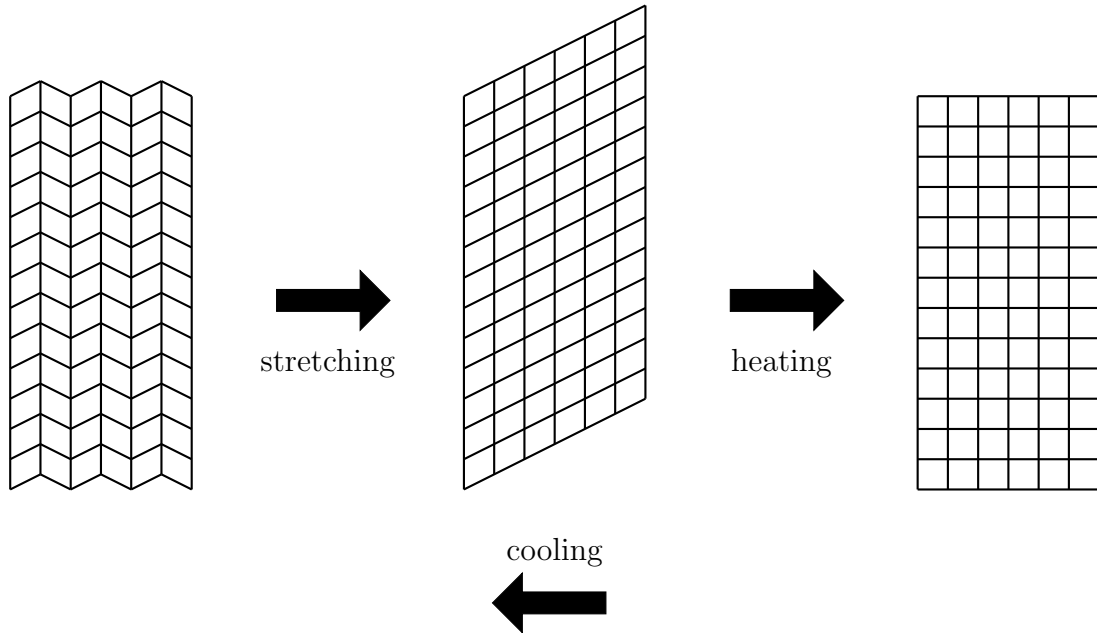


Figure 1.4: Representation of shape memory at a microscopic scale.

- we heat up our sample to a temperature $\theta > \theta_T$. The material undergoes a phase transformation from martensite to austenite. Given the higher symmetry of the austenite lattice, all the different variants of martensite transform back in a unique way, recovering the shape of the sample;
- we cool the material down to $\theta < \theta_T$ until a phase transformation from austenite to martensite occurs. Macroscopically, that means on average, the shape of the sample does not change, due to a phenomenon called *self-accommodation*, which is achieved by finely mixing martensitic variants.

As is easy to imagine, shape-memory alloys have a wide range of applications, from smart actuators to medical stents, energy conversion devices and many others, which make these materials interesting for the automotive, spacecraft, medical and robotic industries.

Hysteresis and reversibility

Major obstacles towards application of shape-memory alloys are the hysteresis and the lack of reversibility of the transformation. Indeed hysteresis can be up to 70°K in NiTi alloys, which is too big for making shape memory efficient in many practical applications. Lack of reversibility in the transformation means that during

every thermal cycle, the material is damaged by plastic effects and by the nucleation of micro-cracks. This leads to the formation of islands in the sample which do not change phase, while the latent heat of the transformation and transition temperature change, and rupture due to fatigue occurs after a few thousand cycles.

An important step towards understanding hysteresis can be found in [84], where the hysteresis is proved to be a function of λ_2 , and is minimised at $\lambda_2 = 1$. Here, λ_2 is the middle eigenvalue of the transformation strain converting the austenite lattice into the martensite lattices (see e.g., (2.8) for cubic to monoclinic transformations). This result seems in good agreement with experimental evidences for many different alloys. It is worth remarking that experiments show a very sensitive dependence of λ_2 on the alloy composition, and that hysteresis increases very rapidly with $|\lambda_2 - 1|$. Unfortunately, up to now, our understanding of reversibility in martensitic transformations is not as good as our understanding of hysteresis. The biggest step in this direction has been achieved in [26], where the authors study some particular conditions of geometric compatibility between phases, called cofactor conditions, and conjecture that these might affect reversibility (see also [51] where a relation between cofactor conditions and reversibility was already speculated). The cofactor conditions include, among others, $\lambda_2 = 1$, and were first introduced in [18], as conditions of degeneracy for the equations of the crystallographic theory of martensite (see Theorem 1.4.1 below).

The recent fabrication announced in [76] of $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$, the first material closely satisfying the cofactor conditions (the relative error is of order 10^{-4}), partially confirms this conjecture. Indeed, this material exhibits ultra-low hysteresis (about 2°K) and does not seem to incur any loss of reversibility after more than 16,000 thermal cycles. Furthermore, the hysteresis loop in this new material remains unchanged after 10^5 cycles of uniaxial compressive loading (see [67]). After $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$, other alloys closely satisfying the cofactor conditions have been fabricated (see [28] and [45]), whose hysteresis curve does not change after 10^7 cycles of uniaxial tension. We refer the reader also to [49] and [50] for two reviews on the topic.

Despite these recent successes in constructing materials undergoing ultra-reversible transformations, a rigorous mathematical explanation of reversibility is not available yet, and the influence of cofactor conditions still needs understanding as well as a rigorous justification.

1.2 Calculus of variations and quasiconvexity

We consider the following problem

$$\text{minimise} \quad I(\mathbf{u}) := \int_{\Omega} f(\mathbf{x}, \mathbf{u}(\mathbf{x}), \nabla \mathbf{u}(\mathbf{x})) \, d\mathbf{x} \quad \text{among } \mathbf{u} \in \mathcal{A}_{\mathbf{g}}, \quad (1.1)$$

where $\Omega \subset \mathbb{R}^N$ is an open, bounded, connected and Lipschitz domain, $f: \mathbb{R}^N \times \mathbb{R}^M \times \mathbb{R}^{M \times N}$ is a smooth non-negative function such that

$$c|\mathbf{A}|^p - C \leq f(\mathbf{x}, \mathbf{s}, \mathbf{A}) \leq C(1 + |\mathbf{A}|^p), \quad \text{for every } \mathbf{x} \in \Omega, \mathbf{s} \in \mathbb{R}^M, \mathbf{A} \in \mathbb{R}^{M \times N}, \quad (1.2)$$

for some $c, C > 0$, $p \in (1, \infty)$, and, for some $\mathbf{g} \in W^{1,p}(\Omega; \mathbb{R}^M)$,

$$\mathcal{A}_{\mathbf{g}} := \{\mathbf{v} \in W^{1,p}(\Omega; \mathbb{R}^M) : \mathbf{v} = \mathbf{g} \text{ on } \partial\Omega\}.$$

Quasiconvexity

The usual approach to prove existence of minimisers for I in $\mathcal{A}_{\mathbf{g}}$ is the direct method of the calculus of variations, based on the following three steps:

- take a minimising sequences $\mathbf{u}_j \in \mathcal{A}_{\mathbf{g}}$ for I ;
- show that the minimising sequence is weakly relatively compact in $\mathcal{A}_{\mathbf{g}}$;
- show that I is sequentially weakly lower semicontinuous.

The second point of the above list can be easily addressed as follows: by (1.2),

$$c\|\mathbf{u}_j\|_{W^{1,p}}^p - C \leq \inf I + \varepsilon \leq I(\mathbf{g}) + \varepsilon \leq C(1 + \|\mathbf{g}\|_{W^{1,p}}^p) \leq \hat{C}.$$

for some $\varepsilon, \hat{C} > 0$. Therefore, by the Banach-Alouglu theorem there exist $\mathbf{u} \in \mathcal{A}_{\mathbf{g}}$ and a subsequence $\mathbf{u}_{j_k}$ of $\mathbf{u}_j$ converging weakly in $W^{1,p}(\Omega; \mathbb{R}^M)$ to $\mathbf{u}$. It is usually more difficult to show that I is sequentially weakly lower semicontinuous. To address this issue, let us introduce the notion of quasiconvexity first introduced by Morrey [59]:

Definition 1.2.1 ([34, Def. 1.5]). *A function $h: \mathbb{R}^{M \times N} \rightarrow \mathbb{R}$ which is Borel measurable and locally bounded is said to be quasiconvex if*

$$h(\mathbf{A}) \leq \frac{1}{\mathcal{L}^N(D)} \int_D h(\mathbf{A} + \nabla \phi(\mathbf{x})) \, d\mathbf{x},$$

for every open set $D \subset \mathbb{R}^N$, $\mathbf{A} \in \mathbb{R}^{M \times N}$ and every $\phi \in W_0^{1,\infty}(D; \mathbb{R}^M)$.

The notion of quasiconvexity is a generalisation of convexity, and the two notions are equivalent if $M = 1$ or $N = 1$. However, the importance of quasiconvexity is given by its strong relationship with sequential weak lower semicontinuity. Indeed, under standard assumptions on f , quasiconvexity is a necessary and sufficient condition:

Theorem 1.2.1 ([34, Thm. 1.13]). *Let $1 \leq p < \infty$, $\Omega \subset \mathbb{R}^N$ bounded open Lipschitz set and*

$$f: \Omega \times \mathbb{R}^M \times \mathbb{R}^{M \times N} \rightarrow \mathbb{R},$$

be a continuous function satisfying

$$0 \leq f(\mathbf{x}, \mathbf{v}, \mathbf{A}) \leq c(1 + |\mathbf{A}|^p),$$

for some $c > 0$. Let

$$I(\mathbf{u}) = \int_{\Omega} f(\mathbf{x}, \mathbf{u}(\mathbf{x}), \nabla \mathbf{u}(\mathbf{x})) \, d\mathbf{x}.$$

Then, I is sequentially weakly lower semicontinuous in $W^{1,p}(\Omega; \mathbb{R}^M)$ if and only if $\mathbf{A} \rightarrow f(\mathbf{x}, \mathbf{v}, \mathbf{A})$ is quasiconvex for every $(\mathbf{x}, \mathbf{v}) \in \Omega \times \mathbb{R}^M$.

There are other generalisations of convexity, which have been introduced to better understand quasiconvexity, in particular polyconvexity [7], and rank-one convexity (see e.g., [34]):

Definition 1.2.2. *We say that $h: \mathbb{R}^{M \times N} \rightarrow \mathbb{R}$ is rank-one convex, if*

$$h(\lambda \mathbf{A} + (1 - \lambda) \mathbf{B}) \leq \lambda h(\mathbf{A}) + (1 - \lambda) h(\mathbf{B}),$$

for every $\lambda \in [0, 1]$, and every $\mathbf{A}, \mathbf{B} \in \mathbb{R}^{M \times N}$ such that $\text{rank}(\mathbf{A} - \mathbf{B}) \leq 1$.

We say that $h: \mathbb{R}^{M \times N} \rightarrow \mathbb{R}$ is polyconvex, if there exists $\hat{h}: \mathbb{R}^{\tau(M,N)} \rightarrow \mathbb{R}$ convex, such that

$$h(\mathbf{A}) = \hat{h}(T(\mathbf{A})), \quad \text{for every } \mathbf{A} \in \mathbb{R}^{M,N},$$

and where $T: \mathbb{R}^{M,N} \rightarrow \mathbb{R}^{\tau(M,N)}$ is such that

$$T(\mathbf{A}) = (\mathbf{A}, \text{adj}_2 \mathbf{A}, \dots, \text{adj}_{\min\{M,N\}} \mathbf{A}).$$

Here, $\text{adj}_s \mathbf{A}$, with $s \in [2, \min\{M, N\}]$, stands for the matrix of all $s \times s$ minors of $\mathbf{A} \in \mathbb{R}^{M,N}$, and

$$\tau(M, N) = \sum_{j=1}^{\min\{M,N\}} \sigma(s), \quad \sigma(s) = \binom{N}{s} \binom{M}{s}.$$

The relations between the different generalisations of convexity are given by the following result:

Proposition 1.2.1 ([34, Thm. 1.7]). *Let $h: \mathbb{R}^{M \times N} \rightarrow \mathbb{R}$ be Borel measurable and locally bounded. Then, if h is quasiconvex, it is also rank-one convex. If h is polyconvex, it is also quasiconvex. The converses are false in general.*

Therefore, an easy way to prove that a function h is not quasiconvex is to show that it is not rank-one convex. Conversely, to show that a function h is quasiconvex, it is sufficient (but not necessary) to show that h is polyconvex.

Lack of quasiconvexity

In certain physically relevant situations, as for the case of martensitic phase transitions discussed in Section 1.3, the energy density f is not quasiconvex, and the question arises as to what happens to minimising sequences. One possible approach is to consider the relaxed energy functional I^{rel} defined as

$$I^{rel}(\mathbf{u}) := \inf \left\{ \liminf_j I(\mathbf{u}_j) : \mathbf{u}_j - \mathbf{u} \rightharpoonup \mathbf{0} \text{ weakly in } W_0^{1,p}(\Omega; \mathbb{R}^M) \right\}.$$

This is always sequentially weakly lower semicontinuous, and hence admits a minimiser. The following result holds.

Theorem 1.2.2 ([34]). *Let $h: \mathbb{R}^{M \times N} \rightarrow \mathbb{R}$ be continuous and satisfy*

$$c|\mathbf{A}|^p - C \leq h(\mathbf{A}) \leq C(1 + |\mathbf{A}|^p), \quad (1.3)$$

for some $c, C > 0$ and $p \in (1, \infty)$. Let also

$$E(\mathbf{u}) = \int_{\Omega} h(\nabla \mathbf{u}(\mathbf{x})) \, d\mathbf{x}.$$

Then,

$$E^{rel}(\mathbf{u}) = \int_{\Omega} h^{qc}(\nabla \mathbf{u}) \, d\mathbf{x},$$

where

$$h^{qc}(\mathbf{A}) := \sup \{ g(\mathbf{A}) : g \leq h \text{ and } g \text{ quasiconvex} \}.$$

Therefore, under the hypotheses of Theorem 1.2.2, relaxation of an energy functional is equivalent to the quasiconvexification of its energy density. However, following this approach, information about minimising sequences is usually lost, and the relaxed problem may have infinitely many minimisers.

Young measures

A way of capturing the oscillations in minimising sequences for non-quasiconvex functionals is given by Young measures, which are parametrised families of probability measures, generated by weakly converging sequences. Before stating a version of the fundamental theorem for Young measures we need to introduce some notation. Below, we denote by $\mathcal{M}(\mathbb{R}^M)$ the space of positive Radon measures with finite mass, and by $C_0(\mathbb{R}^M)$ its predual, the space of continuous functions $f: \mathbb{R}^M \rightarrow \mathbb{R}$ such that $\lim_{|\mathbf{x}| \rightarrow \infty} f(\mathbf{x}) = 0$. We can now introduce the space $L_{w^*}^\infty(\Omega; \mathcal{M}(\mathbb{R}^M))$, which is the space of essentially-bounded measure-valued functions equipped with the weak* topology. We remark that, as $Y := L_{w^*}^\infty(\Omega; \mathcal{M}(\mathbb{R}^M)) = (L^1(\Omega; C_0(\mathbb{R}^M)))^*$, it is a Banach space with the norm

$$\|\nu\|_Y := \operatorname{ess\,sup}_{\mathbf{x} \in \Omega} \|\nu_{\mathbf{x}}\|_{\mathcal{M}}.$$

We are now ready to state the following theorem:

Theorem 1.2.3 ([8, 69]). *Let $\Omega \subset \mathbb{R}^N$ be measurable, $K \in \mathbb{R}^M$ closed, and let $\mathbf{z}_j: \Omega \rightarrow \mathbb{R}^M$ be a sequence of measurable functions such that for any open set $\mathcal{U}$, $K \subset \mathcal{U}$,*

$$\lim_j \mathcal{L}^N \{\mathbf{x} \in \Omega: \mathbf{z}_j(\mathbf{x}) \notin \mathcal{U}\} = 0.$$

Then there exists a subsequence $\mathbf{z}_{j_k}$, and $\nu \in L_{w^}^\infty(\Omega; \mathcal{M}(\mathbb{R}^M))$, such that*

- $\int_{\mathbb{R}^M} d\nu_{\mathbf{x}} \leq 1$, for a.e. $\mathbf{x} \in \Omega$;
- $\operatorname{supp} \nu_{\mathbf{x}} \subset K$ for a.e. $\mathbf{x} \in \Omega$;
- $f(\mathbf{z}_{j_k}) \rightharpoonup \int_{\mathbb{R}^M} f(\mathbf{a}) d\nu_{\mathbf{x}}(\mathbf{a})$ weakly* in $L^\infty(\Omega)$, for every $f \in C_0(\mathbb{R}^M)$.

If further

$$\sup_j \int_{\Omega} \psi(|\mathbf{z}_j|) d\mathbf{x} < \infty,$$

for some continuous non-decreasing $\psi: [0, \infty) \rightarrow [0, \infty)$ such that $\lim_{t \rightarrow \infty} \psi(t) = \infty$, then $\int_{\mathbb{R}^M} d\nu_{\mathbf{x}} = 1$, for a.e. $\mathbf{x} \in \Omega$, and

$$g(\mathbf{z}_{j_k}) \rightharpoonup \int_{\mathbb{R}^M} g(\mathbf{a}) d\nu_{\mathbf{x}}(\mathbf{a}), \quad \text{weakly in } L^1(\Omega),$$

for every continuous $g: \mathbb{R}^M \rightarrow \mathbb{R}$ such that $g(\mathbf{z}_{j_k})$ is weakly relatively compact in $L^1(\Omega)$.

In the situation of the above theorem, we say that the sequence $\mathbf{z}_{j_k}$ generates a Young measure ν . Let us consider the measure $\nu_{\mathbf{x}}^{k,r}$ defined by

$$\nu_{\mathbf{x}}^{k,r}(A) = \frac{\mathcal{L}^M(\{\mathbf{s} \in B_r(\mathbf{x}) \cap \Omega : \mathbf{z}_{j_k}(\mathbf{s}) \in A\})}{\mathcal{L}^M(\Omega \cap B_r(\mathbf{x}))}, \quad (1.4)$$

for every Lebesgue-measurable set $A \subset \mathbb{R}^M$, every $\mathbf{x} \in \Omega$, and where we denoted by B_r the ball of radius r centred at $\mathbf{x} \in \mathbb{R}^N$. It is possible to check (see e.g., [8]) that, for a.e. $\mathbf{x} \in \Omega$, $\lim_{r \rightarrow 0} \lim_{k \rightarrow \infty} \nu_{\mathbf{x}}^{k,r} = \nu_{\mathbf{x}}$, where the convergence is weak* in the sense of measures. The characterisation of Young measures in (1.4) shows how these objects capture the fine oscillations of the sequence $\mathbf{z}_{j_k}$. In other words, $\nu_{\mathbf{x}}$ can be thought as the limiting probability distribution for the values of $\mathbf{z}_{j_k}$ close to $\mathbf{x} \in \Omega$.

The characterisation of Young measures generated by sequences of gradients is less trivial. Indeed, in the general case $M, N > 1$, not all Young measures can be generated by sequences of gradients. The following theorem due to Kinderlehrer and Pedregal [53, 54] relates Gradient Young measures to quasiconvexity:

Theorem 1.2.4 ([69]). *Let $\nu \in L_{w^*}^\infty(\Omega; \mathcal{M}(\mathbb{R}^{M \times N}))$. There exists a sequence $\mathbf{z}_j$ bounded in $W^{1,p}(\Omega; \mathbb{R}^M)$, $p > 1$, such that $\nabla \mathbf{z}_j$ generates ν as $j \rightarrow \infty$ if and only if*

1. $\int_{\mathbb{R}^{M \times N}} \mathbf{A} d\nu_{\mathbf{x}}(\mathbf{A}) = \nabla \mathbf{y}(\mathbf{x})$, for a.e. $\mathbf{x} \in \Omega$, some $\mathbf{y} \in W^{1,p}(\Omega; \mathbb{R}^M)$;
2. $\int_{\Omega} \int_{\mathbb{R}^{M \times N}} |\mathbf{A}|^p d\nu_{\mathbf{x}}(\mathbf{A}) d\mathbf{x}$ if $p \in (1, \infty)$, $\text{supp } \nu_{\mathbf{x}} \subset \hat{K}$, for a.e. $\mathbf{x} \in \Omega$, for some $\hat{K} \subset \mathbb{R}^{M \times N}$ compact, if $p = \infty$;
3. $\int_{\mathbb{R}^{M \times N}} \psi(\mathbf{A}) d\nu_{\mathbf{x}}(\mathbf{A}) \geq \psi(\nabla \mathbf{y}(\mathbf{x}))$, for a.e. $\mathbf{x} \in \Omega$, every quasiconvex and bounded from below $\psi: \mathbb{R}^{M \times N} \rightarrow \mathbb{R}$, such that, $|\psi(\mathbf{A})| \leq c(1 + |\mathbf{A}|^p)$. If $p = \infty$, no growth condition is required on ψ .

A possible way to relax I without losing all the information about the minimising sequences is to consider the problem

$$\text{minimise } I^{GYM} := \int_{\Omega} \int_{\mathbb{R}^{M \times N}} f(\mathbf{x}, \mathbf{y}(\mathbf{x}), \mathbf{A}) d\nu_{\mathbf{x}}(\mathbf{A}) d\mathbf{x}, \quad \begin{array}{l} \text{among } \nu \text{ satisfying 1--3,} \\ \mathbf{y} \in \mathcal{A}_{\mathbf{g}}. \end{array}$$

In this case, by (1.2) we can guarantee the existence of minimisers, and at the same time keep track of the fine oscillations in the minimising sequences.

Existence without weak lower semicontinuity

Following the pioneering work of Nash [66] and Kuiper [56] on isometric embeddings, and the later work of Gromov [44] we can construct minimisers of suitable energy functionals

$$E(\mathbf{u}) := \int_{\Omega} h(\nabla \mathbf{u}) \, d\mathbf{x},$$

for some bounded Lipschitz domain $\Omega \subset \mathbb{R}^N$ and non-negative continuous $h: \mathbb{R}^{M \times N} \rightarrow \mathbb{R}$, by solving a partial differential inclusion problem of the type

$$\text{find } \mathbf{u} \in W^{1,\infty}(\Omega; \mathbb{R}^M) \text{ such that } \begin{cases} \nabla \mathbf{u}(\mathbf{x}) \in \mathcal{K}, & \text{for a.e. } \mathbf{x} \in \Omega \\ \mathbf{u}(\mathbf{x}) = \mathbf{g}(\mathbf{x}), & \text{for all } \mathbf{x} \in \partial\Omega, \end{cases} \quad (1.5)$$

for some $\mathbf{g} \in W^{1,\infty}(\Omega; \mathbb{R}^M)$, where

$$\mathcal{K} := \{\mathbf{F} \in \mathbb{R}^{M \times N} : h(\mathbf{F}) = 0\},$$

and $\mathcal{K}$ is assumed to be compact non-empty. We remark that if $\mathcal{K}$ is big enough to have existence of solutions to (1.5), then no quasiconvexity of the energy density h is required to have existence of minimisers. We follow here the approach of Müller and Šverák [65] (see also [62, 64]), but we refer the reader to the work of Dacorogna and Marcellini (see [34, 35] and references therein) and to [63, 78] for different approaches to the topic.

Let us start by recalling that, given a set $\mathcal{B} \in \mathbb{R}^{M \times N}$, its rank-one convex hull $\mathcal{B}^{rc}$ is defined by

$$\mathcal{B}^{rc} := \left\{ \mathbf{F} \in \mathbb{R}^{M \times N} : h(\mathbf{F}) \leq \sup_{\mathbf{G} \in \mathcal{A}} h(\mathbf{G}), \text{ for each } h: \mathbb{R}^{M \times N} \rightarrow \mathbb{R} \text{ rank-one convex} \right\}.$$

We are now ready to introduce the following definition:

Definition 1.2.3 ([65, Def. 1.2]). *Let $\Sigma \subset \mathbb{R}^{M \times N}$ and $\mathcal{B} \subset \Sigma$. A sequence of sets $U_i \subset \Sigma$ is an in-approximation of $\mathcal{B}$ in Σ , if*

- *the U_i are open in Σ ;*
- *the U_i are uniformly bounded;*
- *$U_i \subset U_{i+1}^{rc}$;*
- *if $\mathbf{F}_i \in U_i$ and $\mathbf{F}_i \rightarrow \mathbf{F}$, then $\mathbf{F} \in \mathcal{B}$.*

Under suitable assumptions, existence of solutions to (1.5) is given by the following theorem:

Theorem 1.2.5 ([65, Thm. 1.3]). *Let $\Sigma := \{\mathbf{F} : \det \mathbf{F} = t\}$, $t \neq 0$, or $\Sigma = \mathbb{R}^{M \times N}$. Let $\mathcal{K} \subset \Sigma$, and U_i be an in-approximation of $\mathcal{K}$ in Σ . Suppose $\mathbf{g} \in C^{2,\alpha}(\Omega; \mathbb{R}^M)$ (or piecewise linear) and that $\nabla \mathbf{g}(\mathbf{x}) \in U_1$ for every $\mathbf{x} \in \Omega$. Then, there exists a solution to problem (1.5).*

Thanks to Theorem 1.2.5, in many cases of relevance for martensitic transformations, we can construct existence of exact minimisers, despite the lack of weak lower semicontinuity of the energy functional. Furthermore, these solutions are minimisers also to the relaxed functionals E^{rel}, E^{GYM} . However, solutions constructed in this way are in general infinitely many, and it is an open problem to select the physically relevant ones.

1.3 Nonlinear elasticity model

In order to describe austenite to martensite phase transitions in crystalline solids, one of the most successful mathematical continuum models is nonlinear elasticity, which has proved to capture many aspects of the physical phenomena such as the formation of twins (see [18]), and to be useful in understanding related behaviour such as the shape-memory effect (see [22]), and, more recently, hysteresis (see [84]). In this and the next subsection, we give a brief overview of the theory following closely [20, 26]. For more details we refer the reader to [18, 19, 23].

In order to understand the nonlinear elasticity model, we need to introduce first the concept of Bravais lattice, which is a defect-free perfectly periodic lattice, as stated in the following definition:

Definition 1.3.1 ([23]). *Let $\{\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3\}$ be three linearly independent vectors. We say that*

$$\mathcal{L}(\mathbf{e}_i, \mathbf{o}) := \left\{ \mathbf{x} \in \mathbb{R}^3 : \mathbf{x} = \mathbf{o} + \sum_{i=1}^3 \nu_i \mathbf{e}_i, \text{ and } \nu_1, \nu_2, \nu_3 \in \mathbb{Z} \right\},$$

is a Bravais lattice through the point $\mathbf{o} \in \mathbb{R}^3$, generated by $\{\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3\}$.

Bravais lattices are often an approximation of real lattices, not only due to the presence of defects, but also due to the lack of periodicity. Indeed, often real lattices can be better approximated as the union of more than one Bravais lattice, called multi-lattices. In what follows, for simplicity, we restrict ourselves to the case of a

single lattice, referring the reader to [23, 72] for more complex situations.

It is possible to show that there are fourteen types of Bravais lattices, which can be divided into the following seven symmetry groups: cubic, tetragonal, orthorhombic, monoclinic, trigonal, hexagonal and triclinic.

We remark that the representation of a Bravais lattice is non-unique. Indeed, given any matrix $\mathbf{M} \in \mathbb{Z}^{3 \times 3}$, $\det \mathbf{M} = 1$, it is possible to show that

$$\mathfrak{L}(\mathbf{e}_i, \mathbf{o}) = \mathfrak{L}(\mathbf{f}_i, \mathbf{o}), \quad \text{where} \quad \mathbf{f}_i = \sum_{j=1}^3 \mathbf{M}_{ij} \mathbf{e}_j, \quad i = 1, 2, 3.$$

Both rotations and simple shears can map a Bravais lattice into itself, but as plastic effects are usually negligible in martensitic transformations, of more interest is the *point group of a lattice* $\mathcal{P}(\mathbf{e}_i)$ defined as

$$\mathcal{P}(\mathbf{e}_i) := \{\mathbf{R} \in SO(3) : \mathfrak{L}(\mathbf{e}_i, \mathbf{o}) = \mathfrak{L}(\mathbf{R}\mathbf{e}_i, \mathbf{o})\}.$$

This is the set of rotations that send a Bravais lattice into itself.

The main ingredient to pass from an atomistic to a continuum theory is the Cauchy-Born hypothesis. To give an idea of this concept, we consider a sample occupying a region $\Omega \subset \mathbb{R}^3$, and we assume that at every $\mathbf{x} \in \Omega$ there exists a Bravais lattice generated by the vectors $\{\mathbf{e}_1^0(\mathbf{x}), \mathbf{e}_2^0(\mathbf{x}), \mathbf{e}_3^0(\mathbf{x})\}$. Let us further assume that our sample gets smoothly deformed, and that the position of the material point initially in $\mathbf{x} \in \Omega$ is now given by $\mathbf{y}(\mathbf{x})$, where $\mathbf{y} \in C^1(\Omega; \mathbb{R}^3)$. The Cauchy-Born hypothesis states that the lattice vectors deform according to the deformation gradient, that is, the Bravais lattice at the point $\mathbf{y}(\mathbf{x})$ after the body has been deformed is generated by $\{\mathbf{e}_1(\mathbf{y}(\mathbf{x})), \mathbf{e}_2(\mathbf{y}(\mathbf{x})), \mathbf{e}_3(\mathbf{y}(\mathbf{x}))\}$, and

$$\mathbf{e}_i(\mathbf{y}(\mathbf{x})) = \nabla \mathbf{y}(\mathbf{x}) \mathbf{e}_i^0(\mathbf{x}).$$

We remark that simple shears may macroscopically deform a body, but locally leave the vectors generating the Bravais lattice unchanged. The Ericksen-Pitteri neighbourhood [38, 71, 72] provides a framework where the Cauchy-Born can be legitimised.

The nonlinear elasticity model is based on the idea of looking at changes in the crystal lattice as elastic deformations in the continuum mechanics framework, and legitimised by the Cauchy-Born hypothesis. Following [18], we hence assume that the deformations minimize a free energy of the type

$$\mathcal{E}(\mathbf{y}, \theta) = \int_{\Omega} \phi(\nabla \mathbf{y}(\mathbf{x}), \theta) \, d\mathbf{x}. \quad (1.6)$$

Here, θ denotes the temperature of the crystal. Three different regimes are distinguishable depending on this parameter: $\theta < \theta_T$ and $\theta > \theta_T$, where respectively martensite and austenite phases minimize the energy, and $\theta = \theta_T$ where these are energetically equivalent. In (5.2), the bounded Lipschitz domain (open and connected) Ω stands for the reference configuration of a single crystal of undistorted austenite at $\theta = \theta_T$ and $\mathbf{y}(\mathbf{x})$ denotes the position of the particle $\mathbf{x} \in \Omega$ after the deformation. Finally, ϕ is the free-energy density, depending on the temperature θ and the deformation gradient $\nabla \mathbf{y}$, satisfying the following properties:

- $\mathcal{D} := \{\mathbf{F} \in \mathbb{R}^{3 \times 3} : \det \mathbf{F} > 0\}$,

$$\phi(\cdot, \theta) : \mathcal{D} \rightarrow \mathbb{R}$$

is a function bounded below by a constant depending on θ for each $\theta > 0$;

- $\phi(\cdot, \theta)$ satisfies frame-indifference, i.e., for all $\mathbf{F} \in \mathcal{D}$ and all rotations $\mathbf{R} \in SO(3)$, $\phi(\mathbf{R}\mathbf{F}, \theta) = \phi(\mathbf{F}, \theta)$. This property reflects the invariance of the free-energy density under rotations;
- ϕ respects lattice symmetries, i.e., $\phi(\mathbf{F}\mathbf{Q}, \theta) = \phi(\mathbf{F}, \theta)$ for all $\mathbf{F} \in \mathcal{D}$ and all rotations $\mathbf{Q}$ in the point group of austenite $\mathcal{P}_a$, usually the group of rotations sending a cube into itself;
- denoting by K_θ the set of minima for the free-energy density at temperature θ , i.e., $K_\theta := \{\mathbf{F} \in \mathcal{D} : \mathbf{F} \in \operatorname{argmin}(\phi(G, \theta))\}$,

$$K_\theta = \begin{cases} \alpha(\theta)SO(3), & \theta > \theta_T \\ SO(3) \cup \bigcup_{i=1}^N SO(3)\mathbf{U}_i(\theta_T), & \theta = \theta_T \\ \bigcup_{i=1}^N SO(3)\mathbf{U}_i(\theta), & \theta < \theta_T. \end{cases} \quad (1.7)$$

Here, $\alpha(\theta)$ is a scalar dilatation coefficient satisfying $\alpha(\theta_T) = 1$, while $\mathbf{U}_i(\theta) \in \mathbb{R}_{Sym+}^{3 \times 3}$ are the N positive definite symmetric matrices corresponding to the transformation from austenite to the N variants of martensite at temperature θ . Here and below $\mathbb{R}_{Sym+}^{3 \times 3}$ represents the set of 3×3 symmetric and positive definite matrices. From now on, we omit the dependence on the temperature in K_θ when $\theta < \theta_T$, and neglect the dependence on θ of the $\mathbf{U}_i$. We remark that $N = \frac{\#\mathcal{P}_a}{\#\mathcal{P}_m}$, where $\mathcal{P}_m$ is the point group of martensite, and that for each $\mathbf{U}_i, \mathbf{U}_j$ there exists $\mathbf{R} \in \mathcal{P}_a$ such that $\mathbf{R}^T \mathbf{U}_j \mathbf{R} = \mathbf{U}_i$, so that $\mathbf{U}_i, \mathbf{U}_j$ share the same eigenvalues.

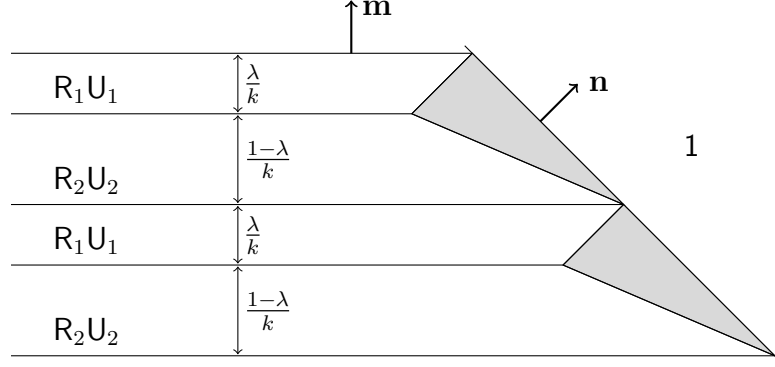


Figure 1.5: Example of minimising sequence $\nabla \mathbf{y}_k$ involving both austenite and martensite. Here, $\mathbf{R}_1 \mathbf{U}_1 - \mathbf{R}_2 \mathbf{U}_2 = \mathbf{b} \otimes \mathbf{m}$, and $(1 - \lambda) \mathbf{R}_1 \mathbf{U}_1 + \lambda \mathbf{R}_2 \mathbf{U}_2 = \mathbf{1} + \mathbf{a} \otimes \mathbf{n}$. Therefore compatibility between phases is achieved on average. As k tends to infinity, the grey region, the only one where the energy is different from zero and the material is not stress-free, vanishes.

1.3.1 Lack of quasiconvexity and relaxation

Based on both experimental evidence and the mathematical complexity of other cases, most results in the literature are related to planar austenite-martensite interfaces $\{\mathbf{x} \in \mathbb{R}^3 : \mathbf{x} \cdot \mathbf{n} = k\}$, with normal $\mathbf{n}$, at $\theta = \theta_T$. In this case, under suitable conditions on the lattice deformation and for some $\mathbf{n} \in \mathbb{R}^3$, it is possible to construct a sequence $\mathbf{y}^j$ such that

$$\nabla \mathbf{y}^j \rightarrow SO(3) \quad \text{in measure} \quad \text{for } \mathbf{x} \cdot \mathbf{n} < k \quad (1.8)$$

$$\nabla \mathbf{y}^j \rightarrow \bigcup_{i=1}^N SO(3) \mathbf{U}_i \quad \text{in measure} \quad \text{for } \mathbf{x} \cdot \mathbf{n} > k. \quad (1.9)$$

Denoting by $\mathcal{L}^3$ three-dimensional Lebesgue measure, we notice that (1.8)–(1.9) imply

$$\lim_{j \rightarrow \infty} \mathcal{L}^3 \{\mathbf{x} \in \Omega : \nabla \mathbf{y}^j(\mathbf{x}) \notin K_{\theta_T}\} = 0,$$

and, under some further hypotheses on ϕ , $\mathbf{y}^j$ is a minimizing sequence for $\mathcal{E}(\cdot, \theta_T)$ (see [18] for more details). Furthermore, since $\mathbf{y}^j$ can be constructed so as to be bounded in $W^{1,\infty}(\Omega, \mathbb{R}^3)$, there exists a subsequence $\mathbf{y}^{k_j}$ and $\mathbf{y} \in W^{1,\infty}(\Omega, \mathbb{R}^3)$ such that $\mathbf{y}^{k_j}$ converges to $\mathbf{y}$ weakly* in the same space. However, the energy functional is not quasiconvex and, in general, the minimum is not attained in the classical sense. Therefore, $\nabla \mathbf{y}$ is not a minimizer for $\mathcal{E}(\cdot, \theta_T)$, but just of its relaxation, $\mathcal{E}^{qc}(\cdot, \theta_T)$. It is worth remarking that the relaxation result in Theorem 1.2.2 does not apply to the context of nonlinear elasticity. Indeed, the energy density ϕ must satisfy $\phi(\mathbf{F}, \cdot) = \infty$

for every $\mathbf{F} \in \mathbb{R}^{3 \times 3} \setminus \mathcal{D}$, and (1.3) is not satisfied. We refer the reader to [30, 55] for relaxation results in case of energy densities which are not locally finite. From a physical point of view, $\nabla \mathbf{y}$ represents the deformation gradient in the sample at a macroscopic scale, an average of the fine microstructures with gradients in

$$K := \bigcup_{i=1}^N SO(3) \mathbf{U}_i.$$

It is important to remark that, in general, macroscopic deformation gradients $\nabla \mathbf{y}$ are not elements of K a.e. in Ω . Instead, we have $\nabla \mathbf{y} \in K^{qc}$ a.e. in Ω , where

$$K^{qc} := \left\{ \mathbf{M} \in \mathbb{R}^{3 \times 3} \mid \begin{array}{l} f(\mathbf{M}) \leq \max_K f, \text{ for all continuous} \\ \text{quasiconvex } f: \mathbb{R}^{3 \times 3} \rightarrow \mathbb{R} \end{array} \right\},$$

is the quasiconvex hull of the set K (see [62]).

On the other hand, a way to deal with the lack of weak lower semicontinuity of the functional is thus to consider the relaxed problem

$$\mathcal{E}^r(\mathbf{y}, \nu_{\mathbf{x}}) = \int_{\Omega} \int_{\mathbb{R}^{3 \times 3}} \phi(\mathbf{F}, \theta_T) d\nu_{\mathbf{x}} d\mathbf{x},$$

where $\nu_{\mathbf{x}}$ is a gradient Young measure with barycentre $\nabla \mathbf{y}$ for some $\mathbf{y} \in W^{1,\infty}(\Omega; \mathbb{R}^3)$. In this way, under mild assumptions on ϕ the problem always admits minimisers, and the gradient Young measure can keep track of the oscillations in the minimising sequences. Thus, $\nu_{\mathbf{x}}$ embeds the information about the microstructures, while the barycentre of $\nu_{\mathbf{x}}$, namely $\nabla \mathbf{y}(\mathbf{x}) := \int_{\mathbb{R}^{3 \times 3}} \mathbf{F} d\nu_{\mathbf{x}}(\mathbf{F})$, is called the macroscopic (or average) deformation gradient, because it is an average of the fine microstructures of $\nabla \mathbf{y}_{k_j}$. However, minimising gradient Young measures for $\mathcal{E}^r$ can sometimes be infinitely many, and is an open problem to select the physical ones.

1.3.2 Hadamard jump conditions

Characterizing the set of possible macroscopic deformations K^{qc} is very important in order to fully understand the nonlinear elasticity model. On the other hand, the set of constant macroscopic gradients $\mathbf{B}$ which can form an interface with austenite, having constant gradient $\mathbf{A}$, is in general smaller than the whole of K^{qc} . Indeed, a Lipschitz function whose gradient is equal to $\mathbf{A}, \mathbf{B}$ a.e. in Ω , with $\mathbf{A}, \mathbf{B} \in \mathbb{R}^{3 \times 3}$ must satisfy a generalized version of the Hadamard jump condition proved in [18] (see Figure 1.6):

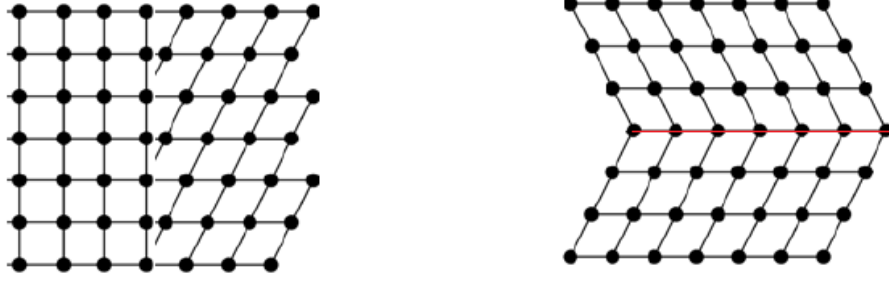


Figure 1.6: Example of incompatible deformation gradients (on the left) and compatible deformation gradients (on the right) at a crystalline length-scale.

Proposition 1.3.1 ([18, Prop. 1]). *Let $\Omega \in \mathbb{R}^3$ be open and connected. Assume $\mathbf{y} \in W^{1,\infty}(\Omega, \mathbb{R}^3)$ satisfies*

$$\nabla \mathbf{y}(\mathbf{x}) = \mathbf{A}, \quad a.e. \ \mathbf{x} \in \Omega_A; \quad \nabla \mathbf{y}(\mathbf{x}) = \mathbf{B}, \quad a.e. \ \mathbf{x} \in \Omega_B,$$

where $\mathbf{A}, \mathbf{B} \in \mathbb{R}^{3 \times 3}$ and Ω_A, Ω_B are disjoint measurable sets such that

$$\Omega_A \cup \Omega_B = \Omega, \quad \mathcal{L}^3(\Omega_A) > 0, \quad \mathcal{L}^3(\Omega_B) > 0.$$

Then,

$$\mathbf{A} - \mathbf{B} = \mathbf{a} \otimes \mathbf{n}, \quad \mathbf{a}, \mathbf{n} \in \mathbb{R}^3, \ |\mathbf{n}| = 1.$$

This result forces a fixed macroscopic gradient of martensite with constant gradient $\mathbf{B}$ to be rank-one connected to austenite, having constant gradient $\mathbf{A}$, across every interface between the two. Furthermore, it implies that, in this case, every phase interface is planar. For these reasons, Proposition 1.3.1 is the background for many results for compatibility between phases.

However, this type of result fails to be true for more general gradients $\nabla \mathbf{y} \in L^\infty(\Omega; \mathbb{R}^{3 \times 3})$. As a matter of fact, as shown in [18], it is possible to construct a Lipschitz function $\mathbf{z} \in W^{1,\infty}(\Omega, \mathbb{R}^3)$ which is constant in the set $\Omega \cap \{\mathbf{x} : \mathbf{x} \cdot \mathbf{n} > c\}$ for some $c \in \mathbb{R}$, $\mathbf{n} \in \mathbb{R}^3$, and whose gradient $\nabla \mathbf{y} \in \{\mathbf{F}_1, \mathbf{F}_2, \mathbf{F}_3, \mathbf{F}_4\}$ in $\Omega \cap \{\mathbf{x} : \mathbf{x} \cdot \mathbf{n} < c\}$ for some matrices $\mathbf{F}_i$, such that $\nabla \mathbf{y}$ is not rank-one almost everywhere in $\Omega \cap \{\mathbf{x} : \mathbf{x} \cdot \mathbf{n} < c\}$. Indeed a fractal behaviour of $\nabla \mathbf{y}$ close to $\mathbf{x} \cdot \mathbf{n} = c$, finely mixing martensitic variants near the interface, allows one to achieve compatibility between incompatible gradients. That is, compatibility is achieved on the average. Possible approaches to recovering the Hadamard jump condition in an average sense can be found in [11], or in Remark 4.5.1 below. Another generalization of Proposition 1.3.1 was proved in [20]

by assuming $\mathbf{y} \in W^{1,\infty}(\Omega, \mathbb{R}^3)$ to be C^1 both in $\overline{\Omega}_A$ and $\overline{\Omega}_B$, with Ω_A, Ω_B two open disjoint subdomains of Ω , separated by a piecewise C^1 , possibly curved, 2-dimensional interface Γ such that $\Omega = \Omega_A \cup \Omega_B \cup \Gamma$. In the case of martensitic transformations, $\nabla \mathbf{y} = \mathbf{1}$ in Ω_A , while in Ω_B $\nabla \mathbf{y}$ represents a continuously varying macroscopic deformation gradient corresponding to a continuously varying martensitic microstructure. This result can be extended to $\mathbf{y} \in H^1(\Omega, \mathbb{R}^3)$ with $\nabla \mathbf{y} \in BV(\Omega; \mathbb{R}^{3 \times 3})$ as done in the two dimensional setting in [37], or more generally in Lemma 4.5.1 below. However, the deep result of [48] states that, in the case where $\mathbf{y} \in W^{1,\infty}(\mathbb{R}^3, \mathbb{R}^3)$, and $\mathbf{y}$ is constant in $\{\mathbf{x} \cdot \mathbf{n} > c\}$, the polyconvex hull of the set $\{\nabla \mathbf{y}(\mathbf{x}) : \mathbf{x} \in \{\mathbf{x} \cdot \mathbf{n} < c\}\}$, which contains $\{\nabla \mathbf{y}(\mathbf{x}) : \mathbf{x} \in \{\mathbf{x} \cdot \mathbf{n} < c\}\}^{qc}$, might not contain a matrix which is rank-one connected to $\mathbf{0}$, the deformation gradient in $\{\mathbf{x} \cdot \mathbf{n} > c\}$.

On the other hand, in order to fully capture the complex microstructures observed in $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$, we are interested in macroscopic deformation gradients that are just in $L^\infty(\Omega, \mathbb{R}^{3 \times 3})$. Therefore, in Section 4.4 we generalize Proposition 1.3.1 to non-constant deformation gradients in $L^\infty(\Omega; \mathbb{R}^{3 \times 3})$ and to curved interfaces. However, given the above mentioned counterexamples of [18, 48], we need to change perspective and introduce some further hypotheses. This is done by introducing the idea of moving mask in Chapter 4.

1.4 Twinning theory and the cofactor conditions

As explained in the previous section, the existence of a constant macroscopic martensitic deformation compatible with austenite, is related to the existence of a matrix $\mathbf{F} \in K^{qc}$ such that $\mathbf{F} = \mathbf{1} + \mathbf{a} \otimes \mathbf{n}$. Conditions on the deformation parameters under which such matrices exist have been first investigated in [18] in the case of two wells, i.e., $N = 2$, and then generalized in [12, 13]. The case where $N = 2$ is the most widely studied, as it is the only one for which an explicit characterization of K^{qc} is known, and turns out to be a fundamental tool to explain a wide range of experimental observations. Therefore, we now focus on the possibility of a pair of martensitic variants forming interfaces with austenite. The notation and results of this section follow closely those in [26].

Let us first recall that given two different variants of martensite, represented by $\mathbf{U}, \mathbf{V} \in \mathbb{R}_{Sym+}^{3 \times 3}$, there exists a rotation $\mathbf{R} \in SO(3)$ satisfying $\mathbf{V} = \mathbf{R}\mathbf{U}\mathbf{R}^T$. A first useful result is the following:

Proposition 1.4.1 ([26, Prop. 12]). *Let $\mathbf{U}, \mathbf{V} \in \mathbb{R}_{Sym^+}^{3 \times 3}$ with $\mathbf{U} \neq \mathbf{V}$. Suppose further that they are compatible in the sense that there is a matrix $\hat{\mathbf{R}} \in SO(3)$ such that*

$$\hat{\mathbf{R}}\mathbf{V} - \mathbf{U} = \mathbf{b} \otimes \mathbf{m}, \quad (1.10)$$

$\mathbf{b}, \mathbf{m} \in \mathbb{R}^3$. Then there is a unit vector $\hat{\mathbf{e}} \in \mathbb{R}^3$ such that

$$\mathbf{V} = (-\mathbf{1} + 2\hat{\mathbf{e}} \otimes \hat{\mathbf{e}})\mathbf{U}(-\mathbf{1} + 2\hat{\mathbf{e}} \otimes \hat{\mathbf{e}}). \quad (1.11)$$

Conversely, if (1.11) is satisfied, then there exist $\hat{\mathbf{R}} \in SO(3)$, $\mathbf{b}, \mathbf{m} \in \mathbb{R}^3$ such that (1.10) holds.

Equation (1.10) is called the compatibility condition for two variants of martensite; the solutions to this equation can be classified into three categories: compound, type I and type II twins. It is possible to prove that (see e.g., [23]), once $\mathbf{U}$ and $\mathbf{V}$ are given and (1.11) holds, the compatibility condition always has two solutions $(\hat{\mathbf{R}}_I, \mathbf{b}_I \otimes \mathbf{m}_I)$ and $(\hat{\mathbf{R}}_{II}, \mathbf{b}_{II} \otimes \mathbf{m}_{II})$. The solutions can be expressed as follows:

$$\text{type I} \quad \mathbf{m}_I = \hat{\mathbf{e}}, \quad \mathbf{b}_I = 2\left(\frac{\mathbf{U}^{-1}\hat{\mathbf{e}}}{|\mathbf{U}^{-1}\hat{\mathbf{e}}|^2} - \mathbf{U}\hat{\mathbf{e}}\right), \quad (1.12)$$

$$\text{type II} \quad \mathbf{m}_{II} = 2\left(\hat{\mathbf{e}} - \frac{\mathbf{U}^2\hat{\mathbf{e}}}{|\mathbf{U}\hat{\mathbf{e}}|^2}\right), \quad \mathbf{b}_{II} = \mathbf{U}\hat{\mathbf{e}}, \quad (1.13)$$

where $\hat{\mathbf{e}}$ is as in (1.11). If $\hat{\mathbf{e}}$ satisfying (1.11) is unique up to change of sign, we call $(\mathbf{U}, \mathbf{b}_I, \mathbf{m}_I)$ and $(\mathbf{U}, \mathbf{b}_{II}, \mathbf{m}_{II})$ respectively the type I and type II twin generated by $\mathbf{U}, \mathbf{V}$. In case there exist two different non-parallel unit vectors satisfying (1.11), the resulting pair of solutions (1.12)–(1.13) are called compound twins. Nonetheless, it is possible to prove (see e.g., [23]) that in the case of compound twins, given two different unit vectors satisfying (1.11), namely $\hat{\mathbf{e}}_1$ and $\hat{\mathbf{e}}_2$, then

$$\mathbf{b}_I^1 \otimes \mathbf{m}_I^1 := 2\left(\frac{\mathbf{U}^{-1}\hat{\mathbf{e}}_1}{|\mathbf{U}^{-1}\hat{\mathbf{e}}_1|^2} - \mathbf{U}\hat{\mathbf{e}}_1\right) \otimes \hat{\mathbf{e}}_1 = \mathbf{U}\hat{\mathbf{e}}_2 \otimes 2\left(\hat{\mathbf{e}}_2 - \frac{\mathbf{U}^2\hat{\mathbf{e}}_2}{|\mathbf{U}\hat{\mathbf{e}}_2|^2}\right) =: \mathbf{b}_{II}^2 \otimes \mathbf{m}_{II}^2.$$

Therefore, there are just two solutions to (1.11), even in the case of compound twins, each of which can be considered as both a type I and a type II twin. Below, however, when we refer to type I or type II solutions of (1.10) we assume implicitly that they are not compound solutions. Furthermore, we sometimes abuse the notation and write that $\mathbf{U}, \mathbf{V}$ generate a compound twin if the solutions of the twinning equations (1.10) are compound twins. The following characterization of compound twins is used below:

Proposition 1.4.2 ([26, Prop. 1]). *Let $\mathbf{U}$ and $\mathbf{V}$ be two different variants of martensite and $\hat{\mathbf{e}}_1$ a unit vector such that (1.11) is satisfied. Then there exists a second unit vector $\hat{\mathbf{e}}_2$ not parallel to $\hat{\mathbf{e}}_1$ satisfying (1.11) if and only if $\hat{\mathbf{e}}_1$ is perpendicular to an eigenvector of $\mathbf{U}$. When this condition is verified, $\hat{\mathbf{e}}_2$ is unique up to change of sign and is perpendicular to both $\hat{\mathbf{e}}_1$ and that eigenvector.*

Let us now consider a simple laminate, i.e., a constant macroscopic gradient $\nabla \mathbf{y}$ equal a.e. to $\lambda \hat{\mathbf{R}}\mathbf{V} + (1 - \lambda)\mathbf{U}$ for some $\lambda \in (0, 1)$ and some rank-one connected $\mathbf{R}\mathbf{U}$, $\mathbf{U} \in K$. Following [18, 26] we focus on the possibility for such $\nabla \mathbf{y}$ to be compatible with austenite. By Proposition 1.3.1, a necessary condition is that $SO(3)$ has a rank-one connection with $\lambda \hat{\mathbf{R}}\mathbf{V} + (1 - \lambda)\mathbf{U}$. The existence of $(\mathbf{R}, \lambda, \mathbf{a} \otimes \mathbf{n})$ solving

$$\mathbf{R}[\lambda \hat{\mathbf{R}}\mathbf{V} + (1 - \lambda)\mathbf{U}] - \mathbf{1} = \mathbf{R}[\lambda(\mathbf{U} + \mathbf{b} \otimes \mathbf{m}) + (1 - \lambda)\mathbf{U}] - \mathbf{1} = \mathbf{a} \otimes \mathbf{n}, \quad (1.14)$$

that is a twinned laminate compatible with austenite, was first studied in [82] and later in [18]. Lattice deformations and parameters of materials that are usually considered in the literature lead to twins with exactly four solutions to equation (1.14). Nonetheless, in some cases the number of solutions can be just zero, one or two, and, under some particular condition on the lattice parameters, as in the case of the material discovered in [76], (1.14) is satisfied for all $\lambda \in [0, 1]$. The following result gives necessary and sufficient conditions for this to hold:

Theorem 1.4.1 ([26, Thm. 2]). *Let $\mathbf{U}, \mathbf{V} \in \mathbb{R}_{Sym+}^{3 \times 3}$ be distinct and such that there exist $\hat{\mathbf{R}} \in SO(3)$ and $\mathbf{b}, \mathbf{m} \in \mathbb{R}^3 \setminus \{\mathbf{0}\}$ satisfying*

$$\hat{\mathbf{R}}\mathbf{V} = \mathbf{U} + \mathbf{b} \otimes \mathbf{m}.$$

*Then, (1.14) has a solution $\mathbf{R} \in SO(3)$, $\mathbf{a}, \mathbf{n} \in \mathbb{R}^3$ for each $\lambda \in [0, 1]$ if and only if the following **cofactor conditions** hold:*

(CC1) *The middle eigenvalue λ_2 of $\mathbf{U}$ satisfies $\lambda_2 = 1$,*

(CC2) $\mathbf{b} \cdot \mathbf{U} \operatorname{cof}(\mathbf{U}^2 - \mathbf{1})\mathbf{m} = 0$,

(CC3) $\operatorname{tr} \mathbf{U}^2 - \det \mathbf{U}^2 - \frac{1}{4}|\mathbf{b}|^2|\mathbf{m}|^2 - 2 \geq 0$.

The condition (CC2) can be rewritten as (see [26, Coroll. 5])

$$(\mathbf{v}_2 \cdot \mathbf{b})(\mathbf{v}_2 \cdot \mathbf{m}) = 0,$$

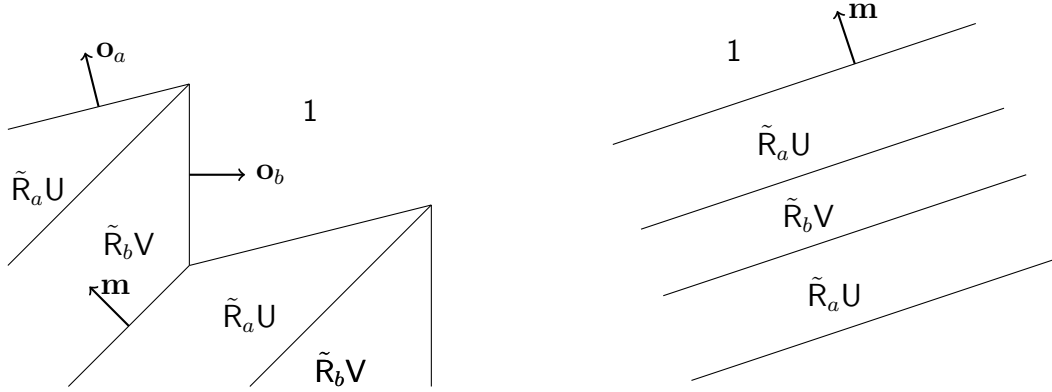


Figure 1.7: Example of triple junctions for type I and type II twins when the cofactor conditions are satisfied (see Theorem 1.4.2 and Theorem 1.4.3).

where $\mathbf{v}_2$ is the eigenvector of $\mathbf{U}$ related to the eigenvalue $\lambda_2 = 1$. Another equivalent formulation of (CC2) for type I/II twins is given by the following chains of equivalences (see [26, Prop. 6])

$$(CC2) \Leftrightarrow (\mathbf{v}_2 \cdot \mathbf{b}_I) = 0, \text{ and } (\mathbf{v}_2 \cdot \mathbf{m}_I) \neq 0 \Leftrightarrow |\mathbf{U}^{-1}\hat{\mathbf{e}}| = 1, \text{ for type I twins} \quad (1.15)$$

$$(CC2) \Leftrightarrow (\mathbf{v}_2 \cdot \mathbf{m}_{II}) = 0 \text{ and } (\mathbf{v}_2 \cdot \mathbf{b}_{II}) \neq 0 \Leftrightarrow |\mathbf{U}\hat{\mathbf{e}}| = 1, \text{ for type II twins} \quad (1.16)$$

where $\hat{\mathbf{e}}$ is given by (1.11).

Remark 1.4.1. Let $(\mathbf{U}, \mathbf{b}, \mathbf{m})$ be as in Theorem 1.4.1 and satisfy (CC1)–(CC3). Then, [26, Coroll. 3] states that the smallest and the largest eigenvalue of $\mathbf{U}$, namely λ_1 and λ_3 , satisfy $\lambda_1 < 1 < \lambda_3$.

We now report two results from [26] related to the cofactor conditions in type I/II twins. These results state the possibility to have exact stress free interface between a martensitic laminate and austenite as shown in Figure 1.7.

Theorem 1.4.2 ([26, Thm. 7]). Let $\mathbf{U}, \mathbf{V} \in \mathbb{R}_{Sym^+}^{3 \times 3}$ be distinct and such that $\hat{\mathbf{R}} \in SO(3)$, $\mathbf{b}_I, \mathbf{m}_I \in \mathbb{R}^3$ is a type I solution to (1.10). Suppose further that $\mathbf{U}, \mathbf{b}_I, \mathbf{m}_I$ satisfy the cofactor conditions. Then, there exist $\mathbf{R}_0 \in SO(3)$, $\mathbf{a}_0 \in \mathbb{R}^3$, $\mathbf{n}_0, \mathbf{n}_1 \in \mathbb{S}^2$ and $\xi \neq 0$ such that

$$\mathbf{R}_0 \mathbf{U} = \mathbf{1} + \mathbf{a}_0 \otimes \mathbf{n}_0, \quad \mathbf{R}_0(\mathbf{U} + \mathbf{b}_I \otimes \mathbf{m}_I) = \mathbf{1} + \mathbf{a}_0 \otimes \xi \mathbf{n}_1. \quad (1.17)$$

Furthermore,

$$\mathbf{R}_0[\mathbf{U} + \lambda \mathbf{b}_I \otimes \mathbf{m}_I] = \mathbf{1} + \mathbf{a}_0 \otimes (\lambda \xi \mathbf{n}_1 + (1 - \lambda) \mathbf{n}_0), \quad \text{for all } \lambda \in [0, 1]. \quad (1.18)$$

The three deformation gradients $\mathbf{1}$, $\mathbf{R}_0\mathbf{U}$, $\mathbf{R}_0\hat{\mathbf{R}}\mathbf{V}$ can form a compatible austenite-martensite triple junction in the sense that

$$\mathbf{R}_0\mathbf{U} - \mathbf{1} = \mathbf{a}_0 \otimes \mathbf{n}_0, \quad \mathbf{R}_0\hat{\mathbf{R}}\mathbf{V} - \mathbf{1} = \mathbf{a}_0 \otimes \xi\mathbf{n}_1, \quad \mathbf{R}_0\hat{\mathbf{R}}\mathbf{V} - \mathbf{R}_0\mathbf{U} = \mathbf{R}_0\mathbf{b}_I \otimes \mathbf{m}_I.$$

There is a constant $c \neq 0$ such that $c\mathbf{m}_I = \xi\mathbf{n}_1 - \mathbf{n}_0$, so that the three vectors $\mathbf{n}_0$, $\mathbf{n}_1$ and $\mathbf{m}_I$ lie in a plane.

Theorem 1.4.3 ([26, Thm. 8]). Let $\mathbf{U}, \mathbf{V} \in \mathbb{R}_{Sym^+}^{3 \times 3}$ be distinct and such that $\hat{\mathbf{R}} \in SO(3)$, $\mathbf{b}_{II}, \mathbf{m}_{II} \in \mathbb{R}^3$ is a type II solution to (1.10). Suppose further that $\mathbf{U}, \mathbf{b}_{II}, \mathbf{m}_{II}$ satisfy the cofactor conditions. Then, there exist $\mathbf{R}_0 \in SO(3)$, $\mathbf{a}_0, \mathbf{a}_1 \in \mathbb{R}^3$, $\mathbf{n}_0 \in \mathbb{S}^2$ and $\xi \neq 0$ such that

$$\mathbf{R}_0\mathbf{U} = \mathbf{1} + \mathbf{a}_0 \otimes \mathbf{n}_0, \quad \mathbf{R}_0(\mathbf{U} + \mathbf{b}_{II} \otimes \mathbf{m}_{II}) = \mathbf{1} + \xi\mathbf{a}_1 \otimes \mathbf{n}_0. \quad (1.19)$$

Furthermore,

$$\mathbf{R}_0[\mathbf{U} + \lambda\mathbf{b}_{II} \otimes \mathbf{m}_{II}] = \mathbf{1} + (\lambda\xi\mathbf{a}_1 + (1 - \lambda)\mathbf{a}_0) \otimes \mathbf{n}_0, \quad \text{for all } \lambda \in [0, 1]. \quad (1.20)$$

The normal $\mathbf{n}_0$ to the austenite-martensite interface is independent of the volume fraction λ and is parallel to the martensite-martensite interface normal $\mathbf{m}_{II}$.

Following Theorem 1.4.2 and Theorem 1.4.3, if a martensitic type I/II twin satisfies the cofactor conditions, then it can form exact phase interfaces, that is without a non stress-free transition layer, between austenite and a martensitic laminate, independently of the volume fractions within the laminate. Therefore, from a theoretical point of view, these materials can convert extremely easily a laminate of martensite into austenite and vice versa, without any elastic stress, and hence without the need to incur plastic deformation. The condition $\lambda_2 = 1$ guarantees already the possibility of having exact stress-free austenite-martensite phase interfaces. However, during martensitic transformations, nucleations occur often at different points of the sample, and, without further compatibility, the single variants of martensite cannot grow further after they have met [47].

1.5 Overview

A better understanding of the microstructures arising in martensitic transformations is an important step to capture the factors influencing reversibility. This

is particularly necessary in materials with six or twelve martensitic variants, like the ultra-reversible ones mentioned in Section 1.1, as the analytical predictions have been mostly developed for the case where just two variants are involved. Furthermore, even in the two variants case, convex integration can lead to infinitely many minimisers for the energy of the system, making the model unhelpful for predictions.

For this reason, the aim of this thesis is to better understand and characterise the microstructures arising in martensitic transformations, with a particular focus on $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$, the first alloy undergoing an ultra-reversible phase transition, and closely satisfying the cofactor conditions. Indeed, it looks striking and unusual that, in this material, microstructures look completely different at every thermal cycle, and are very complex, despite the possibility to reach stress-free phase compatibility at every scale with very simple deformation gradients.

This thesis is divided as follows:

- in Chapter 2 we study further the cofactor conditions from an analytical perspective. We also introduce a physically motivated metric to measure how closely a material satisfies the cofactor conditions, as the two currently used in the literature often give contradictory results, and are not physically grounded. In the same chapter, we show that $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$ not only closely satisfies the cofactor conditions with type II twins, but also some further compatibility between martensitic laminates, which improves the ability of this material to change phase.
- In Chapter 3 we introduce a new system of partial differential equations as a simplified model for the evolution of reversible martensitic transformations under thermal cycling in low hysteresis alloys. The model is developed in the context of nonlinear continuum mechanics, where the developed theory is mostly static, and cannot capture the influence of dynamics on martensitic microstructures. First, we prove existence of weak solutions; secondly, we study the physically relevant limit when the interface energy density vanishes, and the elastic constants tend to infinity. The limit problem provides a framework for the moving mask approximation introduced in Chapter 4. In the last section of Chapter 3 we study the limit equations in a one-dimensional setting. After closing the equations with a constitutive relation between the phase interface velocity and the temperature of the one-dimensional sample, the equations become a two-phase Stefan problem with a kinetic condition at the free boundary. Under

some further assumptions, we show that the phase interface reaches the domain boundary in finite time.

- In Chapter 4 we introduce a moving mask approximation, which is based on experimental observations in $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$ but which appears to be relevant also for thermally induced transformations in many other materials. Roughly, the moving mask approximation reflects the idea of a moving mask (or a moving curtain) continuously uncovering (or covering) a microstructure of martensite. Under this hypothesis, we can prove a new type of Hadamard jump condition, from which we deduce that the deformation gradient must be of the form

$$\mathbf{1} + \mathbf{a}(\mathbf{x}) \otimes \mathbf{n}(\mathbf{x}), \quad (1.21)$$

for almost every $\mathbf{x}$ in the martensite phase. This is useful to better understand the complex microstructures and the formation of curved interfaces between phases in new ultra-low hysteresis alloys such as $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$. In particular, we use the new type of Hadamard jump condition to deduce a rigidity theorem for compound twins, and to study curved phase interfaces in type I laminates satisfying the cofactor conditions. The rigidity result can be extended to general two well problems under some additional, physically relevant assumptions. Furthermore, by assuming the further compatibility between martensitic laminates introduced in Chapter 2, and that is closely satisfied by $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$, we can easily construct infinitely many macroscopic deformation gradients which are non-constant and satisfy (1.21). This is in agreement with experimental observations for $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$ (see [76]), where the microstructures drastically differ at every thermal cycle.

- In the last chapter we study the influence of interface energy in modelling microstructures. In particular, the aim of Chapter 5 is to show that interface energy can be used to select minimising sequences, and minimising gradient Young measures, in the case when these are infinitely many. We do this in a simplified one-dimensional setting for the energy functional

$$\mathcal{E}(u) = \int_0^1 (W(u_x) + u^2) \, dx,$$

where $W \geq 0$ has three wells, that is $W(s) = 0$ if and only if $s = z_1, z_2, z_3 \in \mathbb{R}$, and $z_1 < 0 < z_2 < z_3$. As $\mathcal{E}$ is not quasiconvex, minimisers do not necessarily exist. Furthermore, minimising sequence have gradients taking values close to

the three wells, but without any preference as to which of the three. Minimising sequences thus generate infinitely many gradient Young measures supported on $\{z_1, z_2, z_3\}$ without any preference on the volume fractions. We introduce a regularization

$$\mathcal{E}^\varepsilon = \varepsilon^6 \int_0^1 u_{xx}^2 \, dx + \mathcal{E}(u),$$

of $\mathcal{E}$ with an ε -small penalization of the second derivatives, and we obtain the first and, under some further assumptions, the second Γ -limit of $\mathcal{E}^\varepsilon$ as $\varepsilon \downarrow 0$. The latter selects a unique minimizing gradient Young measure of the former, which is supported just in two wells and not in three. We show that some assumptions are necessary to derive our second Γ -limit, but not to prove that minimizers of $\mathcal{E}^\varepsilon$ generate, as $\varepsilon \downarrow 0$, Young measures supported just in two wells and not in three.

Chapter 2

On the cofactor conditions and further conditions of super-compatibility between phases

The aim of this chapter is to study further the cofactor conditions, with a particular focus on cubic to monoclinic transformations, such as for $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$. The case of cubic to orthorhombic transformations, which is relevant for the new materials in [28, 45], is a special case of cubic to monoclinic, and hence all our results apply. In Section 2.1 we prove that once the cofactor conditions are satisfied by a type I/II twin, they are satisfied by all symmetry related twins. As a consequence, in the cubic to monoclinic case, if a twin satisfies the cofactor conditions, then eleven other different twins enjoy the same property (see Table 2.1). This generalises previous work in [26] (and Table 2 therein) where only three such twins were noted. The presence of multiple twins satisfying the cofactor conditions gives to a material many possibilities to change phase without inducing any elastic stress. We prove that the same pair of martensitic variants cannot satisfy at the same time the cofactor condition as a type I and a type II twin. This result is puzzling because, as discussed below and in Section 2.5, in $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$ the same pair of martensitic variants seems to satisfy very closely the cofactor conditions with both the type I and the type II twin generated by the same pair of martensitic variants. However, being close is a matter of metric, that is, of how we measure the cofactor conditions. Indeed, as the cofactor conditions can never be satisfied exactly by real materials, in Section 2.5 we discuss how these are measured in the literature, and introduce a new metric which we believe to be related to reversibility. We find that in $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$, the stress required to deform austenite in such a way that is exactly compatible (that is with no interface layer) with a laminate of martensite is very small for type II twins, and almost ten

times bigger for type I twins. Therefore, according to our new metric, it seems that $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$ satisfies the cofactor conditions with type II twins much better than with type I twins.

In Section 2.2 we study the possible homogeneous average deformations (also called constant macroscopic deformation gradients in the literature [19]) that can be obtained by finely mixing two unstressed martensitic variants, and we study which are compatible with austenite across a plane. Despite what one might think in a first instance, if the cofactor conditions are satisfied just by the type I (or just by the type II) twinning system the set of average deformation gradients which are compatible with austenite across a plane is of the same dimension as in standard shape-memory alloys.

In Section 2.3 we introduce a new condition of super-compatibility between phases which supplements the cofactor conditions. We call the twins satisfying these conditions *star twins*. If the cofactor conditions are satisfied by $(\mathbf{V}_1, \mathbf{b}_{12}, \mathbf{m}_{12})$, the type I/II twin generated by the two martensitic variants $\mathbf{V}_1, \mathbf{V}_2 \in \mathbb{R}_{\text{Sym}^+}^{3 \times 3}$ and where $\mathbf{b}_{12}, \mathbf{m}_{12} \in \mathbb{R}^3$, then Theorem 7 and Theorem 8 in [26] (reported here as Theorem 1.4.2 and Theorem 1.4.3) imply the existence of $\mathbf{a}_1, \mathbf{a}_2, \mathbf{n}_1, \mathbf{n}_2 \in \mathbb{R}^3$, and $\mathbf{R}_{12}, \hat{\mathbf{R}}_{12} \in SO(3)$ such that

$$\hat{\mathbf{R}}_{12}((1 - \mu)\mathbf{V}_1 + \mu\mathbf{R}_{12}\mathbf{V}_2) = \mathbf{1} + \mathbf{b}_{12} \otimes ((1 - \mu)\mathbf{n}_1 + \mu\mathbf{n}_2), \quad \text{for type I twins,} \quad (2.1)$$

$$\hat{\mathbf{R}}_{12}((1 - \mu)\mathbf{V}_1 + \mu\mathbf{R}_{12}\mathbf{V}_2) = \mathbf{1} + ((1 - \mu)\mathbf{a}_1 + \mu\mathbf{a}_2) \otimes \mathbf{m}_{12}, \quad \text{for type II twins,} \quad (2.2)$$

$$\mathbf{R}_{12}\mathbf{V}_2 - \mathbf{V}_1 = \mathbf{b}_{12} \otimes \mathbf{m}_{12}, \quad (2.3)$$

for every $\mu \in [0, 1]$. For any fixed $\mu_0 \in [0, 1]$ the average deformation gradient

$$\mathbf{F}_0 = \hat{\mathbf{R}}_{12}((1 - \mu_0)\mathbf{V}_1 + \mu_0\mathbf{R}_{12}\mathbf{V}_2), \quad (2.4)$$

is compatible with austenite without an interface layer (cf. Figure 1.7), and can hence easily propagate in austenite. We call such a constructed $\mathbf{F}_0$ an *exactly compatible* laminate generated by $\mathbf{V}_1, \mathbf{V}_2$, of volume fraction μ_0 . In cubic to monoclinic phase transitions, in general, given $\mathbf{F}_0$ as in (2.4), there exists no exactly compatible laminate $\mathbf{F}_1$ generated by the martensitic variants $\mathbf{V}_3, \mathbf{V}_4 \in \mathbb{R}_{\text{Sym}^+}^{3 \times 3}$, of volume fraction $\mu_1 \in [0, 1]$, and such that

$$\text{rank}(\mathbf{F}_0 - \mathbf{F}_1) = 1, \quad (2.5)$$

that is, such that $\mathbf{F}_0$ and $\mathbf{F}_1$ are different but compatible on average. Such compatibility is possible only for specific values of μ_0, μ_1 . However, in general, given $\mathbf{F}_0$ as in

(2.4) there exists at most one exactly compatible laminate F_1 such that (2.5) is satisfied. Exceptions to this fact occur when $(V_1, \mathbf{b}_{12}, \mathbf{m}_{12})$ is a star-twin (or a half-star twin). In this case, there exist values of μ_0 such that, given F_0 as in (2.4), there exist three exactly compatible laminates $F_1, F_2, F_3 \in \mathbb{R}^{3 \times 3}$ (resp. two exactly compatible laminates in the case of half-star twins) with

$$\text{rank}(F_i - F_j) = 1, \quad \text{for every } i, j = 0, 1, 2, 3. \quad (2.6)$$

Furthermore, any three of F_0, F_1, F_2, F_3 are linearly independent (see Figure 2.1). For general cubic to monoclinic transformations there are four deformation parameters determining the transformation strain (cf. a, b, c, d in (2.8)). The cofactor conditions impose two relations between these four parameters, so that there is a two-parameter family that satisfies them. For the existence of star twins a further relation has to be satisfied, reducing the set of possible deformation parameters to a one-dimensional family. The compatibility between different laminates which form an exact interface with austenite is only on average. Nevertheless, this allows three different laminates, nucleated in different regions of the sample, to grow further after they meet, and not to stop due to incompatibility. Also, as emphasised in Remark 2.3.3, in the presence of type II star twins macroscopically curved interfaces whose normal does not lie in a plane are possible between austenite and martensite without an interface layer (cf. also Figure 2.3 and Figure 2.4).

In Section 2.4 we show under which conditions on the eigenvalues of the transformation matrices for cubic to monoclinic phase transitions a twin satisfying the cofactor conditions is actually a star twin. It is striking to notice that in $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$, the smallest eigenvalue is approximately 0.9363 and the largest is 1.0600. In order to have type II star twin, the largest should be 1.0609, so that the error is about $9 \cdot 10^{-4}$, very similar to the approximate error of $6 \cdot 10^{-4}$ which separates λ_2 from 1 in this material. Denoting by U_1^r the measured deformation gradient in $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$ related to the first martensitic variant (see (2.8)), and by U_1^e the same matrix, but allowing a type II star twin, we have

$$U_1^r = \begin{bmatrix} 1.0015 & 0.0073 & 0 \\ 0.0073 & 1.0591 & 0 \\ 0 & 0 & 0.9363 \end{bmatrix}, \quad U_1^e = \begin{bmatrix} 1.0010 & 0.0078 & 0 \\ 0.0078 & 1.0594 & 0 \\ 0 & 0 & 0.9368 \end{bmatrix},$$

which are very close as $\|U_1^r - U_1^e\| \approx 1.1 \cdot 10^{-3}$. As a comparison, the closest matrix to U_1^r , say U_1^c , satisfying the cofactor conditions and describing a cubic to monoclinic phase transformation is such that $\|U_1^r - U_1^c\| \approx 0.9 \cdot 10^{-3}$. Experimental images of

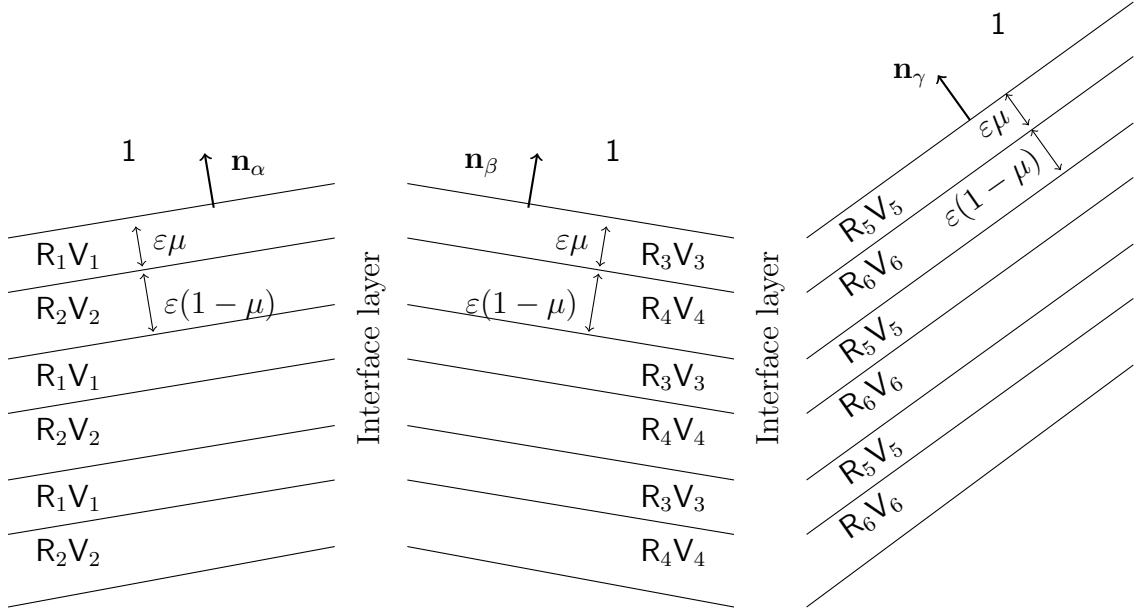


Figure 2.1: Example of average compatibility between three different laminates occurring in type II star twins. Here $R_1, \dots, R_6 \in SO(3)$ and $V_1, \dots, V_6 \in \mathbb{R}_{Sym^+}^{3 \times 3}$ are six deformation gradients corresponding to six martensitic variants (possibly $V_i = V_j$ for $i \neq j$). We have $R_i V_i = 1 + \mathbf{a}_i \otimes \mathbf{n}_\alpha$ for $i = 1, 2$, $R_i V_i = 1 + \mathbf{a}_i \otimes \mathbf{n}_\beta$ for $i = 3, 4$, and $R_i V_i = 1 + \mathbf{a}_i \otimes \mathbf{n}_\gamma$ for $i = 5, 6$, for some $\mathbf{a}_1, \dots, \mathbf{a}_6 \in \mathbb{R}^3$, $\mathbf{n}_\alpha, \mathbf{n}_\beta, \mathbf{n}_\gamma \in \mathbb{S}^2$. Furthermore, there exist $\mu \in (0, 1)$, $\mathbf{a}^* \in \mathbb{R}^3$ such that $\mu R_1 V_1 + (1 - \mu) R_2 V_2 = 1 + \mathbf{a}^* \otimes \mathbf{n}_\alpha$, $\mu R_3 V_3 + (1 - \mu) R_4 V_4 = 1 + \mathbf{a}^* \otimes \mathbf{n}_\beta$ and $\mu R_5 V_5 + (1 - \mu) R_6 V_6 = 1 + \mathbf{a}^* \otimes \mathbf{n}_\gamma$, where $\text{rank}(R_{2i} V_{2i} - R_{2i-1} V_{2i-1}) = 1$ for $i = 1, 2, 3$. In the centre we have a transition layer which makes the two laminates exactly compatible in the limit $\varepsilon \rightarrow 0$. This image is a projection of a plane, but we remark that $\mathbf{n}_\alpha, \mathbf{n}_\beta, \mathbf{n}_\gamma$ are linearly independent.

martensitic microstructures for $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$ show the presence of inexact junctions between three different laminates (see Figure 2.2). The influence of star twins on these microstructures needs however to be confirmed with further experimental investigations. We remark that in $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$ the type I twins are not close to being

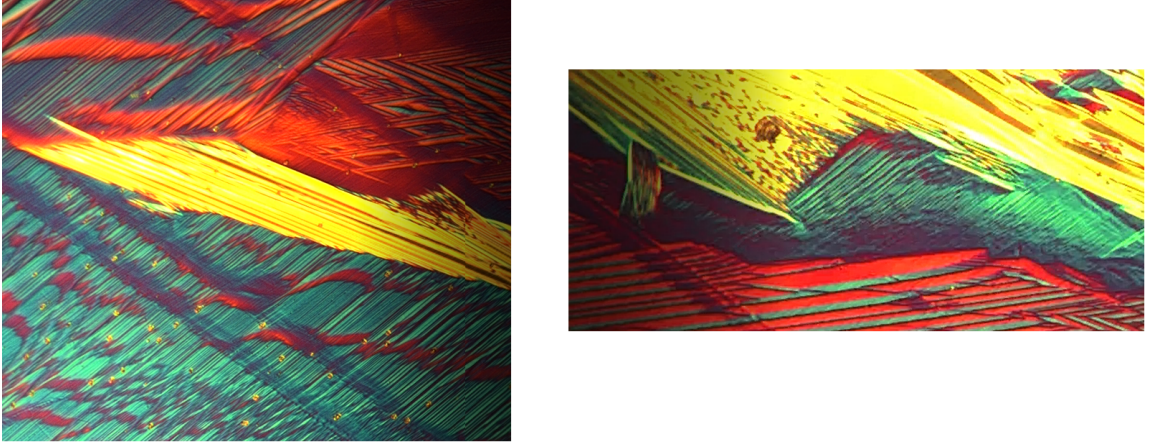


Figure 2.2: Examples of junctions between three different laminates observed in $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$ (cf. Figure 2.1). Courtesy of Xian Chen, Yintao Song and Richard James.

star twins. As proved in Chapter 4, in a first approximation the average deformation gradients in the martensite phase of $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$ are of the form

$$\mathbf{1} + \mathbf{a}(\mathbf{x}) \otimes \mathbf{n}(\mathbf{x}), \quad \text{a.e. } \mathbf{x}. \quad (2.7)$$

The presence of star twins makes it extremely easy to construct very complex average deformation gradients of the form (2.7), and this might explain the presence of such colourful microstructures, as much as the ability of this material to “*perform a much wider and more efficient collection of adjustments of microstructure to environmental changes*”(ref. [50]).

2.1 Martensitic transformations from cubic to monoclinic lattices

We start this section by proving Proposition 2.1.1 below, stating that the presence of a twin system satisfying the cofactor conditions (CC1)–(CC3) often implies the existence of many other twin systems enjoying the same property

Proposition 2.1.1. *Let $U, V \in \mathbb{R}_{Sym^+}^{3 \times 3}$ satisfying (1.11) be such that $R \in SO(3)$, $\mathbf{b} \in \mathbb{R}^3$, $\mathbf{m} \in \mathbb{S}^2$ is a solution of the twinning equation (1.10). Suppose further that $(U, \mathbf{b}, \mathbf{m})$ satisfies the cofactor conditions. Then, for every $Q \in SO(3)$, $QUQ^T, QVQ^T \in \mathbb{R}_{Sym^+}^{3 \times 3}$ are such that $QRQ^T \in SO(3)$, $Q\mathbf{b} \in \mathbb{R}^3$, $Q\mathbf{m} \in \mathbb{S}^2$ is a solution of the twinning equation (1.10) for QUQ^T, QVQ^T , and $(Q^T U Q, Q\mathbf{b}, Q\mathbf{m})$ satisfies the cofactor conditions. Furthermore, if $(R, \mathbf{b}, \mathbf{m})$ is a type I solution (or a type II or a compound solution) to the twinning equation (1.10), then so is $(QRQ^T, Q\mathbf{b}, Q\mathbf{m})$.*

Proof. Multiplying (1.10) by Q on the left, and by Q^T on the right we get

$$QRQ^T QUQ^T - QVQ^T = (Q\mathbf{b}) \otimes Q\mathbf{m}.$$

Furthermore, (1.11) becomes

$$QVQ^T = (-1 + 2(Q\hat{\mathbf{e}} \otimes (Q\hat{\mathbf{e}}))QUQ^T(-1 + 2(Q\hat{\mathbf{e}} \otimes (Q\hat{\mathbf{e}})).$$

Therefore, if $(R, \mathbf{b}, \mathbf{m})$ was a type I (or a type II) solution of the twinning equation (1.10) for U, V , then so is $(QRQ^T, Q\mathbf{b}, Q\mathbf{m})$ for QUQ^T, QVQ^T . If the solutions of (1.10) for U, V are compound twins, then so are the ones for QUQ^T, QVQ^T . Also, it is clear that (CC1) and (CC3) are satisfied.

It just remains to prove that (CC2) holds. But (CC2) can be rewritten as

$$(\lambda_1^2 - 1)(\lambda_3^2 - 1)(\mathbf{b} \cdot \mathbf{v}_2)(\mathbf{m} \cdot \mathbf{v}_2) = 0.$$

Here $\lambda_1 \leq \lambda_2 \leq \lambda_3$ are the eigenvalues of U , and $\mathbf{v}_2$ is the eigenvector of U related to the eigenvalue $\lambda_2 = 1$. But as $QUQ^T Q\mathbf{v}_2 = Q\mathbf{v}_2$, we get that (CC2) reduces to

$$(\lambda_1^2 - 1)(\lambda_3^2 - 1)(Q\mathbf{b} \cdot Q\mathbf{v}_2)(Q\mathbf{m} \cdot Q\mathbf{v}_2) = 0,$$

which concludes the proof. $\square$

In case of cubic to monoclinic transformations, the twelve transformation matrices are given by

$$\begin{aligned} U_1 &= \begin{bmatrix} a & b & 0 \\ b & c & 0 \\ 0 & 0 & d \end{bmatrix}, \quad U_2 = \begin{bmatrix} a & -b & 0 \\ -b & c & 0 \\ 0 & 0 & d \end{bmatrix}, \quad U_3 = \begin{bmatrix} c & b & 0 \\ b & a & 0 \\ 0 & 0 & d \end{bmatrix}, \quad U_4 = \begin{bmatrix} c & -b & 0 \\ -b & a & 0 \\ 0 & 0 & d \end{bmatrix}, \\ U_5 &= \begin{bmatrix} a & 0 & b \\ 0 & d & 0 \\ b & 0 & c \end{bmatrix}, \quad U_6 = \begin{bmatrix} a & 0 & -b \\ 0 & d & 0 \\ -b & 0 & c \end{bmatrix}, \quad U_7 = \begin{bmatrix} c & 0 & b \\ 0 & d & 0 \\ b & 0 & a \end{bmatrix}, \quad U_8 = \begin{bmatrix} c & 0 & -b \\ 0 & d & 0 \\ -b & 0 & a \end{bmatrix}, \\ U_9 &= \begin{bmatrix} d & 0 & 0 \\ 0 & a & b \\ 0 & b & c \end{bmatrix}, \quad U_{10} = \begin{bmatrix} d & 0 & 0 \\ 0 & a & -b \\ 0 & -b & c \end{bmatrix}, \quad U_{11} = \begin{bmatrix} d & 0 & 0 \\ 0 & c & b \\ 0 & b & a \end{bmatrix}, \quad U_{12} = \begin{bmatrix} d & 0 & 0 \\ 0 & c & -b \\ 0 & -b & a \end{bmatrix}. \end{aligned} \tag{2.8}$$

$R(\theta, \mathbf{v})$	type I/II twins (A)	type I/II twins (B)	compound twins (C)
$\pi, (1, 0, 0)$			$(1, 2), (5, 6)$
$\pi, (1, 0, 0)$			$(3, 4), (7, 8)$
$\pi, (0, 1, 0)$			$(1, 2), (11, 12)$
$\pi, (0, 1, 0)$			$(3, 4), (9, 10)$
$\pi, (0, 0, 1)$			$(5, 6), (9, 10)$
$\pi, (0, 0, 1)$			$(11, 12), (7, 8)$
$\pi, (1, 0, 1)$	$(1, 12), (2, 11)$	$(3, 10), (4, 9)$	$(5, 7), (6, 8)$
$\pi, (1, 0, -1)$	$(1, 11), (2, 12)$	$(3, 9), (4, 10)$	$(5, 7), (6, 8)$
$\pi, (1, 1, 0)$	$(5, 10), (6, 9)$	$(8, 11), (7, 12)$	$(1, 3), (2, 4)$
$\pi, (1, -1, 0)$	$(5, 9), (6, 10)$	$(7, 11), (8, 12)$	$(1, 3), (2, 4)$
$\pi, (0, 1, 1)$	$(3, 8), (4, 7)$	$(1, 6), (2, 5)$	$(9, 11), (10, 12)$
$\pi, (0, -1, 1)$	$(3, 7), (4, 8)$	$(1, 5), (2, 6)$	$(9, 11), (10, 12)$
$\frac{\pi}{2}, (0, 1, 0)$	$(1, 12), (2, 11)$	$(3, 10), (4, 9)$	$(5, 8), (6, 7)$
$-\frac{\pi}{2}, (0, 1, 0)$	$(1, 11), (2, 12)$	$(3, 9), (4, 10)$	$(5, 8), (6, 7)$
$\frac{\pi}{2}, (0, 0, 1)$	$(5, 9), (6, 10)$	$(7, 11), (8, 12)$	$(1, 4), (2, 3)$
$-\frac{\pi}{2}, (0, 0, 1)$	$(5, 10), (6, 9)$	$(8, 11), (7, 12)$	$(1, 4), (2, 3)$
$\frac{\pi}{2}, (1, 0, 0)$	$(3, 7), (4, 8)$	$(1, 5), (2, 6)$	$(10, 11), (9, 12)$
$-\frac{\pi}{2}, (1, 0, 0)$	$(3, 8), (4, 7)$	$(1, 6), (2, 5)$	$(10, 11), (9, 12)$

Table 2.1: Table containing all possible twinning systems for cubic to monoclinic transformations.

Following [26, Table 1], in Table 2.1 we listed all the possible twinning systems for cubic to monoclinic transformations, denoting by (i, j) the twins generated by $\mathbf{U}_i, \mathbf{U}_j$. The angle and axis of the rotation $\mathbf{R} \in SO(3)$ such that $\mathbf{U}_j = \mathbf{R}\mathbf{U}_i\mathbf{R}^T$ are given in the first column of the table.

Thanks to Proposition 2.1.1, if a pair of variants in column (A) (or column (B)) generates a twin which satisfies the cofactor conditions, then all the pairs in the same column, that is column (A) (resp. column (B)) satisfy the cofactor conditions for the same type of twin solution. The situation of column (C) is more complex: if a pair of variants generates a twin which satisfies the cofactor conditions, then all possible compound-twins solutions generated by pairs of variants in column (C) satisfy (CC1)–(CC2). However, in general, not all possible compound-twin solutions satisfy (CC3). This is the case for example for

$$\mathbf{U}_1 = \begin{bmatrix} 0.95 & 0.0073 & 0 \\ 0.0073 & 1.2 & 0 \\ 0 & 0 & 1 \end{bmatrix}.$$

In this case, (CC3) is satisfied by both compound solutions generated by $\mathbf{U}_1, \mathbf{U}_2$, but by neither of the two compound solutions generated by $\mathbf{U}_1, \mathbf{U}_3$. Nevertheless, Proposition 2.1.1 implies that if a pair of variants in the upper box of column (C) (or in

the lower box of column (C)) generates a twin satisfying the cofactor conditions, then all other pairs of variants in the same box generate a twin satisfying the cofactor conditions.

The difference from the table in [26] is that we passed from thirteen boxes of this type to four, that is, we have shown that if a twin satisfies the cofactor conditions, there are many different twins enjoying the same property. Furthermore, also thanks to Remark 2.2.1 below, we have clarified the situation of compound twins, which was not completely investigated in [26].

In case of cubic to orthorhombic transformations, the deformation gradients related to the six martensitic variants are

$$\begin{aligned} \tilde{U}_1 &= \begin{bmatrix} a & b & 0 \\ b & a & 0 \\ 0 & 0 & d \end{bmatrix}, \quad \tilde{U}_2 = \begin{bmatrix} a & -b & 0 \\ -b & a & 0 \\ 0 & 0 & d \end{bmatrix}, \quad \tilde{U}_3 = \begin{bmatrix} a & 0 & b \\ 0 & d & 0 \\ b & 0 & a \end{bmatrix}, \\ \tilde{U}_4 &= \begin{bmatrix} a & 0 & -b \\ 0 & d & 0 \\ -b & 0 & a \end{bmatrix}, \quad \tilde{U}_5 = \begin{bmatrix} d & 0 & 0 \\ 0 & a & b \\ 0 & b & a \end{bmatrix}, \quad \tilde{U}_6 = \begin{bmatrix} d & 0 & 0 \\ 0 & a & -b \\ 0 & -b & a \end{bmatrix}. \end{aligned} \quad (2.9)$$

Furthermore, Table 2.1 simplifies to Table 2.2. In the following proposition, we

$R(\theta, \mathbf{v})$	type I/II twins	compound twins
$\pi, (1, 0, 0)$		$(1, 2), (3, 4)$
$\pi, (0, 1, 0)$		$(1, 2), (5, 6)$
$\pi, (0, 0, 1)$		$(3, 4), (5, 6)$
$\pi, (1, 0, 1)$	$(1, 6), (2, 5)$	
$\pi, (1, 0, -1)$	$(1, 5), (2, 6)$	
$\pi, (1, 1, 0)$	$(3, 6), (4, 5)$	
$\pi, (1, -1, 0)$	$(3, 5), (4, 6)$	
$\pi, (0, 1, 1)$	$(1, 4), (2, 3)$	
$\pi, (0, -1, 1)$	$(1, 3), (2, 4)$	

Table 2.2: Table containing all possible twinning systems for cubic to orthorhombic transformations.

prove that in materials undergoing cubic to monoclinic or cubic to orthorhombic transformations, the two martensitic variants cannot satisfy the cofactor conditions both with the type I and with the type II twinning systems.

Proposition 2.1.2. *Let $U_i, U_j \in \mathbb{R}_{Sym+}^{3 \times 3}$ be as in (2.8) for some $i, j = 1, \dots, 12$, and $U_i \neq U_j$. Let us suppose also that $\hat{\mathbf{e}} \in \mathbb{S}^2$ satisfying*

$$U_j = (-1 + 2\hat{\mathbf{e}} \otimes \hat{\mathbf{e}})U_i(-1 + 2\hat{\mathbf{e}} \otimes \hat{\mathbf{e}})$$

is unique up to a change of sign. Then the type I and type II solutions of the twinning equation (1.10) for $\mathbf{U}_i, \mathbf{U}_j$ cannot satisfy (CC1)–(CC2) at the same time.

Proof. By Corollary 2.1.1 (see also Table 2.1) we can restrict ourselves to checking just two cases: $\mathbf{U}_i = \mathbf{U}_1, \mathbf{U}_j = \mathbf{U}_5$ and $\mathbf{U}_i = \mathbf{U}_1, \mathbf{U}_j = \mathbf{U}_{11}$. We focus on the latter, as the former can be deduced similarly.

We now assume that both the type I and the type II twins generated by $\mathbf{U}_1, \mathbf{U}_{11}$ satisfy (CC1)–(CC2), and we aim at a contradiction. Suppose first $d \neq 1$. By (1.16) we have that a type II twin satisfies (CC1)–(CC2) if and only if the middle eigenvalue λ_2 of $\mathbf{U}_1$ satisfies $\lambda_2 = 1$ and $|\mathbf{U}_1 \mathbf{e}| = 1$ with $\mathbf{e} = 2^{-\frac{1}{2}}(1, 0, 1)^T$. That is, we must satisfy at the same time

$$a + c = 1 + \lambda, \quad ac - b^2 = \lambda, \quad (2.10)$$

$$a^2 + b^2 + d^2 = 2. \quad (2.11)$$

Here λ is the eigenvalue of $\mathbf{U}_1$ which is neither 1 nor d . At the same time, if the type I twin generated by $\mathbf{U}_1, \mathbf{U}_{11}$ satisfies (CC2), by (1.16) $|\mathbf{U}_1^{-1} \mathbf{e}| = 1$ and hence

$$\frac{c^2}{\lambda^2} + \frac{b^2}{\lambda^2} + \frac{1}{d^2} = 2. \quad (2.12)$$

From (2.10)–(2.11) we deduce

$$(1 + \lambda)^2 - c(1 + \lambda) = 2 - d^2 + \lambda.$$

On the other hand, (2.10) together with (2.12) imply

$$c(1 + \lambda) = \lambda^2(2 - d^{-2}) + \lambda. \quad (2.13)$$

Collecting the last two identities to get rid of c we get

$$(1 + \lambda)^2 - 2 + d^2 = \lambda^2(2 - d^{-2}) + 2\lambda,$$

which can be rewritten as

$$\frac{\lambda^2}{d^2}(1 - d^2) = (1 - d^2).$$

Having assumed $d \neq 1$, and keeping in account that $\mathbf{U}_1$ is positive definite, we must have $d = \lambda$. However, as the middle eigenvalue of $\mathbf{U}_1$ has to be equal to one (cf. (CC1)), we deduce that $d = \lambda = 1$, which implies $a = 1, c = 1, b = 0$, that is $\mathbf{U}_1 = \mathbf{U}_{11}$. Therefore we reached a contradiction.

Suppose now $d = 1$, then $|\mathbf{U}_1^{-1} \mathbf{e}| = 1, |\mathbf{U}_1 \mathbf{e}| = 1$ reduce to solve

$$a^2 + b^2 = 1, \quad c^2 + b^2 = (ac - b^2)^2.$$

But these imply either $a = -c$, thus contradicting the fact that $\mathbf{U}_1$ is positive definite, or $a = \pm 1, b = 0$, which in turn contradict being positive definite or the fact that $\mathbf{U}_1 \neq \mathbf{U}_{11}$. $\square$

Remark 2.1.1. Following the strategy of Proposition 2.1.2 we can actually prove that in the pure cubic to monoclinic case (that is when $a \neq c$ in (2.8)) it is not possible to satisfy (CC1)–(CC2) with the type I twins (or with the type II twin) of both columns (A) and (B) of Table 2.1 at the same time. In the degenerate case where $a = c$ in (2.8), that is for cubic to orthorhombic transformations, column (A) and column (B) coincide (see Table 2.2).

Despite this result, it is easy to construct examples of matrices $\mathbf{U}_i$ as in (2.8), such that the cofactor conditions are satisfied by the type I twins in column (A) of Table 2.1 (or column (B)) and by the type II twins in column (B) (resp. column (A)) of Table 2.1. It is also possible to construct matrices $\mathbf{U}_i$ as in (2.8), such that $d = 1$, and the cofactor conditions are satisfied by some compound twins in column (C) and by the type I/II twins in column (A)/(B) of Table 2.1. If $d = 1$ and the smallest and the largest eigenvalues of the $\mathbf{U}_i$'s, namely λ_1, λ_3 , satisfy $\lambda_1 = \lambda_3^{-1}$ it is possible to satisfy the cofactor conditions with type I, type II and compound twins at the same time.

2.2 Cofactor conditions for two wells

In this section we study the quasiconvex hull of the set

$$K_{ij} := SO(3)\mathbf{U}_i \cup SO(3)\mathbf{U}_j,$$

where $\mathbf{U}_i \neq \mathbf{U}_j$ are as in (2.8). We recall that the set of matrices K_{ij}^{qc} , that is the quasiconvex hull of K_{ij} , is the set of average deformation gradients which can be achieved by finely and homogeneously mixing the two variants of martensite $\mathbf{U}_i, \mathbf{U}_j$ (see e.g., [19, 62]). If $\mathbf{U}_i, \mathbf{U}_j$ generate a compound twin we show that a necessary condition to satisfy the cofactor conditions is to have $d = 1$. Furthermore, we show that if $\mathbf{U}_i$ has a middle eigenvalue which is equal to one, but $d \neq 1$, then the only constant average deformation gradients in K_{ij}^{qc} which are rank one connected to the identity, and are hence compatible with austenite, are the pure phases, that is $\mathbf{R}_1\mathbf{U}_i, \mathbf{R}_2\mathbf{U}_i, \mathbf{R}_3\mathbf{U}_j, \mathbf{R}_4\mathbf{U}_j$ for some $\mathbf{R}_1, \mathbf{R}_2, \mathbf{R}_3, \mathbf{R}_4$. On the other hand, if $\mathbf{U}_i, \mathbf{U}_j$ generate a type I/II solution of the twinning equation (1.10), say $(\mathbf{R}, \mathbf{b}, \mathbf{m})$, satisfying the cofactor conditions, then there exist two smooth functions $\mathbf{R}_1, \mathbf{R}_2 : [0, 1] \rightarrow SO(3)$ such that the only matrices

in K_{ij}^{qc} rank one connected to the identity are given by $R_k(\mu)(U_i + \mu \mathbf{b} \otimes \mathbf{m})$, with $\mu \in [0, 1]$ and $k = 1, 2$.

2.2.1 Compound twins

We start by recalling the following result that characterizes the quasiconvex hull of a two well problem in some simple case

Lemma 2.2.1 ([36, Thm 2.5.1]). *Let $A, B \in \mathbb{R}_{Sym+}^{3 \times 3}$ with $\det A = \det B = \mathfrak{D} > 0$. Suppose that there exists $\lambda > 0$ and $\mathbf{v} \in \mathbb{S}^2$ such that $A\mathbf{v} = B\mathbf{v} = \lambda\mathbf{v}$. Then, defined $\hat{K}$ as*

$$\hat{K} := SO(3)A \cup SO(3)B,$$

we have

$$\hat{K}^{qc} = \left\{ F \in \mathbb{R}^{3 \times 3} : \det(F) = \mathfrak{D}, F^T F \mathbf{v} = \lambda^2 \mathbf{v}, |\mathbf{F}\mathbf{e}| \leq \max\{|\mathbf{A}\mathbf{e}|, |\mathbf{B}\mathbf{e}|\}, \text{ for each } \mathbf{e} \in \mathbb{S}^2 \right\}.$$

Also, a consequence of the proof of Lemma 4.5.2 in Chapter 4 is the following result

Lemma 2.2.2. *Let $A, B \in \mathbb{R}_{Sym+}^{3 \times 3}$ be such that $A\mathbf{e}_3 = B\mathbf{e}_3 = \lambda\mathbf{e}_3$ for some $\lambda \neq 1, \lambda > 0$. Assume also $\det(A) = \det(B) = \mathfrak{D} > 0$. Then, necessary conditions for the existence of $\mathbf{a} \in \mathbb{R}^3, \mathbf{n} \in \mathbb{S}^2$ satisfying $1 + \mathbf{a} \otimes \mathbf{n} \in K^{qc}$ are*

$$|\mathbf{a}| = \lambda^{-1}|\mathfrak{D} - \lambda^2|, \quad n_3^2 = \frac{\lambda^2(1 - \lambda^2)}{\mathfrak{D}^2 - \lambda^4}$$

and

$$\mathbf{a} = (\mathfrak{D} - \lambda^2)\mathbf{n} - \frac{1}{n_3}(1 - \lambda^2)\mathbf{e}_3.$$

Now let us consider the cubic to monoclinic transformation, and let U_i with $i = 1, \dots, 12$ be as in (2.8). We can prove the following statement

Proposition 2.2.1. *Let $U_1 \in \mathbb{R}_{Sym+}^{3 \times 3}$ be as in (2.8), with $a, c, d > 0, b \geq 0$, and $d \neq 1$. Let also $0 < \lambda_1 \leq \lambda_2 = 1 \leq \lambda_3$ be the eigenvalues of U_1 . Then, for every $i, j \in \{1, \dots, 4\}$ (or, equivalently, $i, j \in \{5, \dots, 8\}$, or $i, j \in \{9, \dots, 12\}$) with $U_i \neq U_j$, there exist exactly four matrices $\mathbf{a}^l \otimes \mathbf{n}^l, l = 1, \dots, 4$ satisfying*

$$1 + \mathbf{a}^l \otimes \mathbf{n}^l \in \left(SO(3)U_i \cup SO(3)U_j \right)^{qc}. \quad (2.14)$$

Furthermore, there exist four rotation matrices $R^l, l = 1, \dots, 4$ such that

$$R^l U_i = 1 + \mathbf{a}^l \otimes \mathbf{n}^l, \quad l = 1, 2 \quad (2.15)$$

$$R^l U_j = 1 + \mathbf{a}^l \otimes \mathbf{n}^l, \quad l = 3, 4. \quad (2.16)$$

Proof. Let us deal with the case $i = 1, j = 4$: the other cases can be treated in a similar way.

On the one hand, given the assumptions on $\mathbf{U}_1, \mathbf{U}_4$, [18, Prop. 4] assures the existence of four matrices $\mathbf{a}^l \otimes \mathbf{n}^l \in \mathbb{R}^{3 \times 3}$, and four rotations $\mathbf{R}^l \in SO(3)$, $l = 1, \dots, 4$, satisfying (2.15)–(2.16).

On the other hand, from Lemma 2.2.2 we know that the first two components of $\mathbf{n}$, namely n_1, n_2 , must lie on a circle

$$n_1^2 + n_2^2 = 1 - n_3^2 = \frac{\mathfrak{D}^2 - d^2}{\mathfrak{D}^2 - d^4}, \quad (2.17)$$

where $\mathfrak{D} = \det \mathbf{U}_1$, and that

$$\mathbf{F} := \mathbf{1} + \mathbf{a} \otimes \mathbf{n} = \mathbf{1} + (\mathfrak{D} - d)\mathbf{n} \otimes \mathbf{n} - \frac{1}{n_3}(1 - d^2)\mathbf{e}_3 \otimes \mathbf{n}. \quad (2.18)$$

Let us now look for unit vectors $\mathbf{v} = (v_1, v_2, 0)$ satisfying $|\mathbf{U}_1 \mathbf{v}| = |\mathbf{U}_4 \mathbf{v}|$, that is

$$(av_1 + bv_2)^2 + (bv_1 + cv_2)^2 = (cv_1 - bv_2)^2 + (-bv_1 + av_2)^2.$$

There are up to a change of sign two such unit vectors $\mathbf{v}^+, \mathbf{v}^-$ satisfying the above identity. Respectively they are such that

$$v_1^+(a - c) = v_2^+(2 - a - c - 2b), \quad v_1^-(a - c) = v_2^-(a + c - 2b - 2), \quad v_1^+ v_1^- = -v_2^+ v_2^-.$$

Here we have used the fact that, as $\lambda_2 = 1$, $ac - b^2 = \lambda$ and $a + c = 1 + \lambda$, with $\lambda = \lambda_1$ if $\lambda_3 = d$ and $\lambda = \lambda_3$ if $\lambda_1 = d$. We remark also that, $v_1^+ v_1^- = -v_2^+ v_2^-$ together with $|\mathbf{v}^+| = |\mathbf{v}^-| = 1$ imply $v_1^- = -v_2^+$ and $v_2^- = v_1^+$. Now, Lemma 2.2.1 implies that a necessary condition for $\mathbf{F}$ to be in K_{ij}^{qc} is

$$|\mathbf{F} \mathbf{v}^+| \leq |\mathbf{U}_1 \mathbf{v}^+|, \quad |\mathbf{F} \mathbf{v}^-| \leq |\mathbf{U}_1 \mathbf{v}^-|.$$

By (2.18), after a few computations these two inequalities become

$$v_1^+(v_1^+ \alpha + v_2^+ \beta) + v_2^+(v_1^+ \beta + v_2^+ \gamma) \leq (av_1^+ + bv_2^+)^2 + (bv_1^+ + cv_2^+)^2, \quad (2.19)$$

$$v_1^-(v_1^- \alpha + v_2^- \beta) + v_2^-(v_1^- \beta + v_2^- \gamma) \leq (av_1^- + bv_2^-)^2 + (bv_1^- + cv_2^-)^2, \quad (2.20)$$

where

$$\alpha = 1 + \frac{n_1^2}{d^2}(\mathfrak{D}^2 - d^4), \quad \gamma = 1 + \frac{n_2^2}{d^2}(\mathfrak{D}^2 - d^4), \quad \beta = \frac{n_1 n_2}{d^2}(\mathfrak{D}^2 - d^4).$$

Summing up the last two inequalities and exploiting the fact that $v_1^- = -v_2^+$, $v_2^- = v_1^+$, $(v_1^+)^2 + (v_2^+)^2 = 1$ we get

$$\alpha + \gamma \leq a^2 + c^2 + 2b^2 = (a + c)^2 - 2\lambda = 1 + \lambda^2.$$

Here, as above, we used the fact that $ac - b^2 = \lambda$ and $a + c = 1 + \lambda$. Now, we notice that (2.17) implies

$$\alpha + \gamma = 2 + (n_1^2 + n_2^2) \frac{1}{d^2} (\mathfrak{D}^2 - d^4) = 2 + \frac{1}{d^2} (\mathfrak{D}^2 - d^2) = 1 + \lambda^2.$$

Thus, (2.19)–(2.20) are actually equalities, that is

$$\begin{aligned} v_1^+(v_1^+\alpha + v_2^+\beta) + v_2^+(v_1^+\beta + v_2^+\gamma) &= (av_1^+ + bv_2^+)^2 + (bv_1^+ + cv_2^+)^2, \\ v_1^-(v_1^-\alpha + v_2^-\beta) + v_2^-(v_1^-\beta + v_2^-\gamma) &= (av_1^i + bv_2^i)^2 + (bv_1^i + cv_2^i)^2, \end{aligned}$$

Exploiting the values of α, β, γ we can rewrite these identities as

$$\begin{aligned} (v_1^+n_1 + v_2^+n_2)^2 &= \frac{d^2}{\mathfrak{D}^2 - d^4} ((av_1^+ + bv_2^+)^2 + (bv_1^+ + cv_2^+)^2 - 1) := r_1^2, \\ (v_2^+n_1 - v_1^+n_2)^2 &= \frac{d^2}{\mathfrak{D}^2 - d^4} ((av_1^+ + bv_2^+)^2 + (bv_1^+ + cv_2^+)^2 - 1) := r_2^2. \end{aligned}$$

Therefore, n_1, n_2 must be at the same time on one of the lines $(v_1^+n_1 + v_2^+n_2) = \pm r_1$ and on one of the lines $(v_2^+n_1 - v_1^+n_2) = \pm r_2$. Therefore there exist a maximum of four couples (n_1^l, n_2^l) such that $|\mathbf{F}\mathbf{v}^+| \leq |\mathbf{U}_1\mathbf{v}^+|$ and $|\mathbf{F}\mathbf{v}^-| \leq |\mathbf{U}_1\mathbf{v}^-|$. As n_3 can have both a positive and a negative sign, we hence found eight $\mathbf{n}$ and, by Lemma 2.2.2 eight $\mathbf{a}$, such that (2.14) is satisfied. These can be expressed as

$$\begin{aligned} &(\mathbf{a}^1, \mathbf{n}^1), \quad (\mathbf{a}^2, \mathbf{n}^2), \quad (\mathbf{a}^3, \mathbf{n}^3), \quad (\mathbf{a}^4, \mathbf{n}^4), \\ &(-\mathbf{a}^1, -\mathbf{n}^1), \quad (-\mathbf{a}^2, -\mathbf{n}^2), \quad (-\mathbf{a}^3, -\mathbf{n}^3), \quad (-\mathbf{a}^4, -\mathbf{n}^4). \end{aligned}$$

Therefore, there exist exactly four matrices $\mathbf{a} \otimes \mathbf{n}$ satisfying (2.14), which are given by (2.15)–(2.16). $\square$

Remark 2.2.1. Consequence of the above result is that a compound twin in a cubic to monoclinic transformation (and hence also in its special cases as the cubic to orthorhombic or the cubic to tetragonal) can satisfy the cofactor conditions if and only if $d = 1$. As shown in Remark 4.5.3 in Chapter 4, if $d = 1$ the set of matrices $\mathbf{F} \in (SO(3)\mathbf{U}_i \cup SO(3)\mathbf{U}_j)^{qc}$ such that $\mathbf{F} = \mathbf{1} + \mathbf{a} \otimes \mathbf{n}$ for some $\mathbf{a} \in \mathbb{R}^3, \mathbf{n} \in \mathbb{S}^2$ can have dimension two if $\mathbf{U}_i, \mathbf{U}_j$ generate a compound twin. Indeed, by Lemma 2.2.1, all matrices in the quasiconvex hull of K have 1 as a singular value.

2.2.2 Type I/II twins

Let now $\mathbf{U}, \mathbf{V} \in \mathbb{R}_{Sym+}^{3 \times 3}$ satisfy (1.11), and $(\mathbf{U}, \mathbf{b}, \mathbf{m}), (\mathbf{U}, \hat{\mathbf{b}}, \hat{\mathbf{m}})$ be the type I/type II solutions of the twinning equation (1.10). Suppose further that either $(\mathbf{U}, \mathbf{b}, \mathbf{m})$ or $(\mathbf{U}, \hat{\mathbf{b}}, \hat{\mathbf{m}})$ satisfies the cofactor conditions. We are interested in which constant average deformation gradients obtained by finely mixing the two martensitic variants $\mathbf{U}$ and $\mathbf{V}$ are compatible with austenite. We can prove the following result.

Proposition 2.2.2. *Suppose $(\mathbf{U}, \mathbf{b}, \mathbf{m}), (\mathbf{U}, \hat{\mathbf{b}}, \hat{\mathbf{m}})$ are the type I/type II (non necessarily in order) solutions of the twinning equation (1.10) for $\mathbf{U}, \mathbf{V} \in \mathbb{R}_{Sym+}^{3 \times 3}$, with $\mathbf{U}, \mathbf{V}$ satisfying (1.11). Suppose further that $(\mathbf{U}, \mathbf{b}, \mathbf{m})$ satisfies the cofactor conditions, while $(\mathbf{U}, \hat{\mathbf{b}}, \hat{\mathbf{m}})$ does not satisfy (CC2). Then, $\mathbf{F} \in (SO(3)\mathbf{U} \cup SO(3)\mathbf{V})^{qc}$ is such that $\mathbf{F} = \mathbf{1} + \mathbf{a} \otimes \mathbf{n}$ for some $\mathbf{a} \in \mathbb{R}^3, \mathbf{n} \in \mathbb{S}^2$ if and only if*

$$\mathbf{F}^T \mathbf{F} = (\mathbf{U} + \mu \mathbf{b} \otimes \mathbf{m})^T (\mathbf{U} + \mu \mathbf{b} \otimes \mathbf{m}),$$

for some $\mu \in [0, 1]$.

Proof. Let us first define

$$\begin{aligned} \mathbf{u}_1 &:= \mathbf{U}^{-1} \mathbf{m} |\mathbf{U}^{-1} \mathbf{m}|^{-1}, \quad \mathbf{u}_3 := \mathbf{b} |\mathbf{b}|^{-1}, \quad \mathbf{u}_2 = \mathbf{u}_1 \times \mathbf{u}_3, \\ \delta &= \frac{1}{2} |\mathbf{b}| |\mathbf{U}^{-1} \mathbf{m}|, \quad \mathbf{L} := \mathbf{U}^{-1} (\mathbf{1} - \delta \mathbf{u}_3 \otimes \mathbf{u}_1). \end{aligned}$$

Thanks to [19, Section 5], we know that, defined

$$K_{\mathbf{UV}} := SO(3)\mathbf{U} \cup SO(3)\mathbf{V},$$

its quasiconvex hull is given by

$$K_{\mathbf{UV}}^{qc} = \{ \mathbf{F} \in \mathbb{R}^{3 \times 3} : \mathbf{F}^T \mathbf{F} = \mathbf{L}^{-T} \mathbf{M}(\alpha, \beta, \gamma) \mathbf{L}^{-1}, 0 < \alpha \leq 1 + \delta^2, 0 < \gamma \leq 1, \alpha\gamma - \beta^2 = 1 \},$$

where

$$\mathbf{M}(\alpha, \beta, \gamma) := \alpha \mathbf{u}_1 \otimes \mathbf{u}_1 + \mathbf{u}_2 \otimes \mathbf{u}_2 + \gamma \mathbf{u}_3 \otimes \mathbf{u}_3 + \beta \mathbf{u}_1 \odot \mathbf{u}_3,$$

and we denoted $\mathbf{u}_1 \odot \mathbf{u}_3 = \mathbf{u}_1 \otimes \mathbf{u}_3 + \mathbf{u}_3 \otimes \mathbf{u}_1$. Let us first consider the scalar function

$$f_1(\alpha, \beta, \gamma) = \det(\mathbf{L}^{-T} \mathbf{M}(\alpha, \beta, \gamma) \mathbf{L}^{-1} - \mathbf{1}).$$

We define $\phi(\alpha, \beta, \gamma) = \lambda_M(\alpha, \beta, \gamma) \lambda_m(\alpha, \beta, \gamma)$, where λ_m, λ_M are respectively the largest and the smallest eigenvalue of $\mathbf{L}^{-T} \mathbf{M}(\alpha, \beta, \gamma) \mathbf{L}^{-1} - \mathbf{1}$. We denote by $\mathcal{A}$ the region

$$\mathcal{A} := \{ (\alpha, \beta, \gamma) \in \mathbb{R}^3 : 0 < \alpha \leq 1 + \delta^2, \quad 0 < \gamma \leq 1, \quad \alpha\gamma - \beta^2 = 1 \}.$$

We are interested in the set

$$\mathcal{B} := \{(\alpha, \beta, \gamma) \in \mathcal{A}: f_1(\alpha, \beta, \gamma) = 0, \quad \phi(\alpha, \beta, \gamma) \leq 0\},$$

which characterises the set of $\mathbf{F} \in (SO(3)\mathbf{U} \cup SO(3)\mathbf{V})^{qc}$ which are rank-one connected to the identity (see e.g., [18, Prop. 4]). We first notice that

$$f_1 = \frac{1}{\det \mathbf{U}^2} \det(\mathbf{M} - \mathbf{L}^T \mathbf{L}),$$

and that

$$g(\beta, \gamma) := \det \mathbf{U}^2 f_1\left(\frac{1 + \beta^2}{\gamma}, \beta, \gamma\right),$$

must be of the form

$$g(\beta, \gamma) = \frac{c_1 \beta^2 + c_2}{\gamma} + c_3 \beta + c_4 \gamma + c_5,$$

We also point out that, by (1.12)–(1.13) together with (1.15)–(1.16), we have

$$\begin{aligned} \mathbf{L}^T \mathbf{L} \mathbf{u}_1 &= \mathbf{u}_1, & \text{if } (\mathbf{U}, \mathbf{b}, \mathbf{m}) \text{ is a type I twin,} \\ \mathbf{L}^T \mathbf{L} \mathbf{u}_3 &= \mathbf{u}_3, & \text{if } (\mathbf{U}, \mathbf{b}, \mathbf{m}) \text{ is a type II twin.} \end{aligned} \tag{2.21}$$

Here we have also used the fact that, by (1.15)–(1.16), $|\mathbf{U}^{-1} \hat{\mathbf{e}}| = 1$ and $|\mathbf{U} \hat{\mathbf{e}}| = 1$, respectively for type I and type II twins satisfying the cofactor conditions, with $\hat{\mathbf{e}}$ being as in (1.11). Let us define $\mathbf{Q} := 2\mathbf{u}_j \otimes \mathbf{u}_j - \mathbf{1}$, with $j = 1$ if $(\mathbf{U}, \mathbf{b}, \mathbf{m})$ is a type I twin, and $j = 3$ if $(\mathbf{U}, \mathbf{b}, \mathbf{m})$ is a type II twin. By (2.21) we have that

$$\begin{aligned} g(\beta, \gamma) &= \det((\mathbf{M}(\gamma^{-1}(1 + \beta^2)), \beta, \gamma) - \mathbf{L}^T \mathbf{L}) \\ &= \det(\mathbf{Q}(\mathbf{M}(\gamma^{-1}(1 + \beta^2)), \beta, \gamma) - \mathbf{L}^T \mathbf{L} \mathbf{Q}) = g(-\beta, \gamma). \end{aligned}$$

For this reason, $g(\beta, \gamma)$ is even in β , and is of the form

$$g(\beta, \gamma) = \frac{c_1 \beta^2 + c_2}{\gamma} + c_4 \gamma + c_5. \tag{2.22}$$

Now, we notice that

$$(\mathbf{U} + \mu \mathbf{b} \otimes \mathbf{m}) \mathbf{L} = (1 + 2\mu \delta \mathbf{u}_3 \otimes \mathbf{u}_1) \mathbf{U} \mathbf{L} = (1 + \delta(2\mu - 1) \mathbf{u}_3 \otimes \mathbf{u}_1). \tag{2.23}$$

Therefore, the fact that $(\mathbf{U}, \mathbf{b}, \mathbf{m})$ satisfies the cofactor conditions implies

$$\frac{1}{\det \mathbf{U}^2} g(\beta, 1) = \det((\mathbf{U} + \mu(\beta) \mathbf{b} \otimes \mathbf{m})^T (\mathbf{U} + \mu(\beta) \mathbf{b} \otimes \mathbf{m}) - \mathbf{1}) = 0, \quad \text{for all } \beta \in [-\delta, \delta],$$

where $\mu(\beta) = \frac{\beta + \delta}{2\delta}$. Thus, (2.22) simplifies to

$$g(\beta, \gamma) = c_2 \left(\frac{1}{\gamma} - 1 \right) + c_4 (\gamma - 1), \tag{2.24}$$

that means, g is constant in β . Now consider the solution of the twinning equation (1.10) other than $(\mathbf{U}, \mathbf{b}, \mathbf{m})$, namely $(\mathbf{U}, \hat{\mathbf{b}}, \hat{\mathbf{m}})$. By [18, Prop. 5] we have

$$\det((\mathbf{U} + \mu \hat{\mathbf{b}} \otimes \hat{\mathbf{m}})^T (\mathbf{U} + \mu \hat{\mathbf{b}} \otimes \hat{\mathbf{m}}) - 1) = c_0 \mu (1 - \mu), \quad (2.25)$$

for some $c_0 \in \mathbb{R}$. Now, we notice that

$$\mathbf{L}^T (\mathbf{U} + \mu \hat{\mathbf{b}} \otimes \hat{\mathbf{m}})^T (\mathbf{U} + \mu \hat{\mathbf{b}} \otimes \hat{\mathbf{m}}) \mathbf{L} = \mathbf{M} \left(1 + \delta^2, \delta(2\mu - 1), \frac{1 + \delta^2(2\mu - 1)^2}{1 + \delta^2} \right), \quad \mu \in [0, 1]. \quad (2.26)$$

This can be shown by using (1.12)–(1.13) and the fact that, by (1.15)–(1.16), $|\mathbf{U}^{-1}\hat{\mathbf{e}}| = 1$ and $|\mathbf{U}\hat{\mathbf{e}}| = 1$, respectively for type I and type II twins satisfying the cofactor conditions, with $\hat{\mathbf{e}}$ being as in (1.11). Thus, setting

$$\hat{\beta}(\mu) := \delta(2\mu - 1), \quad \hat{\gamma}(\mu) := \frac{1 + \delta^2(2\mu - 1)^2}{1 + \delta^2},$$

(2.25) becomes

$$\det((\mathbf{U} + \mu \hat{\mathbf{b}} \otimes \hat{\mathbf{m}})^T (\mathbf{U} + \mu \hat{\mathbf{b}} \otimes \hat{\mathbf{m}}) - 1) = \frac{c_0(1 + \delta^2)}{4\delta^2} (1 - \hat{\gamma}(\mu)).$$

Therefore, by (2.26),

$$g(\hat{\beta}(\mu), \hat{\gamma}(\mu)) = \det \mathbf{U}^2 \frac{c_0(1 + \delta^2)}{4\delta^2} (1 - \hat{\gamma}(\mu)).$$

But recalling that $g(\beta, \gamma)$ is independent of β , a comparison with (2.24) yields

$$g(\beta, \gamma) = \det \mathbf{U}^2 \frac{c_0(1 + \delta^2)}{4\delta^2} (1 - \gamma). \quad (2.27)$$

Here, $c_0 \neq 0$ as we assumed that $(\mathbf{U}, \hat{\mathbf{b}}, \hat{\mathbf{m}})$ does not satisfy (CC2). Therefore, by (2.23) and Theorem 1.4.1, we get that $(\alpha, \beta, \gamma) \in \mathcal{B}$ if and only if $\gamma = 1$, $\beta \in [-1, 1]$ and $\alpha = 1 + \beta^2$, that is (see (2.23)), if and only if

$$\mathbf{L}^{-T} \mathbf{M} \mathbf{L}^{-1} = (\mathbf{U} + \mu \mathbf{b} \otimes \mathbf{m})^T (\mathbf{U} + \mu \mathbf{b} \otimes \mathbf{m}), \quad \mu \in [0, 1],$$

which is the claimed result. $\square$

Remark 2.2.2. From (2.27) we notice that if both $(\mathbf{U}, \mathbf{b}, \mathbf{m})$ and $(\mathbf{U}, \hat{\mathbf{b}}, \hat{\mathbf{m}})$ satisfy the cofactor conditions $f_1 = 0$ everywhere at the interior of $\mathcal{A}$. Therefore, in this case, the set $\mathcal{B}$ is two-dimensional and not one-dimensional as in the case proven in Proposition 2.2.2. Indeed, ϕ is a smooth function of α, β, γ , and is strictly negative on the boundary of $(SO(3)\mathbf{U} \cup SO(3)\mathbf{V})^{qc}$, except for at most two points. Thus, by continuity, there exists an open two-dimensional region contained in $(SO(3)\mathbf{U} \cup SO(3)\mathbf{V})^{qc}$, where ϕ is negative.

Remark 2.2.3. Under the hypotheses of Proposition 2.2.2, the set of matrices in $K_{\mathbf{U},\mathbf{V}}^{qc}$ which are rank-one connected to the identity coincides with two smooth and one-dimensional curves of finite length. The dimension of this set is hence the same as it is in the case of twins not satisfying the cofactor conditions. Indeed, for example, in [12] (see also Lemma 2.2.2 above) it is shown that for

$$\mathbf{U} = \text{diag}(a, b, c), \quad \mathbf{V} = \text{diag}(b, a, c),$$

for some $a, b, c > 0$, $a, b, c \neq 1$, the set of matrices in $K_{\mathbf{U},\mathbf{V}}^{qc}$ which are rank-one connected to the identity coincides with four smooth and one-dimensional curves of finite length. Nonetheless, the difference is in the fact that, when the cofactor conditions are satisfied, the microstructures which can form an interface with austenite are just simple laminates, and not laminates within laminates.

The following corollary is a straightforward consequence of [18, Prop. 4] and Proposition 2.2.2:

Corollary 2.2.1. *Suppose $(\mathbf{U}_i, \mathbf{b}, \mathbf{m})$, $(\mathbf{U}_i, \hat{\mathbf{b}}, \hat{\mathbf{m}})$ are the type I/type II solutions of the twinning equation (1.10) for $\mathbf{U}_i, \mathbf{U}_j \in \mathbb{R}_{Sym+}^{3 \times 3}$ (not necessarily in order), with $\mathbf{U}_i \neq \mathbf{U}_j$ as in (2.8). Suppose further that $(\mathbf{U}_i, \mathbf{b}, \mathbf{m})$ satisfies the cofactor conditions. Then, there exist $\mathbf{R}_1, \mathbf{R}_2: [0, 1] \rightarrow SO(3)$ such that $\mathbf{F} \in K_{ij}^{qc}$ is of the form $\mathbf{F} = \mathbf{1} + \mathbf{a} \otimes \mathbf{n}$ for some $\mathbf{a} \in \mathbb{R}^3, \mathbf{n} \in \mathbb{S}^2$ if and only if*

$$\mathbf{F} = \mathbf{R}_i(\mu)(\mathbf{U} + \mu \mathbf{b} \otimes \mathbf{m}),$$

for some $\mu \in [0, 1], i = 1, 2$.

2.3 Type I and II star twins

Let us first recall the following results from [26]

Theorem 2.3.1 ([26, Thm. 7]). *Let $\mathbf{U}, \mathbf{V} \in \mathbb{R}_{Sym+}^{3 \times 3}$ satisfying (1.11) be such that $\mathbf{R}_I \in SO(3), \mathbf{b}_I \in \mathbb{R}^3, \mathbf{m}_I \in \mathbb{S}^2$ is a type I solution of the twinning equation (1.10). Suppose further that $(\mathbf{U}, \mathbf{b}_I, \mathbf{m}_I)$ satisfies the cofactor conditions. Then there exist $\mathbf{R}_\mathbf{U}, \mathbf{R}_\mathbf{V}$ and $\mathbf{a} \in \mathbb{R}^3, \mathbf{n}_\mathbf{U}, \mathbf{n}_\mathbf{V} \in \mathbb{S}^2$ such that*

$$\mathbf{R}_\mathbf{U} \mathbf{U} = \mathbf{1} + \mathbf{a} \otimes \mathbf{n}_\mathbf{U}, \quad \mathbf{R}_\mathbf{V} \mathbf{V} = \mathbf{1} + \mathbf{a} \otimes \mathbf{n}_\mathbf{V},$$

$\mathbf{R}_\mathbf{U} \mathbf{b}_I \times \mathbf{a} = \mathbf{0}$, $(\mathbf{n}_\mathbf{U} \times \mathbf{n}_\mathbf{V}) \cdot \mathbf{m}_I = 0$, and $\mathbf{R}_\mathbf{V} = \mathbf{R}_\mathbf{U} \mathbf{R}_I$.

Theorem 2.3.2 ([26, Thm. 8]). *Let $\mathbf{U}, \mathbf{V} \in \mathbb{R}_{Sym+}^{3 \times 3}$ satisfying (1.11) be such that $\mathbf{R}_{II} \in SO(3)$, $\mathbf{b}_{II} \in \mathbb{R}^3$, $\mathbf{m}_{II} \in \mathbb{S}^2$ is a type II solution of the twinning equation (1.10). Suppose further that $(\mathbf{U}, \mathbf{b}_{II}, \mathbf{m}_{II})$ satisfies the cofactor conditions. Then there exist $\mathbf{R}_U, \mathbf{R}_V$ and $\mathbf{a}_U, \mathbf{a}_V \in \mathbb{R}^3$ such that*

$$\mathbf{R}_U \mathbf{U} = \mathbf{1} + \mathbf{a}_U \otimes \mathbf{m}_{II}, \quad \mathbf{R}_V \mathbf{V} = \mathbf{1} + \mathbf{a}_V \otimes \mathbf{m}_{II},$$

and $\mathbf{R}_V = \mathbf{R}_U \mathbf{R}_{II}$.

Let us now introduce the definition of type I/II star twin. Below we will make use of the notation of Theorem 2.3.1 and of Theorem 2.3.2.

Definition 2.3.1. *Let $\mathcal{M}$ be a subset of $\mathbb{R}_{Sym+}^{3 \times 3}$. Let $\mathbf{U}, \mathbf{V} \in \mathcal{M}$, $\mathbf{U} \neq \mathbf{V}$, satisfying (1.11) and let $\mathbf{R}_I \in SO(3)$, $\mathbf{b}_I \in \mathbb{R}^3$, $\mathbf{m}_I \in \mathbb{S}^2$ be a type I solution of the twinning equation (1.10) for $\mathbf{U}, \mathbf{V}$. Suppose further that $(\mathbf{U}, \mathbf{b}_I, \mathbf{m}_I)$ satisfies the cofactor conditions. Then we say that $(\mathbf{U}, \mathbf{b}_I, \mathbf{m}_I)$ is a type I star twin generated by $(\mathbf{U}, \mathbf{V})$ if there exist three different rotations $\mathbf{Q}_1, \mathbf{Q}_2, \mathbf{Q}_3 \in SO(3)$, $\mathbf{Q}_i \neq \mathbf{1}$ for each $i = 1, \dots, 3$, and $\mu^* \in (0, 1)$ satisfying*

$$(S1) \quad \mathbf{Q}_i \mathbf{U} \mathbf{Q}_i^T, \mathbf{Q}_i \mathbf{V} \mathbf{Q}_i^T \in \mathcal{M}, \text{ for each } i = 1, \dots, 3,$$

$$(S2) \quad \mathbf{Q}_i (\mu^* \mathbf{n}_U + (1 - \mu^*) \mathbf{n}_V) = \chi_i (\mu^* \mathbf{n}_U + (1 - \mu^*) \mathbf{n}_V) \text{ for each } i = 1, \dots, 3 \text{ and with } \chi_i = \pm 1,$$

$$(S3) \quad ((\mathbf{Q}_i \mathbf{a}) \times (\mathbf{Q}_j \mathbf{a})) \cdot \mathbf{a} \neq 0 \text{ for each } i, j = 1, \dots, 3, i \neq j \text{ and } ((\mathbf{Q}_1 \mathbf{a}) \times (\mathbf{Q}_2 \mathbf{a})) \cdot (\mathbf{Q}_3 \mathbf{a}) \neq 0.$$

If the number of $\mathbf{Q}_i$ satisfying (S1)–(S3) is two and not three we say that $(\mathbf{U}, \mathbf{b}_I, \mathbf{m}_I)$ is a type I half-star twin.

Definition 2.3.2. *Let $\mathcal{M}$ be a subset of $\mathbb{R}_{Sym+}^{3 \times 3}$. Let $\mathbf{U}, \mathbf{V} \in \mathcal{M}$, $\mathbf{U} \neq \mathbf{V}$, and let $\mathbf{R}_{II} \in SO(3)$, $\mathbf{b}_{II} \in \mathbb{R}^3$, $\mathbf{m}_{II} \in \mathbb{S}^2$ be a type II solution of the twinning equation (1.10) for $\mathbf{U}, \mathbf{V}$. Suppose further that $(\mathbf{U}, \mathbf{b}_{II}, \mathbf{m}_{II})$ satisfies the cofactor conditions. Then we say that $(\mathbf{U}, \mathbf{b}_{II}, \mathbf{m}_{II})$ is a type II star twin generated by $(\mathbf{U}, \mathbf{V})$ if there exist three different rotations $\mathbf{Q}_1, \mathbf{Q}_2, \mathbf{Q}_3 \in SO(3)$, $\mathbf{Q}_i \neq \mathbf{1}$ for each $i = 1, \dots, 3$, and $\mu^* \in (0, 1)$ satisfying*

$$(T1) \quad \mathbf{Q}_i \mathbf{U} \mathbf{Q}_i^T, \mathbf{Q}_i \mathbf{V} \mathbf{Q}_i^T \in \mathcal{M}, \text{ for each } i = 1, \dots, 3,$$

$$(T2) \quad \mathbf{Q}_i (\mu^* \mathbf{a}_U + (1 - \mu^*) \mathbf{a}_V) = \chi_i (\mu^* \mathbf{a}_U + (1 - \mu^*) \mathbf{a}_V) \text{ for each } i = 1, \dots, 3 \text{ and with } \chi_i = \pm 1,$$

(T3) $((\mathbf{Q}_i \mathbf{m}_{II}) \times (\mathbf{Q}_j \mathbf{m}_{II})) \cdot \mathbf{m}_{II} \neq 0$ for each $i, j = 1, \dots, 3$, $i \neq j$ and $((\mathbf{Q}_1 \mathbf{m}_{II}) \times (\mathbf{Q}_2 \mathbf{m}_{II})) \cdot (\mathbf{Q}_3 \mathbf{m}_{II}) \neq 0$.

If the number of $\mathbf{Q}_i$ satisfying (T1)–(T3) is two and not three we say that $(\mathbf{U}, \mathbf{b}_{II}, \mathbf{m}_{II})$ is a type II half-star twin.

Remark 2.3.1. It would be possible to generalise (T2) (or similarly (S2)) in the definition of star twins by requiring the existence of $\mu_1^*, \mu_2^* \in [0, 1]$ such that

$$\mathbf{Q}_i(\mu_1^* \mathbf{a}_U + (1 - \mu_1^*) \mathbf{a}_V) = \chi_i(\mu_2^* \mathbf{a}_U + (1 - \mu_2^*) \mathbf{a}_V) \text{ for each } i = 1, \dots, 3 \text{ and with } \chi_i = \pm 1.$$

However, a fine observation of the proof of Theorem 2.4.1 (resp. Theorem 2.4.2) yields that, in the cubic to monoclinic case, the only interesting case is when $\mu_1^* = \mu_2^*$, making the generalisation not relevant to our context.

Remark 2.3.2. Definition 2.3.1 and Definition 2.3.2 both imply the existence of four exactly compatible twins $\mathbf{F}_0, \mathbf{F}_1, \mathbf{F}_2, \mathbf{F}_3$ satisfying (2.6). Conversely, in cubic to monoclinic transformations, given four exactly compatible twins $\mathbf{F}_0, \mathbf{F}_1, \mathbf{F}_2, \mathbf{F}_3$ of type I (or of type II) satisfying (2.6), we have that Proposition 2.1.1, Remark 2.1.1 and Remark 2.3.1 imply the satisfaction of Definition 2.3.1 (resp. Definition 2.3.2).

The following two propositions state that, in the presence of star-twins, the set of homogeneous average deformation gradients which can form stress free phase interfaces with austenite is unusually large:

Proposition 2.3.1. *Let $\mathcal{M}$ be a finite subset of $\mathbb{R}_{Sym+}^{3 \times 3}$, $\mathbf{U}, \mathbf{V} \in \mathcal{M}$ satisfying (1.11), and let $(\mathbf{U}, \mathbf{b}_{II}, \mathbf{m}_{II})$ be a type II half-star twin generated by $(\mathbf{U}, \mathbf{V})$. Then,*

$$\begin{aligned} 1 + (\mu^* \mathbf{a}_U + (1 - \mu^*) \mathbf{a}_V) \otimes \mathbf{m}_{II}, \quad & 1 + (\mu^* \mathbf{a}_U + (1 - \mu^*) \mathbf{a}_V) \otimes \mathbf{Q}_1 \mathbf{m}_{II}, \\ & 1 + (\mu^* \mathbf{a}_U + (1 - \mu^*) \mathbf{a}_V) \otimes \mathbf{Q}_2 \mathbf{m}_{II}, \end{aligned}$$

and all their convex combinations are contained in the set K^{qc} , where

$$K := \bigcup_{\mathbf{M} \in \mathcal{M}} SO(3) \mathbf{M}.$$

If $(\mathbf{U}, \mathbf{b}_{II}, \mathbf{m}_{II})$ is a type II star twin generated by $(\mathbf{U}, \mathbf{V})$, then also

$$1 + (\mu^* \mathbf{a}_U + (1 - \mu^*) \mathbf{a}_V) \otimes \mathbf{Q}_3 \mathbf{m}_{II} \in K^{qc}.$$

Proof. We restrict to the case of half-star twins; the case of star twins follows similarly. From [62] we know that given a set $C \in \mathbb{R}^{3 \times 3}$, the lamination convex hull of C , namely C^{lc} , is contained in C^{qc} . We recall for the benefit of the reader that

$$C^{lc} := \bigcup_{i=1}^{\infty} C^{(i)},$$

where $C^{(1)} = C$ and

$$C^{(i+1)} := C^{(i)} \cup \{\mu A + (1 - \mu)B : A, B \in C^{(i)}, \text{ rank}(A - B) = 1, \mu \in (0, 1)\}.$$

Therefore, let $K := \bigcup_{M \in \mathcal{M}} SO(3)M$. By Theorem 2.3.2 we know that

$$\begin{aligned} R_U U &= 1 + \mathbf{a}_U \otimes \mathbf{m}_{II}, & R_V V &= 1 + \mathbf{a}_V \otimes \mathbf{m}_{II}, \\ Q_1 R_U Q_1^T Q_1 U Q_1^T &= 1 + Q_1 \mathbf{a}_U \otimes Q_1 \mathbf{m}_{II}, & Q_1 R_V Q_1^T Q_1 V Q_1^T &= 1 + Q_1 \mathbf{a}_V \otimes Q_1 \mathbf{m}_{II}, \\ Q_2 R_U Q_2^T Q_2 U Q_2^T &= 1 + Q_2 \mathbf{a}_U \otimes Q_2 \mathbf{m}_{II}, & Q_2 R_V Q_2^T Q_2 V Q_2^T &= 1 + Q_2 \mathbf{a}_V \otimes Q_2 \mathbf{m}_{II}, \end{aligned}$$

are all in K . Thus,

$$\begin{aligned} \mu^* R_U U + (1 - \mu^*) R_V V &= 1 + (\mu^* \mathbf{a}_U + (1 - \mu^*) \mathbf{a}_V) \otimes \mathbf{m}_{II}, \\ \mu^* Q_1 R_U Q_1^T Q_1 U Q_1^T + (1 - \mu^*) Q_1 R_V Q_1^T Q_1 V Q_1^T &= 1 + (\mu^* \mathbf{a}_U + (1 - \mu^*) \mathbf{a}_V) \otimes Q_1 \mathbf{m}_{II}, \\ \mu^* Q_2 R_U Q_2^T Q_2 U Q_2^T + (1 - \mu^*) Q_2 R_V Q_2^T Q_2 V Q_2^T &= 1 + (\mu^* \mathbf{a}_U + (1 - \mu^*) \mathbf{a}_V) \otimes Q_2 \mathbf{m}_{II}, \end{aligned}$$

are in $K^{(2)}$, and all their convex combinations are in $K^{(3)}$. $\square$

Remark 2.3.3. Type II half-star/star twins are interesting also because, by combining laminates, we can easily construct macroscopic curved interfaces between austenite and martensite whose normal does not lie in a plane. Indeed, in the notation of Proposition 2.3.1, let us define

$$\mathbf{a}^* := \mu^* \mathbf{a}_U + (1 - \mu^*) \mathbf{a}_V.$$

Then,

$$G(\mu_1, \mu_2) := 1 + \mathbf{a}^* \otimes (\mu_1(\mu_2 \mathbf{m}_{II} + (1 - \mu_2) Q_1 \mathbf{m}_{II}) + (1 - \mu_1) Q_2 \mathbf{m}_{II})$$

is in the K^{qc} for every $\mu_1, \mu_2 \in [0, 1]$, and $\mathbf{m}_{II}, Q_1 \mathbf{m}_{II}, Q_2 \mathbf{m}_{II}$ span $\mathbb{R}^3$. Therefore, given a smooth bounded domain Ω , one can choose $\mu_1, \mu_2 : \Omega \rightarrow [0, 1]$ to be space dependent, and, provided

$$\mu_1 (\nabla \mu_2 \times (\mathbf{m}_{II} - Q_1 \mathbf{m}_{II})) + \nabla \mu_1 \times (\mu_2 \mathbf{m}_{II} + (1 - \mu_2) Q_1 \mathbf{m}_{II}) - \nabla \mu_1 \times Q_2 \mathbf{m}_{II} = 0,$$

$G(\mu_1(\mathbf{x}), \mu_2(\mathbf{x}))$ is an average deformation gradient. Choosing for example

$$\mu_2(\mathbf{x}) = f(\mathbf{x} \cdot (\mathbf{m}_{II} - \mathbf{Q}_1 \mathbf{m}_{II})),$$

and

$$\mu_1(\mathbf{x}) = c_1 \left(\int_0^{\mathbf{x} \cdot (\mathbf{m}_{II} - \mathbf{Q}_1 \mathbf{m}_{II})} f(s) \, ds + \mathbf{x} \cdot (\mathbf{Q}_1 \mathbf{m}_{II} - \mathbf{Q}_2 \mathbf{m}_{II}) \right) + c_2,$$

for some smooth $f: \mathbb{R} \rightarrow [0, 1]$ and some $c_1, c_2 \in \mathbb{R}$ such that $\mu_2(\mathbf{x}) \in [0, 1]$ for every $\mathbf{x} \in \Omega$, will give an average deformation gradient of martensite, which is compatible with austenite across an interface whose normal does not lie in a plane.

In the same way we can prove

Proposition 2.3.2. *Let $\mathcal{M}$ be a finite subset of $\mathbb{R}_{Sym+}^{3 \times 3}$, $\mathbf{U}, \mathbf{V} \in \mathcal{M}$ and let $(\mathbf{U}, \mathbf{b}_I, \mathbf{m}_I)$ be a type I half-star twin generated by $(\mathbf{U}, \mathbf{V})$. Then*

$$\begin{aligned} 1 + \mathbf{a} \otimes (\mu^* \mathbf{n}_U + (1 - \mu^*) \mathbf{n}_V), \quad 1 + \mathbf{Q}_1 \mathbf{a} \otimes (\mu^* \mathbf{n}_U + (1 - \mu^*) \mathbf{n}_V), \\ 1 + \mathbf{Q}_2 \mathbf{a} \otimes (\mu^* \mathbf{n}_U + (1 - \mu^*) \mathbf{n}_V), \end{aligned}$$

and all their convex combinations are contained in the set K^{qc} , where

$$K := \bigcup_{M \in \mathcal{M}} SO(3)M.$$

If $(\mathbf{U}, \mathbf{b}_I, \mathbf{m}_I)$ is a type I star twin generated by $(\mathbf{U}, \mathbf{V})$, then also

$$1 + \mathbf{Q}_3 \mathbf{a} \otimes (\mu^* \mathbf{n}_U + (1 - \mu^*) \mathbf{n}_V) \in K^{qc}.$$

2.4 Star twins in cubic to monoclinic transformations

In this section we want to find conditions on the eigenvalues of the matrices $\mathbf{U}_i$ in (2.8) so that there exist type I and type II star twins. We start with the following simple lemma

Lemma 2.4.1. *Let $\mathbf{v} \in \mathbb{R}^3$, $\mathbf{v} \neq 0$, then $\mathbf{Q} \in SO(3)$ satisfies $\mathbf{Q}\mathbf{v} = -\mathbf{v}$ if and only if $\mathbf{Q}$ is a rotation of π and axis $\mathbf{u}_R \in \mathbb{S}^2$, that is, $\mathbf{R} = 2\mathbf{u}_R \otimes \mathbf{u}_R - \mathbf{1}$, and $\mathbf{u}_R \cdot \mathbf{v} = 0$. The matrix $\mathbf{Q} \in SO(3)$, $\mathbf{Q} \neq \mathbf{1}$ satisfies $\mathbf{Q}\mathbf{v} = \mathbf{v}$ if and only if the axis of $\mathbf{Q}$ is parallel to $\mathbf{v}$.*

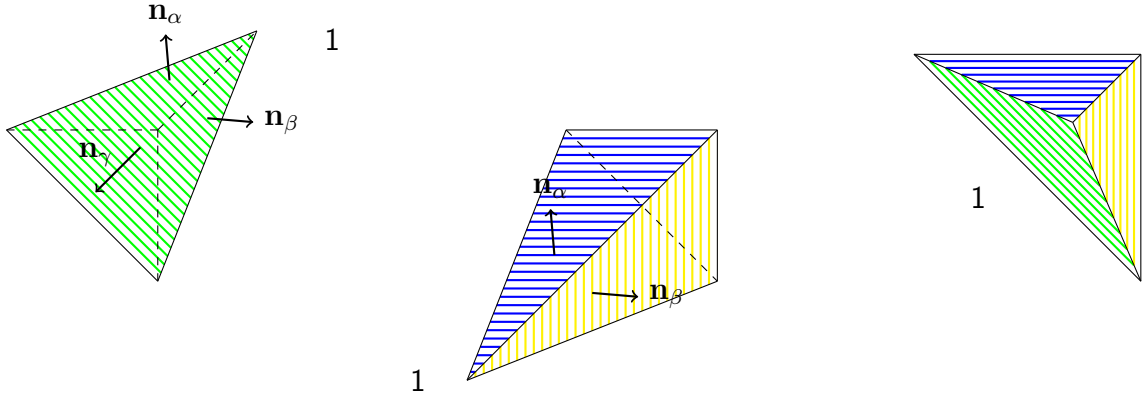


Figure 2.3: Three different views of a pyramidal microstructure constructed with type II half-star twins (also possible with star twins). The polyhedron is divided into a blue, a yellow and a green region, in each of which we have a different exactly compatible laminate, namely $F_1 = \mathbf{1} + \mathbf{a}^* \otimes \mathbf{n}_\alpha$, $F_2 = \mathbf{1} + \mathbf{a}^* \otimes \mathbf{n}_\beta$ and $F_3 = \mathbf{1} + \mathbf{a}^* \otimes \mathbf{n}_\gamma$, for some $\mathbf{a}^* \in \mathbb{R}^3 \setminus \{\mathbf{0}\}$. The three constant average deformation gradients F_1, F_2, F_3 satisfy (2.6), that means that $\mathbf{n}_\alpha, \mathbf{n}_\beta, \mathbf{n}_\gamma$ are linearly independent. At the three faces of the pyramid with normals $\mathbf{n}_\alpha, \mathbf{n}_\beta, \mathbf{n}_\gamma$ the martensitic laminates are compatible with austenite without an interface layer. Such pyramid can be used as a building block for three-dimensional curved interfaces.

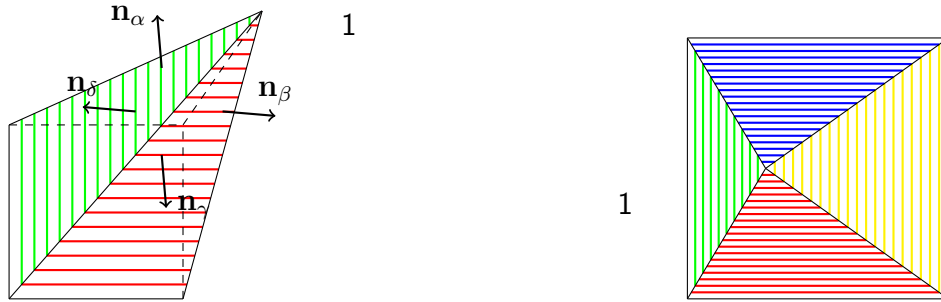


Figure 2.4: Two different views of a pyramidal microstructure constructed with type II star twins. The polyhedron is divided into a blue, a yellow, a red and a green region, in each of which we have a different exactly compatible laminate, namely $F_1 = \mathbf{1} + \mathbf{a}^* \otimes \mathbf{n}_\alpha$, $F_2 = \mathbf{1} + \mathbf{a}^* \otimes \mathbf{n}_\beta$, $F_3 = \mathbf{1} + \mathbf{a}^* \otimes \mathbf{n}_\gamma$ and $F_4 = \mathbf{1} + \mathbf{a}^* \otimes \mathbf{n}_\delta$, for some $\mathbf{a}^* \in \mathbb{R}^3 \setminus \{\mathbf{0}\}$. The four constant average deformation gradients F_1, F_2, F_3, F_4 satisfy (2.6), that means that $\mathbf{n}_\alpha, \mathbf{n}_\beta, \mathbf{n}_\gamma, \mathbf{n}_\delta$ are linearly independent. At the three faces of the pyramid with normals $\mathbf{n}_\alpha, \mathbf{n}_\beta, \mathbf{n}_\gamma, \mathbf{n}_\delta$ the martensitic laminates are compatible with austenite without an interface layer. Such pyramid can be used as a building block for three-dimensional curved interfaces.

Proof. We can write R in terms of its angle of rotation θ and its rotation axis $\mathbf{u}_R$ as

$$R = \cos(\theta)1 + \sin(\theta)[\mathbf{u}_R]_{\times} + (1 - \cos(\theta))\mathbf{u}_R \otimes \mathbf{u}_R,$$

where $[\mathbf{u}_R]_{\times}$ is the cross product matrix of $\mathbf{u}_R$. Therefore, after decomposing $\mathbf{v}$ into a parallel and an orthogonal part to $\mathbf{u}_R$, $\mathbf{v}_{\mathbf{u}_R}$ and $\mathbf{v}_{\mathbf{u}_R}^{\perp}$, we get that $R\mathbf{v} = -\mathbf{v}$ is equivalent to

$$\cos(\theta)(\mathbf{v}_{\mathbf{u}_R} + \mathbf{v}_{\mathbf{u}_R}^{\perp}) + \sin(\theta)(\mathbf{u}_R \times \mathbf{v}_{\mathbf{u}_R}^{\perp}) + (1 - \cos(\theta))\mathbf{u}_R(\mathbf{u}_R \cdot \mathbf{v}_{\mathbf{u}_R}) = -\mathbf{v}_{\mathbf{u}_R} - \mathbf{v}_{\mathbf{u}_R}^{\perp}.$$

Thus, multiplying the equation by $\mathbf{v}_{\mathbf{u}_R}^{\perp}$ we get that either $\theta = \pi$, or $\mathbf{v}_{\mathbf{u}_R}^{\perp} = \mathbf{0}$. But multiplying the equation by $\mathbf{v}_{\mathbf{u}_R}$ leads to $\mathbf{v}_{\mathbf{u}_R} = \mathbf{0}$, and therefore $\theta = \pi$. The second statement follows from the definition of rotation axis. $\square$

The lemma below states that, if a type II twin generated by $\mathbf{U}_i, \mathbf{U}_j$, where (i, j) are a pair in column (A) (resp. (B)) of Table 2.1, is a type I (or a type II) star twin, then all the other type I (resp. type II) twins generated by $\mathbf{U}_k, \mathbf{U}_l$, where (k, l) is a generic pair in column (A) (resp. (B)) of Table 2.1, is a type II star twin.

Lemma 2.4.2. *Let $\mathcal{M} := \cup_{i=1}^{12} \mathbf{U}_i$, where the $\mathbf{U}_i$ are the positive definite matrices, with $a, b, c, d > 0$, given in (2.8). Suppose also that $\mathbf{U}_i, \mathbf{U}_j$, $i \neq j = 1, \dots, 12$, generate a type I/II star twin (resp. a half-star twin). Then, for every $\mathbf{Q} \in \mathcal{P}_{24}$ the type I/II twins generated by $\mathbf{Q}\mathbf{U}_i\mathbf{Q}^T, \mathbf{Q}\mathbf{U}_j\mathbf{Q}^T$ are star twins (resp. a half-star twin).*

Proof. We prove the result for type II star twins, the case of type I star twins follows a similar argument.

The fact that $\mathbf{U}_k, \mathbf{U}_l$, where $\mathbf{U}_k = \mathbf{Q}\mathbf{U}_i\mathbf{Q}^T, \mathbf{U}_l = \mathbf{Q}\mathbf{U}_j\mathbf{Q}^T$ for some $\mathbf{Q} \in \mathcal{P}_{24}$, generate a type II twin satisfying the cofactor conditions is a consequence of Corollary 2.1.1. Let $\mu^* \in (0, 1)$, $\mathbf{Q}_1, \mathbf{Q}_2, \mathbf{Q}_3 \in SO(3)$ be such that $(\mathbf{U}_i, \mathbf{a}_{II}, \mathbf{n}_{II})$, the type II twin generated by $\mathbf{U}_i, \mathbf{U}_j$, satisfies Definition 2.3.2. Let also $\hat{\mathbf{Q}}_m = \mathbf{Q}\mathbf{Q}_m\mathbf{Q}^T$ for every $m = 1, 2, 3$, and $(\mathbf{U}_k, \hat{\mathbf{b}}_{II}, \hat{\mathbf{m}}_{II})$ be the type II twin generated by $\mathbf{U}_k, \mathbf{U}_l$. We claim that $\mu^*, \hat{\mathbf{Q}}_1, \hat{\mathbf{Q}}_2, \hat{\mathbf{Q}}_3 \in SO(3)$ are such that $(\mathbf{U}_k, \hat{\mathbf{b}}_{II}, \hat{\mathbf{m}}_{II})$ satisfies (T1)–(T3) in Definition 2.3.2. Indeed, we recall that $\mathbf{Q}\mathbf{V}\mathbf{Q}^T \in \mathcal{M}$ for every $\mathbf{V} \in \mathcal{M}, \mathbf{Q} \in \mathcal{P}_{24}$ (see e.g., [23]). Thus, as

$$\hat{\mathbf{Q}}_m \mathbf{U}_k \hat{\mathbf{Q}}_m = \mathbf{Q}\mathbf{Q}_m \mathbf{U}_i \mathbf{Q}_m \mathbf{Q}, \quad \hat{\mathbf{Q}}_m \mathbf{U}_l \hat{\mathbf{Q}}_m = \mathbf{Q}\mathbf{Q}_m \mathbf{U}_j \mathbf{Q}_m \mathbf{Q},$$

and given that $\mathbf{U}_i, \mathbf{U}_j$ satisfy (T1), we can write

$$\hat{\mathbf{Q}}_m \mathbf{U}_k \hat{\mathbf{Q}}_m, \hat{\mathbf{Q}}_m \mathbf{U}_l \hat{\mathbf{Q}}_m \in \mathcal{M}, \quad \text{for all } m = 1, 2, 3,$$

that is, (T1). As $\mathbf{a}_{U_k} = \mathbf{Q}\mathbf{a}_{U_i}$, $\mathbf{a}_{U_l} = \mathbf{Q}\mathbf{a}_{U_j}$, we also have

$$\hat{\mathbf{Q}}_m(\mu^* \mathbf{a}_{U_k} + (1 - \mu^*) \mathbf{a}_{U_l}) = \mathbf{Q}(\mu^* \mathbf{a}_{U_i} + (1 - \mu^*) \mathbf{a}_{U_j}) = \mu^* \mathbf{a}_{U_k} + (1 - \mu^*) \mathbf{a}_{U_l},$$

for all $m = 1, 2, 3$, which is (T2). Finally, as $\hat{\mathbf{m}}_{II} = \mathbf{Q}\mathbf{m}_{II}$, we have

$$\begin{aligned} \left((\hat{\mathbf{Q}}_1 \hat{\mathbf{m}}_{II}) \times (\hat{\mathbf{Q}}_2 \hat{\mathbf{m}}_{II}) \right) \cdot (\hat{\mathbf{Q}}_3 \hat{\mathbf{m}}_{II}) &= \left((\mathbf{Q}_1 \mathbf{m}_{II}) \times (\mathbf{Q}_2 \mathbf{m}_{II}) \right) \cdot (\mathbf{Q}_3 \mathbf{m}_{II}) \neq 0, \\ \left((\hat{\mathbf{Q}}_m \hat{\mathbf{m}}_{II}) \times (\hat{\mathbf{Q}}_n \hat{\mathbf{m}}_{II}) \right) \cdot \hat{\mathbf{m}}_{II} &= \left((\mathbf{Q}_m \mathbf{m}_{II}) \times (\mathbf{Q}_n \mathbf{m}_{II}) \right) \cdot \mathbf{m}_{II} \neq 0, \end{aligned}$$

for every $m, n = 1, 2, 3$, which is (T3). $\square$

We are now in the position to prove the following result:

Theorem 2.4.1. *Let $\mathcal{M} := \cup_{i=1}^{12} U_i$, where the U_i are the positive definite matrices, with $a, b, c, d > 0$, given in (2.8). Then, a type II twin generated by U_i, U_j , $i \neq j = 1, \dots, 12$, satisfying the cofactor conditions is a half-star twin if and only if the eigenvalues of U_1 , namely $0 < \lambda_1 \leq \lambda_2 = 1 \leq \lambda_3$ satisfy one of the following conditions*

$$\begin{aligned} 4d^2(d^2 + \lambda_3^2 - 2) &= (d - \lambda_3)^2(1 - d^2), & \text{if } d \neq 1 \text{ and } \lambda_1 = d, \\ 4d^2(d^2 + \lambda_1^2 - 2) &= (d - \lambda_1)^2(1 - d^2), & \text{if } d \neq 1 \text{ and } \lambda_3 = d, \\ \lambda_1^2(5\lambda_3^2 - 1) &= 8\lambda_1\lambda_3 - 5 + \lambda_3^2, & \text{if } d = 1 \text{ and } \lambda_3 > 1. \end{aligned} \quad (2.28)$$

A type II twin generated by U_i, U_j , $i \neq j = 1, \dots, 12$, satisfying the cofactor conditions is a star twin if and only if $a \neq c$ and one of the following holds

$$\begin{aligned} d^2(d^2 + \lambda_3^2 - 2) &= (d - \lambda_3)^2(1 - d^2), & \text{if } d \neq 1 \text{ and } \lambda_1 = d, \\ d^2(d^2 + \lambda_1^2 - 2) &= (d - \lambda_1)^2(1 - d^2), & \text{if } d \neq 1 \text{ and } \lambda_3 = d. \end{aligned} \quad (2.29)$$

Remark 2.4.1. It is possible to solve the equations in (2.28)–(2.29) to get a direct relation between the eigenvalues of the U_i 's. Indeed, (2.28) is equivalent to

$$\begin{aligned} \lambda_1 &= \frac{d - d^3 - d2\sqrt{2}\sqrt{6d^2 - 1 - 3d^4}}{1 - 5d^2}, & \text{if } \lambda_3 = d \in (1, \sqrt{1 + 2^{\frac{1}{2}}3^{-\frac{1}{2}}}), \\ \lambda_1 &= \frac{d - d^3 + d2\sqrt{2}\sqrt{6d^2 - 1 - 3d^4}}{1 - 5d^2}, & \text{if } \lambda_3 = d \in (\sqrt{5^{-1}9}, \sqrt{1 + 2^{\frac{1}{2}}3^{-\frac{1}{2}}}), \\ \lambda_3 &= \frac{d - d^3 - d2\sqrt{2}\sqrt{6d^2 - 1 - 3d^4}}{1 - 5d^2}, & \text{if } \lambda_1 = d \in (\sqrt{1 - 2^{\frac{1}{2}}3^{-\frac{1}{2}}}, 1), \\ \lambda_1 &= \frac{4\lambda_3 \pm \sqrt{5}(\lambda_3^2 - 1)}{5\lambda_3^2 - 1}, & \text{if } \lambda_3 > 1, d = 1, \end{aligned}$$

when $0 < \lambda_1 \leq 1 \leq \lambda_3$. Similarly, the condition (2.29) can be rewritten as

$$\begin{aligned}\lambda_1 &= \frac{d - d^3 - d\sqrt{6d^2 - 2 - 3d^4}}{1 - 2d^2}, & \text{if } \lambda_3 = d \in (1, \sqrt{1 + 3^{-\frac{1}{2}}}), \\ \lambda_1 &= \frac{d - d^3 + d\sqrt{6d^2 - 2 - 3d^4}}{1 - 2d^2}, & \text{if } \lambda_3 = d \in (\sqrt{2^{-1}3}, \sqrt{1 + 3^{-\frac{1}{2}}}), \\ \lambda_3 &= \frac{d - d^3 - d\sqrt{6d^2 - 2 - 3d^4}}{1 - 2d^2}, & \text{if } \lambda_1 = d \in (\sqrt{1 - 3^{-\frac{1}{2}}}, 1),\end{aligned}$$

when $0 < \lambda_1 \leq 1 \leq \lambda_3$.

Proof. We deal with the case where the twins generated by $\mathbf{U}_i, \mathbf{U}_j$, with (i, j) in column (A) of Table 2.1, satisfy the cofactor conditions. The case where (i, j) are in column (B) of Table 2.1 can be treated similarly and leads to the same results. Furthermore, thanks to Lemma 2.4.2 and Proposition 2.1.1 we can carry on the proof by considering just $(\mathbf{U}_1, \mathbf{b}_{II}, \mathbf{m}_{II})$ the type II twin generated by $\mathbf{U}_1, \mathbf{U}_{11}$. We divide the proof in three cases, based on the eigenvalues of $\mathbf{U}_1$ $\lambda_1 < \lambda_2 = 1 < \lambda_3$ (cf. Remark 1.4.1): the cases are $\lambda_2 = 1 > d$, $\lambda_2 = 1 < d$ and $0 < \lambda_1 < d = 1 < \lambda_3$.

Case 1: $0 < d < 1 < \lambda$, $\lambda = ac - b^2$. Here, the eigenvectors of $\mathbf{U}_1$ related to $\lambda_1 = d, \lambda_2 = 1$, and $\lambda_3 = \lambda$ are

$$\mathbf{e}_1 = (0, 0, 1)^T, \quad \mathbf{e}_2 = (-z_2, z_1, 0)^T, \quad \mathbf{e}_3 = (z_1, z_2, 0)^T,$$

while the ones for $\mathbf{U}_{11}$ are given by

$$\mathbf{v}_1 = (1, 0, 0)^T, \quad \mathbf{v}_2 = (0, z_1, -z_2)^T, \quad \mathbf{v}_3 = (0, z_2, z_1)^T.$$

Here, z_1, z_2 are such that $\frac{z_1}{z_2} = \frac{\lambda - c}{b}$, $z_1^2 + z_2^2 = 1$, and we assume without loss of generality that $z_1 > 0$. We can neglect the cases $z_1 = 0$ and $z_2 = 0$, as by (2.10)–(2.11) they would imply $b = 0$ and $\mathbf{U}_1 = \mathbf{U}_{11}$. For the same reason we can neglect the case $z_2 = 0$. Using [18, Prop. 4] we get that

$$\mathbf{R}_1^\pm \mathbf{U}_1 = \mathbf{1} + \mathbf{a}_1^\pm \otimes \mathbf{n}_1^\pm, \quad \mathbf{R}_{11}^\pm \mathbf{U}_{11} = \mathbf{1} + \mathbf{a}_{11}^\pm \otimes \mathbf{n}_{11}^\pm,$$

where

$$\mathbf{a}_1^\pm = \beta_0(-\lambda\eta_1\mathbf{e}_1 \pm d\eta_2\mathbf{e}_3), \quad \mathbf{n}_1^\pm = \eta_1\mathbf{e}_1 \pm \eta_2\mathbf{e}_3, \quad (2.30)$$

$$\mathbf{a}_{11}^\pm = \beta_0(-\lambda\eta_1\mathbf{v}_1 \pm d\eta_2\mathbf{v}_3), \quad \mathbf{n}_{11}^\pm = \eta_1\mathbf{v}_1 \pm \eta_2\mathbf{v}_3, \quad (2.31)$$

and

$$\eta_1 = -\frac{\sqrt{1 - d^2}}{\sqrt{\lambda^2 - d^2}}, \quad \eta_2 = \frac{\sqrt{\lambda^2 - 1}}{\sqrt{\lambda^2 - d^2}}, \quad \beta_0 = \lambda - d.$$

Since the cofactor conditions are satisfied, by Theorem 2.3.2 we know that one of the following identities must hold

$$\mathbf{n}_1^+ \times \mathbf{n}_{11}^+ = 0, \quad \mathbf{n}_1^+ \times \mathbf{n}_{11}^+ = 0, \quad \mathbf{n}_1^- \times \mathbf{n}_{11}^+ = 0, \quad \mathbf{n}_1^+ \times \mathbf{n}_{11}^- = 0.$$

Keeping in mind that (2.10)–(2.11) imply $\eta_1 = -z_1\eta_2$, we get that, in the notation of Theorem 2.3.2, the only possibility up to a sign change is

$$\begin{aligned} \mathbf{a}_{U_1} &= \beta_0(d\eta_1, -z_2d\eta_2, -\lambda\eta_1), & \mathbf{a}_{U_{11}} &= \beta_0(-\lambda\eta_1, -z_2d\eta_2, d\eta_1), \\ \mathbf{m}_{II} &= (\eta_1, -z_2\eta_2, \eta_1), \end{aligned}$$

Therefore,

$$\mu\mathbf{a}_{U_1} + (1 - \mu)\mathbf{a}_{U_{11}} = \eta_1\beta_0\left(\mu d - (1 - \mu)\lambda, d\frac{z_2}{z_1}, -\mu\lambda + (1 - \mu)d\right). \quad (2.32)$$

We are now interested in finding which rotations $\mathbf{R} \in SO(3)$, and which $\mu \in [0, 1]$ are such that (T1)–(T3) are satisfied. First, we claim that $\mathbf{R} \in \mathcal{P}_{24}$ and $\mathbf{R}$ is a rotation whose angle and axis given in the first column of Table 2.1. Indeed, given $\mathbf{R} \in SO(3)$, by Proposition 2.1.1 $(\mathbf{R}\mathbf{U}_1\mathbf{R}^T, \mathbf{R}\mathbf{b}_{II}, \mathbf{R}\mathbf{m}_{II})$ is a type II twin generated by $\mathbf{R}\mathbf{U}_1\mathbf{R}^T, \mathbf{R}\mathbf{U}_{11}\mathbf{R}^T$ and satisfies the cofactor conditions. Therefore, (T1) together with Remark 2.1.1 imply that $\mathbf{R}\mathbf{U}_1\mathbf{R}^T, \mathbf{R}\mathbf{U}_{11}\mathbf{R}^T$ must be a pair in column (A) of Table 2.1 (or in the second column of Table 2.2 in the degenerate case where $a = c$). As a consequence, there exists $\mathbf{Q} \in \mathcal{P}_{24}$, with $\mathbf{Q}$ being a rotation of angle and axis given in the first column of Table 2.1, such that $\mathbf{Q}\mathbf{U}_1\mathbf{Q}^T = \mathbf{R}\mathbf{U}_1\mathbf{R}^T$ and $\mathbf{Q}\mathbf{U}_{11}\mathbf{Q}^T = \mathbf{R}\mathbf{U}_{11}\mathbf{R}^T$. Thus, $\mathbf{U}_1 = \mathbf{Q}^T\mathbf{R}\mathbf{U}_1\mathbf{R}^T\mathbf{Q}$, $\mathbf{U}_{11} = \mathbf{Q}^T\mathbf{R}\mathbf{U}_{11}\mathbf{R}^T\mathbf{Q}$, and Remark 1.4.1 together with Lemma 2.4.1, imply that $\mathbf{Q}^T\mathbf{R} = \mathbf{1}$ or $\mathbf{Q}^T\mathbf{R} = (2\mathbf{v} \otimes \mathbf{v} - \mathbf{1})$ where $\mathbf{v} \in \mathbb{S}^2$ must be at the same time an eigenvector for $\mathbf{U}_1$ and for $\mathbf{U}_{11}$. Under our hypotheses, there exists no $\mathbf{v} \in \mathbb{S}^2$ being at the same time an eigenvector for $\mathbf{U}_1$ and for $\mathbf{U}_{11}$. Thus $\mathbf{Q} = \mathbf{R}$ concluding the proof of the claim.

Now, by Lemma 2.4.1 we have to check for which μ , and which rotation axis $\mathbf{u} \in \mathbb{S}^2$ among the ones given in the first column of Table 2.1, $\mu\mathbf{a}_{U_1} + (1 - \mu)\mathbf{a}_{U_{11}}$ is parallel or perpendicular to $\mathbf{u}$. Table 2.3 below illustrates the possible cases. As a result, we have two different rotation axes $\mathbf{u}$ which are either parallel or perpendicular to $\mu\mathbf{a}_{U_1} + (1 - \mu)\mathbf{a}_{U_{11}}$ if and only if

1. $\mu = \frac{1}{2}$ and $\frac{z_2}{z_1} = \pm \frac{d-\lambda}{2d}$,
2. $\mu = \frac{d}{\lambda+d}$ and $\frac{z_2}{z_1} = \pm \frac{d-\lambda}{d}$,
3. $\mu = \frac{\lambda}{\lambda+d}$ and $\frac{z_2}{z_1} = \pm \frac{d-\lambda}{d}$.

Rotation axis $\mathbf{e}$	Parallel	Orthogonal
(1, 0, 0)	never	$\mu = \frac{\lambda}{\lambda+d}$
(0, 1, 0)	$d = \lambda$ (never)	never
(0, 0, 1)	never	$\mu = \frac{d}{\lambda+d}$
(1, 0, 1)	never	$\lambda = d$ (never)
(1, 0, -1)	never	$\mu = \frac{1}{2}$ (but $\mathbf{R}\mathbf{n}_{II} = -\mathbf{n}_{II}$)
(1, 1, 0)	$\mu = \frac{d}{\lambda+d}$ and $\frac{z_2}{z_1} = \frac{d-\lambda}{d}$	$d\frac{z_2}{z_1} = (1-\mu)\lambda - \mu d$
(1, -1, 0)	$\mu = \frac{d}{\lambda+d}$ and $\frac{z_2}{z_1} = \frac{\lambda-d}{d}$	$d\frac{z_2}{z_1} = -(1-\mu)\lambda + \mu d$
(0, 1, 1)	$\mu = \frac{\lambda}{\lambda+d}$ and $\frac{z_2}{z_1} = \frac{d-\lambda}{d}$	$d\frac{z_2}{z_1} = -(1-\mu)d + \mu\lambda$
(0, -1, 1)	$\mu = \frac{\lambda}{\lambda+d}$ and $\frac{z_2}{z_1} = \frac{\lambda-d}{d}$	$d\frac{z_2}{z_1} = (1-\mu)d - \mu\lambda$

Table 2.3: Conditions on μ to have $(2\mathbf{e} \otimes \mathbf{e} - \mathbf{1})(\mu\mathbf{a}_{U_1} + (1-\mu)\mathbf{a}_{U_{11}}) = \pm\mu\mathbf{a}_{U_1} + (1-\mu)\mathbf{a}_{U_{11}}$, under the assumption that $d < 1 < \lambda$.

We remark that, in the last two cases, the rotation axes which are either parallel or perpendicular to $\mu\mathbf{a}_{U_1} + (1-\mu)\mathbf{a}_{U_{11}}$ are actually three and not two. As $(U_1, \mathbf{b}_{II}, \mathbf{m}_{II})$ satisfies the cofactor conditions as a type II twin, from (2.10)–(2.11) we deduce that

$$b = \frac{\sqrt{-d^4 + d^2(3 - \lambda^2) + (\lambda^2 - 2)}}{1 + \lambda}, \quad \lambda - c = \frac{1 - d^2}{1 + \lambda},$$

which implies

$$\frac{z_2}{z_1} = \frac{\sqrt{-d^4 + d^2(3 - \lambda^2) + (\lambda^2 - 2)}}{1 - d^2}.$$

After some computations we finally deduce that, under the hypothesis $d < 1 < \lambda$, there exist at least two different rotations such that (T1)–(T2) are satisfied if and only if

$$\begin{aligned} \frac{z_2}{z_1} = \pm \frac{d - \lambda}{2d} &\iff d^2 4(d^2 + \lambda^2 - 2) = (d - \lambda)^2(1 - d^2), \\ \frac{z_2}{z_1} = \pm \frac{d - \lambda}{d} &\iff d^2(d^2 + \lambda^2 - 2) = (d - \lambda)^2(1 - d^2). \end{aligned}$$

It remains to check that (T3) is satisfied by the $\mathbf{Q} \in \mathcal{P}_{24}$ such that $\mathbf{Q}(\mu\mathbf{a}_{U_1} + (1-\mu)\mathbf{a}_{U_{11}}) = \pm\mu\mathbf{a}_{U_1} + (1-\mu)\mathbf{a}_{U_{11}}$. Under our assumptions, there always exist two such $\mathbf{Q}$ for type II half-star twins. For type II star twins this happens if and only if $\eta_1 \neq \pm z_2 \eta_2$. But, recalling that $\eta_1 = -z_1 \eta_2$, we need $z_1^2 \neq z_2^2$. As $U_1 \mathbf{e}_2 = \mathbf{e}_2$ implies $\frac{z_1^2}{z_2^2} = \frac{(a-1)^2}{(c-1)^2}$, we hence deduce that (T3) is satisfied by type II star twins if and only if $a \neq c$ or $a + c \neq 2$. But the latter case never occurs as by (2.10) $a + c = 2$ implies $\lambda = 1$, contradicting our assumption that $\lambda > 1$.

Case 2: $0 < \lambda < 1 < d$, $\lambda = ac - b^2$. Here, the eigenvectors related to $\lambda, 1, d$ for $\mathbf{U}_1$ and $\mathbf{U}_{11}$ are respectively $\mathbf{e}_i$ and $\mathbf{v}_i$, given by

$$\begin{aligned}\mathbf{e}_1 &= (z_1, z_2, 0)^T, & \mathbf{e}_2 &= (-z_2, z_1, 0)^T, & \mathbf{e}_3 &= (0, 0, 1)^T, \\ \mathbf{v}_1 &= (0, z_2, z_1)^T, & \mathbf{v}_2 &= (0, z_1, -z_2)^T, & \mathbf{v}_3 &= (1, 0, 0)^T.\end{aligned}$$

where, again, we assume without loss of generality that $z_1 > 0$. Here

$$\mathbf{R}_1^\pm \mathbf{U}_1 = \mathbf{1} + \mathbf{a}_1^\pm \otimes \mathbf{n}_1^\pm, \quad \mathbf{R}_{11}^\pm \mathbf{U}_{11} = \mathbf{1} + \mathbf{a}_{11}^\pm \otimes \mathbf{n}_{11}^\pm,$$

where, $\mathbf{n}_1^\pm, \mathbf{n}_{11}^\pm$ are still given by (2.30)–(2.31), and

$$\mathbf{a}_1^\pm = \beta_0(-d\eta_1\mathbf{e}_1 \pm \lambda\eta_2\mathbf{e}_3), \quad \mathbf{a}_{11}^\pm = \beta_0(-d\eta_1\mathbf{v}_1 \pm \lambda\eta_2\mathbf{v}_3),$$

but this time

$$\eta_1 = -\frac{\sqrt{1-\lambda^2}}{\sqrt{d^2-\lambda^2}}, \quad \eta_2 = \frac{\sqrt{d^2-1}}{\sqrt{d^2-\lambda^2}}, \quad \beta_0 = d - \lambda.$$

From these, we can again find

$$\begin{aligned}\mathbf{a}_{\mathbf{U}_1} &= \beta_0(d\eta_2, -z_2d\eta_1, -\lambda\eta_2), & \mathbf{a}_{\mathbf{U}_{11}} &= \beta_0(-\lambda\eta_2, -z_2d\eta_1, d\eta_2), \\ \mathbf{m}_{II} &= (-\eta_2, z_2\eta_1, -\eta_2),\end{aligned}$$

and $z_1\eta_1 = -\eta_2$. Thus,

$$\mu\mathbf{a}_{\mathbf{U}_1} + (1-\mu)\mathbf{a}_{\mathbf{U}_{11}} = \eta_2\beta_0\left(\mu d - (1-\mu)\lambda, d\frac{z_2}{z_1}, -\mu\lambda + (1-\mu)d\right),$$

which is the same as (2.32). We can hence argue as in the first case, and deduce that, if $d > 1$, $(\mathbf{U}_1, \mathbf{U}_{11})$ generate a half-star twin if and only if

$$d^2 4(d^2 + \lambda^2 - 2) = (d - \lambda)^2(1 - d^2).$$

In the same way, $(\mathbf{U}_1, \mathbf{U}_{11})$ generates a star twin if and only if $a \neq c$, and

$$d^2(d^2 + \lambda^2 - 2) = (d - \lambda)^2(1 - d^2).$$

Case 3: $0 < \lambda_1 < d = 1 < \lambda_3$, where $\lambda_1\lambda_3 = ac - b^2$ and $\lambda_1 + \lambda_3 = a + c$. Repeating the above arguments, this time we deduce that

$$\begin{aligned}\mathbf{a}_{\mathbf{U}_1} &= \beta_0\eta_1(-z_1(\lambda_1 + \lambda_3), (\lambda_1 z_1^2 z_2^{-1} - \lambda_3 z_2), 0), \\ \mathbf{a}_{\mathbf{U}_{11}} &= \beta_0\eta_1(0, (\lambda_1 z_1^2 z_2^{-1} - \lambda_3 z_2), -z_1(\lambda_1 + \lambda_3)), \\ \mathbf{m}_{II} &= (0, \eta_1 z_2^{-1}, 0),\end{aligned}$$

where

$$\eta_1 = -\frac{\sqrt{1-\lambda_1^2}}{\sqrt{\lambda_3^2-\lambda_1^2}}, \quad \eta_2 = \frac{\sqrt{\lambda_3^2-1}}{\sqrt{\lambda_3^2-\lambda_1^2}}, \quad \beta_0 = \lambda_3 - \lambda_1,$$

and $z_1\eta_1 = z_2\eta_2$. Thus,

$$\mu \mathbf{a}_{\mathbf{U}_1} + (1-\mu) \mathbf{a}_{\mathbf{U}_{11}} = -\eta_1 \beta_0 \left(\mu z_1(\lambda_1 + \lambda_3), (\lambda_3 z_2 - \lambda_1 z_1^2 z_2^{-1}), (1-\mu) z_1(\lambda_1 + \lambda_3) \right).$$

Define

$$\beta_1 = z_1(\lambda_1 + \lambda_3), \quad \beta_2 = (\lambda_3 z_2 - \lambda_1 z_1^2 z_2^{-1}).$$

In this case, Table 2.3 becomes Table 2.4. Here we neglected all the rotations whose

Rotation axis $\mathbf{e}$	Parallel	Orthogonal
$(1, 1, 0)$	$\mu = 1$ and $\beta_1 = \beta_2$	$\mu\beta_1 = -\beta_2$
$(1, -1, 0)$	$\mu = 1$ and $\beta_1 = -\beta_2$	$\mu\beta_1 = \beta_2$
$(0, 1, 1)$	$\mu = 0$ and $\beta_1 = \beta_2$	$(1-\mu)\beta_1 = -\beta_2$
$(0, -1, 1)$	$\mu = 0$ and $\beta_1 = -\beta_2$	$(1-\mu)\beta_1 = \beta_2$

Table 2.4: Conditions on μ to have $(2\mathbf{e} \otimes \mathbf{e} - \mathbf{1})(\mu \mathbf{a}_{\mathbf{U}_1} + (1-\mu) \mathbf{a}_{\mathbf{U}_{11}}) = \mu \mathbf{a}_{\mathbf{U}_1} + (1-\mu) \mathbf{a}_{\mathbf{U}_{11}}$, under the assumption that $\lambda_1 < 1 < \lambda_3$.

axes, due to Lemma 2.4.1, do not satisfy (T3). Thus, the only option to form a type II half-star twin is

$$\mu = \frac{1}{2} \quad \text{and} \quad \frac{1}{2}\beta_1 = \pm\beta_2.$$

As $0 < \lambda_1, \lambda_3$ and $\lambda_1, \lambda_3 \neq 1$, $\frac{1}{2}\beta_1 = \pm\beta_2$ simplifies to

$$\lambda_1^2(5\lambda_3^2 - 1) - 8\lambda_1\lambda_3 - \lambda_3^2 + 5 = 0,$$

and the solutions satisfying $\lambda_3 > 1 > \lambda_1$ are

$$\lambda_1 = \frac{4\lambda_3 \pm \sqrt{5}(\lambda_3^2 - 1)}{5\lambda_3^2 - 1}, \quad \lambda_3 > 1.$$

□

In a similar way, we can prove

Theorem 2.4.2. *Let $\mathcal{M} := \cup_{i=1}^{12} \mathbf{U}_i$, where the $\mathbf{U}_i$ are the positive definite matrices, with $a, b, c, d > 0$, given in (2.8). Then, a type I twin generated by $\mathbf{U}_i, \mathbf{U}_j$, $i \neq j$*

$j = 1, \dots, 12$, satisfying the cofactor conditions is a I half-star twin if and only the eigenvalues of $\mathbf{U}_1$, namely $\lambda_1 \leq \lambda_2 = 1 \leq \lambda_3$ satisfy one of the following conditions

$$\begin{aligned} 4(2d^2\lambda_3^2 - \lambda_3^2 - d^2) &= (d - \lambda_3)^2(1 - d^2), & \text{if } d \neq 1 \text{ and } d = \lambda_1, \\ 4(2d^2\lambda_1^2 - \lambda_1^2 - d^2) &= (d - \lambda_1)^2(1 - d^2), & \text{if } d \neq 1 \text{ and } d = \lambda_3, \\ \lambda_1^2(5\lambda_3^2 - 1) &= 8\lambda_1\lambda_3 - 5 + \lambda_3^2, & \text{if } d = 1 \text{ and } \lambda_3 > 1. \end{aligned} \quad (2.33)$$

A type I twin generated by $\mathbf{U}_i, \mathbf{U}_j, i \neq j = 1, \dots, 12$, satisfying the cofactor conditions is a I star twin if and only if $a \neq c$ and one of the following holds

$$\begin{aligned} (2d^2\lambda_3^2 - \lambda_3^2 - d^2) &= (d - \lambda_3)^2(1 - d^2), & \text{if } d \neq 1 \text{ and } d = \lambda_1, \\ (2d^2\lambda_1^2 - \lambda_1^2 - d^2) &= (d - \lambda_1)^2(1 - d^2), & \text{if } d \neq 1 \text{ and } d = \lambda_3. \end{aligned} \quad (2.34)$$

Remark 2.4.2. It is possible to solve the equations in (2.33)–(2.34) to get a direct relation between the eigenvalues of the $\mathbf{U}_i$'s. Indeed, (2.33) is equivalent to

$$\begin{aligned} \lambda_1 &= \frac{d^3 - d + d2\sqrt{2}\sqrt{6d^2 - 3 - d^4}}{9d^2 - 5}, & \text{if } \lambda_3 = d \in (1, \sqrt{3 + \sqrt{6}}), \\ \lambda_1 &= \frac{d^3 - d - d2\sqrt{2}\sqrt{6d^2 - 3 - d^4}}{9d^2 - 5}, & \text{if } \lambda_3 = d \in (\sqrt{5}, \sqrt{3 + \sqrt{6}}), \\ \lambda_3 &= \frac{d^3 - d + d2\sqrt{2}\sqrt{6d^2 - 3 - d^4}}{9d^2 - 5}, & \text{if } \lambda_1 = d \in (\sqrt{3 - \sqrt{6}}, 1), \\ \lambda_3 &= \frac{d^3 - d - d2\sqrt{2}\sqrt{6d^2 - 3 - d^4}}{9d^2 - 5}, & \text{if } \lambda_1 = d \in (\sqrt{3 - \sqrt{6}}, \frac{\sqrt{5}}{3}), \\ \lambda_1 &= \frac{4\lambda_3 \pm \sqrt{5}(\lambda_3^2 - 1)}{5\lambda_3^2 - 1}, & \text{if } \lambda_3 > 1, d = 1, \end{aligned}$$

when $0 < \lambda_1 \leq 1 \leq \lambda_3$. Similarly, the condition (2.29) can be rewritten as

$$\begin{aligned} \lambda_1 &= \frac{d^3 - d + d\sqrt{6d^2 - 3 - 2d^4}}{3d^2 - 2}, & \text{if } \lambda_3 = d \in (1, \sqrt{\frac{(3 + \sqrt{3})}{2}}), \\ \lambda_1 &= \frac{d^3 - d - d\sqrt{6d^2 - 3 - 2d^4}}{3d^2 - 2}, & \text{if } \lambda_3 = d \in (\sqrt{2}, \sqrt{\frac{(3 + \sqrt{3})}{2}}), \\ \lambda_3 &= \frac{d^3 - d + d\sqrt{6d^2 - 3 - 2d^4}}{3d^2 - 2}, & \text{if } \lambda_1 = d \in (\sqrt{\frac{(3 - \sqrt{3})}{2}}, 1), \\ \lambda_3 &= \frac{d^3 - d - d\sqrt{6d^2 - 3 - 2d^4}}{3d^2 - 2}, & \text{if } \lambda_1 = d \in (\sqrt{\frac{(3 - \sqrt{3})}{2}}, \sqrt{\frac{2}{3}}), \end{aligned}$$

$0 < \lambda_1 \leq 1 \leq \lambda_3$.

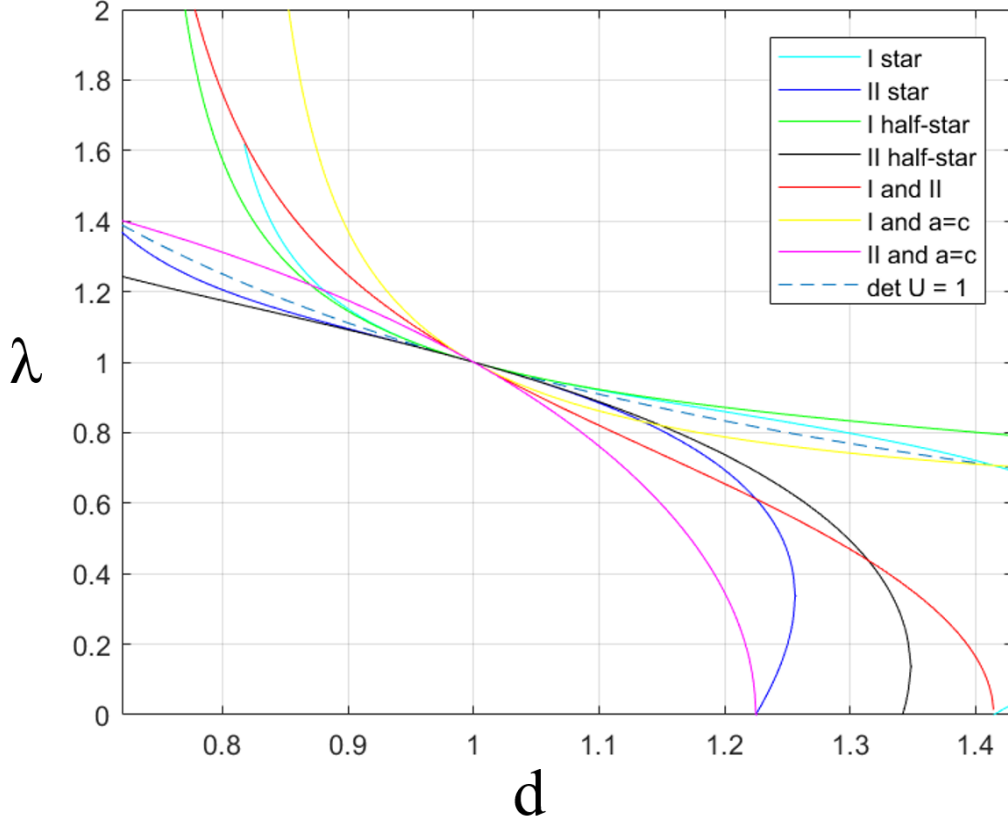


Figure 2.5: Given matrices $U_i \in \mathbb{R}_{Sym^+}^{3 \times 3}$ as in (2.8) describing cubic to monoclinic phase transitions, with eigenvalues $1, \lambda, d$, we plot in cyan and in blue the curves of lattice parameters such that respectively type I star and type II star twins exist. In green and in black the curves of lattice parameters such that respectively type I half-star and type II half-star twins exist. In red the curve of lattice parameters such that both, type I and type II twins exist (see Remark 2.1.1). In magenta and in yellow we plot the curves where $a = c$ (cubic to orthorhombic transformation), and the cofactor conditions are satisfied for type II and type I twin respectively. Finally, the dashed line is the curve where $\det U_i = 1$.

2.5 How closely does a material satisfy the cofactor conditions?

In practice the cofactor conditions are never satisfied exactly. For this reason, an interesting question for application is how closely must a material satisfy the cofactor conditions, in order to behave as if these were satisfied exactly. The problem is complex not only from the point of view of establishing a threshold, but especially in terms of choosing the right metric. There are various ways to measure how closely a material satisfies the cofactor conditions, but, up to our knowledge, just two have been used in the literature up till now. Provided (CC3) holds, the first, more intuitive, way is to check (CC1)-(CC2) directly, that is, to see how close the numbers $|\lambda_2 - 1|$, $|\mathbf{b} \cdot \mathbf{U} \operatorname{cof}(\mathbf{U}^2 - \mathbf{1})\mathbf{m}|$ are to zero. This is, for example, the way the cofactor conditions are measured in [28, 45]. However, there is no physical motivation behind these quantities. On the other hand, provided (CC3) holds, a second way to measure how closely the cofactor conditions are satisfied is to measure how small are the quantities $|\lambda_2 - 1|$, and $||\mathbf{U}\hat{\mathbf{e}}| - 1|$, $||\mathbf{U}^{-1}\hat{\mathbf{e}}| - 1|$, respectively in the case of type II and type I twins. Here $\hat{\mathbf{e}}$ is as in Proposition 1.11. This approach has been adopted for example in [76] and relies on (1.15)–(1.16). However, if we apply these two different metrics to $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$, we obtain contradictory results. Indeed, while following the first approach we get that the lowest value of $|\mathbf{b} \cdot \mathbf{U} \operatorname{cof}(\mathbf{U}^2 - \mathbf{1})\mathbf{m}|$ among the possible twin systems is approximately equal to $4.11 \cdot 10^{-5}$ for type I twins, and to $3.76 \cdot 10^{-5}$ for type II twins, while we have that the second approach leads to lowest values of $||\mathbf{U}^{-1}\hat{\mathbf{e}}| - 1| = 8.1 \cdot 10^{-3}$ and $||\mathbf{U}\hat{\mathbf{e}}| - 1| = 4.15 \cdot 10^{-4}$. Therefore, while the first approach entails that the cofactor conditions are closely satisfied by both type I and type II twins, the second states that $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$ satisfies the cofactor conditions with type II twins much better than with type I twins. Furthermore, for the completely different alloy $\text{Ti}_{23}\text{Nb}_3\text{Al}_{74}$ (see [47]), where the experimentally observed microstructures look very different from those in $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$, we have $\lambda_2 = 0.999996$ and values for $|\mathbf{b} \cdot \mathbf{U} \operatorname{cof}(\mathbf{U}^2 - \mathbf{1})\mathbf{m}|$ of $4.3754 \cdot 10^{-5}$ and $3.7524 \cdot 10^{-5}$ respectively for type I and type II twins. These values would lead to the conclusion that $\text{Ti}_{23}\text{Nb}_3\text{Al}_{74}$ satisfies the cofactor conditions as closely as $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$, even if their behaviours are quite different. However, in $\text{Ti}_{23}\text{Nb}_3\text{Al}_{74}$ we get $||\mathbf{U}^{-1}\hat{\mathbf{e}}| - 1| = 9.9 \cdot 10^{-3}$ and $||\mathbf{U}\hat{\mathbf{e}}| - 1| = 8.3 \cdot 10^{-3}$, which seem to confirm that the cofactor conditions are not so closely satisfied in this material.

Following the results of Section 2.2.1 (see Remark 2.2.1), it seems reasonable for compound twins to measure the quantity $|d - 1|$, under the condition that d is the middle eigenvalue. For type I and type II, we want to use a different strategy, which is based on critical shear stresses, which generate plastic effects. Theorem 2.3.1 and Theorem 2.3.2 tell us that if a material satisfies the cofactor conditions, then there exist triple junctions. These allow a lot of flexibility in the microstructures, as one can create a laminate with an arbitrary volume fraction, which is compatible with austenite without an interface layer. For these reasons, the presence of triple junctions might play an important role in the reversibility of the transformations. We want thus to measure how close a certain twin is to create triple junctions. Under some simplifying assumptions, below we quantify the shear stress necessary to deform austenite in such a way that triple junctions are possible. If this shear stress is small, then the material behaves as if the cofactor conditions were satisfied; if this shear stress is large, then triple junctions are energetically too expensive to be observed. We then apply the new metric to $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$ and to $\text{Ti}_{23}\text{Nb}_3\text{Al}_{74}$. We deduce that, in the former material, the cofactor conditions are better satisfied by type II twins than by type I twins; in the latter material our metric seems to confirm that triple junctions require enough elastic energy to induce plastic deformation.

Below we consider $\mathbf{U}, \mathbf{V} \in \mathbb{R}_{Sym}^{3 \times 3}$ satisfying (1.11), $\mathbf{U} \neq \mathbf{V}$, and assume that $(\mathbf{U}, \mathbf{b}_I, \mathbf{m}_I), (\mathbf{U}, \mathbf{b}_{II}, \mathbf{m}_{II})$ are respectively the type I and type II twins generated by $\mathbf{U}, \mathbf{V}$. We now want to look for $\epsilon \in \mathbb{R}^{3 \times 3}$, such that

$$\tilde{\mathbf{R}}_a \mathbf{U} = \mathbf{1} + \epsilon + \mathbf{c}_a \otimes \mathbf{o}_a, \quad \tilde{\mathbf{R}}_b \mathbf{V} = \mathbf{1} + \epsilon + \mathbf{c}_b \otimes \mathbf{o}_b, \quad \tilde{\mathbf{R}}_b \mathbf{V} - \tilde{\mathbf{R}}_a \mathbf{U} = \tilde{\mathbf{R}}_a \mathbf{b} \otimes \mathbf{m}, \quad (2.35)$$

for some $\tilde{\mathbf{R}}_a, \tilde{\mathbf{R}}_b \in SO(3)$, $\mathbf{c}_a, \mathbf{c}_b \in \mathbb{R}^3$, $\mathbf{o}_a, \mathbf{o}_b \in \mathbb{S}^2$ and where either $\mathbf{b} = \mathbf{b}_I$ and $\mathbf{m} = \mathbf{m}_I$ or $\mathbf{b} = \mathbf{b}_{II}$ and $\mathbf{m} = \mathbf{m}_{II}$. Here and below we also assume that $(\mathbf{U}, \mathbf{b}, \mathbf{m})$ satisfies 1.4.1. Physically, we want to find an elastic deformation of austenite, which allows for triple junctions, and hence is compatible with a laminate without an interface layer (see Figure 2.6). This is an approximation of what happens in reality, where also martensite is deformed, but the above assumptions allow us both to stick to simple computations, and to get qualitative results. Since we are interested in the case where we are close to satisfying the cofactor conditions, by Theorem 2.3.1 and Theorem 2.3.2 we expect this deformation to be small, and we hence stick to the context of linear elasticity, where we can express the deformation gradient of the elastically perturbed austenite as $\mathbf{1} + \epsilon$. Furthermore, as ϵ is expected to be small, we will also adopt the

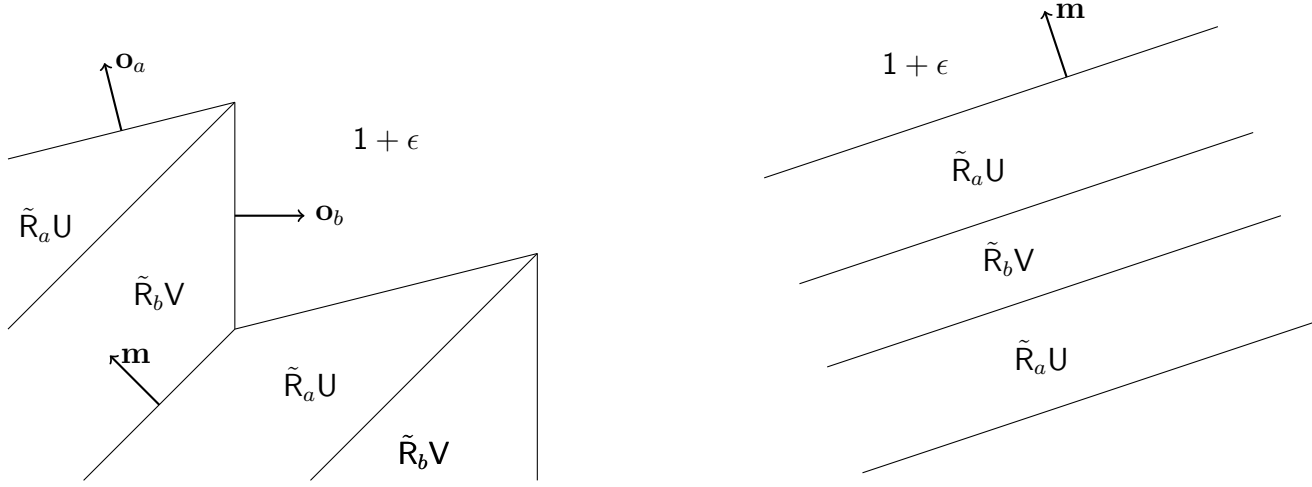


Figure 2.6: Example of triple junctions for type I and type II twins with perturbed austenite.

following approximation

$$\epsilon + \epsilon^T \approx \epsilon + \epsilon^T + \epsilon^T \epsilon = (\tilde{R}_a U - \mathbf{c}_a \otimes \mathbf{o}_a)^T (\tilde{R}_a U - \mathbf{c}_a \otimes \mathbf{o}_a) - 1. \quad (2.36)$$

We remark that the conditions in (2.35) imply

$$\mathbf{c}_b \otimes \mathbf{o}_b - \mathbf{c}_a \otimes \mathbf{o}_a = \tilde{R}_a \mathbf{b} \otimes \mathbf{m}$$

that is, either $\mathbf{o}_a = \pm \mathbf{o}_b = \pm \mathbf{m}$, or $\mathbf{c}_a \parallel \mathbf{c}_b \parallel \tilde{R}_a \mathbf{b}$. Therefore, the only possibilities are:

$$\mathbf{c}_a \otimes \mathbf{o}_a = -\tilde{R}_a \mathbf{b} \otimes \mathbf{o}, \quad (2.37)$$

$$\mathbf{c}_a \otimes \mathbf{o}_a = -\tilde{R}_a \mathbf{c} \otimes \mathbf{m}, \quad (2.38)$$

for some $\mathbf{c}, \mathbf{o} \in \mathbb{R}^3$.

We now want to minimise the energy of $\epsilon + \epsilon^T$ with respect to $\mathbf{c}, \mathbf{o}$ depending on the case. For simplicity, following the approach in [84], we consider just the shear component of the energy, that is we consider the energy of $\epsilon + \epsilon^T$ to be given by

$$\frac{G}{2} \|\epsilon + \epsilon^T\|^2, \quad (2.39)$$

where G is the shear modulus for the material. The following two lemmas are useful in what follows, as they allow to find the energy minimisers, under the assumption that our elastic deformation is small enough.

Lemma 2.5.1. *Let $\mathbf{U}, \mathbf{V} \in \mathbb{R}_{Sym+}^{3 \times 3}$ satisfy (1.11), and let $(\mathbf{U}, \mathbf{b}, \mathbf{m})$ be a type I or a type II solution of the twinning equation (1.10) for $\mathbf{U}, \mathbf{V}$. Suppose further that the minimum and the maximum eigenvalues of $\mathbf{U}$, denoted respectively by λ_1 and λ_3 , satisfy $2\lambda_1^2 - \lambda_3^2 > 0$, and define*

$$\mathbf{C}(\mathbf{c}) = (\mathbf{U} + \mathbf{c} \otimes \mathbf{m})^T (\mathbf{U} + \mathbf{c} \otimes \mathbf{m}) - 1.$$

Assume also

$$\min_{\mathbf{c} \in \mathbb{R}^3} \|\mathbf{C}(\mathbf{c})\|^2 < \min\{(2\lambda_1^2 - \lambda_3^2)^2, 1\}.$$

Then,

$$\|\mathbf{C}(\hat{\mathbf{c}})\|^2 = \min_{\mathbf{c} \in \mathbb{R}^3} \|\mathbf{C}(\mathbf{c})\|^2,$$

if and only if $\hat{\mathbf{c}} = \hat{\mathbf{c}}^\pm = \pm \frac{\mathbf{U}^{-1}\mathbf{m}}{|\mathbf{U}^{-1}\mathbf{m}|} - \mathbf{U}\mathbf{m}$. The eigenvalues of $\mathbf{C}_ := \mathbf{C}(\hat{\mathbf{c}}^+) = \mathbf{C}(\hat{\mathbf{c}}^-)$ are*

$$\frac{1}{2} \left(\text{tr}(\mathbf{C}_*) - \sqrt{2 \text{tr}(\mathbf{C}_*^2) - (\text{tr} \mathbf{C}_*)^2} \right), \quad 0, \quad \frac{1}{2} \left(\text{tr}(\mathbf{C}_*) + \sqrt{2 \text{tr}(\mathbf{C}_*^2) - (\text{tr} \mathbf{C}_*)^2} \right).$$

Finally, $\|\mathbf{C}_\|^2 = 0$ if and only if $(\mathbf{U}, \mathbf{b}, \mathbf{m})$ is a type II twin satisfying the cofactor conditions.*

Proof. We first notice that $\|\mathbf{C}(\mathbf{c})\|^2$ is a smooth function of $\mathbf{c}$, so its minimizer $\mathbf{c}$ must satisfy the equation

$$(\mathbf{U} + \mathbf{c} \otimes \mathbf{m})(\mathbf{U}^2 - 1 + |\mathbf{c}|^2 \mathbf{m} \otimes \mathbf{m} + \mathbf{U}\mathbf{c} \odot \mathbf{m})\mathbf{m} = \mathbf{0}, \quad (2.40)$$

where we denoted by $\odot$ the symmetrised tensor product $\mathbf{u} \odot \mathbf{w} := \mathbf{u} \otimes \mathbf{w} + \mathbf{w} \otimes \mathbf{u}$. We introduce a change of variable $\mathbf{v} = \mathbf{c} + \mathbf{U}\mathbf{m}$, so that, after rearranging the terms, (2.40) becomes

$$(\mathbf{U}^2 - 1 - \mathbf{U}\mathbf{m} \otimes \mathbf{U}\mathbf{m})\mathbf{v} = -|\mathbf{v}|^2 \mathbf{v}. \quad (2.41)$$

Therefore, $\mathbf{v}$ is either zero or is an eigenvector of

$$\mathbf{D} := (\mathbf{U}^2 - 1 - \mathbf{U}\mathbf{m} \otimes \mathbf{U}\mathbf{m}),$$

and the related eigenvalue must be equal to $-|\mathbf{v}|^2$. We first notice that -1 is an eigenvalue for $\mathbf{D}$ related to the eigenvector $\mathbf{w}_1 := \frac{\mathbf{U}^{-1}\mathbf{m}}{|\mathbf{U}^{-1}\mathbf{m}|}$. Let $\mathbf{w}_2, \mathbf{w}_3$ be the other two eigenvectors of $\mathbf{D}$. We have

$$\mathbf{D} = -\mathbf{w}_1 \otimes \mathbf{w}_1 + \bar{\lambda}_2 \mathbf{w}_2 \otimes \mathbf{w}_2 + \bar{\lambda}_3 \mathbf{w}_3 \otimes \mathbf{w}_3,$$

for some $\bar{\lambda}_2, \bar{\lambda}_3 \in \mathbb{R}, \bar{\lambda}_2 \leq \bar{\lambda}_3$. Let now $\alpha_1, \alpha_2, \alpha_3 \in [0, 1]$ be such that

$$\frac{\mathbf{U}\mathbf{m}}{|\mathbf{U}\mathbf{m}|} = \alpha_1 \mathbf{w}_1 + \alpha_2 \mathbf{w}_2 + \alpha_3 \mathbf{w}_3.$$

The following identity must hold

$$\alpha_1^2 + \alpha_2^2 + \alpha_3^2 = 1. \quad (2.42)$$

Furthermore,

$$\alpha_1 = \left| \mathbf{w}_1 \cdot \frac{\mathbf{U}\mathbf{m}}{|\mathbf{U}\mathbf{m}|} \right| = \frac{1}{|\mathbf{U}\mathbf{m}||\mathbf{U}^{-1}\mathbf{m}|} \geq \frac{\lambda_1}{\lambda_3}, \quad (2.43)$$

where λ_1, λ_3 are respectively the minimum and the maximum eigenvalue of $\mathbf{U}$. Thus,

$$\begin{aligned} (\lambda_3^2 - 1) &\geq \max_{\mathbf{w} \in \mathbb{S}^2, \mathbf{w} \perp \mathbf{w}_1} \mathbf{D}\mathbf{w} \cdot \mathbf{w} = \bar{\lambda}_3 \geq \bar{\lambda}_2 = \min_{\mathbf{w} \in \mathbb{S}^2, \mathbf{w} \perp \mathbf{w}_1} \mathbf{D}\mathbf{w} \cdot \mathbf{w} \\ &\geq (\lambda_1^2 - 1) - \max_{\mathbf{w} \in \mathbb{S}^2, \mathbf{w} \perp \mathbf{w}_1} |\mathbf{U}\mathbf{m} \cdot \mathbf{w}|^2. \end{aligned} \quad (2.44)$$

The last term on the right hand side can be estimated by,

$$-|\mathbf{U}\mathbf{m} \cdot \mathbf{w}|^2 \geq -|\mathbf{U}\mathbf{m}|((\alpha_2\mathbf{w}_2 + \alpha_3\mathbf{w}_3) \cdot \mathbf{w})^2 \geq -\lambda_3^2(\alpha_2^2 + \alpha_3^2) \geq \lambda_1^2 - \lambda_3^2, \quad (2.45)$$

where we made use of (2.42)–(2.43) in the last inequality. Thus, collecting the inequalities in (2.44)–(2.45), we prove

$$\bar{\lambda}_2 \geq 2\lambda_1^2 - 1 - \lambda_3^2. \quad (2.46)$$

In this way, we have

$$\|\mathbf{C}(\mathbf{v} - \mathbf{U}\mathbf{m})\|^2 \geq \max_{\mathbf{w} \in \mathbb{S}^2} |\mathbf{C}(\mathbf{v} - \mathbf{U}\mathbf{m})\mathbf{w} \cdot \mathbf{w}|^2 \geq |\mathbf{C}(\mathbf{v} - \mathbf{U}\mathbf{m})\mathbf{m} \cdot \mathbf{m}|^2 = (1 - |\mathbf{v}|^2)^2.$$

Now, (2.41) implies that, if $\mathbf{v} \neq \mathbf{w}_1$, $|\mathbf{v}|^2$ is equal to $-\bar{\lambda}_2, -\bar{\lambda}_3, 0$. Hence, by (2.44)–(2.46), together with the assumption $2\lambda_1^2 - \lambda_3^2 > 0$, we can conclude

$$\|\mathbf{C}(\mathbf{v} - \mathbf{U}\mathbf{m})\|^2 \geq \min\{(2\lambda_1^2 - \lambda_3^2)^2, 1\}.$$

We have hence proved that, if $\min_{\mathbf{c} \in \mathbb{R}^3} \|\mathbf{C}(\mathbf{c})\|^2 < \min\{(2\lambda_1^2 - \lambda_3^2)^2, 1\}$, the only minimizers are $\hat{\mathbf{c}}^\pm = \pm\mathbf{w}_1 - \mathbf{U}\mathbf{m}$. In this case, we obtain

$$\mathbf{C}_* := \mathbf{C}(\pm\mathbf{w}_1 - \mathbf{U}\mathbf{m}) = \mathbf{U}^2 - \mathbf{1} + (1 + |\mathbf{U}\mathbf{m}|^2)\mathbf{m} \otimes \mathbf{m} - \mathbf{U}^2\mathbf{m} \odot \mathbf{m}.$$

Clearly, $\mathbf{m}$ is an eigenvector of $\mathbf{C}_*$ related to the null eigenvalue. Thus, the two remaining eigenvalues are given by the solutions ρ of

$$-\rho^2 + \text{tr}(\mathbf{C}_*)\rho - \frac{1}{2}((\text{tr}(\mathbf{C}_*))^2 - \text{tr}(\mathbf{C}_*^2)) = 0.$$

We now claim that, if $(\mathbf{U}, \mathbf{b}, \mathbf{m})$ satisfies the cofactor conditions as a type II twin, then $\mathbf{C}_* = 0$. Indeed, by (1.16) we have that $\mathbf{m} \cdot \mathbf{v}_2 = 0$, where $\mathbf{v}_2$ is such that $\mathbf{U}\mathbf{v}_2 = \mathbf{v}_2$.

Then, clearly $\mathbf{w}_1 \cdot \mathbf{v}_2 = 0$, and $\mathbf{v}_2$ is an eigenvector for $\mathbf{C}_*$ related to a null eigenvalue. Therefore, as $\mathbf{C}_*$ is symmetric, we just need to show that $\text{tr } \mathbf{C}_* = \text{tr } \mathbf{U}^2 - 2 - |\mathbf{U}\mathbf{m}|^2 = 0$. But if the cofactor conditions hold as a type II twin, then, thanks to Theorem 2.3.2 and [18, Prop. 4] we deduce that

$$|\mathbf{U}\mathbf{m}|^2 = (\lambda_1^2 + \lambda_3^2 - 1),$$

that is $\text{tr } \mathbf{C}_* = 0$, and thus $\mathbf{C}_* = 0$. Conversely, let $\mathbf{C}_* = 0$, and define $\tilde{\mathbf{R}}_a \in \mathbb{R}^{3 \times 3}$ by

$$\tilde{\mathbf{R}}_a^T = \mathbf{U} + \hat{\mathbf{c}}^\pm \otimes \mathbf{m}.$$

First, by the fact that $\mathbf{C}_* = 0$, we have $\tilde{\mathbf{R}}_a \tilde{\mathbf{R}}_a^T = \mathbf{1}$, thus $\tilde{\mathbf{R}}_a^T$ is an orthogonal matrix, and $\det \tilde{\mathbf{R}}_a^T = \pm 1$. Furthermore,

$$\det \tilde{\mathbf{R}}_a = \det \mathbf{U}(1 + \mathbf{U}^{-1} \hat{\mathbf{c}}^\pm \cdot \mathbf{m}) = \pm |\mathbf{U}^{-1} \mathbf{m}| \det \mathbf{U}.$$

Therefore, if we choose $\mathbf{c}_a = \hat{\mathbf{c}}^+$, $\tilde{\mathbf{R}}_a \in SO(3)$, we have

$$\tilde{\mathbf{R}}_a \mathbf{U} = \mathbf{1} - \tilde{\mathbf{R}}_a \mathbf{c}_a \otimes \mathbf{m}, \quad \tilde{\mathbf{R}}_b \mathbf{V} - \tilde{\mathbf{R}}_a \mathbf{U} = \tilde{\mathbf{R}}_a \mathbf{b} \otimes \mathbf{m}, \quad \tilde{\mathbf{R}}_b \mathbf{V} = \mathbf{1} - \tilde{\mathbf{R}}_a (\mathbf{c}_a + \mathbf{b}) \otimes \mathbf{m},$$

for some $\tilde{\mathbf{R}}_b \in SO(3)$. We can hence construct laminates

$$\tilde{\mathbf{R}}_a (\mathbf{U} + \mu \mathbf{b} \otimes \mathbf{m}) = \mu \tilde{\mathbf{R}}_b \mathbf{V} + (1 - \mu) \tilde{\mathbf{R}}_a \mathbf{U} = \mathbf{1} + \tilde{\mathbf{R}}_a (\mathbf{c}_a + \mu \mathbf{b}) \otimes \mathbf{m},$$

which are rank-one connected to $\mathbf{1}$ for arbitrary volume fractions $\mu \in [0, 1]$. As a consequence of Theorem 1.4.1 $(\mathbf{U}, \mathbf{b}, \mathbf{m})$ satisfies the cofactor conditions. Furthermore, as $\mathbf{C}_* = 0$, $\mathbf{v}_2 := \mathbf{m} \times \mathbf{U}^2 \mathbf{m}$ satisfies $\mathbf{U} \mathbf{v}_2 = \mathbf{v}_2$, and is thus the eigenvector corresponding to the eigenvalue of $\mathbf{U}$ equal to one. Therefore, (1.15) finally implies that $(\mathbf{U}, \mathbf{b}, \mathbf{m})$ cannot be a type I twin, and, by assumption, must hence be a type II twin. $\square$

In a similar way, we can prove the following result

Lemma 2.5.2. *Let $\mathbf{E}(\mathbf{o}) = (\mathbf{U} + \mathbf{b} \otimes \mathbf{o})^T (\mathbf{U} + \mathbf{b} \otimes \mathbf{o}) - \mathbf{1}$. Let also*

$$\min_{\mathbf{o} \in \mathbb{R}^3} \|\mathbf{E}(\mathbf{o})\|^2 < 1.$$

Then,

$$\|\mathbf{E}(\hat{\mathbf{o}}^\pm)\|^2 = \min_{\mathbf{o} \in \mathbb{R}^3} \|\mathbf{E}(\mathbf{o})\|^2,$$

if and only if $\hat{\mathbf{o}}^\pm = \pm \frac{\mathbf{U}^{-1} \mathbf{b}}{|\mathbf{U}^{-1} \mathbf{b}| |\mathbf{b}|} - \frac{\mathbf{U} \mathbf{b}}{|\mathbf{b}|^2}$. The eigenvalues of $\mathbf{E}_ = \mathbf{E}(\hat{\mathbf{o}}^+) = \mathbf{E}(\hat{\mathbf{o}}^-)$ are*

$$\frac{1}{2} \left(\text{tr}(\mathbf{E}_*) - \sqrt{2 \text{tr}(\mathbf{E}_*^2) - (\text{tr } \mathbf{E}_*)^2} \right), \quad 0, \quad \frac{1}{2} \left(\text{tr}(\mathbf{E}_*) + \sqrt{2 \text{tr}(\mathbf{E}_*^2) - (\text{tr } \mathbf{E}_*)^2} \right).$$

Finally, $\|\mathbf{E}(\hat{\mathbf{o}})\|^2 = 0$ if and only if $(\mathbf{U}, \mathbf{b}, \mathbf{m})$ is a type I twin satisfying the cofactor conditions.

Proof. Again, $\|\mathbf{E}(\mathbf{o})\|^2$ is a smooth function of $\mathbf{o}$, so its minimizer $\mathbf{o}$ must be a stationary point, that is

$$(\mathbf{U}^2 - 1 + |\mathbf{b}|^2 \mathbf{o} \otimes \mathbf{o} + \mathbf{U}\mathbf{b} \odot \mathbf{o})(\mathbf{U}\mathbf{b} + |\mathbf{b}|^2 \mathbf{o}) = \mathbf{0}.$$

The change of variable $\mathbf{v} = |\mathbf{b}|^2 \mathbf{o} + \mathbf{U}\mathbf{b}$ thus entails

$$\left(\mathbf{U}^2 - 1 - \frac{1}{|\mathbf{b}|^2} \mathbf{U}\mathbf{b} \otimes \mathbf{U}\mathbf{b} \right) \mathbf{v} = -\frac{|\mathbf{v}|^2}{|\mathbf{b}|^2} \mathbf{v}.$$

As in the case of Lemma 2.5.1, $\mathbf{v}$ is either zero or is an eigenvector of

$$\mathbf{F} := \mathbf{U}^2 - 1 - \frac{1}{|\mathbf{b}|^2} \mathbf{U}\mathbf{b} \otimes \mathbf{U}\mathbf{b},$$

and the related eigenvalue must be equal to $-\frac{|\mathbf{v}|^2}{|\mathbf{b}|^2}$. As in the previous case, -1 is an eigenvalue for $\mathbf{F}$ related to the eigenvector $\mathbf{w}_1 := \frac{\mathbf{U}^{-1}\mathbf{b}}{|\mathbf{U}^{-1}\mathbf{b}|}$. If we choose $\mathbf{v} \neq \pm \mathbf{w}_1$, then it must hold $\mathbf{v} \cdot \mathbf{w}_1 = 0$, and therefore

$$\mathbf{E}\left(\pm \frac{\mathbf{v}}{|\mathbf{b}|^2} - \frac{\mathbf{U}\mathbf{b}}{|\mathbf{b}|^2}\right) = \mathbf{U}^2 - 1 - \frac{1}{|\mathbf{b}|^2} \mathbf{U}\mathbf{b} \otimes \mathbf{U}\mathbf{b} + \frac{1}{|\mathbf{b}|^2} \mathbf{v} \otimes \mathbf{v},$$

has -1 as eigenvalue related to the eigenvector $\mathbf{w}_1$. In this case, hence,

$$\left\| \mathbf{E}\left(\pm \frac{\mathbf{v}}{|\mathbf{b}|^2} - \frac{\mathbf{U}\mathbf{b}}{|\mathbf{b}|^2}\right) \right\|^2 \geq \max_{\mathbf{w} \in \mathbb{S}^2} \left| \mathbf{E}\left(\pm \frac{\mathbf{v}}{|\mathbf{b}|^2} - \frac{\mathbf{U}\mathbf{b}}{|\mathbf{b}|^2}\right) \mathbf{w} \cdot \mathbf{w} \right|^2 \geq \left| \mathbf{E}\left(\pm \frac{\mathbf{v}}{|\mathbf{b}|^2} - \frac{\mathbf{U}\mathbf{b}}{|\mathbf{b}|^2}\right) \mathbf{w}_1 \cdot \mathbf{w}_1 \right|^2 \geq 1.$$

Therefore, $\|\mathbf{E}(\mathbf{n})\| < 1$ if and only if $\mathbf{v} = \pm |\mathbf{b}|^{-1} \mathbf{w}_1$. On the other hand,

$$\mathbf{E}_* := \mathbf{E}\left(\pm \frac{\mathbf{w}_1}{|\mathbf{b}|} - \frac{\mathbf{U}\mathbf{b}}{|\mathbf{b}|^2}\right) = \mathbf{U}^2 - 1 - \frac{1}{|\mathbf{b}|^2} \mathbf{U}\mathbf{b} \otimes \mathbf{U}\mathbf{b} + \mathbf{w}_1 \otimes \mathbf{w}_1,$$

has always 0 as an eigenvalue related to the eigenvector $\mathbf{w}_1$. Therefore the other two eigenvalues are given by the solutions ρ of

$$-\rho^2 + \text{tr}(\mathbf{E}_*)\rho - \frac{1}{2}((\text{tr}(\mathbf{E}_*))^2 - \text{tr}(\mathbf{E}_*^2)) = 0.$$

We now want to prove that, if $(\mathbf{U}, \mathbf{b}, \mathbf{m})$ satisfies the cofactor conditions as a type I twin, then $\mathbf{E}_* = \mathbf{0}$. By (1.15) we have that $\mathbf{b} \cdot \mathbf{v}_2 = 0$, where $\mathbf{v}_2$ is such that $\mathbf{U}\mathbf{v}_2 = \mathbf{v}_2$. So that $\mathbf{v}_2$ is also an eigenvector of $\mathbf{E}_*$, and $\mathbf{E}_* \mathbf{v}_2 = \mathbf{0}$, as much as $\mathbf{w}_1 \cdot \mathbf{v}_2 = 0$. Therefore, we just need to check that $\text{tr} \mathbf{E}_* = 0$, that is, $|\mathbf{b}|^{-2} |\mathbf{U}\mathbf{b}|^2 = (\lambda_1^2 + \lambda_3^2 - 1)$. But again, by Theorem 2.3.1, we know that $|\mathbf{b}|^{-1} \mathbf{R}_0 \mathbf{b} = |\mathbf{a}|^{-1} \mathbf{a}$, and

$$|\mathbf{b}|^{-1} \mathbf{U}\mathbf{b} = |\mathbf{b}|^{-1} \mathbf{U} \mathbf{R}_0^T \mathbf{R}_0 \mathbf{b} = |\mathbf{a}|^{-1} (1 + \mathbf{n}_0 \otimes \mathbf{a}) \mathbf{a}.$$

Therefore, exploiting [18, Prop. 4] we thus deduce

$$\frac{|\mathbf{U}\mathbf{b}|^2}{|\mathbf{b}|^2} = \left\| \frac{\mathbf{a}}{\|\mathbf{a}\|} + \|\mathbf{a}\|\mathbf{n}_U \right\|^2 = (\lambda_1^2 + \lambda_3^2 - 1),$$

as required. Conversely, let $\mathbf{E}_* = 0$, and define $\tilde{\mathbf{R}}_a \in \mathbb{R}^{3 \times 3}$ by

$$\tilde{\mathbf{R}}_a^T = \mathbf{U} + \mathbf{b} \otimes \hat{\mathbf{o}}^\pm.$$

As $\mathbf{E}_* = 0$, we have $\tilde{\mathbf{R}}_a \tilde{\mathbf{R}}_a^T = 1$, thus $\tilde{\mathbf{R}}_a^T$ is an orthogonal matrix, and $\det \tilde{\mathbf{R}}_a^T = \pm 1$. Furthermore,

$$\det \tilde{\mathbf{R}}_a = \det \mathbf{U}(1 + \mathbf{U}^{-1}\mathbf{b} \cdot \hat{\mathbf{o}}^\pm) = \pm |\mathbf{U}^{-1}\mathbf{b}| \det \mathbf{U}.$$

Therefore, if we choose $\mathbf{o}_a = \hat{\mathbf{o}}^+$, $\tilde{\mathbf{R}}_a \in SO(3)$, we have

$$\tilde{\mathbf{R}}_a \mathbf{U} = 1 - \tilde{\mathbf{R}}_a \mathbf{b} \otimes \mathbf{o}_a, \quad \tilde{\mathbf{R}}_b \mathbf{V} - \tilde{\mathbf{R}}_a \mathbf{U} = \tilde{\mathbf{R}}_a \mathbf{b} \otimes \mathbf{m}, \quad \tilde{\mathbf{R}}_b \mathbf{V} = 1 - \tilde{\mathbf{R}}_a \mathbf{b} \otimes (\mathbf{o}_a + \mathbf{m}),$$

for some $\tilde{\mathbf{R}}_b \in SO(3)$. We can hence construct laminates compatible with 1 for arbitrary volume fractions, and hence, by Theorem 1.4.1 the cofactor conditions must be satisfied. Furthermore, as $\mathbf{E}_* = 0$, $\mathbf{v}_2 := (\mathbf{U}\mathbf{b}) \times (\mathbf{U}^{-1}\mathbf{b})$ satisfies $\mathbf{U}\mathbf{v}_2 = \mathbf{v}_2$, and hence $\mathbf{v}_2$ is the eigenvector related to the eigenvalue of $\mathbf{U}$ equal to one. Therefore, $\mathbf{v}_2 \cdot \mathbf{b} = \mathbf{v}_2 \cdot \mathbf{U}\mathbf{b} = 0$, and (1.15) finally implies that $(\mathbf{U}, \mathbf{b}, \mathbf{m})$ cannot be a type I twin. As a consequence, by hypotheses, $(\mathbf{U}, \mathbf{b}, \mathbf{m})$ must be a type I twin. $\square$

By putting together (2.36) and (2.39), together with Lemma 2.5.1 and Lemma 2.5.2 we have that, if we are close to satisfy the cofactor conditions, the stress induced by ϵ is given by

$$\begin{aligned} \sigma(\epsilon) &= G\mathbf{R}_{SI} \operatorname{diag}(0, \omega_1, \omega_2) \mathbf{R}_{SI}^T, & \text{for type I,} \\ \sigma(\epsilon) &= G\mathbf{R}_{SII} \operatorname{diag}(0, \theta_1, \theta_2) \mathbf{R}_{SII}^T, & \text{for type II,} \end{aligned}$$

for some $\mathbf{R}_{SI}, \mathbf{R}_{SII} \in SO(3)$, and where

$$\begin{aligned} \omega_1 &= \frac{1}{2} \left(\operatorname{tr} \mathbf{E}_* - \sqrt{2 \operatorname{tr}(\mathbf{E}_*^2) - (\operatorname{tr} \mathbf{E}_*)^2} \right), & \omega_2 &= \frac{1}{2} \left(\operatorname{tr}(\mathbf{E}_*) + \sqrt{2 \operatorname{tr}(\mathbf{E}_*^2) - (\operatorname{tr} \mathbf{E}_*)^2} \right), \\ \theta_1 &= \frac{1}{2} \left(\operatorname{tr}(\mathbf{C}_*) - \sqrt{2 \operatorname{tr}(\mathbf{C}_*^2) - (\operatorname{tr} \mathbf{C}_*)^2} \right), & \theta_2 &= \frac{1}{2} \left(\operatorname{tr}(\mathbf{C}_*) + \sqrt{2 \operatorname{tr}(\mathbf{C}_*^2) - (\operatorname{tr} \mathbf{C}_*)^2} \right). \end{aligned}$$

Here, as above,

$$\begin{aligned} \mathbf{E}_* &:= \mathbf{U}^2 - 1 - \frac{1}{|\mathbf{b}|^2} \mathbf{U}\mathbf{b} \otimes \mathbf{U}\mathbf{b} + \frac{\mathbf{U}^{-1}\mathbf{b}}{|\mathbf{U}^{-1}\mathbf{b}|} \otimes \frac{\mathbf{U}^{-1}\mathbf{b}}{|\mathbf{U}^{-1}\mathbf{b}|}, \\ \mathbf{C}_* &:= \mathbf{U}^2 - 1 + (1 + |\mathbf{U}\mathbf{m}|^2) \mathbf{m} \otimes \mathbf{m} - \mathbf{U}^2 \mathbf{m} \odot \mathbf{m}. \end{aligned}$$

If we think now of reversibility as the lack of plastic effects during the phase transition, we can say that we are close to the cofactor conditions if our principal stresses satisfy some yield criterion. Adopting for simplicity Tresca's yield criterion, we can say that we are closely satisfying the cofactor conditions if

$$\begin{aligned}\frac{G}{2} \max\{|\omega_1 - \omega_2|, |\omega_1|, |\omega_2|\} &\leq \sigma_C, & \text{for type I,} \\ \frac{G}{2} \max\{|\theta_1 - \theta_2|, |\theta_1|, |\theta_2|\} &\leq \sigma_C, & \text{for type II,}\end{aligned}$$

where σ_C is the shear yield stress. In the absence of experimental values for G and σ_C we can measure how closely the cofactor conditions are satisfied in a certain alloy by simply computing

$$\begin{aligned}\max\{|\omega_1 - \omega_2|, |\omega_1|, |\omega_2|\}, & \quad \text{for type I twins,} \\ \max\{|\theta_1 - \theta_2|, |\theta_1|, |\theta_2|\}, & \quad \text{for type II twins.}\end{aligned}$$

Lowest values among the possible twin systems for $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$ are 0.0173 and 0.0021 respectively for type I and type II twins. These values seem to confirm that the satisfaction of the cofactor conditions in this material is much closer with some type II twins than with type I twins, in agreement with Proposition 2.1.2. Values for $\text{Ti}_{23}\text{Nb}_3\text{Al}_{74}$ are 0.0267 and 0.0225 respectively for type I and type II twins, confirming that triple junctions require high elastic energy in this material.

Chapter 3

A model for the evolution of highly reversible martensitic transformations

3.1 Introduction

Based on experimental observations for $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$, in Chapter 4.3 we introduce the moving mask approximation. This, under suitable assumptions, enables characterisation of the macroscopic deformation gradients for martensite as the ones of the form

$$1 + \mathbf{a}(\mathbf{x}) \otimes \mathbf{n}(\mathbf{x}), \quad \text{a.e. } \mathbf{x},$$

for some $\mathbf{a}, \mathbf{n} \in L^\infty(\Omega; \mathbb{R}^3)$. This result, which would not have been achievable in a static framework, turned out to be in accordance with experiments on $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$ [27, 76], and appears relevant for many other alloys.

In order to frame the moving mask approximation in a dynamical model in the context of continuum mechanics, with the aim of better understanding the complex microstructures arising in $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$, in this chapter we introduce a simplified model to describe the evolution of thermally induced martensitic phase transitions with ultra-low hysteresis, and hence taking place in a small range of temperature values. This is the case, for example, for $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$ where the experimental values of hysteresis are of 2°K , in contrast to values up to 70°K in NiTi alloys. Our model results in a system of partial differential equations, for which we prove existence of weak solutions, and that we relate to the moving mask approximation of Chapter 4.3. The structure of this chapter is the following: in Section 3.2, we recall the equations

describing the balance of momentum and the balance of energy in continuum mechanics. Indeed, following the work in [1, 2], it seems relevant for the evolution of the austenite-martensite interface to take into account thermodynamic effects. In Section 3.3, we introduce the assumptions on the free-energy density of the system, and the boundary conditions, which reflect the experimental setup for $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$ (see [76]). Furthermore, we justify the linearization of the equation for the conservation of energy around the critical temperature. Existence of solutions to the resulting system of partial differential equations is proved in Section 3.4. This proof requires the application of a version of the div-curl lemma (see e.g., [31]) to show that the solutions to a suitably defined approximated problem converge to weak solutions of our problem, when the regularisation vanishes. In Section 3.5 we study the limit of the solutions as the elastic constants tend to infinity and the interface energy density goes to zero. This approximation is physically relevant and has been already adopted, for example in [15] and [21]. Here we prove that the deformation gradients $\nabla \mathbf{y}$ generate in the limit a gradient Young measure $\nu_{\mathbf{x},t}$ which is supported on the phases, and whose evolution cannot be deduced from the conservation of momentum alone. On the other hand, the equation for the conservation of energy becomes a heat equation with a source term of the type $\frac{d}{dt} \int_{\mathbb{R}^{3 \times 3}} \eta_1(\mathbf{A}) d\nu_{\mathbf{x},t}(\mathbf{A})$, where $\eta_1 - \log\left(\frac{\theta}{\theta_T}\right)$ is the entropy of the system, and η_1 is supposed to be a smooth function, assuming different values in austenite and martensite. In Section 3.6 we recall the moving mask approximation introduced in Chapter 4.3, and explain how it can be framed in the context of our limit problem. In this section, we also devise a possible strategy to make the connection rigorous.

In the last section, we study the propagation of a laminate in a long bar, where we can simplify the problem, under some further assumptions, to a one-dimensional one. The resulting equation is an easier version of the equations introduced in [2], that are underdetermined, and should be closed with a constitutive relation for the velocity of the phase boundary. After introducing a constitutive relation between the velocity and the temperature of the phase boundary, our problem becomes a two-phase Stefan problem with kinetic condition at the free boundary (see e.g., [80, 81]). We prove that, under our assumptions, the position of the austenite-martensite phase interface is a monotone function of time, which reaches the domain boundary in finite time. This simple, one-dimensional case is an example of solution satisfying the moving mask approximation to the limit problem obtained in Section 3.5.

3.2 Dynamics of martensitic phase transitions in continuum solid mechanics

As mentioned in the previous section, the statics of martensite to austenite phase transitions can be successfully modelled in the framework of nonlinear elasticity. For these reasons, many authors have tried to study the evolution of martensitic phase transitions within the same theory, where the conservation of linear momentum can be expressed in Lagrangian coordinates as

$$\rho_0 \ddot{\mathbf{y}} = \text{Div } \boldsymbol{\sigma}. \quad (3.1)$$

Again, here $\mathbf{y}(\mathbf{x}, t)$ is the position at time t of the material point $\mathbf{x} \in \Omega$, $\rho_0(\mathbf{x})$ the density of the body in its reference configuration, and $\boldsymbol{\sigma}$ the Piola-Kirchhoff stress tensor, whose dependence on $\mathbf{y}$ and the deformation gradient $\mathbf{F} = \nabla \mathbf{y}$ is given in terms of the free energy ϕ of the material through the relation

$$\boldsymbol{\sigma}(\mathbf{F}, \theta) = \frac{\partial \phi}{\partial \mathbf{F}}(\mathbf{F}, \theta). \quad (3.2)$$

However, as ϕ is in general not rank-one convex, the operator $\text{Div}(\boldsymbol{\sigma}(\nabla \mathbf{y}, \theta))$ is not elliptic, making the problem extremely complex from a mathematical point of view, and the related initial value problem formally ill posed. Furthermore, the total energy is conserved in the model for smooth solutions, which is not coherent with physical observations, where heat is absorbed or released during the process. Therefore, one should seek solutions with shocks that, as explained for example in [1], dissipate energy, but which are mathematically very difficult to study. Also studied in the literature are viscoelastic models. In the case of viscoelasticity of rate type the stress tensor has the form

$$\boldsymbol{\sigma} = \boldsymbol{\sigma}(\nabla \mathbf{y}, \nabla \dot{\mathbf{y}}, \theta),$$

in (3.1), or in its quasi-static variant where $\ddot{\mathbf{y}}$ is assumed to be negligible. However, many difficulties arise in dimension greater than one due to the fact that frame-indifference prohibits a linear dependence on $\nabla \dot{\mathbf{y}}$ (see e.g., [74]). For further discussion, we refer the reader to [9, 17, 10, 70, 74] and references therein.

On the other hand, the phase transition generates heat, and it is therefore reasonable, when modelling the evolution for the austenite to martensite phase transition, to take into account thermal effects. The equations of continuum mechanics expressing the balance of linear momentum and energy in the absence of heat supply and

neglecting body forces are given by (see e.g., [9, 79])

$$\begin{aligned} \frac{\partial}{\partial t} \left(\rho_0 \frac{1}{2} |\dot{\mathbf{y}}|^2 + U \right) - \text{Div}(\mathbf{C}(\nabla \mathbf{y}, \theta) \nabla \theta) - \text{Div}(\dot{\mathbf{y}} \sigma(\nabla \mathbf{y}, \theta)) &= 0, \\ \rho_0 \ddot{\mathbf{y}} - \text{Div} \sigma(\nabla \mathbf{y}, \theta) &= \mathbf{0}. \end{aligned}$$

Here, U is the internal energy of the system, and we also made use of the Fourier law for the heat flux $\mathbf{q}$, that is $\mathbf{q} = -\mathbf{C} \nabla \theta$. For simplicity $\mathbf{C}$ will be considered to be a constant, symmetric, positive definite matrix in $\mathbb{R}^{3 \times 3}$ with smallest eigenvalue equal to c_m and ρ_0 is assumed constant. Thus, after exploiting the relation $U = \phi + \theta \eta$ relating the internal energy to the entropy η and the free energy ϕ , the above equations become

$$\begin{aligned} \rho_0 \theta \dot{\eta}(\nabla \mathbf{y}, \theta) - \text{Div}(\mathbf{C} \nabla \theta) &= 0, \\ \rho_0 \ddot{\mathbf{y}} - \text{Div} \sigma(\nabla \mathbf{y}, \theta) &= \mathbf{0}, \end{aligned} \tag{3.3}$$

where we also made use of

$$\eta(\nabla \mathbf{y}, \theta) = -\frac{\partial \phi}{\partial \theta}(\nabla \mathbf{y}, \theta). \tag{3.4}$$

3.3 Hypotheses, boundary conditions and approximations

The aim of this section is to introduce some hypotheses on the free-energy density ϕ . Given the fact that we are restricting ourselves to elastic deformations, and that elastic moduli are in general considerable in metals, minimisers of the free energy lie extremely close to the wells. Therefore, it is a reasonable approximation to adopt in Section 3.5 the approach of [21], and let ϕ grow to infinity outside the wells, thus neglecting the influence of the shape of $\phi(\mathbf{F})$ for $\mathbf{F} \notin K \cup SO(3)$ on our results.

Following the approach of [52] and references therein, we split the energy contributions into a term, ϕ_3 below, responsible for the increase of energy with temperature, and terms, namely ϕ_1, ϕ_2 , describing the energy thermo-mechanical effects. As hysteresis is smaller than 5 °K in materials undergoing ultra-reversible martensitic transformations such as $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$, we can neglect the effects of thermal expansion, and we also neglect the heat generated by elastic deformations which are not a change of phase, which, as mentioned above, are usually observed to be small. For $k \geq 1$ we thus consider

$$\phi^k(\mathbf{F}, \theta) = \phi_1(\mathbf{F}, \theta) + k\phi_2(\mathbf{F}) + \phi_3(\theta), \tag{3.5}$$

where:

- $\phi_2: \mathbb{R}^{3 \times 3} \rightarrow [0, \infty)$ is smooth and satisfies

$$\phi_2(\mathbf{F}) = 0 \Leftrightarrow \mathbf{F} \in K \cup SO(3), \quad c_2(|\mathbf{F}|^p + 1) \geq \phi_2(\mathbf{F}) \geq c_0|\mathbf{F}|^p - c_1,$$

with $c_0, c_1, c_2 > 0$ and $p \in (2, \infty)$. When $k \gg 1$, the ϕ_2 term is responsible for the growth of the energy outside the wells, and k plays the role of an elastic constant;

- $\phi_3(\theta) = \gamma\theta\left(1 - \log\left(\frac{\theta}{\theta_T}\right)\right)$, arises naturally by assuming that the specific heat $\gamma > 0$ is constant, and hence independent of the phase and of the temperature. Indeed, ϕ^k must satisfy $\gamma = -\theta \frac{\partial^2 \phi^k}{\partial \theta^2}$ (see [2]);
- $\rho_0 \phi_1 = \bar{\phi}_1(\mathbf{F}) - \theta \eta_1(\mathbf{F})$, where $\bar{\phi}_1, \eta_1: \mathbb{R}^{3 \times 3} \rightarrow \mathbb{R}$ are smooth and such that

$$\begin{aligned} \min_{\mathbb{R}^{3 \times 3}} \phi_1(\cdot, \theta) &= \min \left\{ 0, -\alpha \left(1 - \frac{\theta}{\theta_T} \right) \right\}, \\ \bar{\phi}_1(\mathbf{F}) &= -\alpha, \quad \text{and} \quad \eta_1(\mathbf{F}) = -\frac{\alpha}{\theta_T}, \quad \text{for every } \mathbf{F} \in K, \\ \bar{\phi}_1(\mathbf{F}) &= 0 \quad \text{and} \quad \eta_1(\mathbf{F}) = 0, \quad \text{for every } \mathbf{F} \in SO(3), \end{aligned}$$

for every $\mathbf{F} \in \mathbb{R}^{3 \times 3}$, for some positive constants α, c_a, c_b and $q \in (1, p)$. This term is responsible for the phase transition, that is, this term is responsible for changing the global minimizers of ϕ^k from $SO(3)$ above θ_c to K below θ_c . The behaviour is chosen to be linear in θ because of the small hysteresis;

- the following growth conditions hold

$$\begin{aligned} |\eta_1(\mathbf{F})| + |\bar{\phi}_1(\mathbf{F})| + \left| \frac{\partial \phi_2}{\partial \mathbf{F}}(\mathbf{F}) \right| + \left| \frac{\partial \bar{\phi}_1}{\partial \mathbf{F}}(\mathbf{F}) \right| &\leq c_d(1 + |\mathbf{F}|^r), \\ \left| \frac{\partial \eta_1}{\partial \mathbf{F}}(\mathbf{F}) \right| &\leq c_f(1 + |\mathbf{F}|^{\tilde{r}}), \end{aligned} \tag{3.6}$$

for some positive constants c_d, c_f , $r < p$, $\tilde{r} < \frac{5}{6}p$ and for all $\mathbf{F} \in \mathbb{R}^{3 \times 3}$.

The choice of our free energy is coherent with the one adopted in [52] and references therein, and works particularly well for computations; other choices for the free energy reflecting the change of energetically preferable phase across θ_T would be possible. We finally add to the energy of the system also a term of the type $\frac{1}{k}|\Delta \mathbf{y}|^2$, which penalizes interfaces between martensitic variants. A term of this type is often kept in account when modelling martensitic transformations (see e.g., [18, 33, 37]) and makes the model more accurate. This term is going to disappear from the equations when we send $k \rightarrow \infty$, in the next section. We remark that we could choose $\frac{c_\beta}{k^\beta}$ for some $c_\beta, \beta > 0$ in front of $|\Delta \mathbf{y}|^2$ instead of $\frac{1}{k}$ without affecting the results below.

Furthermore we can replace $|\Delta \mathbf{y}|^2$ with $|\nabla \nabla \mathbf{y}|^2$. Indeed, as these terms differ only by a null Lagrangian, the equations governing the evolution are the same.

Remark 3.3.1. In the definition of the energy we did not take into account the constraint $\phi^k(\mathbf{F}, \cdot) \rightarrow \infty$ as $\det \mathbf{F} \rightarrow 0$, which is usually introduced to avoid interpenetration. This constraint is extremely difficult to handle from a mathematical point of view, as, in this case, there exists no $c > 0$ such that $\left| \frac{\partial \phi^k(\mathbf{F}, \cdot)}{\partial \mathbf{F}} \right| \leq c(1 + |\phi^k(\mathbf{F}, \cdot)|)$. That is, a finite energy deformation map $\mathbf{y} \in W^{1,p}(\Omega; \mathbb{R}^3)$ might be such that $\frac{\partial \phi^k(\nabla \mathbf{y}, \cdot)}{\partial \mathbf{F}} \notin L^1(\Omega; \mathbb{R}^{3 \times 3})$, and there is, up to our knowledge, no reference in the literature on how to control this term. We refer the reader to Section 2.4 in [9]. Nonetheless, the interpenetration constraint holds in the limit $k \rightarrow \infty$, as shown in Section 3.5.

Under the above assumptions, and after adding the term responsible for the interface energy, (3.3) becomes

$$\begin{cases} \rho_0 \ddot{\mathbf{y}} - \operatorname{Div} \boldsymbol{\sigma}^k(\nabla \mathbf{y}, \theta) + \frac{1}{k} \Delta^2 \mathbf{y} &= \mathbf{0}, \\ \gamma \rho_0 \dot{\theta} - \operatorname{Div}(\mathbf{C} \nabla \theta) + \theta \dot{\eta}_1(\nabla \mathbf{y}) &= 0, \end{cases} \quad (3.7)$$

where

$$\boldsymbol{\sigma}^k(\mathbf{F}, \theta) := \frac{\partial \phi^k}{\partial \mathbf{F}}(\mathbf{F}, \theta), \quad \eta_1(\mathbf{F}) := -\frac{\partial \phi_1}{\partial \theta}(\mathbf{F}).$$

We supplement the problem with initial and boundary conditions for θ and for $\mathbf{y}$ reflecting the experimental setup for $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$ in [76]. The initial conditions $\theta_0(\mathbf{x}), \mathbf{y}_a(\mathbf{x}), \mathbf{y}_b(\mathbf{x})$ respectively represent from a physical point of view the temperature, the position and the velocity at the point $\mathbf{x}$ at $t = 0$. The boundary conditions for the equation describing the conservation of momentum reflect the lack of external forces imposed at the boundary of the sample. We consider $\Omega \subset \mathbb{R}^3$ to be bounded and smooth, $\partial\Omega_D \subset \partial\Omega$ to be open and non-empty in $\partial\Omega$, $\partial\Omega_N = \partial\Omega \setminus \partial\Omega_D$, and set

$$\begin{cases} \theta = \theta_B, & \text{on } \partial\Omega_D \times (0, T), \\ \mathbf{C} \nabla \theta \cdot \mathbf{m} = 0, & \text{on } \partial\Omega_N \times (0, T), \\ (\boldsymbol{\sigma}^k(\nabla \mathbf{y}, \theta) + \nabla \Delta \mathbf{y}) \cdot \mathbf{m} = \mathbf{0}, & \text{on } \partial\Omega \times (0, T), \\ \Delta \mathbf{y} = \mathbf{0}, & \text{on } \partial\Omega \times (0, T), \\ \theta|_{t=0} = \theta_0(\mathbf{x}), & \text{in } \Omega, \\ \mathbf{y}|_{t=0} = \mathbf{y}_a(\mathbf{x}), & \text{in } \Omega, \\ \dot{\mathbf{y}}|_{t=0} = \mathbf{y}_b(\mathbf{x}), & \text{in } \Omega, \end{cases} \quad (3.8)$$

where $\theta_B, \theta_0, \mathbf{y}_b, \mathbf{y}_a$ are functions in suitable function spaces, and $\mathbf{m}$ is the unit outer normal to Ω . In a consistent way with the experiments in [76] on $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$, for

which $\theta_T \approx 240^\circ\text{K}$ and the hysteresis is about 2°K , we can assume $\|\theta_B - \theta_T\|_{L^\infty} + \|\theta_T - \theta_0\|_{L^\infty} \ll \theta_T$. In this way, we can formally justify the replacement of the θ in front of $\dot{\eta}_1$ with θ_T , thus rewriting (3.7) as

$$\begin{cases} \rho_0 \ddot{\mathbf{y}} - \text{Div } \boldsymbol{\sigma}^k(\nabla \mathbf{y}, \theta) + \frac{1}{k} \Delta^2 \mathbf{y} &= \mathbf{0}, \\ \gamma \rho_0 \dot{\theta} - \text{Div}(\mathbf{C} \nabla \theta) + \theta_T \dot{\eta}_1(\nabla \mathbf{y}) &= 0. \end{cases} \quad (3.9)$$

Under suitable assumptions on η_1 , a maximum principle such as the one in Proposition 3.7.1 could be proved to rigorously justify this approximation.

Below c, c_T represent two positive constants whose value is dependent just on the parameters, and, in the case of c_T , on the final time T and the initial data. Their value may change from line to line or even within the same line, but is always independent of n, k, ε, M, N . When there is a dependence on k or on M , this constant will be denoted by $c_{T,k}, c_{T,M}$ respectively.

Below, we will denote by $\mathbf{A} : \mathbf{B}$ the Frobenius product between matrices $\mathbf{A}, \mathbf{B} \in \mathbb{R}^{3 \times 3}$ defined as $\mathbf{A} : \mathbf{B} = \text{tr}(\mathbf{A}^T \mathbf{B})$.

3.4 Existence of weak solutions

Existence of solutions in a suitable sense to (3.8)–(3.9) is given by Theorem 3.4.1 below. We will denote by H_D^1 the space of functions $\xi \in H^1(\Omega)$ such that $\xi = 0$ on $\partial\Omega_D$, by H_D^{-1} its dual, and by $L_w^s(\Omega)$ the $L^s(\Omega)$ spaces endowed with the weak topology. Finally, we will denote by $\mathbf{V}$ the Hilbert space

$$\mathbf{V} := \{\mathbf{u} \in L^2(\Omega; \mathbb{R}^3) : \|\mathbf{u}\|_{\mathbf{V}} < \infty\}, \quad \|\mathbf{u}\|_{\mathbf{V}}^2 := \int_{\Omega} (|\mathbf{u}|^2 + |\nabla \mathbf{u}|^2 + |\Delta \mathbf{u}|^2) \, d\mathbf{x},$$

endowed with the scalar product $(\cdot, \cdot)_{\mathbf{V}}$ defined by

$$(\mathbf{u}, \mathbf{w})_{\mathbf{V}} := \int_{\Omega} (\mathbf{u} \cdot \mathbf{w} + \nabla \mathbf{u} : \nabla \mathbf{w} + \Delta \mathbf{u} \cdot \Delta \mathbf{w}) \, d\mathbf{x}.$$

Given the boundary conditions for our problem, we are not able to apply the Miranda-Talenti inequality, and therefore the embedding $\mathbf{V} \hookrightarrow H^1(\Omega; \mathbb{R}^3)$ might not be compact. The following lemma gives us sufficient conditions for this to hold:

Lemma 3.4.1. *Let $\mathbf{u}_j \in \mathbf{V}$ be such that $\|\mathbf{u}_j\|_{\mathbf{V}} \leq c$ and such that $|\nabla \mathbf{u}_j|^2$ is uniformly integrable. Then, there exists $\mathbf{u} \in \mathbf{V}$ and a non-relabelled subsequence $\mathbf{u}_j$ converging weakly in $\mathbf{V}$ and strongly in $H^1(\Omega; \mathbb{R}^3)$ to $\mathbf{u}$. Furthermore, the space $\mathbf{V} \cap W^{1,p}(\Omega; \mathbb{R}^3)$ is compactly embedded in $H^1(\Omega; \mathbb{R}^3)$ for every $p > 2$.*

Proof. By Banach-Alaoglu, we get the existence of $\mathbf{u} \in \mathbf{V}$ and of a non-relabelled subsequence $\mathbf{u}_j$ converging weakly in $\mathbf{V}$. By Sobolev embeddings, $\mathbf{u}_j \rightarrow \mathbf{u}$ also strongly in $L^2(\Omega; \mathbb{R}^3)$. Now, we notice that, on the one hand, $\nabla \times \nabla(\mathbf{u} - \mathbf{u}_j) = 0$ in the sense of distributions. On the other hand $\|\operatorname{Div} \nabla(\mathbf{u} - \mathbf{u}_j)\|_2 \leq c$. Therefore, an application of a version of the div-curl lemma (see e.g., [31]) entails

$$\lim_j \|\nabla \mathbf{u}_j - \nabla \mathbf{u}\|_2^2 = \lim_j \int_{\Omega} \nabla(\mathbf{u} - \mathbf{u}_j) : \nabla(\mathbf{u} - \mathbf{u}_j) \, d\mathbf{x} = 0.$$

Now, we recall that for a sequence $\mathbf{u}_j \in W^{1,p}(\Omega; \mathbb{R}^3)$ with $\|\mathbf{u}_j\|_{W^{1,p}} \leq c$, for some $c > 0$ $p > 2$, $|\nabla \mathbf{u}_j|^2$ is uniformly integrable. Therefore, thanks to the first statement of the lemma, we deduce that $\mathbf{V} \cap W^{1,p}(\Omega; \mathbb{R}^3)$ is compactly embedded in $H^1(\Omega; \mathbb{R}^3)$. $\square$

We are now ready to state the existence theorem:

Theorem 3.4.1. *Let $\theta_B \in H^1(\Omega) \cap L^{\frac{p}{p-r}}(\Omega)$, $\theta_0 \in L^2(\Omega; \mathbb{R}_+)$, $\mathbf{y}_b \in L^2(\Omega; \mathbb{R}^3)$, $\mathbf{y}_a \in W^{1,p}(\Omega; \mathbb{R}^3) \cap \mathbf{V}$. Then, there exist $(\mathbf{y}_k, \theta_k)$ such that for every $T > 0$*

$$\begin{aligned} \mathbf{y}_k &\in L^\infty(0, T; W^{1,p}(\Omega; \mathbb{R}^3) \cap \mathbf{V}) \cap W^{1,\infty}(0, T; L^2(\Omega; \mathbb{R}^3)), \\ \theta_k &\in L^2(0, T; H^1(\Omega)) \cap L^\infty(0, T; L^2(\Omega)), \end{aligned}$$

and

$$\begin{aligned} &\int_0^T \int_{\Omega} \zeta \mathbf{C} \nabla \theta_k \cdot \nabla \omega \, d\mathbf{x} dt \\ &= \int_0^T \int_{\Omega} (\theta_T \eta_1(\nabla \mathbf{y}_k) + \gamma \rho_0 \theta_k) \dot{\zeta} \omega \, d\mathbf{x} dt + \zeta(0) \int_{\Omega} \omega (\gamma \rho_0 \theta_0 + \theta_T \eta_1(\nabla \mathbf{y}_a)) \, d\mathbf{x}, \end{aligned} \quad (3.10)$$

$$\int_0^T \int_{\Omega} \left(\frac{\xi}{k} \Delta \mathbf{y} \cdot \Delta \boldsymbol{\psi} + \xi \nabla \boldsymbol{\psi} : \frac{\partial \phi^k}{\partial \mathbf{F}}(\nabla \mathbf{y}_k, \theta_k) - \rho_0 \dot{\xi} \dot{\mathbf{y}}_k \cdot \boldsymbol{\psi} \right) d\mathbf{x} dt = \rho_0 \xi(0) \int_{\Omega} \mathbf{y}_1 \cdot \boldsymbol{\psi} \, d\mathbf{x}, \quad (3.11)$$

for all $\xi, \zeta \in C_c^1(-T, T)$, $\boldsymbol{\psi} \in C^1(\Omega; \mathbb{R}^3) \cap \mathbf{V}$, and all $\omega \in C^1(\Omega) \cap H_D^1$. Furthermore, $\theta_k = \theta_B$ on $\partial\Omega_D$ in the sense of traces for a.e. $t \in (0, T)$,

$$\dot{\mathbf{y}}_k \in C([0, T]; L_w^2(\Omega; \mathbb{R}^3)), \quad \theta_k + \eta(\nabla \mathbf{y}_k) \in C([0, T]; L_w^{\min\{2, p/r\}}(\Omega))$$

and

$$\begin{aligned} &\left(\|\theta_k\|_2^2 + \|\dot{\mathbf{y}}_k\|_2^2 + \|\mathbf{y}_k\|_{W^{1,p}}^p + \frac{1}{k} \|\Delta \mathbf{y}_k\|_2^2 \right)(t) + \int_0^t \|\nabla \theta_k\|_2^2(\tau) \, d\tau + k \int_{\Omega} \phi_2(\nabla \mathbf{y}_k) \, d\mathbf{x} \\ &\leq C_T + c \left| \int_{\Omega} \mathbf{y}_a(\mathbf{x}) \, d\mathbf{x} \right|^p + \frac{c}{k} \|\Delta \mathbf{y}_a\|_2^2 + ck \int_{\Omega} \phi_2(\nabla \mathbf{y}_a) \, d\mathbf{x}, \end{aligned} \quad (3.12)$$

for a.e $t \in (0, T)$, for some positive C_T depending on the parameters and on $\|\theta_B\|_{\frac{p}{p-r}}$, $\|\theta_B\|_{H^1}$, $\|\theta_0\|_2$, $\|\mathbf{y}_b\|_2$ and T only, and for some c depending solely on the parameters.

Proof. Let $\{\omega_j\}_{j \in \mathbb{N}} \subset H_D^1(\Omega)$ and $\{\psi_j\}_{j \in \mathbb{N}} \subset H^2(\Omega; \mathbb{R}^3)$ be bases of $H^1(\Omega)$ and $H^2(\Omega; \mathbb{R}^3)$ respectively. We assume that the first is constructed by taking the eigenvectors of the Laplacian with homogeneous Dirichlet boundary conditions on $\partial\Omega_D$ and homogeneous Neumann on $\partial\Omega_N$. The second one is constructed by taking the eigenvectors of the compact, symmetric, and linear operator $A: H^2(\Omega; \mathbb{R}^3) \rightarrow (H^2(\Omega; \mathbb{R}^3))^*$, where

$$\langle A\mathbf{u}, \mathbf{w} \rangle = (\mathbf{u}, \mathbf{w})_{H^2},$$

for each $\mathbf{u}, \mathbf{w} \in H^2(\Omega; \mathbb{R}^3)$. Here $(\cdot, \cdot)_{H^2}$ is the scalar product in the space $H^2(\Omega; \mathbb{R}^3)$. We can define P_n and $\mathbf{P}_n$ to be respectively the $L^2(\Omega)$ and the $L^2(\Omega; \mathbb{R}^3)$ projectors on the finite subspaces $\{\omega_j\}_{j=1, \dots, n}$ and $\{\psi_j\}_{j=1, \dots, n}$. Let us consider k to be fixed, we drop the subscript k to simplify notation, and we will suppose without loss of generality that $\rho_0 = \theta_T = \gamma = 1$. Furthermore, we start by assuming that $\mathbf{y}_a \in H^2(\Omega; \mathbb{R}^3) \cap W^{1,p}(\Omega; \mathbb{R}^3)$. The strategy of the proof is the following: we show existence of solutions via the Galerkin method to an approximated problem depending on parameters $M, N \in \mathbb{N}$. Then we send first M and then N to infinity, and show that solutions to the approximated problems converge to solutions to the original problem. Finally, we recover the result for general $\mathbf{y}_a \in \mathbf{V} \cap W^{1,p}(\Omega; \mathbb{R}^3)$.

For every $n \geq 1$ let us consider the following functions

$$\theta_n(\mathbf{x}, t) = \theta_B + \sum_{j=1}^n b_j^{(n)}(t) \omega_j(\mathbf{x}), \quad \mathbf{y}_n(\mathbf{x}, t) = \sum_{j=1}^n c_j^{(n)}(t) \psi_j(\mathbf{x}),$$

satisfying

$$\int_{\Omega} (\dot{\theta}_n \omega_j + \mathbf{C} \nabla \theta_n \cdot \nabla \omega_j) \, d\mathbf{x} = - \int_{\Omega} \dot{\eta}_1(H_M(\nabla \mathbf{y}_n)) \omega_j \, d\mathbf{x}, \quad (3.13)$$

$$\begin{aligned} & \int_{\Omega} \left(\frac{\partial \phi^k}{\partial \mathbf{F}}(H_M(\nabla \mathbf{y}_n), \theta_n) \frac{\partial H_M(\nabla \mathbf{y}_n)}{\partial \mathbf{F}} : \nabla \psi_i + \frac{1}{N} (\mathbf{y}_n, \psi_i)_{H^2} \right) \, d\mathbf{x} \\ &= - \int_{\Omega} \left(\frac{1}{k} \Delta \psi_i \cdot \Delta \mathbf{y}_n + \ddot{\mathbf{y}}_n \cdot \psi_i \right) \, d\mathbf{x}, \end{aligned} \quad (3.14)$$

$$\theta_n(\mathbf{x}, 0) = P_n(\theta_0(\mathbf{x}) - \theta_B) + \theta_B, \quad \mathbf{y}_n(\mathbf{x}, 0) = \mathbf{P}_n \mathbf{y}_a(\mathbf{x}), \quad \dot{\mathbf{y}}_n(\mathbf{x}, 0) = \mathbf{P}_n \dot{\mathbf{y}}_b(\mathbf{x}),$$

for all $i, j = 1, \dots, n$, for almost every $\mathbf{x} \in \Omega$, for $M, N \in \mathbb{N}$ fixed, and where H_M is a smooth function such that

$$H_M(\mathbf{F}) = \begin{cases} \mathbf{F}, & \text{if } |\mathbf{F}| \leq M, \\ (M+1) \frac{\mathbf{F}}{|\mathbf{F}|}, & \text{if } |\mathbf{F}| \geq M+1, \end{cases} \quad (3.15)$$

and $\left| \frac{\partial H_M(\mathbf{F})}{\partial \mathbf{F}} \right| \leq 2$. We remark that $\frac{\partial H_M(\mathbf{F})}{\partial \mathbf{F}}$ is a fourth order tensor, and $\frac{\partial \phi^k(H_M(\mathbf{F}))}{\partial \mathbf{F}} \frac{\partial H_M(\mathbf{F})}{\partial \mathbf{F}}$ should read as

$$\left(\frac{\partial \phi^k(H_M(\mathbf{G}))}{\partial \mathbf{F}} \frac{\partial H_M(\mathbf{G})}{\partial \mathbf{F}} \right)_{lm} = \sum_{i,j=1}^3 \frac{\partial \phi^k(\mathbf{H})}{\partial \mathbf{H}_{ij}} \Big|_{\mathbf{H}=H_M(\mathbf{G})} \frac{\partial (H_M(\mathbf{F}))_{ij}}{\partial \mathbf{F}_{lm}} \Big|_{\mathbf{F}=\mathbf{G}} = \frac{\partial \phi^k(H_M(\mathbf{F}))}{\partial \mathbf{F}_{lm}} \Big|_{\mathbf{F}=\mathbf{G}}.$$

The above system of differential equations in $\mathbf{b}^{(n)}, \mathbf{c}^{(n)}$ admits locally in time a smooth solution by standard ODE theory. Define

$$\hat{\theta}_n := \theta_n - \theta_B \text{ and } \hat{\theta}_{0,n} := P_n(\theta_0(\mathbf{x}) - \theta_B).$$

After multiplying (3.13) tested against ω_j by $b_j^{(n)}$, we can sum over j and get

$$\frac{1}{2} \frac{d}{dt} \|\hat{\theta}_n\|_2^2 + c_m \|\nabla \hat{\theta}_n\|_2^2 \leq - \int_{\Omega} \dot{\eta}_1(H_M(\nabla \mathbf{y}_n)) \hat{\theta}_n \, d\mathbf{x} - \int_{\Omega} \nabla \theta_B \cdot \nabla \hat{\theta}_n. \quad (3.16)$$

Here we have also used the fact that c_m is the minimum eigenvalue of $\mathbf{C}$ to bound $\mathbf{C} \nabla \hat{\theta}_n \cdot \nabla \hat{\theta}_n$ from below with $c_m |\nabla \hat{\theta}_n|^2$. Thanks to the Cauchy-Schwarz and the Young inequalities, (3.16) becomes

$$\frac{1}{2} \frac{d}{dt} \|\hat{\theta}_n\|_2^2 + \frac{c_m}{2} \|\nabla \hat{\theta}_n\|_2^2 \leq - \int_{\Omega} \hat{\theta}_n \dot{\eta}_1(H_M(\nabla \mathbf{y}_n)) \, d\mathbf{x} + c \|\nabla \theta_B\|_2^2.$$

On the other hand, testing (3.14) with ψ_i , multiplying it by $\dot{c}_i^{(n)}$ and summing over i from 1 to n we get

$$\begin{aligned} & \frac{1}{2} \frac{d}{dt} \left(\|\dot{\mathbf{y}}_n\|_2^2 + \frac{1}{k} \|\Delta \mathbf{y}_n\|_2^2 + \frac{1}{N} \|\mathbf{y}_n\|_{H^2}^2 \right) + \frac{d}{dt} \int_{\Omega} \left(k \phi_2(H_M(\nabla \mathbf{y}_n)) + \bar{\phi}_1(H_M(\nabla \mathbf{y}_n)) \right) d\mathbf{x} \\ &= \int_{\Omega} \theta_n \dot{\eta}_1(H_M(\nabla \mathbf{y}_n)) \, d\mathbf{x}. \end{aligned}$$

Therefore, summing the last two inequalities and integrating in time between 0 and

$t \in (0, T)$ we get

$$\begin{aligned}
& \frac{1}{2} \left(\|\dot{\mathbf{y}}_n\|_2^2 + \|\hat{\theta}_n\|_2^2 + \frac{1}{k} \|\Delta \mathbf{y}_n\|_2^2 + \frac{1}{N} \|\mathbf{y}_n\|_{H^2}^2 \right) (t) + \frac{c_m}{2} \int_0^t \|\nabla \hat{\theta}_n\|_2^2 d\tau \\
& + \int_{\Omega} \left(k\phi_2(H_M(\nabla \mathbf{y}_n(t))) + \bar{\phi}_1(H_M(\nabla \mathbf{y}_n(t))) \right) d\mathbf{x} \\
& \leq \int_{\Omega} \left(k\phi_2(H_M(\nabla \mathbf{P}_n \mathbf{y}_a)) + \bar{\phi}_1(H_M(\nabla \mathbf{P}_n \mathbf{y}_a)) \right) d\mathbf{x} \tag{3.17} \\
& + \int_{\Omega} \eta_1(H_M(\nabla \mathbf{y}_n(t))) \theta_B d\mathbf{x} - \int_{\Omega} \eta_1(H_M(\nabla \mathbf{P}_n \mathbf{y}_a)) \theta_B d\mathbf{x} + ct \|\nabla \theta_B\|_2^2 \\
& + \frac{1}{2} \left(\|\mathbf{P}_n \mathbf{y}_b\|_2^2 + \|\hat{\theta}_{0,n}\|_2^2 + \frac{1}{k} \|\Delta \mathbf{P}_n \mathbf{y}_a\|_2^2 + \frac{1}{N} \|\mathbf{P}_n \mathbf{y}_a\|_{H^2}^2 \right).
\end{aligned}$$

Now, we notice that by means of the assumptions on η we can deduce

$$\begin{aligned}
\left| \int_{\Omega} \eta_1(H_M(\nabla \mathbf{P}_n \mathbf{y}_a)) \theta_B d\mathbf{x} \right| & \leq c \|\theta_B\|_{\frac{p}{p-r}} (1 + \|H_M(\nabla \mathbf{P}_n \mathbf{y}_a)\|_p^r) \\
& \leq c \left(1 + \|\theta_B\|_{\frac{p}{p-r}}^{\frac{p}{p-r}} \right) + \frac{c_0}{4} \|H_M(\nabla \mathbf{P}_n \mathbf{y}_a)\|_p^p \\
& \leq c \left(1 + \|\theta_B\|_{\frac{p}{p-r}}^{\frac{p}{p-r}} \right) + \frac{k}{4} \phi_2(H_M(\nabla \mathbf{P}_n \mathbf{y}_a)).
\end{aligned}$$

In the same way,

$$\begin{aligned}
\left| \int_{\Omega} \eta_1(H_M(\nabla \mathbf{y}_n(t))) \theta_B d\mathbf{x} \right| & \leq c \left(1 + \|\theta_B\|_{\frac{p}{p-r}}^{\frac{p}{p-r}} \right) + \frac{k}{4} \int_{\Omega} \phi_2(H_M(\nabla \mathbf{y}_n(t))) d\mathbf{x}, \\
\left| \int_{\Omega} \bar{\phi}_1(H_M(\nabla \mathbf{P}_n \mathbf{y}_a)) d\mathbf{x} \right| & \leq c + \frac{k}{4} \int_{\Omega} \phi_2(H_M(\nabla \mathbf{P}_n \mathbf{y}_a)) d\mathbf{x}, \\
\left| \int_{\Omega} \bar{\phi}_1(H_M(\nabla \mathbf{y}_n(t))) d\mathbf{x} \right| & \leq c + \frac{k}{4} \int_{\Omega} \phi_2(H_M(\nabla \mathbf{y}_n(t))) d\mathbf{x}.
\end{aligned}$$

Thus, (3.17) becomes

$$\begin{aligned}
& \left(\frac{1}{k} \|\Delta \mathbf{y}_n\|_2^2 + \|\dot{\mathbf{y}}_n\|_2^2 + \|\hat{\theta}_n\|_2^2 + \frac{1}{N} \|\mathbf{y}_n\|_{H^2}^2 \right) (t) \\
& + c_m \int_0^t \|\nabla \hat{\theta}_n\|_2^2 d\tau + k \int_{\Omega} \phi_2(H_M(\nabla \mathbf{y}_n(t))) d\mathbf{x} \\
& \leq \frac{1}{k} \|\Delta \mathbf{P}_n \mathbf{y}_a\|_2^2 + \|\mathbf{P}_n \mathbf{y}_b\|_2^2 + \|\hat{\theta}_{0,n}\|_2^2 + ct \|\nabla \theta_B\|_2^2 + c \left(1 + \|\theta_B\|_{\infty}^{\frac{p}{p-r}} \right) \\
& + 3k \int_{\Omega} \phi_2(H_M(\nabla \mathbf{P}_n \mathbf{y}_a)) d\mathbf{x} + \frac{1}{N} \|\mathbf{P}_n \mathbf{y}_a\|_{H^2}^2. \tag{3.18}
\end{aligned}$$

We remark that as $\mathbf{P}_n \mathbf{y}_a \rightarrow \mathbf{y}_a$ strongly in $H^2(\Omega; \mathbb{R}^3)$, and $\mathbf{P}_n \mathbf{y}_b \rightarrow \mathbf{y}_b$ strongly in $L^2(\Omega; \mathbb{R}^3)$, the sequences $\|\mathbf{P}_n \mathbf{y}_a\|_{H^2}$, $\|\mathbf{P}_n \mathbf{y}_b\|_2$ are bounded. Therefore, by the boundedness of H_M (cf. (3.15)), and thanks to the fact that ϕ_2 is non-negative, we deduce that for every $T > 0$ there exists a constant $\tilde{C} = \tilde{C}(T, M, k, N)$, independent of n , such that

$$\frac{1}{k} \|\Delta \mathbf{y}_n\|_2^2(t) + \|\dot{\mathbf{y}}_n\|_2^2(t) + \|\hat{\theta}_n\|_2^2(t) + c_m \int_0^t \|\nabla \hat{\theta}_n\|_2^2 d\tau + \frac{1}{N} \|\mathbf{y}_n\|_{H^2}^2(t) \leq \tilde{C},$$

for a.e. $t \in (0, T)$. Hence, there exists

$$\begin{aligned} y_M &\in L^\infty(0, T; H^2(\Omega; \mathbb{R}^3)) \cap W^{1,\infty}(0, T; L^2(\Omega; \mathbb{R}^3)), \\ \theta_M &\in L^2(0, T; H^1(\Omega)) \cap L^\infty(0, T; L^2(\Omega)), \end{aligned}$$

such that, up to a subsequence,

$$\theta_n \rightharpoonup \theta_M \quad \text{weakly in } L^2(0, T; H^1(\Omega)), \quad \text{weakly}^* \text{ in } L^\infty(0, T; L^2(\Omega)), \quad (3.19)$$

$$\mathbf{y}_n \rightharpoonup \mathbf{y}_M \quad \text{weakly}^* \text{ in } L^\infty(0, T; H^2(\Omega; \mathbb{R}^3)) \cap W^{1,\infty}(0, T; L^2(\Omega; \mathbb{R}^3)). \quad (3.20)$$

Furthermore, (3.20) together with a version of the Aubin-Lions lemma (see e.g., [24, Thm. II.5.16]) imply

$$\mathbf{y}_n \rightarrow \mathbf{y}_M, \quad \text{strongly in } C([0, T]; H^1(\Omega; \mathbb{R}^3)), \quad (3.21)$$

that is, up to a further subsequence, $\nabla \mathbf{y}_n \rightarrow \nabla \mathbf{y}_M$ for each $t \in (0, T)$, a.e. $\mathbf{x} \in \Omega$. Thus, as H_M is continuous and bounded, by Vitali's convergence theorem $H_M(\nabla \mathbf{y}_n) \rightarrow H_M(\nabla \mathbf{y}_M)$ strongly in $L^{\tilde{q}}(0, T; L^{\tilde{q}}(\Omega))$ for each $\tilde{q} \in [1, \infty)$.

We can thus multiply (3.13) and (3.14) by $\zeta, \xi \in C_c^1(-T, T)$ and integrate in time between 0 and T . After an integration by parts these become

$$\begin{aligned} &\int_0^T \int_\Omega (\zeta C \nabla \theta_n \cdot \nabla \omega_j - \omega_j \dot{\zeta} (\eta_1(H_M(\nabla \mathbf{y}_n)) + \theta_n)) \, d\mathbf{x} \, dt \\ &= \zeta(0) \int_\Omega (\eta_1(H_M(\nabla \mathbf{P}_n \mathbf{y}_a)) + \theta_{0,n}) \omega_j \, d\mathbf{x}, \end{aligned} \quad (3.22)$$

$$\begin{aligned} &\int_0^T \int_\Omega \left(\frac{\xi}{k} \Delta \mathbf{y}_n \cdot \Delta \boldsymbol{\psi}_i + \xi \frac{\partial \phi^k}{\partial \mathbf{F}}(H_M(\nabla \mathbf{y}_n), \theta_n) \frac{\partial H_M}{\partial \mathbf{F}}(\nabla \mathbf{y}_n) : \nabla \boldsymbol{\psi}_i \right) \, d\mathbf{x} \, dt \\ &= \int_0^T \int_\Omega \left(\xi \dot{\mathbf{y}}_n \cdot \boldsymbol{\psi}_i - \frac{\xi}{N} (\boldsymbol{\psi}_i, \mathbf{y}_n)_{H^2} \right) \, d\mathbf{x} \, dt + \xi(0) \int_\Omega \mathbf{P}_n \mathbf{y}_b \cdot \boldsymbol{\psi}_i \, d\mathbf{x}, \end{aligned} \quad (3.23)$$

for each $i, j = 1, \dots, n$. Passage to the limit and recovering the equations for every $\omega \in H_D^1(\Omega) \cap C^1(\Omega)$, $\boldsymbol{\psi} \in C^1(\Omega; \mathbb{R}^3) \cap H^2(\Omega; \mathbb{R}^3)$ is standard. Thus, (3.22)–(3.23)

become

$$\begin{aligned} & \int_0^T \int_{\Omega} (\zeta \mathbf{C} \nabla \theta_M \cdot \nabla \omega - \omega \dot{\zeta} (\eta_1(H_M(\nabla \mathbf{y}_M)) + \theta_M)) \, d\mathbf{x} \, dt \\ &= \zeta(0) \int_{\Omega} (\eta_1(H_M(\nabla \mathbf{y}_a)) + \theta_0) \omega \, d\mathbf{x}, \end{aligned} \quad (3.24)$$

$$\begin{aligned} & \int_0^T \int_{\Omega} \left(\frac{\xi}{k} \Delta \mathbf{y}_M \cdot \Delta \boldsymbol{\psi} + \xi \frac{\partial \phi^k}{\partial \mathbf{F}}(H_M(\nabla \mathbf{y}_M), \theta_M) \frac{\partial H_M}{\partial \mathbf{F}}(\nabla \mathbf{y}_M) : \nabla \boldsymbol{\psi} \right) \, d\mathbf{x} \, dt \\ &= \int_0^T \int_{\Omega} \left(\dot{\xi} \dot{\mathbf{y}}_M \cdot \boldsymbol{\psi} - \frac{\xi}{N} (\boldsymbol{\psi}_i, \mathbf{y}_n)_{H^2} \right) \, d\mathbf{x} \, dt + \xi(0) \int_{\Omega} \mathbf{y}_b \cdot \boldsymbol{\psi} \, d\mathbf{x}, \end{aligned} \quad (3.25)$$

for every $\omega \in H_D^1(\Omega) \cap C^1(\Omega)$, $\boldsymbol{\psi} \in C^1(\Omega; \mathbb{R}^3) \cap \mathbf{H}^2(\Omega; \mathbb{R}^3)$, $\zeta, \xi \in C_c^1(-T, T)$. A passage to the limit as $n \rightarrow \infty$ in (3.18) and the weak lower semicontinuity of the norms leads to

$$\begin{aligned} & \frac{1}{N} \|\mathbf{y}_M\|_{H^2}^2(t) + \frac{1}{k} \|\Delta \mathbf{y}_M\|_2^2(t) + \|\dot{\mathbf{y}}_M\|_2^2(t) + \|\hat{\theta}_M\|_2^2 \\ & \quad + c_m \int_0^t \|\nabla \hat{\theta}_M\|_2^2 \, d\tau + k \int_{\Omega} \phi_2(H_M(\nabla \mathbf{y}_M(t))) \, d\mathbf{x} \\ & \leq c_T + \frac{1}{N} \|\mathbf{y}_a\|_{H^2}^2 + \frac{1}{k} \|\Delta \mathbf{y}_a\|_2^2 + 3k \int_{\Omega} \phi_2(H_M(\nabla \mathbf{y}_a)) \, d\mathbf{x}, \end{aligned} \quad (3.26)$$

for a.e. $t \in (0, T)$. Now, given the fact that $\mathbf{y}_a \in W^{1,p}(\Omega; \mathbb{R}^3)$, we have that $H_M(\nabla \mathbf{y}_a)$ converges strongly to $\nabla \mathbf{y}_a$ as $M \rightarrow \infty$ by dominated convergence in L^p , and hence, by (3.6), $\phi_2(H_M(\nabla \mathbf{y}_a))$ converges strongly in L^1 to $\phi_2(\nabla \mathbf{y}_a)$. Therefore, inequality (3.26) yields

$$\begin{aligned} & \|\mathbf{y}_M\|_{L^\infty(0,T;H^2)} + \|\dot{\mathbf{y}}_M\|_{L^\infty(0,T;L^2)} + \|\hat{\theta}_M\|_{L^\infty(0,T;L^2)} + \|\nabla \hat{\theta}_M\|_{L^2(0,T;L^2)} \\ & \quad + \operatorname{ess\,sup}_{t \in (0,T)} \int_{\Omega} \phi_2(H_M(\nabla \mathbf{y}_M(t))) \, d\mathbf{x} \leq c_{T,k,N}. \end{aligned} \quad (3.27)$$

We can hence deduce the existence of

$$\begin{aligned} & \mathbf{y}_N \in L^\infty(0, T; H^2(\Omega; \mathbb{R}^3)) \cap W^{1,\infty}(0, T; L^2(\Omega; \mathbb{R}^3)), \\ & \theta_N \in L^2(0, T; H^1(\Omega)) \cap L^\infty(0, T; L^2(\Omega)), \end{aligned}$$

such that, up to a subsequence, $\theta_M \rightarrow \theta_N$ and $\mathbf{y}_M \rightarrow \mathbf{y}_N$ as $M \rightarrow \infty$ in the sense of (3.19)–(3.21). Furthermore, by (3.27) and (3.6) we also have

$$\|H_M(\mathbf{y}_M)\|_{L^\infty(0,T;L^p)} \leq c_{T,k,N}. \quad (3.28)$$

Therefore, as $H_M(\nabla \mathbf{y}_M) \rightarrow \nabla \mathbf{y}_N$ for a.e. $(x, t) \in \Omega \times (0, T)$, by Vitali's theorem we deduce

$$H_M(\nabla \mathbf{y}_M) \rightarrow \nabla \mathbf{y}_N, \quad \text{strongly in } L^{\bar{q}}(\Omega \times (0, T)), \quad (3.29)$$

for every $\bar{q} \in (1, p)$. We remark that Fatou's lemma and the fact that $H_M(\nabla \mathbf{y}_M) \rightarrow \nabla \mathbf{y}_N$ for a.e. $(\mathbf{x}, t) \in \Omega \times (0, T)$, imply

$$\int_{\Omega} \phi_2(\nabla \mathbf{y}_N(\mathbf{x}, t)) \, d\mathbf{x} \leq \liminf_{M \rightarrow \infty} \int_{\Omega} \phi_2(H_M(\nabla \mathbf{y}_M(\mathbf{x}, t))) \, d\mathbf{x}, \quad \text{for a.e. } t \in (0, T).$$

This inequality, together with the lower semicontinuity of the norms, and the convergences (3.19)–(3.21) entail

$$\begin{aligned} & \left(\frac{1}{N} \|\mathbf{y}_N\|_{H^2}^2 + \frac{1}{k} \|\Delta \mathbf{y}_N\|_2^2 + \|\dot{\mathbf{y}}_N\|_2^2 + \|\theta_N\|_2^2 \right)(t) \\ & \quad + c_m \int_0^t \|\nabla \theta_N\|_2^2 \, d\tau + k \int_{\Omega} \phi_2(\nabla \mathbf{y}_N(t)) \, d\mathbf{x} \\ & \leq C_T^* + \frac{1}{N} \|\mathbf{y}_a\|_{H^2}^2 + \frac{1}{k} \|\Delta \mathbf{y}_a\|_2^2 + 3k \int_{\Omega} \phi_2(\nabla \mathbf{y}_a) \, d\mathbf{x}, \end{aligned} \quad (3.30)$$

for a.e. $t \in (0, T)$, for some $C^* > 0$ depending on $\|\theta_B\|_{H^1}$, $\|\theta_B\|_{\frac{p}{p-r}}$, $\|\mathbf{y}_b\|_2$, $\|\theta_0\|_2$ and T only. By (3.6), together with (3.30), we thus have also

$$\mathbf{y}_N \in L^\infty(0, T; W^{1,p}(\Omega; \mathbb{R}^3)).$$

The assumptions in (3.6), and the bound on $\frac{\partial H_M}{\partial \mathbf{F}}$, together with (3.29) imply also that

$$\frac{\partial \phi_2}{\partial \mathbf{F}}(H_M(\nabla \mathbf{y}_M)) \frac{\partial H_M}{\partial \mathbf{F}}(\nabla \mathbf{y}_M) \rightarrow \frac{\partial \phi_2}{\partial \mathbf{F}}(\nabla \mathbf{y}_N), \quad \text{strongly in } L^1(\Omega; \mathbb{R}^{3 \times 3}), \quad (3.31)$$

$$\frac{\partial \bar{\phi}_1}{\partial \mathbf{F}}(H_M(\nabla \mathbf{y}_M)) \frac{\partial H_M}{\partial \mathbf{F}}(\nabla \mathbf{y}_M) \rightarrow \frac{\partial \bar{\phi}_1}{\partial \mathbf{F}}(\nabla \mathbf{y}_N), \quad \text{strongly in } L^1(\Omega; \mathbb{R}^{3 \times 3}). \quad (3.32)$$

Hence, passage to the limit as $M \rightarrow \infty$ in (3.24)–(3.25) is standard thanks to (3.19)–(3.21) together with (3.31)–(3.32). The only difficulty is to show that, defined

$$R_M := \frac{\partial \eta_1}{\partial \mathbf{F}}(H_M(\nabla \mathbf{y}_M)) \frac{\partial H_M}{\partial \mathbf{F}}(\nabla \mathbf{y}_M),$$

we have

$$\begin{aligned} & \left| \int_0^T \int_{\Omega} \xi \nabla \boldsymbol{\psi} : \left(\frac{\partial \eta_1}{\partial \mathbf{F}}(\nabla \mathbf{y}_N) \theta_N - R_M \theta_M \right) \, d\mathbf{x} \, dt \right| \\ & = \left| \int_0^T \int_{\Omega} \xi \nabla \boldsymbol{\psi} : \left(\frac{\partial \eta_1}{\partial \mathbf{F}}(\nabla \mathbf{y}_N) (\theta_N - \theta_M) + \theta_M \left(\frac{\partial \eta_1}{\partial \mathbf{F}}(\nabla \mathbf{y}_N) - R_M \right) \right) \, d\mathbf{x} \, dt \right| \rightarrow 0, \end{aligned}$$

as $M \rightarrow \infty$, for every $\boldsymbol{\psi} \in C^1(\Omega; \mathbb{R}^3) \cap H^1(\Omega; \mathbb{R}^3)$, every $\xi \in C_c^1(-T, T)$. On the one hand,

$$\left| \int_0^T \int_{\Omega} \xi \nabla \boldsymbol{\psi} : \left(\frac{\partial \eta_1}{\partial \mathbf{F}}(\nabla \mathbf{y}_N) (\theta_N - \theta_M) \right) \, d\mathbf{x} \, dt \right| \rightarrow 0, \quad \text{as } M \rightarrow \infty,$$

because of the fact that $\theta_M \rightarrow \theta_N$ weakly in $L^2(0, T; L^6(\Omega))$ as $M \rightarrow \infty$, and that (3.6) together with $\nabla \mathbf{y}_N \in L^\infty(0, T; L^p(\Omega; \mathbb{R}^{3 \times 3}))$ imply $\frac{\partial \eta_1}{\partial \mathbf{F}}(\nabla \mathbf{y}_N) \in L^\infty(0, T; L^{\frac{6}{5}}(\Omega))$. On the other hand,

$$\begin{aligned} & \left| \int_0^T \int_\Omega \xi \nabla \psi : \left(\theta_M \left(R_M - \frac{\partial \eta_1}{\partial \mathbf{F}}(\nabla \mathbf{y}_N) \right) \right) d\mathbf{x} dt \right| \\ & \leq \|\theta_M\|_{L^2(0, T; L^6)} \|\psi\|_{C^1} \|\xi\|_{C^1} \left\| \frac{\partial \eta_1}{\partial \mathbf{F}}(\nabla \mathbf{y}_N) - R_M \right\|_{L^2(0, T; L^{\frac{6}{5}})} \rightarrow 0 \end{aligned}$$

because of (3.29). In the limit, (3.24)–(3.25) thus become

$$\int_0^T \int_\Omega (\zeta \mathbf{C} \nabla \theta_N \cdot \nabla \omega - \omega \dot{\zeta} (\eta_1(\nabla \mathbf{y}_N) + \theta_N)) d\mathbf{x} dt = \zeta(0) \int_\Omega (\eta_1(\nabla \mathbf{y}_a) + \theta_0) \omega d\mathbf{x}, \quad (3.33)$$

$$\begin{aligned} & \int_0^T \int_\Omega \left(\frac{\xi}{k} \Delta \mathbf{y}_N \cdot \Delta \psi + \xi \nabla \psi : \frac{\partial \phi^k}{\partial \mathbf{F}}(\nabla \mathbf{y}_N, \theta_N) \right) d\mathbf{x} dt \\ & = \int_0^T \int_\Omega \left(\dot{\xi} \dot{\mathbf{y}}_N \cdot \psi - \frac{\xi}{N} (\psi, \mathbf{y}_N)_{H^2} \right) d\mathbf{x} dt + \xi(0) \int_\Omega \mathbf{y}_b \cdot \psi d\mathbf{x}, \end{aligned} \quad (3.34)$$

for every $\omega \in H_D^1(\Omega) \cap C^1(\Omega)$, $\psi \in C^1(\Omega; \mathbb{R}^3) \cap H^2(\Omega; \mathbb{R}^3)$, $\zeta, \xi \in C_c^1(-T, T)$. Now, we want to send N to infinity. To this end, we first suppose $N \geq N_*$, where $N_* = N_*(\|\mathbf{y}_a\|_{H^2})$ is the smallest integer such that $N_*^{-1} \|\mathbf{y}_a\|_{H^2}^2 \leq 1$. In this case, by (3.30) we have

$$\|\nabla \mathbf{y}_N\|_p^p(t) - c \leq \int_\Omega \phi_2(\nabla \mathbf{y}_N) d\mathbf{x} \leq c_{T,k}, \quad \text{a.e. } t \in (0, T). \quad (3.35)$$

Therefore, (3.35) together with (3.30) entail

$$\left(\frac{1}{N} \|\mathbf{y}_N\|_{H^2}^2 + \|\Delta \mathbf{y}_N\|_2^2 + \|\dot{\mathbf{y}}_N\|_2^2 + \|\nabla \mathbf{y}_N\|_p^p + \|\theta_N\|_2^2 \right)(t) + c_m \int_0^t \|\nabla \theta_N\|_2^2 d\tau \leq c_{T,k}, \quad (3.36)$$

for a.e. $t \in (0, T)$. Let $\{\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3\}$ be a Cartesian coordinate system for $\mathbb{R}^3$. We can choose $\psi = \mathbf{e}_i$ with $i = 1, \dots, 3$ in (3.34) and deduce that

$$\frac{d^2}{dt^2} \int_\Omega \mathbf{y}_N(t) \cdot \mathbf{e}_i d\mathbf{x} = -\frac{1}{N} \int_\Omega \mathbf{y}(t) \cdot \mathbf{e}_i d\mathbf{x},$$

for a.e. $t \in (0, T)$. This implies,

$$\int_\Omega \mathbf{y}_N(t) \cdot \mathbf{e}_i d\mathbf{x} = \cos\left(\frac{t}{\sqrt{N}}\right) \int_\Omega \mathbf{y}_a \cdot \mathbf{e}_i d\mathbf{x} + \sqrt{N} \sin\left(\frac{t}{\sqrt{N}}\right) \int_\Omega \mathbf{y}_b \cdot \mathbf{e}_i d\mathbf{x},$$

and therefore, for every T finite, we deduce that

$$\left| \int_\Omega \mathbf{y}_N(t) \cdot \mathbf{e}_i d\mathbf{x} \right| \leq \left| \int_\Omega \mathbf{y}_a \cdot \mathbf{e}_i d\mathbf{x} \right| + T \left| \int_\Omega \mathbf{y}_b \cdot \mathbf{e}_i d\mathbf{x} \right| \leq c_T, \quad (3.37)$$

where, again, c_T is independent of N, k . Therefore, an application of the Poincaré inequality, together with (3.35) leads also to

$$\|\mathbf{y}_N\|_{W^{1,p}}(t) \leq c_{T,k}, \quad \text{a.e. } t \in (0, T). \quad (3.38)$$

Therefore, from (3.36)–(3.38) we deduce the existence of

$$\begin{aligned} \mathbf{y} &\in L^\infty(0, T; \mathbf{V} \cap W^{1,p}(\Omega; \mathbb{R}^3)) \cap W^{1,\infty}(0, T; L^2(\Omega; \mathbb{R}^3)), \\ \theta &\in L^2(0, T; H^1(\Omega)) \cap L^\infty(0, T; L^2(\Omega)), \end{aligned}$$

and a non-relabelled subsequence, such that

$$\theta_N \rightharpoonup \theta \quad \text{weakly in } L^2(0, T; H^1(\Omega)), \quad \text{weakly}^* \text{ in } L^\infty(0, T; L^2(\Omega)), \quad (3.39)$$

$$\mathbf{y}_N \rightharpoonup \mathbf{y} \quad \text{weakly}^* \text{ in } L^\infty(0, T; \mathbf{V} \cap W^{1,p}(\Omega; \mathbb{R}^3)) \cap W^{1,\infty}(0, T; L^2(\Omega; \mathbb{R}^3)). \quad (3.40)$$

Furthermore, from the fact that $p > 2$, and Lemma 3.4.1 we known that the Banach space $\mathbf{V} \cap W^{1,p}(\Omega; \mathbb{R}^3)$ is compact in $H^1(\Omega; \mathbb{R}^3)$. Hence, by (3.39)–(3.40) and the Aubin-Lions lemma, $\mathbf{y}_N$ converges to $\mathbf{y}$ also in the sense of (3.21). This, together with (3.36), (3.38) and Vitali's theorem, lead

$$\nabla \mathbf{y}_N \rightarrow \nabla \mathbf{y}, \quad \text{strongly in } L^{\bar{q}}(\Omega \times (0, T)), \quad (3.41)$$

for every $\bar{q} < p$. We can thus pass again to the limit in (3.30), and deduce

$$\begin{aligned} \left(\frac{1}{k} \|\Delta \mathbf{y}\|_2^2 + \|\dot{\mathbf{y}}\|_2^2 + \|\theta\|_2^2 \right)(t) + c_m \int_0^t \|\nabla \theta\|_2^2 d\tau + k \int_\Omega \phi_2(\nabla \mathbf{y}(t)) d\mathbf{x} \\ \leq C_T^* + \frac{1}{k} \|\Delta \mathbf{y}_a\|_2^2 + 3k \int_\Omega \phi_2(\nabla \mathbf{y}_a) d\mathbf{x}, \end{aligned}$$

a.e. $t \in (0, T)$. Now the Poincaré inequality, together with (3.35) lead to

$$\int_\Omega \phi_2(\nabla \mathbf{y}) d\mathbf{x} \geq \|\nabla \mathbf{y}\|_p^{p-c} \geq c_0 \|\mathbf{y}\|_{W^{1,p}}^{p-c-c} \left| \int_\Omega \mathbf{y} d\mathbf{x} \right|^p \geq c_0 \|\mathbf{y}\|_{W^{1,p}}^{p-C_T^*-c} \left| \int_\Omega \mathbf{y}_a d\mathbf{x} \right|^p,$$

where we made use of (3.37) in the last inequality. Combining the last two estimates we thus deduce (3.12). Furthermore, by arguing as above, we can use (3.6) and (3.41) to pass to the limit as $N \rightarrow \infty$ in (3.33)–(3.34) and deduce (3.10)–(3.11). Here, the only additional difficulty is to pass to the limit in the term $(\mathbf{y}_N, \boldsymbol{\psi})_{H^2}$, which, thanks to (3.36) can be treated as follows

$$\frac{1}{N} |(\mathbf{y}_N, \boldsymbol{\psi})_{H^2}| \leq \frac{1}{N} \|\mathbf{y}_N\|_{H^2} \|\boldsymbol{\psi}\|_{H^2} \leq \frac{\sqrt{c_{T,k}}}{\sqrt{N}} \|\boldsymbol{\psi}\|_{H^2} \rightarrow 0, \quad \text{as } N \rightarrow \infty, \text{ a.e. } t \in (0, T).$$

Now, the existence result for a generic initial datum $\mathbf{y}_a \in \mathbf{V} \cap W^{1,p}(\Omega; \mathbb{R}^3)$ can be recovered as follows: we take a sequence of $\mathbf{y}_a^{(j)} \in H^2(\Omega; \mathbb{R}^3) \cap W^{1,p}(\Omega; \mathbb{R}^3)$ converging strongly in $\mathbf{V} \cap W^{1,p}(\Omega; \mathbb{R}^3)$ to $\mathbf{y}_a$. In this case,

$$\int_{\Omega} \phi_2(\nabla \mathbf{y}_a^{(j)}) \leq c, \quad \int_{\Omega} \phi_2(\nabla \mathbf{y}_a^{(j)}) \rightarrow \int_{\Omega} \phi_2(\nabla \mathbf{y}_a), \quad (3.42)$$

so thanks to (3.12), the sequence of solutions $(\theta_j, \mathbf{y}_j)$ related to $\mathbf{y}_a^{(j)}$ satisfies

$$\left(\|\Delta \mathbf{y}_j\|_2^2 + \|\dot{\mathbf{y}}_j\|_2^2 + \|\nabla \mathbf{y}_j\|_{W^{1,p}}^p + \|\theta_j\|_2^2 \right)(t) + c_m \int_0^t \|\nabla \theta_j\|_2^2 d\tau \leq c_{T,k}, \quad \text{a.e. } t \in (0, T). \quad (3.43)$$

Therefore we can deduce the existence of

$$\begin{aligned} \mathbf{y} &\in L^\infty(0, T; \mathbf{V} \cap W^{1,p}(\Omega; \mathbb{R}^3)) \cap W^{1,\infty}(0, T; L^2(\Omega; \mathbb{R}^3)), \\ \theta &\in L^2(0, T; H^1(\Omega)) \cap L^\infty(0, T; L^2(\Omega)), \end{aligned}$$

such that, up to a non relabelled subsequence, $(\theta_j, \mathbf{y}_j)$ converge to $(\theta, \mathbf{y})$ in the sense of (3.39)–(3.41). Passage to the limit to recover that $(\theta, \mathbf{y})$ satisfy (3.12), (3.10)–(3.11) follows the steps above.

Finally, by comparison in (3.11) we also know that

$$\ddot{\mathbf{y}} \in L^2(0, T; (W^{1,\bar{p}}(\Omega; \mathbb{R}^3) \cap \mathbf{V})^*).$$

Hence, by [24, Prop. II.5.11] and [24, Lemma II.5.9] we deduce $\dot{\mathbf{y}} \in C([0, T]; L_w^2(\Omega; \mathbb{R}^3))$. On the other hand, (3.10) implies also $\dot{\theta} + \eta_1(\nabla \mathbf{y}) \in L^2(0, T; H_D^{-1}(\Omega))$. As by (3.6) and $\mathbf{y} \in L^\infty(0, T; W^{1,p}(\Omega; \mathbb{R}^3))$,

$$\operatorname{ess\,sup}_{t \in (0, T)} \int_{\Omega} |\eta(\nabla \mathbf{y}_n(t))|^{\frac{p}{r}} d\mathbf{x} \leq c + \operatorname{ess\,sup}_{t \in (0, T)} \int_{\Omega} |\nabla \mathbf{y}_n|^p d\mathbf{x} \leq c_{T,k},$$

it must hold $\eta_1(\nabla \mathbf{y}) \in L^\infty(0, T; L^{p/r}(\Omega))$. Therefore, again by [24, Prop. II.5.11] and [24, Lemma II.5.9] we thus deduce that

$$\theta + \eta_1(\nabla \mathbf{y}) \in C([0, T]; L_w^{\min\{2, p/r\}}(\Omega)).$$

□

3.5 Convergence as $k \rightarrow +\infty$ to a limiting constrained theory

As mentioned above, in this section we send k to ∞ in (3.10)–(3.11), and obtain in the limit a constrained theory for the deformation gradient $\nabla \mathbf{y}$. This is equivalent to assuming that elastic constants tend to infinity, which, as remarked in [15], is usually a reasonable approximation when studying martensitic phase transitions with no external (or at least small) load. In this way we capture the essential behaviour of a generic free energy satisfying the properties listed in Section 1.3, and neglect at the same time all aspects depending on the growth of the energy density. In the limit, the weak formulation of the energy conservation equation (3.10) becomes a heat equation with a heat source at the austenite-martensite phase interface. This heat source is proportional to the difference of entropy between martensite and austenite, that is to the latent heat, is concentrated at the phase boundary, which might be sharp or diffuse, and depends also on the velocity of the phase transition. Where the transition is from austenite to martensite, the heat source has a positive sign, while its sign is negative where the transition is from martensite to austenite. The deformation gradient $\nabla \mathbf{y}_k$ generates as k tends to infinity a Gradient Young Measure supported on $K \cup SO(3)$. As shown in Remark 3.5.1 the evolution of the obtained Gradient Young Measure cannot be deduced from the conservation of momentum alone.

In what follows we will make use of Young Measures, for which we refer the reader to [8, 21, 69] and references therein. Below, the space of Young Measures is denoted by $L_{w*}^\infty(\Omega \times (0, T); \mathcal{M}_1(\mathbb{R}^{3 \times 3}))$, where $\mathcal{M}_1(\mathbb{R}^{3 \times 3})$ is the space of probability measures on $\mathbb{R}^{3 \times 3}$ and where the subscript w^* stands for the fact that we endow this space with the weak* topology.

Proposition 3.5.1. *Let $T > 0$, let $(\theta_k, \mathbf{y}_k)$ be a solution of (3.10)–(3.11) given by Theorem 3.4.1 and such that $\mathbf{y}|_{t=0} = \mathbf{y}_{a,k} \in W^{1,p}(\Omega, \mathbb{R}^3) \cap \mathbf{V}$. Assume also $\mathbf{y}_{a,k}$ to be such that*

$$\int_{\Omega} \phi_2(\nabla \mathbf{y}_{a,k}) \leq ck^{-1}, \quad \left| \int_{\Omega} \mathbf{y}_{a,k} \, d\mathbf{x} \right| \leq c, \quad \|\Delta \mathbf{y}_{a,k}\|_2 \leq ck$$

for every $k \geq 1$, and for some positive constant c independent of k . Then there exist

$$\begin{aligned} \mathbf{y} &\in L^\infty(0, T; W^{1,p}(\Omega; \mathbb{R}^3)) \cap W^{1,\infty}(0, T; L^2(\Omega; \mathbb{R}^3)), \\ \theta &\in L^2(0, T; H^1(\Omega)) \cap L^\infty(0, T; L^2(\Omega)), \end{aligned} \tag{3.44}$$

and $\nu_{\mathbf{x},t} \in L_{w^*}^\infty(\Omega \times (0, T); \mathcal{M}_1(\mathbb{R}^{3 \times 3}))$, $\nu_{\mathbf{x},0} \in L_{w^*}^\infty(\Omega; \mathcal{M}_1(\mathbb{R}^{3 \times 3}))$ such that, up to a subsequence,

$$\begin{aligned} \mathbf{y}_k &\rightharpoonup \mathbf{y}, & \text{weakly}^* \text{ in } L^\infty(0, T; W^{1,p}(\Omega; \mathbb{R}^{3 \times 3})), \\ \theta_k &\rightharpoonup \theta, & \text{weakly in } L^2(0, T; H^1(\Omega)), \\ \delta_{\nabla \mathbf{y}_k(\mathbf{x}, t)} &\rightharpoonup \nu_{\mathbf{x}, t}, & \text{weakly}^* \text{ in } L_{w^*}^\infty(\Omega \times (0, T); \mathcal{M}_1(\mathbb{R}^{3 \times 3})), \\ \delta_{\nabla \mathbf{y}_{a,k}(\mathbf{x})} &\rightharpoonup \nu_{\mathbf{x}, 0}, & \text{weakly}^* \text{ in } L_{w^*}^\infty(\Omega; \mathcal{M}_1(\mathbb{R}^{3 \times 3})). \end{aligned} \quad (3.45)$$

Furthermore, $\nu_{\mathbf{x}, t}, \nu_{\mathbf{x}, 0}$ satisfy

$$\begin{aligned} \text{supp } \nu_{\mathbf{x}, t} &\subset K \cup SO(3), & \text{for a.e. } (\mathbf{x}, t) \in \Omega \times (0, T), \\ \text{supp } \nu_{\mathbf{x}, 0} &\subset K \cup SO(3), & \text{for a.e. } \mathbf{x} \in \Omega, \end{aligned} \quad (3.46)$$

and $(\theta, \nabla \mathbf{y}, \nu_{\mathbf{x}, t})$ are satisfying the following weak formulation of the equation governing the conservation of energy

$$\begin{aligned} \int_0^T \int_\Omega \zeta C \nabla \theta \cdot \nabla \omega \, d\mathbf{x} &= \int_0^T \int_\Omega \left(\theta_T \int_{\mathbb{R}^{3 \times 3}} \eta_1(\mathbf{A}) \, d\nu_{\mathbf{x}, t}(\mathbf{A}) + \gamma \rho_0 \theta \right) \dot{\zeta} \omega \, d\mathbf{x} dt \\ &+ \zeta(0) \int_\Omega \omega \left(\theta_T \int_{\mathbb{R}^{3 \times 3}} \eta_1(\mathbf{A}) \, d\nu_{\mathbf{x}, 0}(\mathbf{A}) + \gamma \rho_0 \theta_0 \right) d\mathbf{x}, \end{aligned} \quad (3.47)$$

for each $\zeta \in C_c^1(-T, T)$, $\omega \in C^1(\Omega) \cap H_D^1$, and

$$\int_{\mathbb{R}^{3 \times 3}} \mathbf{A} \, d\nu_{\mathbf{x}, t}(\mathbf{A}) = \nabla \mathbf{y}(\mathbf{x}, t), \quad \text{a.e. } (\mathbf{x}, t) \in \Omega \times (0, T). \quad (3.48)$$

Also $\theta = \theta_B$ on $\partial\Omega_D$ for a.e. $t \in (0, T)$ and $\dot{\theta} + \frac{\partial}{\partial t} \int_{\mathbb{R}^{3 \times 3}} \eta(\mathbf{A}) \nu_{\mathbf{x}, t}(\mathbf{A}) \in L^2(0, T; H_D^{-1})$.

Proof. Thanks to the hypotheses on $\mathbf{y}_{a,k}$ we can deduce the existence of a positive constant c_T such that (3.12) becomes

$$\begin{aligned} \left(\|\dot{\mathbf{y}}_k\|_2^2 + \|\theta_k\|_2^2 + \|\mathbf{y}_k\|_{W^{1,p}}^p + \frac{1}{k} \|\Delta \mathbf{y}_k\|_2^2 \right)(t) \\ + \int_0^t \|\nabla \theta_k\|_2^2(\tau) \, d\tau + k \int_\Omega \phi_2(\nabla \mathbf{y}_k(\mathbf{x}, t)) \, d\mathbf{x} \leq c_T, \end{aligned}$$

for a.e. $t \in (0, T)$. Therefore, we can deduce the existence of $(\mathbf{y}, \theta)$ satisfying (3.44), and of a converging subsequence satisfying (3.45)₁–(3.45)₂.

Now, for fixed $\varepsilon > 0$ we have that the above inequality implies

$$\begin{aligned} \frac{c_T T}{k} &\geq \int_0^T \int_\Omega \phi_2(\nabla \mathbf{y}_k) \, d\mathbf{x} \, dt \\ &\geq c_\varepsilon \mathcal{L}^4 \{ (\mathbf{x}, t) \in \Omega \times (0, T) : |\nabla \mathbf{y}_k(\mathbf{x}, t) - \mathbf{F}| \geq \varepsilon, \forall \mathbf{F} \in K \cup SO(3) \}, \end{aligned}$$

for some constant $c_\varepsilon > 0$ depending on ε and on the continuous function ϕ_2 . Therefore,

$$\lim_{k \rightarrow \infty} \mathcal{L}^4 \{(\mathbf{x}, t) \in \Omega \times (0, T) : |\nabla \mathbf{y}_k(\mathbf{x}, t) - \mathbf{F}| \geq \varepsilon, \forall \mathbf{F} \in K \cup SO(3)\} = 0,$$

which is convergence in measure of $\nabla \mathbf{y}_k$ to $K \cup SO(3)$ on $\Omega \times (0, T)$. Thus, the fundamental theorem of Young measures (see e.g., Theorem 1.2.3) assures the existence of a non-relabelled subsequence and of a family of parametrized probability measures $\nu_{\mathbf{x},t} \in L_{w^*}^\infty((0, T) \times \Omega; \mathcal{M}_1(\mathbb{R}^{3 \times 3}))$ such that

$$\text{supp } \nu_{\mathbf{x},t} \subset K \cup SO(3), \quad \text{for a.e. } (\mathbf{x}, t) \in \Omega \times (0, T),$$

and

$$\lim_n \int_0^T \int_\Omega h(\nabla \mathbf{y}_k(\mathbf{x}, t)) \zeta(\mathbf{x}, t) \, d\mathbf{x} \, dt = \int_0^T \int_\Omega \int_{\mathbb{R}^{3 \times 3}} h(\mathbf{A}) \, d\nu_{\mathbf{x},t}(\mathbf{A}) \zeta(\mathbf{x}, t) \, d\mathbf{x} \, dt,$$

for every $h: \mathbb{R}^{3 \times 3} \rightarrow \mathbb{R}$ of growth strictly less than p , and every $\zeta \in L^\infty((0, T) \times \Omega)$. Choosing $h(\mathbf{A}) = \mathbf{A}$, by (3.20) we immediately get that (3.48) holds. Furthermore,

$$\eta_1(\nabla \mathbf{y}_k) \rightharpoonup \int_{\mathbb{R}^{3 \times 3}} \eta_1(\mathbf{A}) \, d\nu_{\mathbf{x},t}(\mathbf{A}), \quad \text{in } L^q(\Omega \times (0, T)),$$

for every $q \in [1, \frac{p}{r})$. A similar argument entails that $\mathbf{y}_{a,k}$ generates $\nu_{\mathbf{x},0}$ and that

$$\begin{aligned} \text{supp } \nu_{\mathbf{x},0} &\subset K \cup SO(3), \quad \text{for a.e. } \mathbf{x} \in \Omega, \\ \eta_1(\nabla \mathbf{y}_{a,k}) &\rightharpoonup \int_{\mathbb{R}^{3 \times 3}} \eta_1(\mathbf{A}) \, d\nu_{\mathbf{x},0}(\mathbf{A}), \quad \text{in } L^q(\Omega), \end{aligned}$$

for every $q \in [1, \frac{p}{r})$.

We can thus pass to the limit in (3.10) and get (3.47). Finally, by comparison, we deduce $\dot{\theta} + \frac{\partial}{\partial t} \int_{\mathbb{R}^{3 \times 3}} \eta(\mathbf{A}) \nu_{\mathbf{x},t}(\mathbf{A}) \in L^2(0, T; H_D^{-1})$. $\square$

Remark 3.5.1. Under the assumptions of Proposition 3.5.1, when passing to the limit as $k \rightarrow \infty$, we notice that the equation for the conservation of momentum, that is (3.11), reduces to

$$\int_\Omega \int_{\mathbb{R}^{3 \times 3}} \frac{\partial \phi_2}{\partial \mathbf{F}}(\mathbf{A}) \, d\nu_{\mathbf{x},t}(\mathbf{A}) \, d\mathbf{x} : \nabla \boldsymbol{\psi} \, d\mathbf{x} = 0, \quad \forall \boldsymbol{\psi} \in C^1(\Omega; \mathbb{R}^3), \quad \text{a.e. } t \in (0, T). \quad (3.49)$$

Indeed, (3.11) with $\xi \in C_c^1(0, T)$ becomes

$$\begin{aligned} &\int_0^T \int_\Omega \xi \frac{\partial \phi_2}{\partial \mathbf{F}}(\nabla \mathbf{y}_k) \, d\mathbf{x} : \nabla \boldsymbol{\psi} \, d\mathbf{x} \, dt \\ &= -\frac{1}{k} \int_0^T \int_\Omega \left(\xi \frac{\partial \phi_1}{\partial \mathbf{F}}(\nabla \mathbf{y}_k, \theta) : \nabla \boldsymbol{\psi} - \rho_0 \dot{\xi} \dot{\mathbf{y}}_k \cdot \boldsymbol{\psi} + \frac{1}{k} \Delta \mathbf{y}_k \cdot \Delta \boldsymbol{\psi} \right) \, d\mathbf{x} \, dt. \end{aligned} \quad (3.50)$$

By the Cauchy-Schwarz and the Hölder inequalities, thanks to (3.6) we get

$$\begin{aligned}
& \int_0^T \int_{\Omega} \left(\xi \frac{\partial \phi_1}{\partial \mathbf{F}}(\nabla \mathbf{y}_k, \theta) : \nabla \boldsymbol{\psi} \right) d\mathbf{x} dt \\
& \leq \hat{c}_T (1 + \|\nabla \mathbf{y}_k\|_{L^\infty(0,T;L^p)}^p + \|\nabla \mathbf{y}_k\|_{L^\infty(0,T;L^p)}^{\tilde{r}} \|\theta_k\|_{L^2(0,T;L^6)}) \|\xi\|_{C^1} \|\boldsymbol{\psi}\|_{C^1} \\
& \leq \hat{C}_T \|\xi\|_{C^1} \|\boldsymbol{\psi}\|_{C^1},
\end{aligned}$$

where $c_T, \hat{C}_T$ are positive constants independent of k , and where we also made use of the bounds on $(\theta_k, \nabla \mathbf{y}_k)$ in the proof of Proposition 3.5.1. On the other hand,

$$\begin{aligned}
& \int_0^T \int_{\Omega} \rho_0 \xi \dot{\mathbf{y}}_k \cdot \boldsymbol{\psi} d\mathbf{x} dt \leq c_T \rho_0 \|\dot{\mathbf{y}}_k\|_{L^\infty(0,T;L^2)} \|\boldsymbol{\psi}\|_2 \|\xi\|_{C^1} \leq \tilde{c}_T \|\xi\|_{C^1} \|\boldsymbol{\psi}\|_{C^1}, \\
& \int_0^T \int_{\Omega} \frac{1}{k} \xi \Delta \mathbf{y}_k \cdot \Delta \boldsymbol{\psi} d\mathbf{x} dt \leq \frac{c_T}{k} \|\Delta \mathbf{y}_k\|_{L^\infty(0,T;L^2)} \|\xi\|_{C^1} \|\boldsymbol{\psi}\|_{H^2} \leq \tilde{c}_T \|\xi\|_{C^1} \|\boldsymbol{\psi}\|_{H^2}
\end{aligned}$$

for some positive constant $\tilde{c}_T$ independent of k . A passage to the limit in (3.50), together with (3.6) and the fact that $\nabla \mathbf{y}_k$ generates $\nu_{\mathbf{x},t}$ (cf. (3.45)) lead to (3.49). In conclusion, (3.9)₂ does not determine the evolution for $\nu_{\mathbf{x},t}$ after taking the limit $k \rightarrow \infty$. Actually, given that ϕ_2 is smooth, and hence $\frac{\partial \phi_2}{\partial \mathbf{F}}(\mathbf{F}) = \mathbf{0}$ for each $\mathbf{F} \in SO(3) \cup K$, (3.49) does not even add any further information to (3.46).

Remark 3.5.2. After taking the limit $k \rightarrow \infty$, the only information on the time evolution of $\mathbf{y}$ is embedded in $\frac{\partial}{\partial t} \int_{\mathbb{R}^{3 \times 3}} \eta_1(\mathbf{A}) d\nu_{\mathbf{x},t}$, and in $\mathbf{y} \in W^{1,\infty}(0,T;L^2(\Omega)) \cap L^\infty(0,T;W^{1,\infty}(\Omega))$. Indeed, we point out that, by (3.46), $\int_{\mathbb{R}^{3 \times 3}} \eta_1(\mathbf{A}) d\nu_{\mathbf{x},t} \in L^\infty(\Omega \times (0,T))$. Thus,

$$\int_{\mathbb{R}^{3 \times 3}} \eta_1(\mathbf{A}) d\nu_{\mathbf{x},t} + \theta \in H^1(0,T;H_D^{-1}) \cap L^\infty((0,T);L^2(\Omega)).$$

By [24, Lemma II.5.9] this also belongs to $C(0,T;L_w^2(\Omega))$, from which we can make sense of the initial conditions. Furthermore, we also have that [24, Lemma II.5.9] together with (3.44) imply that $\mathbf{y} \in C(0,T;W_w^{1,\infty}(\Omega))$.

Remark 3.5.3. The gradient Young measure generated by $\nabla \mathbf{y}_k$ as $k \rightarrow \infty$ is in general non-trivial, that is not of the form $\delta_{\nabla \mathbf{y}(\mathbf{x},t)}$ for a.e. $(\mathbf{x},t) \in \Omega \times (0,T)$. Indeed, we know from the static theory (see Section 1.3) that martensitic microstructures may arise in order to achieve compatibility between phases, and hence, as $k \rightarrow \infty$ faster and faster oscillations may occur in $\nabla \mathbf{y}_k$, generating non-trivial gradient Young measures.

Remark 3.5.4. If we assume as in Chapter 4 or in Section 3.7 below that there exist $\Omega_A(t), \Omega_M(t) \subset \Omega$ open such that

$$\Omega_A(t) \cap \Omega_M(t) = \emptyset, \quad \mathcal{L}^3(\Omega \setminus (\Omega_A(t) \cup \Omega_M(t))) = 0, \quad \text{a.e. } t \in (0, T), \quad (3.51)$$

and

$$\begin{aligned} \nu_{\mathbf{x},t}(SO(3)) &= 1, & \text{a.e. } \mathbf{x} \in \Omega_A(t), \text{ a.e. } t \in (0, T), \\ \nu_{\mathbf{x},t}(K) &= 1, & \text{a.e. } \mathbf{x} \in \Omega_M(t), \text{ a.e. } t \in (0, T). \end{aligned}$$

then

$$\int_{\mathbb{R}^{3 \times 3}} \eta_1(\mathbf{A}) \, d\nu_{\mathbf{x},t} = -\alpha \chi_{\Omega_m},$$

where by χ_{Ω_m} we denoted the indicator function on Ω_M . The formula for differentiation of integrals on time dependent domains implies

$$\langle \dot{\chi}_{\Omega_M}(\nabla \mathbf{y}), \psi \rangle = \frac{d}{dt} \int_{\Omega_M} \psi \, d\mathbf{x} = \int_{\Gamma(t)} (\mathbf{v} \cdot \mathbf{n}) \psi \, d\mathcal{H}^2, \quad \forall \psi \in C_0^\infty(\Omega), \quad (3.52)$$

provided $\Omega_A(t), \Omega_M(t)$ and $\mathbf{v} \cdot \mathbf{n}$ are smooth enough (see e.g., [42]). Here $\Gamma(t) := \Omega \setminus (\Omega_A \cup \Omega_M)(t)$ is a surface separating $\Omega_A(t)$ from $\Omega_M(t)$, $\mathbf{n}$ denotes the outer normal to Ω_M and $\mathbf{v}(\mathbf{s})$ is the velocity of the interface at the point $\mathbf{s} \in \Gamma(t)$ at time t . By $\langle \cdot, \cdot \rangle$ we denoted the duality pairing between a distribution and a test function. A version of (4.4) under the moving mask approximation can be found in Corollary 4.4.4 in Chapter 4. It thus appears clear from (4.4) that the model needs to be closed with some relation between $\mathbf{v}$ and $\nabla \mathbf{y}$, θ_B , and possibly θ , as much as it requires an initial condition for $\Omega_M(t)$.

3.6 Connections with the moving mask assumption

In this section we draw connections between the results of Section 3.5, and the moving mask assumptions introduced in Chapter 4. Let us start by recalling the *moving mask* assumptions:

- (MM1) the phases are separated, that is at almost every point we have either austenite or martensite but not both, so that phase interfaces between austenite and martensite are sharp (cf. (3.51));

(MM2) during the phase transition, the deformation gradient remains equal to the identity in the austenite region. This is the case, for example, when the austenite region is connected;

(MM3) microstructures do not change after the transformation has happened;

(MM4) the phase interface moves continuously. More precisely, for almost every point $\mathbf{x}$ in the domain, there exists a time when $\mathbf{x}$ is contained in the phase interface.

As proved in Section 3.5, in the limit as k tends to ∞ , the system (3.9) reduces to finding a function θ and a time-dependant gradient Young measure $\nu_{\mathbf{x},t}$ satisfying

$$\text{supp } \nu_{\mathbf{x},t} \subset K \cup SO(3), \quad \text{a.e. } (\mathbf{x}, t) \in \Omega \times (0, T), \quad (3.53)$$

and the weak formulation in (3.47) of

$$\gamma \rho_0 \theta_t - \text{Div}(\mathbf{C} \nabla \theta) = -\theta_T \frac{\partial}{\partial t} \int_{\mathbb{R}^{3 \times 3}} \eta_1(\mathbf{A}) \, d\nu_{\mathbf{x},t}(\mathbf{A}). \quad (3.54)$$

On the one hand, it is clear that, given an arbitrary parametrised family of measures $\nu_{\mathbf{x},t} \in L^\infty(\Omega \times (0, T); \mathcal{M}_1(\mathbb{R}^{3 \times 3}))$ satisfying (3.46), (3.48), $\frac{\partial}{\partial t} \int_{\mathbb{R}^{3 \times 3}} \eta_1(\mathbf{A}) \, d\nu_{\mathbf{x},t}(\mathbf{A}) \in L^2(0, T; H_D^{-1})$ and (MM1)–(MM4), then one can find the related temperature field $\theta(\mathbf{x}, t) \in L^2(0, T; H_D^1)$ by solving the weak form of the heat equation (3.54). This means that sufficiently regular moving mask solutions are solutions to (3.53)–(3.54). On the other hand, as pointed out in Remark 3.5.4, the system (3.53)–(3.54) is under-determined, and should be closed with some constitutive relation between the velocity of the phase interface and the temperature, $\nu_{\mathbf{x},t}$, and possibly other quantities. This is done, for example, in Section 3.7 in a one dimensional context. However, we conjecture that (MM1)–(MM4) can be deduced to hold in the limit as $k \rightarrow \infty$ for $(\theta_k, \mathbf{y}_k)$, weak solutions to (3.8)–(3.9) provided by Theorem 3.4.1.

In order to justify (MM1) one should prove that ,

$$\mathcal{L}^4(\{(\mathbf{x}, t) : \nu_{\mathbf{x},t}(K) \in \{0, 1\}\}) = 0,$$

where $\nu_{\mathbf{x},t}$ is as in Proposition 3.5.1. We recall that (see e.g., [8, 62, 69]), defined $\nu_{\mathbf{x},t}^{k,\delta}$ as

$$\nu_{\mathbf{x},t}^{k,\delta}(\mathcal{B}) = \frac{\mathcal{L}^4(\{(\mathbf{z}, s) \in B((\mathbf{x}, t); \delta) \cap (\Omega \times (0, T)) : \nabla \mathbf{y}_k(\mathbf{z}, s) \in \mathcal{B}\})}{\mathcal{L}^4(B((\mathbf{x}, t); \delta) \cap (\Omega \times (0, T)))},$$

for every $\mathcal{B} \subset \mathbb{R}^{3 \times 3}$ Borel, for every $(\mathbf{x}, t) \in \Omega \times (0, T)$, and where $B((\mathbf{x}, t); \delta)$ is the open ball of radius δ centred at $(\mathbf{x}, t)$, we have

$$\lim_{\delta \rightarrow 0} \lim_{k \rightarrow \infty} \nu_{\mathbf{x},t}^{k,\delta} = \nu_{\mathbf{x},t},$$

for a.e. $(\mathbf{x}, t) \in \Omega \times (0, T)$, and where the convergence is weak* in the sense of measures. Here, $\nabla \mathbf{y}_k$ is the sequence generating $\nu_{\mathbf{x}, t}$ as in (3.45). Therefore, to rigorously prove (MM1), one could prove that for almost every $(\mathbf{x}, t) \in \Omega \times (0, T)$, for every Borel set $\mathcal{B} \subset \mathbb{R}^{3 \times 3}$, and for every $\varepsilon > 0$ there exist $\delta_0 > 0$ such that for every $\delta \in (0, \delta_0)$

$$\lim_{k \rightarrow \infty} \frac{\mathcal{L}^4(\{(z, s) \in B((\mathbf{x}, t); \delta) \cap (\Omega \times (0, T)) : \nabla \mathbf{y}_k(z, s) \in \mathcal{B}\})}{\mathcal{L}^4(B((\mathbf{x}, t); \delta) \cap (\Omega \times (0, T)))} \in (0, \varepsilon) \cup (1 - \varepsilon, 1 + \varepsilon).$$

This result should be easier to prove in the case of materials satisfying the cofactor conditions where martensitic laminates can form exact austenite-martensite interfaces (see [26]). Furthermore, one should restrict to initial data $\mathbf{y}_{a, k}$ generating a gradient Young measure $\nu_{\mathbf{x}, 0}$ such that $\nu_{\mathbf{x}, 0}(K) \in \{0, 1\}$ for a.e. $\mathbf{x} \in \Omega$.

Regarding (MM2), it is a consequence of a well known result by Reshetnyak (see [73], but also [19, 62]) that a parametrised measure $\mu_{\mathbf{x}}$ with $\text{supp } \mu_{\mathbf{x}} \subset SO(3)$ for a.e. $\mathbf{x}$ in an open connected subset $\mathcal{A} \subset \Omega$ must be of the form $\mu_{\mathbf{x}} = \delta_{\mathbf{R}}$ for a.e. $\mathbf{x} \in \mathcal{A}$, some $\mathbf{R} \in SO(3)$, and where $\delta_{\mathbf{R}}$ denotes a Dirac delta at $\mathbf{R}$. So the problem reduces to show that the set $\{\mathbf{x} : \nu_{\mathbf{x}, t}(SO(3)) = 1\}$ is connected, or, more in general, indecomposable (see [37, 62]).

A context where it might be easier to show that (MM1)–(MM2) hold is when $|\det(\mathbf{F}) - 1| > \delta$ for every $\mathbf{F} \in K$, and for some $\delta > 0$. Indeed, in this case, we expect for k large enough to just have nucleations at the corners of the sample (see [21, 22]), and the presence of martensitic islands in austenite (or viceversa) to be energetically too expensive.

Proving (MM3) rigorously from the solutions to (3.8)–(3.9) in the limit as k tends to infinity might be very difficult in general. Assuming for simplicity (MM1), we now give an intuitive explanation of why (MM3) should hold based on rigidity, which is well known to play an important role in vectorial problems. As the map $\mathbf{y} \in W^{1, \infty}(\Omega; \mathbb{R}^3)$, with $\nabla \mathbf{y}$ satisfying (3.48), is Lipschitz continuous, its gradient cannot be arbitrary (see e.g., [18, 19, 62]). Therefore, changing the macroscopic deformation gradient $\nabla \mathbf{y}$ in a subset of $\{\mathbf{x} : \nu_{\mathbf{x}, t}(K) = 1\}$ of positive measure, might require changing $\nabla \mathbf{y}$ in a much larger set, leading to discontinuities in t of $\mathbf{y}$ on a subset Ω of positive measure. But such a jump in t would contradict the fact that, by Proposition 3.5.1, $\mathbf{y} \in W^{1, \infty}(0, T; L^2(\Omega; \mathbb{R}^3))$.

The continuous movement of the phase interface, that is assumption (MM4), is partially a consequence of the observations in Remark 3.5.2. Indeed, let us assume that

(MM1) and (MM3) hold, and, for simplicity, that

$$\min_{\mathbf{F} \in K^{qc}} \min_{\mathbf{R} \in SO(3)} |\mathbf{F} - \mathbf{R}| \geq \varepsilon,$$

for some $\varepsilon > 0$. Here we are taking $SO(3)$ and not $SO(3)^{qc}$ as these two sets coincide (see e.g., [19]). In this case, suppose (MM4) is not satisfied. Then there exist $t_0 \in (0, T)$, and a measurable set $\mathcal{A}$, with $\mathcal{L}^3(\mathcal{A}) > 0$, such that the macroscopic deformation gradient $\nabla \mathbf{y}$ satisfies $\nabla \mathbf{y}(\mathbf{x}, s) \in SO(3)$ for a.e. $(\mathbf{x}, s) \in \Omega \times (0, t_0)$, and $\nabla \mathbf{y}(\mathbf{x}, s) \in K$ for a.e. $(\mathbf{x}, s) \in \Omega \times [t_0, t_0 + \delta]$, for some $\delta > 0$. Given the fact that we are assuming that (MM3) also holds, there exist $\boldsymbol{\psi} \in W^{1,\infty}(\Omega; \mathbb{R}^3)$, independent of time, such that $\nabla \boldsymbol{\psi} = \nabla \mathbf{y}(\mathbf{x}, t) - \nabla \mathbf{y}(\mathbf{x}, s)$, for every s, t satisfying $s < t_0 < t < t_0 + \delta$. In this case,

$$\int_{\Omega} (\nabla \mathbf{y}(\mathbf{x}, t) - \nabla \mathbf{y}(\mathbf{x}, s)) : \nabla \boldsymbol{\psi} \, d\mathbf{x} = \int_{\Omega} |\nabla \mathbf{y}(\mathbf{x}, t) - \nabla \mathbf{y}(\mathbf{x}, s)|^2 \, d\mathbf{x} \geq \varepsilon^2 \mathcal{L}^3(\mathcal{A}) > 0,$$

for every $s < t$, thus contradicting the fact that, by Remark 3.5.2, $\mathbf{y}$ is in the space $C([0, T]; W_w^{1,\infty}(\Omega; \mathbb{R}^3))$, which implies

$$\int_{\Omega} (\nabla \mathbf{y}(\mathbf{x}, t) - \nabla \mathbf{y}(\mathbf{x}, s)) : \nabla \boldsymbol{\psi} \, d\mathbf{x} \rightarrow 0, \quad \text{as } s \rightarrow t.$$

3.7 The evolution of a simple laminate in a one-dimensional context

The aim of this section is to introduce some hypotheses that allow us to construct solutions to the limit problem (3.46)–(3.47) by solving a one-dimensional heat equation with a measure valued source depending nonlinearly on the unknown. This models the evolution of the phase boundary between a laminate of martensite and austenite and is a simplified version of the model in [2, 52]. The resulting solutions are an example of solutions to (3.53)–(3.54) satisfying (MM1)–(MM4).

3.7.1 An example: the simple laminate

Let us suppose that an austenite to martensite phase transition takes place in a parallelepiped with a circular base, that is, in cylindrical coordinates, $\Omega_C := (0, 2\pi) \times (0, r) \times (0, L)$ for some $0 < r \ll L$. Suppose further that there exists a single phase interface $\Gamma(t) = \{\mathbf{x} \in \Omega_C : \mathbf{x} \cdot \mathbf{e}_3 = u(t)\}$, $u(t) \in (0, L)$ perpendicular to $\mathbf{e}_3$ for every

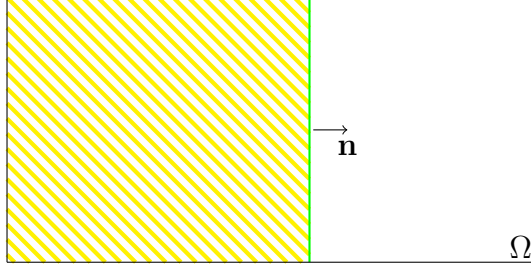


Figure 3.1: Picture of a section of Ω_C , on the left $\Omega_M(t)$ where $\nabla \mathbf{y} = \mathbf{1} + \mathbf{a} \otimes \mathbf{n}$, on the right Ω_A , where $\nabla \mathbf{y} = \mathbf{1}$.

t , and that

$$\begin{aligned} \nu_{\mathbf{x},t} &= \lambda \delta_{\mathbf{A}} + (1 - \lambda) \delta_{\mathbf{B}}, & \text{a.e. in } \Omega_M(t), \text{ a.e. } t > 0, \\ \nu_{\mathbf{x},t} &= \delta_{\mathbf{1}}, & \text{a.e. in } \Omega_A(t), \text{ a.e. } t > 0, \end{aligned}$$

with $\mathbf{A}, \mathbf{B} \in K$, $\lambda \in [0, 1]$ fixed, and

$$\Omega_M(t) = \{\mathbf{x} \in \Omega_C : \mathbf{x} \cdot \mathbf{e}_3 < u(t)\}, \quad \Omega_A(t) = \{\mathbf{x} \in \Omega_C : \mathbf{x} \cdot \mathbf{e}_3 > u(t)\}.$$

Here and below $\delta_{\mathbf{F}}$ is a Dirac measure at $\mathbf{F}$. By (3.48), the macroscopic deformation gradient, that is an average of the fine microstructures, is given by the barycentre of the Young measure, that is, in our case

$$\begin{aligned} \nabla \mathbf{y} &= \int_{\mathbb{R}^{3 \times 3}} \mathbf{F} d\nu_{\mathbf{x},t}(\mathbf{F}) = \lambda \mathbf{A} + (1 - \lambda) \mathbf{B}, & \text{a.e. in } \Omega_M(t), \text{ a.e. } t > 0, \\ \nabla \mathbf{y} &= \int_{\mathbb{R}^{3 \times 3}} \mathbf{F} d\nu_{\mathbf{x},t}(\mathbf{F}) = \mathbf{1}, & \text{a.e. in } \Omega_A(t), \text{ a.e. } t > 0. \end{aligned}$$

However, as shown in [18], $\nabla \mathbf{y}$ is a gradient of a continuous deformation map $\mathbf{y} : \Omega_C \rightarrow \mathbb{R}^3$ if and only if there exist $\mathbf{a} \in \mathbb{R}^3$ such that

$$\lambda \mathbf{A} + (1 - \lambda) \mathbf{B} = \mathbf{1} + \mathbf{a} \otimes \mathbf{e}_3.$$

We can hence deduce that

$$\mathbf{y} = \begin{cases} \mathbf{x} + \mathbf{a}(\mathbf{x} \cdot \mathbf{e}_3) + \mathbf{c}_1(t), & \text{in } \Omega_M(t), \\ \mathbf{x} + \mathbf{c}_2(t), & \text{in } \Omega_A(t), \end{cases}$$

for some time dependent vectors $\mathbf{c}_1, \mathbf{c}_2 : [0, \infty) \rightarrow \mathbb{R}^3$ such that $\mathbf{c}_1(t) - \mathbf{c}_2(t) = \mathbf{a}u(t)$. In an equivalent way, knowing $\mathbf{c}_1, \mathbf{c}_2$ we can determine the position of the phase boundary given by those $\mathbf{x}$ satisfying $(\mathbf{x} \cdot \mathbf{e}_3) = \frac{|\mathbf{c}_2(t) - \mathbf{c}_1(t)|}{|\mathbf{a}|}$. For simplicity we choose constant Dirichlet boundary conditions $\theta_B \neq \theta_T$ for $x_3 = 0, L$, and $\frac{\partial \theta}{\partial \mathbf{m}} = 0$ on the

other face, having normal $\mathbf{m}$. In this case, the equation governing the conservation of energy simplifies to a one-dimensional equation. We define $s \in (0, L)$ to be the coordinate in direction $\mathbf{e}_3$, we assume the phase interface to be located at $s = u_0$ when $t = 0$, and assume without loss of generality that $\Omega_M(0) = \{s : s < u_0\}$. In this case, the classical formulation of (3.47) reduces to

$$\begin{cases} \rho_0 \gamma \dot{\theta} - K \theta_{ss} = \alpha \dot{\chi}_{\Omega_M}(t), & \text{in } \Omega \times (0, T), \\ \chi_{\Omega_M}(s, t) = \chi_{\{s < u(t)\}}, & \text{in } (0, T), \\ \theta(0, t) = \theta(L, t) = \theta_B, & \text{in } (0, T), \\ \theta|_{t=0} = \theta_0, & \text{in } (0, L), \\ u(0) = u_0, \end{cases} \quad (3.55)$$

where we denoted by $\chi_{\mathcal{B}}$ the indicator function on the Borel set $\mathcal{B}$. Without loss of generality, below we assume $\alpha = \gamma = \rho_0 = 1$ to simplify notation. This system of equations describes the evolution of a sharp phase interface at $u(t)$. However, this is underdetermined, and needs to be closed with a constitutive relation for the interface speed $v = \frac{d}{dt} \frac{|\mathbf{c}_2(t) - \mathbf{c}_1(t)|}{|\mathbf{a}|} = \dot{u}$ depending on the temperature at $u(t)$ and possibly other variables (see also Remark 3.5.4). In order to respect the physics of the problem we want the martensite domain to expand when $\theta < \theta_T$ and to shrink when $\theta > \theta_T$. Mathematically, this means that we assume

$$\begin{cases} v(\theta) > 0, & \text{if } \theta < \theta_T, \\ v(\theta) < 0, & \text{if } \theta > \theta_T. \end{cases} \quad (3.56)$$

Imposing

$$\dot{u}(t) = v(\theta(u(t), t)) \quad (3.57)$$

makes the problem a two-phase Stefan problem with a kinetic condition at the free boundary (see e.g. [80, 81, 83]). The one-dimensional two-phase Stefan problem is usually closed by imposing that the temperature θ is equal to θ_T at the phase interface located in $s = u(t)$. However, in martensitic transformations nucleation usually happens strictly below the critical temperature due to hysteresis (see e.g., [84]). Given the high thermal conductivity of metals, the high-temperature phase, that is austenite, as much as the phase interface, can be at a temperature strictly lower than the critical-one. This is usually called undercooling in the literature of the Stefan problem (see e.g. [81]). Therefore, the condition $\theta = \theta_T$ at $s = u(t)$ is inaccurate, and is better replaced by (3.57) (see also [2, 52]). Existence of a weak solution to system (3.55) has been studied in [80], while existence and uniqueness of suitably

defined classical solutions is achieved in [83], under the assumption that $\mathbf{v}$ is linear. The aim of the next section is to show that, if $\min\{0, \theta_B\} \leq \theta_0(s) \leq \max\{0, \theta_B\}$ for a.e. $s \in (0, L)$, then the position of the phase interface u is a monotone function of t , and reaches the domain boundary in a finite amount of time.

3.7.2 Behaviour of solutions to the 1-dimensional problem

For simplicity, we define a rescaled temperature $\bar{\theta} = \theta - \theta_B$, the rescaled critical temperature $\theta_c := \theta_T - \theta_B$, and below we consider the following definition of weak solution

Definition 3.7.1. *Let $T > 0$, and $\Omega := (0, L)$. We say that $(\bar{\theta}, u)$ is a weak solution to problem (3.55) if*

$$\begin{aligned} u &\in H^1(0, T), \\ \bar{\theta} &\in L^2(0, T; H_0^1(\Omega)) \cap H^1(0, T; H^{-1}(\Omega)), \\ \bar{\theta} + \theta_B &> 0, \quad \text{a.e. } x \in \Omega, t \in (0, T), \end{aligned}$$

and

$$\langle \dot{\bar{\theta}}, \psi \rangle + \int_{\Omega} K \bar{\theta}_s \psi_s \, dx = \langle \dot{u} \delta_u, \psi \rangle, \quad \text{a.e. } t \in (0, T), \forall \psi \in H_0^1(\Omega), \quad (3.58)$$

$$u(t) = u(0) + \int_0^t \chi_{\{u(\tau) \in \Omega\}} v(\bar{\theta}(u(\tau), \tau)) \, d\tau, \quad \forall t \in (0, T). \quad (3.59)$$

Furthermore, $\bar{\theta}(0, x) = \bar{\theta}_0 := \theta_0 - \theta_B$ almost everywhere in Ω .

Remark 3.7.1. For the sake of coherence, we multiplied by $\chi_{\{t: u(t) \in (0, L)\}}$ the velocity v of the interface, in order to avoid the non-physical situation where the interface moves out of the sample. This does not affect the $H^1(0, T)$ regularity of u .

The following result concerning weak solutions holds

Theorem 3.7.1. *Let v be Lipschitz, $\bar{\theta}_0 \in L^2(\Omega)$ and $u(0) \in \Omega$. Then there exists $(\bar{\theta}, u)$ weak solution in the sense of Definition 3.7.1 to (3.55).*

We refer to [80] for the proof of this result that can be carried out via a Galerkin approximation.

The following proposition gives some good behaviour of the solutions dependent only on (3.56):

Proposition 3.7.1. *Let $(\bar{\theta}, u)$ be a weak solution as in Definition 3.7.1, θ_0 be measurable and let v satisfy (3.56). Then, if $\theta_c \leq 0$ and $\theta_c \leq \bar{\theta}_0 \leq 0$ a.e. in Ω , we have*

$$\theta_c \leq \bar{\theta}(s, t) \leq 0, \quad \forall s \in \Omega, \text{ a.e. } t \in (0, T). \quad (3.60)$$

If $\theta_c \geq 0$ and $\theta_c \geq \bar{\theta}_0 \geq 0$ a.e. in Ω , we have

$$\theta_c \geq \bar{\theta}(s, t) \geq 0, \quad \forall s \in \Omega, \text{ a.e. } t \in (0, T). \quad (3.61)$$

Proof. The proof of a similar statement can be found in [80], but we report here the proof for the sake of completeness. We start with the proof of (3.60). Let us choose as a test function ψ in (3.58) $\psi = (\bar{\theta} - \theta_c)^- := -\max(0, \theta_c - \bar{\theta})$. In this case, by using (3.56) we have

$$\frac{1}{2} \frac{d}{dt} \|(\bar{\theta} - \theta_c)^-\|_2^2 + \|\nabla(\bar{\theta} - \theta_c)^-\|_2^2 \leq v(\bar{\theta}(u(t), t))(\bar{\theta}(u(t), t) - \theta_c)^- \leq 0, \quad \text{a.e. } t \in (0, T).$$

An integration in time, using the hypothesis on $\bar{\theta}_0$ leads to $\|(\bar{\theta} - \theta_c)^-\|_{H^1} = 0$ for a.e. $t \in (0, T)$, that is $\bar{\theta}(s, t) \geq \theta_c$ everywhere in Ω , for almost every $t \in [0, T]$. Let us now choose in (3.58) $\psi = \bar{\theta}^+ := \max(0, \bar{\theta})$, and notice that (3.56) together $\theta_c \leq \bar{\theta}$ implies $v \leq 0$. Therefore,

$$\frac{d}{dt} \|\bar{\theta}^+\|_2^2 + \|\nabla \bar{\theta}^+\|_2^2 \leq v(\bar{\theta}(u(t), t))\bar{\theta}^+(u(t), t) \leq 0, \quad \text{a.e. } t \in (0, T).$$

This, together with the fact that $\|\bar{\theta}_0^+\|_2 = 0$, proves the first statement. The second statement can be proved in an analogous way by first testing against $(\bar{\theta} - \theta_c)^+$ and then against $\bar{\theta}^-$. $\square$

Remark 3.7.2. The above proposition states that, under hypothesis (3.56), a weak solution is bounded between θ_B and θ_T . Therefore, if our boundary condition θ_B is close to θ_T , then the temperature solutions $\bar{\theta}$ have small L^∞ norm. Furthermore, this implies that the position of the interface is monotone, in the sense that if the boundary of the martensite sub-domain is expanding (resp. shrinking) then, as long as the boundary conditions do not change, the domain keeps expanding (resp. shrinking).

We can prove further regularity for θ , as stated in the following

Proposition 3.7.2. *Let $\bar{\theta}_0 \in H_0^1(\Omega)$, $0 \leq \bar{\theta}_0(s) \leq \theta_c$ (or $0 \geq \bar{\theta}_0(s) \geq \theta_c$) for each $s \in [0, 1]$, $u_0 \in \Omega$ and $v \in W_{loc}^{1,\infty}(\mathbb{R})$ satisfy (3.56). Then, for every $T > 0$,*

$$\dot{\bar{\theta}} \in L^2(0, T; L^2(\Omega)), \quad \bar{\theta} \in L^2(0, T; W^{1,\infty}(\Omega)) \cap L^\infty(0, T; H_0^1(\Omega)).$$

Furthermore, for a.e. $t \in (0, T)$ $\bar{\theta}_s$ is continuous in $[0, u(t)) \cap (u(t), L]$.

The proof of this result relies on

Lemma 3.7.1. *Let $v \in W_{loc}^{1,\infty}(\mathbb{R})$. Let $(u, \bar{\theta}_1)$ and $(u, \bar{\theta}_2)$ be two weak solutions to Problem (3.55) sharing the same $u(t) \in H^1(0, T)$ and the same initial datum $\bar{\theta}_0 \in L^\infty(\Omega)$ such that $\bar{\theta}_0 \in [\min\{0, \theta_c\}, \max\{0, \theta_c\}]$ a.e. in Ω . Then, $\bar{\theta}_1 = \bar{\theta}_2$.*

Proof. Let $\theta := \bar{\theta}_1 - \bar{\theta}_2$. Then, it satisfies

$$\langle \theta_t, \psi \rangle + \int_{\Omega} \theta_s \psi_s \, ds = (v(\bar{\theta}_1(u(t), t)) - v(\bar{\theta}_2(u(t), t)))\psi(u(t)),$$

for every $\psi \in H_0^1(0, L)$, a.e. $t \in (0, T)$. By testing this equation with $\psi = \theta$ and using the fact that v is Lipschitz we thus get

$$\frac{1}{2} \frac{d}{dt} \|\theta\|_2^2 + \|\theta_s\|_2^2 \leq c\theta^2(u(t), t).$$

Now, [75, Cor. 27], implies

$$|\theta(u(t), t)| \leq c\|\theta\|_{W^{r,2}}(t),$$

with $r = \frac{1}{2} + \eta$ and $\eta > 0$ arbitrary. An interpolation inequality thus leads to

$$\frac{1}{2} \frac{d}{dt} \|\theta\|_2^2 + \|\theta_s\|_2^2 \leq c(\|\theta\|_2^2 + \|\theta_s\|_2^{2r} \|\theta\|_2^{2(1-r)}) \leq c\|\theta\|_2^2 + \frac{1}{2} \|\theta_s\|_2^2.$$

The Gronwall lemma and the fact that $\theta(0, s) = 0$ for a.e. $s \in \Omega$ conclude the proof. $\square$

Proof of Prop. 3.7.2. Let us fix a solution $(\bar{\theta}, u)$ among the possible solutions to Problem (3.55). We drop the bar over $\bar{\theta}$, and we put all the constants equal to one to simplify notation. Let us also define $\delta_{u(t)}^\varepsilon(s) = \delta_0^\varepsilon(s - u(t))$ as the convolution of the Dirac measure centred in $u(t)$ with a smooth mollifier. $\delta_{u(t)}^\varepsilon$ is positive and converges weakly* in the sense of measures to $\delta_{u(t)}$ for every fixed t as $\varepsilon \rightarrow 0$ (see [4] for details). Furthermore, Fubini's Theorem yields $\|\delta_{u(t)}^\varepsilon\|_{\mathcal{M}} \leq c$, where c is independent of $u(t)$ and ε . For every fixed $\varepsilon > 0$ the approximate problem

$$\langle \dot{\theta}_\varepsilon, \psi \rangle + \int_{\Omega} \theta_{\varepsilon,s} \psi_s \, ds = \int_{\Omega} v(\theta_\varepsilon(s, t)) \delta_{u(t)}^\varepsilon(s) \psi(s) \, ds, \quad \forall \psi \in H_0^1, \text{ a.e. } t \in (0, T) \quad (3.62)$$

$$\theta_\varepsilon(s, 0) = \theta_0(s), \quad \text{a.e. } s \in \Omega \quad (3.63)$$

admits a solution θ_ε . The following bound can be obtained thanks to Proposition 3.7.1 by choosing $\psi = \theta_\varepsilon$ in (3.62)

$$\|\theta_\varepsilon\|_{L^\infty(0,T;L^2)} + \|\theta_{\varepsilon,s}\|_{L^2(0,T;L^2)} + \|\dot{\theta}_\varepsilon\|_{L^2(0,T;H^{-1})} \leq c_T, \quad (3.64)$$

so that, by Lemma 3.7.1, we may assume that

$$\theta_\varepsilon \rightharpoonup \theta \quad \text{weakly in } L^2(0, T; H_0^1) \cap H^1(0, T; H^{-1}), \quad (3.65)$$

$$\theta_\varepsilon \rightarrow \theta \quad \text{strongly in } L^2(0, T; C[0, L]). \quad (3.66)$$

We recall that here $\dot{u}(t) = v(\theta(u(t), t))\chi_{\{u(t) \in (0, L)\}}$, as the equation for u has not been approximated. Besides, by bootstrapping and using standard regularity for the heat equation (see e.g., [39]) we also have $\theta_{\varepsilon, ss} \in L^2(0, T; H^2(\Omega))$ and $\theta_{\varepsilon, t} \in L^2(0, T; L^2(\Omega))$. Finally, an argument similar to that used to prove Proposition 3.7.1 entails $\max_{s \in [0, L]} |\theta_\varepsilon(s, t)| \leq \theta_c$ for a.e. $t \in (0, T)$. We can hence choose $\psi = \dot{\theta}_\varepsilon$ in (3.62) and get

$$\begin{aligned} \|\dot{\theta}_\varepsilon\|_2^2 + \frac{1}{2} \frac{d}{dt} \|\theta_{\varepsilon, s}\|_2^2 &= \int_\Omega v(\theta_\varepsilon) \delta_{u(t)}^\varepsilon \dot{\theta}_\varepsilon \, ds \\ &= \frac{d}{dt} \int_\Omega V(\theta_\varepsilon) \delta_{u(t)}^\varepsilon \, ds - \dot{u} \int_\Omega V(\theta_\varepsilon) (\delta_{u(t)}^\varepsilon)_{,s} \, ds \\ &= \frac{d}{dt} \int_\Omega V(\theta_\varepsilon) \delta_{u(t)}^\varepsilon \, ds + \dot{u} \int_\Omega \theta_{\varepsilon, s} v(\theta_\varepsilon) \delta_{u(t)}^\varepsilon \, ds, \end{aligned} \quad (3.67)$$

where $V(r) = \int_0^r v(\tau) \, d\tau$, and we integrated by parts in the last term. In addition, exploiting the additional regularity of θ_ε , we notice that

$$|\theta_{\varepsilon, s}(x, t)| \leq c(\|\theta_{\varepsilon, s}\|_1 + \|\theta_{\varepsilon, ss}\|_1) \leq c\left(\|\theta_{\varepsilon, s}\|_2 + \|\dot{\theta}_\varepsilon\|_1 + \int_\Omega |v(\theta_\varepsilon(s, t)) \delta_{u(t)}^\varepsilon(s)| \, ds\right).$$

As $\delta_{u(t)}^\varepsilon$ is positive, and v is Lipschitz,

$$\int_\Omega |v(\theta_\varepsilon) \delta_{u(t)}^\varepsilon| \, ds \leq \int_\Omega \delta_{u(t)}^\varepsilon |v(\theta_\varepsilon)| \, ds \leq c(1 + \|\theta_\varepsilon\|_{C[0, L]}) \int_\Omega \delta_{u(t)}^\varepsilon \, ds \leq c(1 + \|\theta_\varepsilon\|_{C[0, L]}),$$

which yields

$$\max_{x \in [0, L]} |\theta_{\varepsilon, s}(x, t)| \leq c(\|\theta_{\varepsilon, s}\|_2 + \|\dot{\theta}_\varepsilon\|_2 + 1 + \|\theta_\varepsilon\|_{C[0, L]}).$$

Therefore,

$$\begin{aligned} \left| \dot{u} \int_\Omega \theta_{\varepsilon, s} v(\theta_\varepsilon) \delta_{u(t)}^\varepsilon \, ds \right| &\leq c |\dot{u}| (\|\theta_{\varepsilon, s}\|_2 + \|\dot{\theta}_\varepsilon\|_2 + 1 + \|\theta_\varepsilon\|_{C[0, L]}) (1 + \|\theta_{\varepsilon, s}\|_2) \\ &\leq c + c |\dot{u}|^2 (1 + \|\theta_{\varepsilon, s}\|_2^2) + \frac{1}{4} (\|\theta_{\varepsilon, s}\|_2^2 + \|\dot{\theta}_\varepsilon\|_2^2 + \|\theta_{\varepsilon, s}\|_2^2). \end{aligned} \quad (3.68)$$

Let us now consider a sequence of $\varepsilon_j > 0$ such that $\varepsilon_j \rightarrow 0$. We remark that (3.66) implies $\theta_{\varepsilon_j} \rightarrow \theta$ strongly in $L^2(0, T; C[0, L])$ and hence in $C[0, L]$ for almost every $t \in (0, T)$. Furthermore, as mentioned above, an argument as in the proof of Proposition

3.7.1 tells us that $\|\theta_{\varepsilon_j}(\cdot, t)\|_{C[0,L]} \leq \theta_c$ for almost every t . We denote by J_T the set of points in $(0, T)$ where the above uniform convergence holds and where the uniform bound for $\theta_{\varepsilon_j}(\cdot, t)$ holds for every ε_j . We point out that this set is measurable and has full measure. After exploiting (5.54) in (3.67), and integrating the resulting inequality in time between 0 and $t \in J_T$, we obtain

$$\begin{aligned} & \int_0^t \|\dot{\theta}_{\varepsilon_j}\|_2^2(\tau) d\tau + \|\theta_{\varepsilon_j,s}\|_2^2(t) \\ & \leq c + 2 \int_{\Omega} V(\theta_{\varepsilon_j}(s, t)) \delta_{u(t)}^{\varepsilon_j}(s) ds - 2 \int_{\Omega} V(\theta_{\varepsilon_j}(s, 0)) \delta_{u(0)}^{\varepsilon_j}(s) ds \quad (3.69) \\ & \quad + 2\|\theta_{\varepsilon_j,s}\|_2^2(0) + c \int_0^t |\dot{u}|^2(1 + \|\theta_{\varepsilon_j,s}\|_2^2) d\tau + \int_0^t \|\theta_{\varepsilon_j,s}\|_2^2 d\tau. \end{aligned}$$

Now, by means of the regularity on the initial condition we easily deduce

$$-2 \int_{\Omega} V(\theta_{\varepsilon_j}(s, 0)) \delta_{u(0)}^{\varepsilon_j}(s) ds + \|\theta_{\varepsilon_j,s}\|_2^2(0) \leq c.$$

and, as $V \in C^1(\mathbb{R})$ and $t \in J_T$,

$$2 \int_{\Omega} V(\theta_{\varepsilon_j}(s, t)) \delta_{u(t)}^{\varepsilon_j}(s) ds \leq 2\|\delta_{u(t)}^{\varepsilon_j}\|_{\mathcal{M}} \|V(\theta_{\varepsilon_j})(\cdot, t)\|_{C[0,L]} \leq c.$$

Finally, we know that (3.64) holds, and that, thanks to Proposition 3.7.1 and Remark 3.7.2, $\dot{u} \in L^\infty(0, T)$. Therefore,

$$\int_0^t |\dot{u}|^2(1 + \|\theta_{\varepsilon_j,s}\|_2^2) d\tau \leq c,$$

and (3.69) becomes

$$\int_0^t \|\dot{\theta}_{\varepsilon_j}(\tau)\|_2^2 d\tau + \|\theta_{\varepsilon_j,s}\|_2^2(t) \leq c, \quad \text{a.e. } t \in (0, T),$$

where c depends neither on $t \in (0, T)$, nor on ε_j . We can thus pass to the limit as $\varepsilon \rightarrow 0$, and obtain

$$\theta \in L^2(0, T; L^2(\Omega)), \quad \theta_s \in L^\infty(0, T; L^2(\Omega)).$$

By comparison, this also implies that $\theta_{ss} + \dot{u}\delta_{u(t)} \in L^2(0, T; L^2(\Omega))$, that is

$$\theta_s + \dot{u}\chi_{\{s < u(t)\}} \in L^2(0, T; C[0, L]),$$

which yields

$$\theta \in L^2(0, T; W^{1,\infty}(\Omega)).$$

□

Remark 3.7.3. The regularity proved in Proposition 3.7.2 entails uniqueness thanks to an argument as that in [83] for classical solutions. We remark that without having $\bar{\theta} \in L^2(0, T; W^{1,\infty}(\Omega))$ uniqueness of the equation $\dot{u} = v(\bar{\theta}(u(t), t))\chi_{\{u(t) \in (0, L)\}}$ would be hopeless. Nonetheless, we point out that even if $\bar{\theta}_0 \notin H_0^1$, the proof of Proposition 3.7.2 entails that an instantaneous regularization takes place, and that in this case $\bar{\theta} \in L^2(\rho, T; W^{1,\infty}(\Omega))$, for each $\rho > 0$.

As a corollary to the above regularity result we can prove that the austenite-martensite interface reaches the domain boundary in finite time, and that $\bar{\theta} \rightarrow 0$ as $t \rightarrow \infty$.

Theorem 3.7.2. *Let $(\bar{\theta}, u)$ be a weak solution in the sense of Definition 3.7.1, let $u_0 \in (0, L)$ and let $\bar{\theta}_0 \in H_0^1(\Omega)$ be such that $0 \leq \bar{\theta}_0(s) \leq \theta_c$ (or $0 \geq \bar{\theta}_0(s) \geq -\theta_c$) for every $s \in [0, L]$. Let also $v \in W_{loc}^{1,\infty}(\mathbb{R})$ satisfy (3.56). Then there exists $t^* > 0$ such that $u(t^*) = L$ (resp. $u(t^*) = 0$). Furthermore, $\|\bar{\theta}\|_2(t) \rightarrow 0$ as $t \rightarrow \infty$.*

Proof. We assume without loss of generality that $\theta_c > 0$ and prove the theorem for the case where $0 \leq \bar{\theta}_0(s) \leq \theta_c$ for every $s \in \Omega$. The other case can be proved similarly. We remark that thanks to Proposition 3.7.1 and (3.56) we have $0 \leq \dot{u}(t) \leq \max_{s \in [0, \theta_c]} v(s)$ for almost every $t \geq 0$, so that u is non-decreasing. Therefore $\lim_{t \rightarrow \infty} u(t)$ exists. Suppose now that there exists no $t^* < \infty$ such that $u(t^*) = L$. Then, for every $\varepsilon > 0$ there exists t_0 such that

$$\int_{t_0}^{\infty} \dot{u}(\tau) d\tau \leq \varepsilon.$$

Let us now test (3.58) with $\bar{\theta}$ to obtain

$$\frac{1}{2} \frac{d}{dt} \|\bar{\theta}\|_2^2(t) + \|\bar{\theta}_s\|_2^2(t) = \dot{u}(t) \bar{\theta}(u(t), t), \quad \text{a.e. } t \geq 0.$$

Thanks to the Sobolev inequality $\|\psi\|_{C[0, L]} \leq c \|\psi_s\|_2$, which holds for every $\psi \in H_0^1(\Omega)$, by Young's inequality we thus deduce

$$\frac{d}{d\tau} \|\bar{\theta}\|_2^2(t_0 + \tau) + \|\bar{\theta}_s\|_2^2(t_0 + \tau) \leq c |\dot{u}|^2(t_0 + \tau) \leq \hat{c} |\dot{u}|(t_0 + \tau), \quad \text{a.e. } \tau \geq 0,$$

for some constants $c, \hat{c} > 0$ and where we made use of $|\dot{u}| \leq c$. Using once again Sobolev embeddings we have

$$\frac{d}{d\tau} y(\tau) + c_P y(\tau) \leq \hat{c} |\dot{u}|(t_0 + \tau), \quad \text{a.e. } \tau \geq 0,$$

for some $c_P > 0$ and where we defined $y(\tau) := \|\bar{\theta}\|_2^2(t_0 + \tau)$. Using the comparison principle for ordinary differential equations we get that

$$y(\tau) \leq y(0) e^{-c_P \tau} + e^{-c_P \tau} \int_0^\tau e^{c_P \hat{\tau}} \hat{c} |\dot{u}|(t_0 + \hat{\tau}) d\hat{\tau}.$$

As $\|\bar{\theta}\|_2$ is a continuous function of time and as $\|\bar{\theta}\|_2 \leq L^{\frac{1}{2}}\theta_c$ for a.e. $t \geq 0$, we have $y(0) \leq L\theta_c^2$. Let us now choose τ such that $L\theta_c^2 e^{-c_0\tau} \leq \varepsilon$, and such that $\bar{\theta}(u(t_0 + \tau), t_0 + \tau) \geq \frac{3\theta_c}{4}$. This is possible because, as $u \in (0, L)$ for each $t \geq 0$, $\dot{u} = 0$ if and only if $\bar{\theta}(u(t), t) = \theta_c$, and because $v(s) > 0$ if $0 \leq s < \theta_c$ (cf. (3.56)). Indeed

$$\varepsilon \geq \int_{t_0}^{\infty} \dot{u}(\tau) d\tau \geq \min\left\{v(s) : s \in \left[0, \frac{3\theta_c}{4}\right]\right\} \mathcal{L}\left(\left\{t \in [t_0, \infty) : \bar{\theta}(u(t), t) \in \left[0, \frac{3\theta_c}{4}\right]\right\}\right),$$

that is $\bar{\theta}(u(t), t) \notin \left[\frac{3\theta_c}{4}, \theta_c\right]$ only for a finite amount of time. Thus, as we are assuming that the interface position stays in $(0, L)$ for infinite time, such a value of τ exists. We obtain

$$y(\tau) \leq \varepsilon + \hat{c} \int_0^{\tau} |\dot{u}(t_0 + \hat{\tau})| d\hat{\tau} \leq (1 + \hat{c})\varepsilon. \quad (3.70)$$

On the other hand, there exists $\tilde{c} > 0$ such that for almost every $t \geq 0$ we have

$$|\bar{\theta}(u(t), t) - \bar{\theta}(x(t), t)| \leq \tilde{c}\|\bar{\theta}_s\|_2(t)|u(t) - x(t)|^{\frac{1}{2}}. \quad (3.71)$$

Proposition 3.7.2 also implies that $\bar{\theta} \in L^\infty(0, T; H_0^1(\Omega))$, and we can therefore assume without loss of generality that $\|\bar{\theta}_s\|_2(t_0 + \tau) \leq \hat{\alpha}$, for some constant $\hat{\alpha} > 0$. Let us now consider

$$\mathcal{D}_\tau := \left\{s \in (0, u(t_0 + \tau)) : \bar{\theta}(s, t_0 + \tau) = \frac{\theta_c}{4}\right\}.$$

This set is closed, and non-empty given the fact $\bar{\theta}(\cdot, t_0 + \tau)$ is continuous, that $\bar{\theta}(0, t_0 + \tau) = 0$ and that $\bar{\theta}(u(t_0 + \tau), t_0 + \tau) > \frac{3\theta_c}{4}$. Let us hence choose $x(\tau) = \max \mathcal{D}_\tau$ so that (3.71) becomes

$$\frac{\theta_c}{2} = |\bar{\theta}(u(t_0 + \tau), t_0 + \tau) - \bar{\theta}(x(\tau), t_0 + \tau)| \leq \tilde{c}\hat{\alpha}|u(t_0 + \tau) - x(\tau)|^{\frac{1}{2}}.$$

This yields

$$\|\bar{\theta}\|_2^2(t_0 + \tau) = \int_{\Omega} |\bar{\theta}(s, t_0 + \tau)|^2 ds \geq \frac{\theta_c^2}{16}|u(t_0 + \tau) - x(\tau)| \geq \frac{\theta_c^4}{64\tilde{c}\hat{\alpha}}.$$

Combining this inequality with (3.70), by choosing ε small enough, we are led to a contradiction. Therefore, for every $t \geq t^*$ (3.58) simplifies to a standard homogeneous heat equation with initial datum in $H_0^1(\Omega)$, for which convergence to the only equilibrium, namely 0, is a well-known result (see e.g., [39]). $\square$

Chapter 4

Analysis of a moving mask approximation for martensitic transformations

4.1 Introduction

The aim of this chapter is to study from a mathematical point of view the complex microstructures arising during the austenite to martensite phase transition in ultra-low hysteresis alloys such as $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$ (see [76]). As remarked in [76], it is intriguing and unusual that martensitic microstructures in $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$ are drastically different in consecutive thermally induced transformation cycles, this being partially motivated by the fact that the cofactor conditions are close to being satisfied by both some type I and some type II twins, which can all form zero energy interfaces with austenite.

The aim of this chapter is to further study these microstructures, and to identify a common characterization for all of them. To this end we start from the following observation: from the dynamical point of view, it looks as if in every thermally induced phase transformation cycle in $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$ there was a mask moving across the domain covering and uncovering martensite microstructures. This is equivalent to saying that the martensitic microstructures do not change after the phase transition has happened, which seems to be a particularly legitimate approximation in materials satisfying the cofactor conditions, where no interface layer is needed between phases. But this hypothesis makes sense also in many other materials as long as one considers macroscopic deformations.

As a first step, we frame the moving mask approximation in the context of nonlinear elasticity, where phase changes are interpreted as elastic deformations. In

particular, we start from a simplified model for the dynamics of martensitic transformations which was introduced in Chapter 3. In this context, one can describe the evolution of the phase interface as a moving shock wave. The model equations have to be closed with some constitutive relation for the phase boundary velocity. This should depend on the local temperature of the sample, the thermal boundary conditions, and on the local martensitic microstructures, and is therefore not easy to determine. Nonetheless, by deriving a new type of Hadamard jump condition which is based on the moving mask assumption, we prove that every martensitic microstructure arising from the evolution and satisfying this assumption must be of the form

$$\nabla \mathbf{y}(\mathbf{x}) = \mathbf{1} + \mathbf{a}(\mathbf{x}) \otimes \mathbf{n}(\mathbf{x}), \quad \text{a.e.}, \quad (4.1)$$

where $\mathbf{n}(\mathbf{x})$ is, up to a change of sign, the phase interface normal at $\mathbf{x}$ when the point $\mathbf{x}$ is on the interface. The above result does not follow directly from the assumption that the deformation gradient is unchanged after the phase transition, and is not a direct consequence of previously known Hadamard jump conditions if $\mathbf{y}$ is just Lipschitz as in our case. We refer the interested reader to Section 1.3 for a brief review on known versions of the Hadamard jump conditions and on why they do not imply (4.1). Subsequent experiments (see [27]) have measured $\|\text{cof}(\nabla \mathbf{y} - \mathbf{1})\|$ in a sample of $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$. The measured values are of the order of 10^{-4} , which seem to be small enough to partially confirm the validity of (4.1) and of the moving mask approximation.

Under suitable hypotheses, we prove also that $\nabla \cdot \mathbf{a} = 0$. This result relies on a Hadamard jump condition for strains in $BV(\Omega)$ that is proved in Section 4.5 and generalizes that in [37]. As a consequence of the fact that $\nabla \mathbf{y} = \mathbf{1} + \mathbf{a} \otimes \mathbf{n}$ almost everywhere in the martensitic microstructures, we can prove a rigidity theorem for compound twins, a rigidity theorem for general two wells problem, and a result that allows for a better understanding of the nature of curved austenite-martensite interfaces for type I twins satisfying the cofactor conditions.

The dynamics of the phase transition in $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$ are very complex, and far from being completely understood. As in the static case, a major obstacle towards a good understanding of the phenomenon remains the lack of a characterization of the quasiconvex hull of the set of possible deformation gradients. Nonetheless, our results provide an interesting set of tools that can be used to understand further the complex microstructures arising in martensitic phase transitions.

Further investigation on why martensitic microstructures are so different in different thermal cycles in $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$ is carried out in Chapter 2. Indeed, in Chapter

2 we show that this material satisfies some further conditions of compatibility that makes the set of possible macroscopic deformation gradients of the form (4.1) unusually large.

The plan for the chapter is the following: in Section 4.2 we recall the definition and the main properties of k -rectifiable sets. In Section 4.3 we recall results from Chapter 3 and in this context we introduce the mathematical definition of moving mask approximation. In Section 4.4 we prove a dynamic variant of the Hadamard jump condition for Lipschitz functions and curved interfaces. As explained above, the results rely on the hypothesis that the deformation gradient remains constant in time at a point of the domain, once the phase transition has occurred. The last two sections are devoted to proving some results on moving austenite-martensite interfaces and on possible microstructures that can be explained using our model.

4.2 Some preliminaries on k -rectifiable sets

In this section we recall some standard results on Lipschitz functions and k -rectifiable sets from [4, 41, 58] (see also [3] for properties of level sets of Lipschitz functions). We denote by $\mathcal{H}^k$ the k -dimensional Hausdorff measure, and write $\mathcal{H}^k \llcorner E$ for its restriction to an $\mathcal{H}^k$ measurable subset E . $C_c(\mathbb{R}^d)$ stands for the space of continuous functions with compact support in $\mathbb{R}^d$, while $B^d(\mathbf{x}, r)$ denotes the d -dimensional ball centred at x , with radius r and of volume $\omega_d r^d$. We start with the following definitions:

Definition 4.2.1. *A Lipschitz k -graph $\mathcal{G}$ is a set of points in $\mathbb{R}^d$ with $d > k$ such that there exists an open and connected set $\omega \subset \mathbb{R}^k$, a Lipschitz map $\psi: \omega \rightarrow \mathbb{R}^{d-k}$, and a rotation $Q \in SO(d)$ satisfying*

$$\mathcal{G} := \{Q\mathbf{x}, \mathbf{x} = (\mathbf{x}', \psi(\mathbf{x}')), \mathbf{x}' \in \omega\}.$$

Definition 4.2.2. *Let $E \subset \mathbb{R}^d$ be an $\mathcal{H}^k$ -measurable set satisfying $\mathcal{H}^k(E) < \infty$. We say that E is k -rectifiable if there exist countably many Lipschitz mappings $\mathbf{f}_i: \mathbb{R}^k \rightarrow \mathbb{R}^d$ such that*

$$\mathcal{H}^k\left(E \setminus \bigcup_{i=1}^{\infty} \mathbf{f}_i(\mathbb{R}^k)\right) = 0.$$

An equivalent characterization for such sets is given by the following result:

Proposition 4.2.1 ([4, Prop. 2.76]). *Any $\mathcal{H}^k$ -measurable set E is countably $\mathcal{H}^k$ -rectifiable if and only if there exist countably many Lipschitz k -graphs $\mathcal{G}_i \subset \mathbb{R}^N$, such that*

$$\mathcal{H}^k\left(E \setminus \bigcup_{i=1}^{\infty} \mathcal{G}_i\right) = 0.$$

In what follows, a particular case of [41, Theorem 3.2.22] is also used:

Theorem 4.2.1. *Let $\Omega \subset \mathbb{R}^3$ be open, bounded and connected, $\mathcal{Z} \subset \mathbb{R}^3$ be a 1-rectifiable set and $\mathbf{f}: \Omega \rightarrow \mathcal{Z}$ a Lipschitz function. Define the 1-dimensional Jacobian of $\mathbf{f}$ by:*

$$J_1 \mathbf{f} := \sqrt{\sum_{i,j} (\nabla \mathbf{f})_{ij}^2}.$$

Then:

- *for $\mathcal{L}^3$ almost every $\mathbf{x} \in \Omega$, either $J_1 \mathbf{f}(\mathbf{x}) = 0$, or the image of $\nabla \mathbf{f}(\mathbf{x})$ is a 1-dimensional vector space, i.e., $\text{rank } \nabla \mathbf{f}(\mathbf{x}) \leq 1$, $\mathcal{L}^3$ -almost everywhere in Ω ;*
- *for $\mathcal{H}^1$ almost all $\boldsymbol{\xi} \in \mathcal{Z}$, $\mathbf{f}^{-1}(\boldsymbol{\xi})$ is 2-rectifiable;*
- *for every integrable $g: \Omega \rightarrow [-\infty, \infty]$*

$$\int_{\Omega} g(\mathbf{x}) J_1 \mathbf{f}(\mathbf{x}) \, d\mathbf{x} = \int_{\mathcal{Z}} \int_{\mathbf{f}^{-1}(\boldsymbol{\xi}) \cap \Omega} g(\mathbf{s}) \, d\mathcal{H}^2(\mathbf{s}) \, d\mathcal{H}^1(\boldsymbol{\xi}) \quad (4.2)$$

4.3 The moving mask assumption

In Chapter 3 we introduced a continuum model for the evolution of martensitic transformations. After passing to the limit in which the elastic constants tend to infinity and the interface energy density tends to zero, we deduced that the deformation gradients and the temperature field generate in the limit a Young measure $\nu_{\mathbf{x},t}$ (see as a reference [62, 69]) and a function θ satisfying in a suitable sense

$$\begin{aligned} \rho_0 \theta_t - d \Delta \theta &= -\theta_T \frac{\partial}{\partial t} \int_{\mathbb{R}^{3 \times 3}} \eta_1(\mathbf{A}) \, d\nu_{\mathbf{x},t}(\mathbf{A}), & \text{a.e. in } \Omega \times (0, T), \\ \text{supp } \nu_{\mathbf{x},t} &\subset SO(3) \cup K, & \text{a.e. in } \Omega \times (0, T), \end{aligned}$$

complemented with some initial and boundary conditions. Here, $\rho_0 > 0$ is the density of the body, $d > 0$ is a diffusivity coefficient which is supposed to be constant, and η_1 is a smooth function such that

$$\eta_1(\mathbf{F}) = 0, \quad \text{for all } \mathbf{F} \in SO(3), \quad \eta_1(\mathbf{F}) = -\frac{\alpha}{\theta_T}, \quad \text{for all } \mathbf{F} \in K,$$

for some constant $\alpha > 0$ representing the latent heat of the transformation. This system of equations is underdetermined, and should therefore be closed with some constitutive relation between $\int_{\mathbb{R}^{3 \times 3}} \eta_1(\mathbf{A}) \, d\nu_{\mathbf{x},t}(\mathbf{A})$, and θ and $\nu_{\mathbf{x},t}$. Nonetheless, we aim to characterise solutions independently of the constitutive relation. In order to do this we introduce some hypotheses on the solutions, that are based on experimental observation and that together we call the *moving mask* approximation, defined precisely in Definition 4.3.1 below, using the following ingredients:

- the phases are separated, that is there exist open sets $\Omega_A(t), \Omega_M(t) \subset \Omega$ such that

$$\Omega_A(t) \cap \Omega_M(t) = \emptyset, \quad \mathcal{L}^3(\Omega \setminus (\Omega_A(t) \cup \Omega_M(t))) = 0, \quad \text{a.e. } t \in (0, T),$$

and

$$\begin{aligned} \nu_{\mathbf{x},t}(SO(3)) &= 1, & \text{a.e. } \mathbf{x} \in \Omega_A(t), \text{ a.e. } t \in (0, T), \\ \nu_{\mathbf{x},t}(K) &= 1, & \text{a.e. } \mathbf{x} \in \Omega_M(t), \text{ a.e. } t \in (0, T). \end{aligned}$$

The domain can hence be divided for almost every $t \in (0, T)$ into two regions, the region with martensite $\Omega_M(t)$, and the region with austenite $\Omega_A(t)$. Thus,

$$\int_{\mathbb{R}^{3 \times 3}} \eta_1(\mathbf{A}) \, d\nu_{\mathbf{x},t}(\mathbf{A}) = -\frac{\alpha}{\theta_T} \chi_{\Omega_M}(\mathbf{x}, t), \quad (4.3)$$

where $\chi_{\Omega_M}(\mathbf{x}, t)$ is the characteristic function of $\Omega_M(t)$. The austenite-martensite phase boundary is sharp in this case. In terms of macroscopic deformation gradients this reads

$$\begin{aligned} \nabla \mathbf{y}(\mathbf{x}, t) &\in K^{qc}, & \text{a.e. in } \Omega_M(t), \text{ a.e. } t \in (0, T), \\ \nabla \mathbf{y}(\mathbf{x}, t) &\in SO(3), & \text{a.e. in } \Omega_A(t), \text{ a.e. } t \in (0, T), \end{aligned}$$

- during the phase transition, the macroscopic deformation gradient remains equal to a constant rotation in the austenite region. This is the case, for example, when the austenite region is connected;
- the phase interface moves continuously. More precisely, for almost every point $\mathbf{x}$ in the domain, there exists a time when $\mathbf{x}$ is contained in the phase interface (see also Section 3.6);
- the microstructures do not change after the transformation has happened. This assumption makes particular sense in the context of materials satisfying the

cofactor conditions, where austenite and finely twinned martensite can be exactly compatible across interfaces, and even more in $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$ where the phase transition has very low thermal hysteresis and thermal expansion is hence negligible.

As remarked in the introduction, this construction reflects the idea of a moving mask that uncovers a martensitic microstructure, as can be seen in the video of [76]. Mathematically we can define the moving mask approximation as follows:

Definition 4.3.1. *We say that $\nabla \mathbf{y} \in L^\infty(\Omega; \mathbb{R}^{3 \times 3})$ satisfies the moving mask approximation if*

- *for each $t \in [0, T]$ there exist $\Omega_M(t), \Omega_A(t) \subset \Omega$ disjoint and open, such that*

$$\mathcal{L}^3(\Omega \setminus (\Omega_A(t) \cup \Omega_M(t))) = 0;$$

- *either*

$$\Omega_A(t_2) \subset \Omega_A(t_1), \quad \text{for all } 0 \leq t_1 \leq t_2 \leq T,$$

or

$$\Omega_A(t_1) \subset \Omega_A(t_2), \quad \text{for all } 0 \leq t_1 \leq t_2 \leq T;$$

- *for a.e. $\mathbf{x} \in \Omega$ there exists $t = t(\mathbf{x}) \in [0, T]$ such that $\mathbf{x} \in \overline{\Omega}_A(t) \cap \overline{\Omega}_M(t)$;*
- *there exists $\mathbf{Q} \in SO(3)$ such that for every $t \in [0, T]$ the map $\mathbf{y}_M(\cdot, t)$ satisfying*

$$\nabla \mathbf{y}_M(\mathbf{x}, t) = \begin{cases} \nabla \mathbf{y}(\mathbf{x}), & \text{a.e. in } \Omega_M(t) \\ \mathbf{Q}, & \text{a.e. in } \Omega_A(t) \end{cases}$$

is in $W^{1,\infty}(\Omega; \mathbb{R}^3)$.

Remark 4.3.1. We note that, in the case $\Omega_M(s) \subset \Omega_M(t)$ for each $s, t \in [0, T]$ with $s < t$, we have

$$\bigcap_{t \in [0, T]} \Omega_A(t) = \emptyset, \quad \bigcup_{t \in [0, T]} \Omega_M(t) = \Omega.$$

Remark 4.3.2. If we assume (4.3), then the formula for differentiation of integrals on time dependent domains implies that

$$\begin{aligned} \left\langle \frac{\partial}{\partial t} \int_{\mathbb{R}^{3 \times 3}} \eta_1(\mathbf{A}) \, d\nu_{\mathbf{x}, t}(\mathbf{A}), \psi \right\rangle &= \langle \dot{\chi}_{\Omega_M}(\nabla \mathbf{y}), \psi \rangle = \frac{d}{dt} \int_{\Omega_M} \psi \, d\mathbf{x} \\ &= \int_{\Gamma(t)} (\mathbf{v} \cdot \mathbf{n}) \psi \, d\mathcal{H}^2, \quad \forall \psi \in C_0^\infty(\Omega), \end{aligned} \tag{4.4}$$

provided $\Omega_A(t), \Omega_M(t)$ and $\mathbf{v} \cdot \mathbf{n}$ are smooth enough (see e.g., [42]). Here $\Gamma(t) := \Omega \setminus (\Omega_A \cup \Omega_M)(t)$ is a surface separating $\Omega_A(t)$ from $\Omega_M(t)$, $\mathbf{n}$ denotes the outer normal to Ω_M and $\mathbf{v}(\mathbf{s})$ is the velocity of the interface at the point $\mathbf{s} \in \Gamma(t)$ at time t . By $\langle \cdot, \cdot \rangle$ we denoted the duality pairing between a distribution and a test function. A version of (4.4) in a case of interest can be found in Corollary 4.4.4 below. It thus appears clear from (4.4) that the model needs to be closed with some constitutive relation for $\mathbf{v} \cdot \mathbf{n}$ in terms of $\nabla \mathbf{y}, \theta$, and possibly other quantities as the curvature of the phase interface.

4.4 Generalized Hadamard conditions

In this section we restrict our attention to deformation gradients $\nabla \mathbf{y}$ that satisfy the moving mask approximation as stated in Definition 4.3.1. In order to say something more about solutions under these assumptions, we prove below a variant of the Hadamard jump condition reflecting this hypothesis. In what follows, we restrict, without loss of generality, to the case $\Omega_M(s) \subset \Omega_M(t)$ for every $s < t$. As before, below $\Omega \subset \mathbb{R}^3$ is an open bounded connected set with Lipschitz boundary.

Proposition 4.4.1. *Let $\Gamma(t)$ be a family of parallel planes perpendicular to $\mathbf{n} \in \mathbb{S}^2$,*

$$\Gamma(t) := \{\mathbf{x} \in \mathbb{R}^3 : \mathbf{x} \cdot \mathbf{n} = h(t)\},$$

for some non-decreasing function $h \in C([0, T])$ satisfying

$$h(0) = \inf_{\mathbf{x} \in \Omega} \mathbf{x} \cdot \mathbf{n}, \quad h(T) = \sup_{\mathbf{x} \in \Omega} \mathbf{x} \cdot \mathbf{n}.$$

For $t \in [0, T]$ define

$$\Omega_M(t) := \Omega \cap \{\mathbf{x} \cdot \mathbf{n} < h(t)\}, \quad \Omega_A(t) := \Omega \cap \{\mathbf{x} \cdot \mathbf{n} > h(t)\}.$$

Let $\mathbf{z} \in W^{1,\infty}(\Omega; \mathbb{R}^3)$ be such that $\mathbf{Z} = \mathbf{Z}(\mathbf{x}, t)$ satisfying

$$\nabla \mathbf{Z}(\mathbf{x}, t) = \begin{cases} \nabla \mathbf{z}(\mathbf{x}), & \text{a.e. in } \Omega_M(t) \\ 0, & \text{a.e. in } \Omega_A(t) \end{cases} \quad (4.5)$$

is in $W^{1,\infty}(\Omega; \mathbb{R}^3)$ for a.e. $t \in (0, T)$. Then,

1. *there exists $\mathbf{a} \in L^\infty(\Omega; \mathbb{R}^3)$ such that*

$$\nabla \mathbf{z}(\mathbf{x}) = \mathbf{a}(\mathbf{x}) \otimes \mathbf{n}, \quad \text{a.e. } \mathbf{x} \in \Omega.$$

2. if $\Omega \cap \Gamma(t)$ is connected for every $t \in (0, T)$ then $\mathbf{z} = \mathbf{f}(\mathbf{x} \cdot \mathbf{n})$ for some $\mathbf{f} \in W^{1,\infty}((0, T); \mathbb{R}^3)$.

Proof. By rotating the system of coordinates we can assume without loss of generality that $\mathbf{n} = \mathbf{e}_3$. Let us consider the set $\mathcal{B}_1 \subset \Omega$ of points where $\mathbf{z}$ is differentiable, and the set $\mathcal{B}_2$ of points $\mathbf{x} \in \Omega$ such that there exists $t^* \in (0, T)$ for which $\mathbf{x} \in \Gamma(t^*)$ and $\mathbf{Z}(\cdot, t^*) \in W^{1,\infty}(\Omega; \mathbb{R}^3)$. By continuity of h we have that $\mathcal{L}^3(\Omega \setminus (\mathcal{B}_1 \cap \mathcal{B}_2)) = 0$. Let us thus consider a generic point $\mathbf{x} \in \mathcal{B}_1 \cap \mathcal{B}_2$, and notice that, since Ω is open, there exists $r > 0$ such that the ball $B(\mathbf{x}, r) \subset \Omega$. By (4.5), $\mathbf{Z}(\cdot, t^*)$ must be constant in each connected component of $\Omega_A(t)$. In particular, as $\mathbf{Z}(\cdot, t^*) \in W^{1,\infty}(\Omega, \mathbb{R}^3)$ is continuous, it must be a constant on $\Gamma(t^*) \cap B(\mathbf{x}, r)$. At the same time, continuity of $\mathbf{z}$ and $\mathbf{Z}(\cdot, t^*)$ implies also $\mathbf{z}(\mathbf{x}) = \mathbf{Z}(\mathbf{x}, t^*)$ for every $\mathbf{x} \in \Gamma(t^*) \cap B(\mathbf{x}, r)$. Therefore, the function $\mathbf{z}(\mathbf{x})$ must be constant on $\Gamma(t^*) \cap B(\mathbf{x}, r)$. This implies,

$$\frac{\partial \mathbf{z}}{\partial x_i}(\mathbf{x}) = 0, \quad i = 1, 2,$$

which is the first statement. On the other hand, if $\Omega \cap \Gamma(t)$ is connected for a.e. $t \in (0, T)$, then $\mathbf{z}(\mathbf{x})$ is constant on $\Gamma(t^*) \cap \Omega$ for a.e. $t \in (0, T)$ and hence $\mathbf{z} = \mathbf{z}(x_3)$. This concludes the proof. $\square$

Remark 4.4.1. We could replace the hypothesis concerning the connectedness of $\Omega_A(t)$ by assuming that $\mathbf{Z}(\mathbf{x}, t)$ is equal to a constant in $\Omega_A(t)$ for a.e. $t \in [0, T]$. Both these assumptions are automatically satisfied if Ω is convex.

Remark 4.4.2. In the case $\mathbf{z} = \mathbf{f}(\mathbf{x} \cdot \mathbf{n})$, and $\Gamma(t)$ is a single plane, the phase interfaces must coincide with the level sets of $\mathbf{f}$. Therefore, given an experimentally measured martensitic macroscopic deformation gradient, under the assumption that it satisfies the moving mask approximation, and is of the form $\mathbf{1} + \mathbf{a}(\mathbf{x} \cdot \mathbf{n}) \otimes \mathbf{n}$, for some $\mathbf{a} \in \mathbb{R}^3, \mathbf{n} \in \mathbb{S}^2$, we can reconstruct the position of austenite-martensite phase interfaces, by taking the level sets of $\mathbf{f}(\mathbf{x} \cdot \mathbf{n}) = \int_0^{\mathbf{x} \cdot \mathbf{n}} \mathbf{a}(s) ds$. Furthermore, in the case $\mathbf{z} = \mathbf{f}(\mathbf{x} \cdot \mathbf{n})$, and $\Gamma(t)$ is a single plane, the discontinuities in the macroscopic deformation gradient can occur only across the planes $\mathbf{x} \cdot \mathbf{n} = \text{constant}$. This is, for example, the case for type II twins satisfying the cofactor conditions, for which we refer the reader to Proposition 4.6.2.

Proposition 4.4.1 can be partially generalized to the case where $\Gamma(t)$ is a family of curved interfaces. As a first step, we need to introduce the concept of moving interfaces for our problem, generalizing the previous requirements on planar such interfaces.

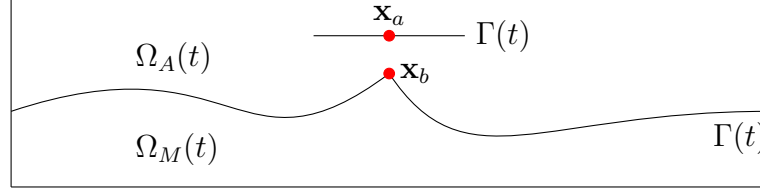


Figure 4.1: Points which are not regular: $\mathbf{x}_a$ is in a smooth k -graph contained in $\Gamma(t)$, but is not separating Ω_A from Ω_M . $\Gamma(t)$ does not coincide with a Lipschitz function differentiable at $\mathbf{x}_b$.

Definition 4.4.1. We say that $\Gamma(t) \subset \Omega$ is a family of moving interfaces in Ω if:

- (a) there exist two families of open disjoint sets $\Omega_M(t), \Omega_A(t) \subset \Omega$ and a bounded open interval $I_T := [0, T]$ such that for every t in I_T ,

$$\Omega = \Omega_M(t) \cup \Omega_A(t) \cup \Gamma(t) \quad \text{and} \quad \Gamma(t) \cap \Omega_M(t) = \Gamma(t) \cap \Omega_A(t) = \emptyset.$$

Furthermore, $\Omega_M(t)$ is non-decreasing in t , i.e.,

$$\Omega_M(t) \subset \Omega_M(s), \quad \Omega_A(s) \subset \Omega_A(t), \quad \forall t < s \in I_T;$$

- (b) the set

$$\mathcal{B} := \left\{ \mathbf{x} \in \Omega \left| \begin{array}{l} \mathbf{x} \in \Gamma(t^*), t^* \in I_T \text{ and there exist } \mathcal{U}_{\mathbf{x}} \subset \Omega \\ \text{open and connected, } \mathbf{x} \in \mathcal{U}_{\mathbf{x}}, \text{ and a Lipschitz} \\ \text{2-graph } \mathcal{G}_{\mathbf{x}} \text{ differentiable at } \mathbf{x} \text{ such that} \\ \mathcal{G}_{\mathbf{x}} \cap \mathcal{U}_{\mathbf{x}} \subset \Gamma(t^*) \cap \mathcal{U}_{\mathbf{x}} \subset \overline{\Omega}_M(t^*) \cap \overline{\Omega}_A(t^*) \cap \mathcal{U}_{\mathbf{x}} \end{array} \right. \right\}$$

is measurable and $\mathcal{L}^3(\Omega \setminus \mathcal{B}) = 0$.

Points in $\mathcal{B}$ are called regular points for $\Gamma(t)$.

At this point we can also introduce the concept of a regular moving mask approximation:

Definition 4.4.2. We say that $\mathbf{y} \in W^{1,\infty}(\Omega; \mathbb{R}^3)$ satisfies a regular moving mask approximation if it satisfies the moving mask approximation and

$$\Gamma(t) = \Omega \setminus (\Omega_A(t) \cup \Omega_M(t))$$

is a family of moving interfaces in Ω , where Ω_A, Ω_M are as in Definition 4.3.1.

Remark 4.4.3. The requirement $\Gamma(t^*) \cap \mathcal{U} \subset \overline{\Omega}_A(t^*) \cap \overline{\Omega}_M(t^*) \cap \mathcal{U}$ in Definition 4.4.1((b)) is mainly to guarantee that the set where an interface is cutting either Ω_A or Ω_M and not separating one from the other is small (see e.g., the point $\mathbf{x}_a$ in Figure 4.1). In this way, families of moving interfaces satisfying the separation condition may also describe further nucleations in the interior of Ω_A during the phase transition.

Below, we say that a curve $\mathbf{c}: [t_0, t_1] \rightarrow \mathbb{R}^3$, for some $t_0, t_1 \in \mathbb{R}$, is simple if $\mathbf{c}(s) \neq \mathbf{c}(t)$ for each $s, t \in [t_0, t_1]$. The following theorem generalizes Proposition 4.4.1 to curved interfaces:

Theorem 4.4.1. *Let $\Gamma(t)$ be a family of moving interfaces in Ω . Assume $\mathbf{z} \in W^{1,\infty}(\Omega; \mathbb{R}^3)$ is such that the function $\mathbf{Z} = \mathbf{Z}(\mathbf{x}, t)$ satisfying*

$$\begin{cases} \nabla \mathbf{z}(\mathbf{x}), & \text{a.e. in } \Omega_M(t) \\ 0, & \text{a.e. in } \Omega_A(t) \end{cases} \quad (4.6)$$

with $\Omega_A(t), \Omega_M(t)$ as in Definition 4.4.1, is in $W^{1,\infty}(\Omega; \mathbb{R}^3)$ for every $t \in I_T$. Then, there exist $\mathbf{a} \in L^\infty(\Omega; \mathbb{R}^3)$, $\mathbf{n} \in L^\infty(\Omega; \mathbb{S}^2)$ such that

$$\nabla \mathbf{z}(\mathbf{x}) = \mathbf{a}(\mathbf{x}) \otimes \mathbf{n}(\mathbf{x}), \quad \text{a.e. } \mathbf{x} \in \Omega. \quad (4.7)$$

Conversely, let $\mathbf{z} \in W^{1,\infty}(\Omega; \mathbb{R}^3)$ be such that (4.7) is satisfied and $\mathbf{z}(\Omega)$ is contained in the image of an absolutely continuous simple curve $\mathbf{c}: I_T \rightarrow \mathbb{R}^3$ of finite length. Then, if $|\mathbf{a}| > 0$ a.e. in Ω , there exists a family of moving interfaces in Ω , and a $\mathbf{Z} = \mathbf{Z}(\mathbf{x}, t)$ in $W^{1,\infty}(\Omega; \mathbb{R}^3)$ satisfying (4.6) for every $t \in I_T$. Furthermore, $\mathcal{L}^3(\Omega \setminus (\Omega_A(t) \cup \Omega_M(t))) = 0$ for every $t \in I_T$.

Remark 4.4.4. The assumptions on the image of $\mathbf{z}$ in Theorem 4.4.1 are motivated by the following observation: if $\mathbf{z} \in C^1(\Omega; \mathbb{R}^3)$, and $\nabla \mathbf{z}$ is rank-one everywhere in Ω , then the constant rank theorem implies that, around every $\mathbf{x} \in \Omega$, the image of $\mathbf{z}$ is a simple absolutely continuous curve. However, the set $\mathbf{z}(\Omega)$ can a priori show branching and other complex structures even in the regular case (e.g., if Ω is non-convex). For the sake of clarity of the proof, in this chapter we restrict ourselves to the easier case where $\mathbf{z}(\Omega)$ is a simple absolutely continuous curve. Nonetheless, it is possible to generalise the second implication in Theorem 4.4.1 to maps $\mathbf{z}: \Omega \rightarrow \mathbb{R}^3$ whose image satisfies

- $\mathbf{z}(\Omega)$ is 1-rectifiable;
- for $\mathcal{H}^1$ -a.e. $\boldsymbol{\xi} \in \mathbf{z}(\Omega)$ there exist an open ball $\mathcal{D}_{\boldsymbol{\xi}} \subset \mathbb{R}^3$ such that $\mathcal{D}_{\boldsymbol{\xi}} \cap \mathbf{z}(\Omega)$ is a simple curve of finite length which is absolutely continuous.

Remark 4.4.5. Assume that the moving mask assumption holds, and that we can reconstruct $\mathbf{z}$ from experimental observations. Then, provided the image of $\mathbf{z}$ satisfies the stated assumptions, the second part of Theorem 4.4.1 gives a useful tool to reconstruct phase interfaces during the phase transformation.

Lemma 4.4.1. *Let $\mathbf{f} \in L^\infty(\Omega; \mathbb{R}^{3 \times 3})$ be such that $\mathbf{f}(\mathbf{x}) = (\mathbf{b} \otimes \mathbf{m})(\mathbf{x})$ for a.e. $\mathbf{x} \in \Omega$. Then, there exist $\mathbf{a}, \mathbf{n} \in L^\infty(\Omega; \mathbb{R}^3)$ such that $\mathbf{f}(\mathbf{x}) = \mathbf{a}(\mathbf{x}) \otimes \mathbf{n}(\mathbf{x})$ and $|\mathbf{n}(\mathbf{x})| = 1$ for a.e. $\mathbf{x} \in \Omega$.*

Proof. This is just a matter of measurability of $\mathbf{a}, \mathbf{n}$. As $\mathbf{f}$ is measurable, so is $\mathbf{f}^T \mathbf{f} = (|\mathbf{b}|^2 \mathbf{m} \otimes \mathbf{m})$, so is its trace $\text{tr}(\mathbf{f}^T \mathbf{f}) = |\mathbf{b}|^2 |\mathbf{m}|^2$ and so is the function

$$\mathbf{g} := \begin{cases} |\mathbf{b}|^{-2} |\mathbf{m}|^{-2} \mathbf{f}^T \mathbf{f}, & \text{if } |\mathbf{b}|^2 |\mathbf{m}|^2 \neq 0, \\ 0, & \text{otherwise.} \end{cases}$$

Therefore, we define $\Omega_1 := \{\mathbf{x} \in \Omega : g_{11}(\mathbf{x}) \neq 0\}$ and

$$\mathbf{n}(\mathbf{x}) = \left((g_{11}(\mathbf{x}))^{\frac{1}{2}}, g_{12}(\mathbf{x})(g_{11}(\mathbf{x}))^{-\frac{1}{2}}, g_{13}(\mathbf{x})(g_{11}(\mathbf{x}))^{-\frac{1}{2}} \right)^T,$$

for almost every $\mathbf{x} \in \Omega_1$. This is actually possible because $g_{ii} \geq 0$ a.e. in Ω , for $i = 1, 2, 3$. Define also

$$\Omega_2 := \{\mathbf{x} \in \Omega \setminus \Omega_1 : g_{22} \neq 0\}, \quad \Omega_3 := \{\mathbf{x} \in \Omega \setminus (\Omega_1 \cup \Omega_2) : g_{33} \neq 0\},$$

and define $\mathbf{n}$ in Ω_2, Ω_3 respectively by

$$\begin{aligned} \mathbf{n} &= \left(g_{21}(g_{22})^{-\frac{1}{2}}, (g_{22})^{\frac{1}{2}}, g_{32}(g_{22})^{-\frac{1}{2}} \right)^T, \\ \mathbf{n} &= \left(g_{31}(g_{33})^{-\frac{1}{2}}, g_{32}(g_{33})^{-\frac{1}{2}}, (g_{33})^{\frac{1}{2}} \right)^T. \end{aligned}$$

Therefore, choosing $\mathbf{n}$ arbitrarily and such that $|\mathbf{n}| = 1$ in the set where $\mathbf{g} = 0$, we have constructed $\mathbf{n} \in L^\infty(\Omega; \mathbb{R}^3)$ as desired. Defining $\mathbf{a} := \mathbf{f}\mathbf{n}$, we thus conclude the proof. $\square$

Proof of Theorem 4.4.1. We first prove (4.7).

Let $N_{\mathbf{z}}$ be the set where $\mathbf{z}$ is not differentiable and remark that, by the hypotheses, $N := N_{\mathbf{z}} \cup (\Omega \setminus \mathcal{B})$ is an $\mathcal{L}^3$ -negligible set. Let $\mathbf{x}_0 \in \Omega \setminus N$ and take $\mathcal{U}_{\mathbf{x}_0}$ to be a neighbourhood of $\mathbf{x}_0$ as in Definition 4.4.1 ((b)). By taking a smaller connected neighbourhood of $\mathbf{x}_0$, which we still denote by $\mathcal{U}_{\mathbf{x}_0}$, we can assume that $\mathcal{G}_{\mathbf{x}_0} \cap \mathcal{U}_{\mathbf{x}_0}$ is connected. As $\nabla \mathbf{Z}(\mathbf{x}, t^*) = 0$ a.e. in $\Omega_A(t^*)$, the continuity of $\mathbf{Z}(\mathbf{x}, t^*)$ implies that $\mathbf{Z}(\mathbf{x}, t^*)$ must be constant on every connected component of $\overline{\Omega}_A(t^*)$. Since Definition 4.4.1 ((b)) implies $\mathcal{G}_{\mathbf{x}_0} \cap \mathcal{U}_{\mathbf{x}_0} \subset \overline{\Omega}_A(t^*)$, we must have $\mathbf{Z}(\cdot, t^*) = \hat{\mathbf{c}}$ for some $\hat{\mathbf{c}} \in \mathbb{R}^3$ on

$\mathcal{G}_{\mathbf{x}_0} \cap \mathcal{U}_{\mathbf{x}_0}$. On the other hand, continuity of $\mathbf{z}, \mathbf{Z}(\cdot, t^*)$ together with (4.6) imply that on every connected component of $\bar{\Omega}_M(t^*)$ $\mathbf{z} = \mathbf{Z}(\cdot, t^*) + \bar{\mathbf{c}}$ for some $\bar{\mathbf{c}} \in \mathbb{R}^3$ depending on the connected component. Therefore, as by Definition 4.4.1 ((b)) $\mathcal{G}_{\mathbf{x}_0} \cap \mathcal{U}_{\mathbf{x}_0} \subset \bar{\Omega}_M(t^*)$, the fact that $\mathbf{Z}(\cdot, t^*)$ is constant on $\mathcal{G}_{\mathbf{x}_0} \cap \mathcal{U}_{\mathbf{x}_0}$ implies that so must be $\mathbf{z}$.

Now, as $\mathcal{G}_{\mathbf{x}_0}$ is a Lipschitz 2-graph, we can find a Lipschitz change of coordinates $\psi: \mathcal{U}_{\mathbf{x}_0} \rightarrow V$ such that

$$\psi(\mathcal{G}_{\mathbf{x}_0} \cap \mathcal{U}_{\mathbf{x}_0}) = \{\mathbf{x} \in \mathbb{R}^3: \mathbf{x} \cdot \mathbf{n}(\mathbf{x}_0) = c_\Gamma\} \cap V,$$

for some $V \subset \mathbb{R}^3$, $c_\Gamma \in \mathbb{R}$ and where $\mathbf{n}(\mathbf{x}_0)$ is the normal vector to $\mathcal{G}_{\mathbf{x}_0}$ at $\mathbf{x}_0$ pointing outwards from $\Omega_M(t^*)$. Let us denote $\bar{\psi}(\mathbf{x}) = \bar{\mathbf{x}}$ for every $\mathbf{x} \in \mathcal{U}_{\mathbf{x}_0}$. We define $\bar{\mathbf{z}}$ as $\bar{\mathbf{z}}(\bar{\mathbf{x}}) = \mathbf{z}(\psi^{-1}(\bar{\mathbf{x}}))$, and assuming without loss of generality that $\mathbf{n}(\mathbf{x}) = \mathbf{e}_3$, we get that

$$\bar{\mathbf{z}}(\bar{\mathbf{x}}_0) = \bar{\mathbf{z}}(\bar{\mathbf{x}}_0 + s\mathbf{e}_i) = \hat{\mathbf{c}}, \quad i = 1, 2,$$

for each s such that $\bar{\mathbf{x}}_0 + s\mathbf{e}_i \in V$. This is due to the fact that $\mathbf{z}(\mathbf{x}) = \hat{\mathbf{c}}$ for every $\mathbf{x} \in \mathcal{U}_{\mathbf{x}_0} \cap \mathcal{G}_{\mathbf{x}_0}$. Therefore,

$$\frac{\partial \bar{\mathbf{z}}}{\partial \bar{x}_i}(\bar{\mathbf{x}}) = 0, \quad i = 1, 2.$$

On the other hand, as $\mathbf{x}_0 \in \Omega \setminus N$, we have

$$\nabla_{\mathbf{x}} \mathbf{z}(\mathbf{x}_0) = \nabla_{\bar{\mathbf{x}}} \bar{\mathbf{z}}(\psi(\mathbf{x}_0)) \nabla_{\mathbf{x}} \psi(\mathbf{x}_0).$$

Since $\mathbf{x}_0$ is a regular point, ψ can be chosen to be differentiable in $\mathbf{x}_0$, and therefore $\bar{\mathbf{x}}_0$ is a point of differentiability for $\bar{\mathbf{z}}$. Therefore, there exists $\mathbf{a} \in \mathbb{R}^3$ such that

$$\nabla_{\bar{\mathbf{x}}} \bar{\mathbf{z}}(\bar{\mathbf{x}}_0) = \mathbf{a} \otimes \mathbf{n}(\mathbf{x}_0).$$

By putting together the last two identities and using the fact that N is negligible we finally deduce (4.7). Measurability of $\mathbf{a}, \mathbf{n}$ follows from Lemma 4.4.1.

We now prove the second statement. We first remark that since $\mathbf{c}$ is absolutely continuous and of finite length, it belongs also to $W^{1,1}(I_T; \mathbb{R}^3)$ and there exists $\mathbf{d} \in W^{1,\infty}(I_T^*; \mathbb{R}^3)$ for some $I_T^* \subset \mathbb{R}$ such that $\mathbf{c}(I_T) = \mathbf{d}(I_T^*)$ (see e.g., [3] and references therein). Therefore $\mathbf{c}(I_T)$ is 1-rectifiable. By Theorem 4.2.1, $\Gamma(t) := \mathbf{z}^{-1}(\mathbf{c}(t)) \cap \Omega$ are 2-rectifiable surfaces for almost every t , and $\mathbf{z}$ is equal to a constant on them. Defining

$$\begin{aligned} \Omega_M(t) &:= \{\mathbf{x} \in \Omega: \exists s \in [0, t) \text{ such that } \mathbf{x} \in \mathbf{z}^{-1}(\mathbf{c}(s))\}, \\ \Omega_A(t) &:= \{\mathbf{x} \in \Omega: \mathbf{x} \notin \mathbf{z}^{-1}(\mathbf{c}(s)), \forall s \in [0, t]\}, \end{aligned} \tag{4.8}$$

it is easy to see that Definition 4.4.1((a)) is satisfied. We only need to check that $\Omega_A(t), \Omega_M(t)$ are open. To this end, let us fix $\hat{\mathbf{x}} \in \Omega_M(t)$, and let us denote by $s_{\hat{\mathbf{x}}} \in [0, t)$ the point such that $\mathbf{z}(\hat{\mathbf{x}}) = \mathbf{c}(s_{\hat{\mathbf{x}}})$. As $\mathbf{z}$ is Lipschitz, we can define $R := \frac{1}{2} \|\nabla \mathbf{z}\|_{L^\infty}^{-1} |\mathbf{c}(t) - \mathbf{z}(\hat{\mathbf{x}})|$, so that

$$|\mathbf{z}(\mathbf{x}) - \mathbf{z}(\hat{\mathbf{x}})| \leq \|\nabla \mathbf{z}\|_{L^\infty} |\mathbf{x} - \hat{\mathbf{x}}| \leq \frac{1}{2} |\mathbf{c}(t) - \mathbf{z}(\hat{\mathbf{x}})| = \frac{1}{2} |\mathbf{c}(t) - \mathbf{c}(s_{\hat{\mathbf{x}}})|, \quad (4.9)$$

for all $\mathbf{x} \in B_R(\hat{\mathbf{x}})$. Suppose now that in $B_R(\hat{\mathbf{x}})$ there exists a point $\mathbf{x}_0$ such that $\mathbf{z}(\mathbf{x}_0) = \mathbf{c}(t_0)$ for some $t_0 \geq t$. Then the segment connecting $\mathbf{x}_0$ to $\hat{\mathbf{x}}$ is still contained in $B_R(\hat{\mathbf{x}})$ and its image through $\mathbf{z}$ must be a connected part of the image of $\mathbf{c}$. But as $\mathbf{c}$ is a simple curve, this implies that there exists $\mathbf{x}_1 \in B_R(\hat{\mathbf{x}})$ such that $\mathbf{z}(\mathbf{x}_1) = \mathbf{c}(t)$, which is in contradiction with (4.9). Therefore, for every $\hat{\mathbf{x}} \in \Omega_M(t)$ there exists an open ball centred at $\hat{\mathbf{x}}$ contained in $\Omega_M(t)$, and therefore $\Omega_M(t)$ is open. The same argument can be used to show that also $\Omega_A(t)$ is open for each t . Clearly, $\Gamma(t)$ is sequentially closed in Ω and $\Omega_M(t) \cup \Gamma(t), \Omega_A(t) \cup \Gamma(t)$ are closed in Ω as well. In this way we have also shown that $\mathbf{Z}(x, t)$ defined as

$$\mathbf{Z}(x, t) = \begin{cases} \mathbf{z}(x), & \text{in } \Omega_M(t) \\ \mathbf{c}(t), & \text{in } \Omega_A(t) \end{cases}$$

is in $W^{1,\infty}(\Omega; \mathbb{R}^3)$ for every $t \in I_T$.

Since $\Gamma(t)$ is 2-rectifiable for almost every t , in order to show that

$$\mathcal{L}^3(\Omega \setminus (\Omega_A(t) \cup \Omega_M(t))) = 0$$

for every $t \in I_T$ we just need to prove that

$$\mathcal{C} := \{\mathbf{x} \in \Omega : \mathbf{x} \in \Gamma(t), t \in I_T, \Gamma(t) \text{ is not 2-rectifiable}\}$$

has null $\mathcal{L}^3$ measure. By Theorem 4.2.1, $\mathcal{C}$ is the preimage through $\mathbf{z}$, which is continuous, of a set of measure zero, and is hence measurable. Now, we notice that by choosing g to be the indicator function on $\mathcal{C}$ in the coarea formula (4.2), and identifying $\mathcal{Z}$ with the support of $\mathbf{c}$, we have

$$0 \leq \int_{\Omega} g(\mathbf{x}) |\mathbf{a}| \, d\mathbf{x} = \int_{\mathcal{Z}} \int_{\mathbf{z}^{-1}(\boldsymbol{\xi})} g(\mathbf{s}) \, d\mathcal{H}^2(\mathbf{s}) \, d\mathcal{H}^1(\boldsymbol{\xi}) = 0 \quad (4.10)$$

as, by Theorem 4.2.1, this can just happen for a set of measure zero in $\mathcal{Z}$. This leads to $\mathcal{L}^3(\mathcal{C}) = 0$.

Now we claim that for every point $\mathbf{x} \in \mathcal{D}$, with

$$\mathcal{D} := \{\mathbf{x} \in \Omega: \nabla \mathbf{z}(\mathbf{x}) \text{ exists, and } \nabla \mathbf{z}(\mathbf{x}) \neq \mathbf{0}\},$$

there exist a Lipschitz 2-graph $\mathcal{G}_{\mathbf{x}}$ which is differentiable at $\mathbf{x}$, an open neighbourhood $\mathcal{U}_{\mathbf{x}}$ and a $t^* \in I_T$ satisfying $\mathcal{G}_{\mathbf{x}} \cap \mathcal{U}_{\mathbf{x}} \subset \Gamma(t^*) \cap \mathcal{U}_{\mathbf{x}}$. In order to do that, we would need a generalised version of the constant rank theorem. However we were not able to find a version of it in the literature suitable to our application. We hence strongly exploit the structure of the image of $\mathbf{z}$ and a weak version of the implicit function theorem. Let us consider a generic $\hat{\mathbf{x}} \in \mathcal{D}$ and suppose, without loss of generality, that $a_1(\hat{\mathbf{x}}) \neq 0$ and $\mathbf{n}(\hat{\mathbf{x}}) = \mathbf{e}_3$. In this case, a version of the implicit function theorem as the one in [46, Thm. E] gives the existence of a connected neighbourhood $\mathcal{N}$ of $(\hat{x}_1, \hat{x}_2)$, and of a function $\psi: \mathcal{N} \rightarrow \mathbb{R}$, such that $\psi(\hat{x}_1, \hat{x}_2) = \hat{x}_3$, and $z_1(x_1, x_2, \psi(x_1, x_2)) = z_1(\hat{\mathbf{x}})$ for every $(x_1, x_2) \in \mathcal{N}$. Furthermore ψ is differentiable in $(\hat{x}_1, \hat{x}_2)$ and hence continuous and Lipschitz in $\mathcal{N}$, and $\nabla \psi(\hat{x}_1, \hat{x}_2) = \mathbf{0}$.

Fixed $\varepsilon = \frac{|a_1(\hat{\mathbf{x}})|}{2}$, the fact that $\mathbf{z}$ is differentiable in $\hat{\mathbf{x}}$ implies the existence of $\delta > 0$ such that

$$z_1(\hat{\mathbf{x}} + \rho \mathbf{e}_3) - z_1(\hat{\mathbf{x}}) = \rho a_1 + r\rho, \quad \forall |\rho| < \delta,$$

and where $|r| < \varepsilon$. Therefore,

$$\begin{aligned} z_1(\hat{\mathbf{x}} + \rho \mathbf{e}_3) &> z_1(\hat{\mathbf{x}}), & \text{if } a_1(\hat{\mathbf{x}})\delta > a_1(\hat{\mathbf{x}})\rho > 0, \\ z_1(\hat{\mathbf{x}} + \rho \mathbf{e}_3) &< z_1(\hat{\mathbf{x}}), & \text{if } -a_1(\hat{\mathbf{x}})\delta < a_1(\hat{\mathbf{x}})\rho < 0, \end{aligned} \tag{4.11}$$

for all $|\rho| < \delta$. This implies the existence of $h > 0$ and $\mathbf{c}(t^* + h), \mathbf{c}(t^* - h)$ in $\mathbf{z}(\Omega)$ such that

$$c_1(t^* + h) > c_1(t^*) > c_1(t^* - h).$$

Here, $t^* \in I_T$ is such that $\mathbf{z}(\hat{\mathbf{x}}) = \mathbf{c}(t^*)$. Furthermore, since $\mathbf{z}$ is Lipschitz, the dependence of $c(t(\mathbf{x})) := \mathbf{z}(\mathbf{x})$ is continuous. This together with (4.11) and the fact that $\mathbf{c}$ is simple, imply that the unique path connecting $\mathbf{c}(t^* \pm h)$ to $\mathbf{c}(t^*)$ must be such that $c_1(t^* + s) > c_1(t^*) > c_1(t^* - s)$ either for every $s \in (0, h)$ or for every $s \in (-h, 0)$. Suppose now the existence of $(x_1, x_2) \in \mathcal{N}$ such that $\mathbf{z}(x_1, x_2, \psi(x_1, x_2)) \neq \mathbf{c}(t^*)$. By continuity of $c(t(\mathbf{x}))$ there exist $(\tilde{x}_1, \tilde{x}_2) \in \mathcal{N}$ such that $\mathbf{z}(\tilde{x}_1, \tilde{x}_2, \psi(\tilde{x}_1, \tilde{x}_2)) = \mathbf{c}(t^* + s)$ for some s with $|s| < h$. Thus, at the same time we should have $c_1(t^* + s) = c_1(t^*)$ because we are on a level set for z_1 , and $c_1(t^* + s) \neq c_1(t^*)$, which leads to a contradiction. We hence showed that c_1 is constant implies also that c_2, c_3 are constants, that is $\mathbf{z}(x_1, x_2, \psi(x_1, x_2)) = \mathbf{z}(\hat{\mathbf{x}}) = \mathbf{c}(t^*)$ for every $(x_1, x_2) \in \mathcal{N}$. This concludes the proof of the claim.

It remains to prove that for all $\mathbf{x} \in \mathcal{D}$, it holds $\Gamma(t^*) \cap \mathcal{U}_{\mathbf{x}} \subset \overline{\Omega}_L(t^*) \cap \overline{\Omega}_R(t^*) \cap \mathcal{U}_{\mathbf{x}}$, where again $t^* \in I_T$ is such that $\mathbf{x} \in \Gamma(t^*)$. Suppose first that there exists a neighbourhood $\mathcal{U}$ of $\mathbf{x}_0 \in \Gamma(t^*)$ for some $t^* \in I_T$ such that

$$\Omega_A(t^*) \cap \mathcal{U} = \emptyset \quad \text{or} \quad \Omega_M(t^*) \cap \mathcal{U} = \emptyset. \quad (4.12)$$

This is $\mathbf{z}(\mathbf{x}_0) = \mathbf{c}(t^*)$, and $\mathbf{z}(\mathbf{x}) = \mathbf{c}(s(\mathbf{x}))$ with $s(\mathbf{x}) > t^*$ or $s(\mathbf{x}) < t^*$ for every $\mathbf{x} \in \mathcal{U}$. We want to prove that either $\mathbf{z}$ is not differentiable in $\mathbf{x}_0$, or $\nabla \mathbf{z}(\mathbf{x}_0) = \mathbf{0}$. Suppose not, then there exists $\beta_j \neq 0$ and a unit vector $\mathbf{v}_j$ such that $\nabla \mathbf{z}_j(\mathbf{x}_0) \cdot \mathbf{v}_j = \beta_j$ for some $j = 1, 2, 3$. Observe also that the differentiability of $\mathbf{z}$ implies the existence of $\delta_j \in (0, 1)$ such that

$$\mathbf{z}_j(\mathbf{x}_0 + \alpha \mathbf{v}_j) - \mathbf{z}_j(\mathbf{x}_0) - \alpha \beta_j = \alpha r_j(\alpha \mathbf{v}_j), \quad \forall \alpha: |\alpha| < \delta_j, \quad (4.13)$$

for some continuous functions r_j bounded in modulus by $\frac{\beta_j}{2}$. This implies that $\mathbf{z}_j(\mathbf{x}_0 + \alpha \mathbf{v}_j) - \mathbf{z}_j(\mathbf{x}_0)$ has the same sign as $\alpha \beta_j$. Therefore, as $\mathbf{c}$ is simple, there exists an interval (t_{δ_j}, t^*) (or (t^*, t_{δ_j})) where $c_j(t) - c_j(t^*)$ is both strictly positive and strictly negative for every $t \in (t_{\delta_j}, t^*)$ (or in (t^*, t_{δ_j})), thus leading to a contradiction. Therefore, if $\mathbf{z}$ is differentiable at $\mathbf{x}_0 \in \Gamma(t^*)$ and $\nabla \mathbf{z}(\mathbf{x}_0) \neq \mathbf{0}$, then $\mathbf{x}_0 \in \overline{\Omega}_A(t^*) \cap \overline{\Omega}_M(t^*)$.

By the coarea formula (4.2) with g chosen to be the characteristic function of $\mathcal{D}$, we notice that

$$0 = \int_{\mathcal{Z}} \int_{f^{-1}(\boldsymbol{\xi}) \cap \Omega} g(\mathbf{s}) \, d\mathcal{H}^2(\mathbf{s}) \, d\mathcal{H}^1(\boldsymbol{\xi}),$$

and we deduce that $\mathbf{z}$ is differentiable with $\nabla \mathbf{z} \neq \mathbf{0}$ for $\mathcal{H}^2$ -almost every $\mathbf{x} \in \Gamma(t)$ for almost every $t \in I_T$. Let us call J_T the subset of I_T such that $\mathbf{x} \in \mathcal{D}$ $\mathcal{H}^2$ -almost everywhere in $\Gamma(t)$ for every $t \in J_T$. By arguing as above for the set $\mathcal{C}$, the coarea formula implies that the set of $\mathbf{x} \in \Omega$ such that $\mathbf{z}(\mathbf{x}) \in \mathbf{c}(I_T \setminus J_T)$ has measure zero. We can hence focus without loss of generality on $\Gamma(t^*)$ for some $t^* \in J_T$. Suppose now that there does not exist a neighbourhood of $\mathbf{x}_s \in \Gamma(t^*)$ such that $\Gamma(t^*) \cap \mathcal{U} \subset \overline{\Omega}_A(t^*) \cap \overline{\Omega}_M(t^*)$. In this case, as $\Gamma(t^*)$ is closed in Ω , there exists a neighbourhood $\mathcal{U}_s$ of $\mathbf{x}_s$ satisfying (4.12). However, as $\nabla \mathbf{z}$ exists and is non null $\mathcal{H}^2$ almost everywhere on $\Gamma(t^*)$, there exists $\mathbf{x}_a \in \Gamma(t^*) \cap \mathcal{U}_s$ satisfying this properties, and which must hence be in $\overline{\Omega}_A(t^*) \cap \overline{\Omega}_M(t^*)$, thus leading to a contradiction. We have therefore proved that the constructed family of moving interfaces $\Gamma(t)$ satisfies the condition in Definition 4.4.1 ((b)), which concludes the proof. $\square$

The following corollaries are straightforward consequences of the above theorem:

Corollary 4.4.1. *Let $\mathbf{y} \in W^{1,\infty}(\Omega; \mathbb{R}^3)$ satisfy a regular moving mask approximation. Then, there exist $\mathbf{a} \in L^\infty(\Omega; \mathbb{R}^3)$, $\mathbf{n} \in L^\infty(\Omega; \mathbb{S}^2)$ such that*

$$\nabla \mathbf{y} = \mathbf{Q} + \mathbf{a}(\mathbf{x}) \otimes \mathbf{n}(\mathbf{x}), \quad \text{a.e. } \mathbf{x} \in \Omega. \quad (4.14)$$

Conversely, if $\mathbf{y} \in W^{1,\infty}(\Omega; \mathbb{R}^3)$ satisfies (4.14) for some $\mathbf{Q} \in SO(3)$, $\mathbf{a} \in L^\infty(\Omega; \mathbb{R}^3)$, $\mathbf{n} \in L^\infty(\Omega; \mathbb{S}^2)$ and $\mathbf{z}(\Omega)$, with $\mathbf{z}(\mathbf{x}) = \mathbf{y}(\mathbf{x}) - \mathbf{Q}\mathbf{x}$, is contained in an absolutely continuous simple curve of finite length, then $\mathbf{y}$ satisfies a regular moving mask approximation.

Corollary 4.4.2. *Let $T > 0$, $\Gamma(t)$ be a family of moving interfaces and $\bar{\Omega}_A(t)$, $\bar{\Omega}_M(t)$ be connected for every $t \in (0, T)$. Then, $\mathbf{Z} \in W^{1,\infty}(\Omega; \mathbb{R}^3)$ for each $t \in (0, T)$ satisfying (4.6) is equal to*

$$\mathbf{Z}(x, t) = \begin{cases} \mathbf{z}(x) + \mathbf{c}_1(t), & \text{in } \Omega_M(t), \\ \mathbf{c}_2(t), & \text{in } \Omega_A(t), \end{cases}$$

for some $\mathbf{z} \in W^{1,\infty}(\Omega; \mathbb{R}^3)$ such that $\nabla \mathbf{z} = \mathbf{a} \otimes \mathbf{n}$ almost everywhere.

Proof. The statement follows directly from Theorem 4.4.1, the fact that $Z(x, t)$ is continuous in $\mathbf{x}$ for each t and the hypothesis that $\bar{\Omega}_A(t)$, $\bar{\Omega}_M(t)$ are connected. $\square$

Corollary 4.4.2 hence implies that, under the above hypotheses, for each $t \in I_T$, the interface $\Gamma(t)$ is a subset of $\{\mathbf{x} \in \Omega: \mathbf{z}(\mathbf{x}) = \mathbf{c}_2(t) - \mathbf{c}_1(t)\}$, that is of a level set of $\mathbf{z}$. This also means that the image of $\mathbf{z}$ is a one-dimensional curve. However $\Gamma(t)$ does not need to coincide with the family of moving interfaces constructed in the proof of Theorem 4.4.1, even if $\mathbf{c}_2 - \mathbf{c}_1$ is absolutely continuous, simple and of finite length. Indeed, in the proof of Theorem 4.4.1 the phase interfaces must be constructed as subsets of level sets for $\mathbf{z}$, but the construction of $\Omega_A(t)$, $\Omega_M(t)$ is arbitrary and could be done differently. For example one could swap the definition of $\Omega_A(t)$ and $\Omega_M(t)$ in (4.8), or replace $s \in [0, t)$ and $s \in [0, t]$ in the definition of Ω_M and Ω_A respectively with $|s - t_0| < t$ and $s \in I_T \setminus [t_0 - t, t_0 + t]$ for some $t_0 \in I_T$, thus getting a different family of moving interfaces.

The next corollary gives some information about the interface velocity. We define the normal velocity of $\Gamma(t^*)$ at time $t^* \in I_T$ and at $\mathbf{x} \in \Gamma(t^*)$, namely $(\mathbf{v} \cdot \mathbf{n})(\mathbf{x}, t^*)$, as $\dot{\gamma}(t^*) \cdot \mathbf{n}(\mathbf{x}, t^*)$, where $\mathbf{n}(\mathbf{x}, t^*)$ is the unit normal to $\Gamma(t^*)$ at $\mathbf{x} \in \Gamma(t^*)$, and $\gamma(t)$ is a generic absolutely continuous path differentiable at t^* such that $\gamma(t) \in \Gamma(t)$ for each $t \in I_T$ and $\gamma(t^*) = \mathbf{x}$. Clearly $(\mathbf{v} \cdot \mathbf{n})(\mathbf{x}, t^*)$ is well defined if its value is independent of the choice of γ among the admissible paths, and if $\mathbf{n}(\mathbf{x}, t^*)$ is well defined.

Corollary 4.4.3. *Let $\mathbf{z} \in W^{1,\infty}(\Omega, \mathbb{R}^3)$ satisfy (4.7), $|a| > 0$ a.e. in Ω , and $\mathbf{z}(\Omega)$ be contained in the image of an absolutely continuous simple curve $\mathbf{c} : I_T \rightarrow \mathbb{R}^3$ of finite length. Assume further that $\Gamma(t)$ is the family of moving interfaces constructed in the proof of Theorem 4.4.1. Then, the normal velocity of $\Gamma(t)$ at a point $\mathbf{x} \in \Gamma(t)$, denoted $(\mathbf{v} \cdot \mathbf{n})(\mathbf{x}, t)$, satisfies*

$$\mathbf{a}(\mathbf{x})(\mathbf{v} \cdot \mathbf{n})(\mathbf{x}, t) = \dot{\mathbf{c}}(t), \quad \text{a.e. } t \in (0, T), \mathcal{H}^2\text{-a.e. } \mathbf{x} \in \Gamma(t). \quad (4.15)$$

Proof. By the coarea formula (4.2) with g chosen to be the characteristic function of the set where z is not differentiable and $|a| > 0$, we notice that

$$0 = \int_0^T \int_{\mathbf{z}^{-1}(\mathbf{c}(t)) \cap \Omega} g(\mathbf{s}) d\mathcal{H}^2(\mathbf{s}) |\dot{\mathbf{c}}(t)| dt.$$

As the argument in the integral is non negative, we deduce that $\mathbf{z}$ is differentiable and $|a| > 0$ for $\mathcal{H}^2$ -almost every $\mathbf{x} \in \Gamma(t)$ for almost every $t \in I_T$. As showed in the proof of Theorem 4.4.1 $\mathbf{n}$ is well defined for all these $\mathbf{x}$. Let us consider $\gamma(t)$, an absolutely continuous path in Ω such that $\gamma(t) \in \Gamma(t)$ for every $t \in I_T$. We have

$$\mathbf{z}(\gamma(t)) = \mathbf{c}(t), \quad \forall t \in (t_0, t_1),$$

for some $0 \leq t_0 < t_1 \leq T$. Taking the time derivative of this identity we get

$$\dot{\mathbf{c}}(t) = \nabla \mathbf{z}(\gamma(t)) \dot{\gamma}(t) = \mathbf{a}(\gamma(t))(\dot{\gamma}(t) \cdot \mathbf{n}(\gamma(t), t)) = \mathbf{a}(\mathbf{x})(\dot{\gamma}(t) \cdot \mathbf{n}(\mathbf{x}, t)),$$

which is the claimed result, as $(\dot{\gamma}(t) \cdot \mathbf{n}(\mathbf{x}, t))$ is independent of γ chosen for a.e. $t \in I_T$, a.e. $\mathbf{x} \in \Gamma(t)$. $\square$

Remark 4.4.6. An important consequence of the above corollary is that $\frac{\mathbf{a}(\mathbf{x})}{|\mathbf{a}(\mathbf{x})|}$ is constant $\mathcal{H}^2$ almost everywhere on $\Gamma(t)$ for almost every t . At the same time, there might be jumps in $|\mathbf{a}(\mathbf{x})|$ along a single interface and jumps for $\frac{\mathbf{a}(\mathbf{x})}{|\mathbf{a}(\mathbf{x})|}$ across interfaces.

Remark 4.4.7. If we assume the determinant of $\nabla \mathbf{y} = \mathbf{1} + \nabla \mathbf{z} = \mathbf{1} + \mathbf{a} \otimes \mathbf{n}$ to be a positive constant $\mathfrak{D}$ almost everywhere in Ω , than we can deduce that on almost all interfaces $|\mathbf{a}(\mathbf{x})|$ can jump if and only if there is a jump in $\mathbf{n}(\mathbf{x})$. Indeed, this is a direct consequence of the following two facts: the first is that the direction of $\mathbf{a}(\mathbf{x})$ is fixed on almost all interfaces, the second is that, $\det(\nabla \mathbf{y}) = \mathfrak{D}$ a.e. in Ω implies $\mathbf{a} \cdot \mathbf{n} = \mathfrak{D} - 1$ a.e. in Ω , and hence, by the coarea formula (see the argument in the proof of Corollary 4.4.3), $\mathbf{a} \cdot \mathbf{n} = \mathfrak{D} - 1$, $\mathcal{H}^2$ -almost everywhere on $\Gamma(t)$ for almost all t .

A different perspective on the velocity of $\Gamma(t)$ is given by

Corollary 4.4.4. *Let $\mathbf{z} \in W^{1,\infty}(\Omega, \mathbb{R}^3)$ satisfy (4.7), $|\mathbf{a}| > 0$ a.e. in Ω , and $\mathbf{z}(\Omega)$ be contained in the image of an absolutely continuous simple curve $\mathbf{c} : I_T \rightarrow \mathbb{R}^3$ of finite length. Assume further that $\Gamma(t)$ is the family of moving interfaces constructed in the proof of Theorem 4.4.1. Then,*

$$\langle \dot{\chi}_{\Omega_M}, \xi \rangle = |\dot{\mathbf{c}}(t)| \int_{\Gamma(t)} \frac{\xi(\mathbf{s})}{|\mathbf{a}(\mathbf{s})|} d\mathcal{H}^2(\mathbf{s}), \quad \forall \xi \in C^0(\Omega), \text{ a.e. } t \in (0, T).$$

Proof. We first notice that, by the coarea formula,

$$\begin{aligned} \int_{\Omega} \chi_{\Omega_M}(\mathbf{x}, t) \xi(\mathbf{x}) d\mathbf{x} &= \int_0^T \int_{\mathbf{z}^{-1}(\mathbf{c}(\tau)) \cap \Omega_A(t)} \frac{\xi(\mathbf{s})}{|\mathbf{a}(\mathbf{s})|} d\mathcal{H}^2(\mathbf{s}) |\dot{\mathbf{c}}(\tau)| d\tau \\ &= \int_0^t \int_{\mathbf{z}^{-1}(\mathbf{c}(\tau)) \cap \Omega} \frac{\xi(\mathbf{s})}{|\mathbf{a}(\mathbf{s})|} d\mathcal{H}^2(\mathbf{s}) |\dot{\mathbf{c}}(\tau)| d\tau. \end{aligned}$$

Therefore,

$$\frac{d}{dt} \int_{\Omega} \chi_{\Omega_M}(\mathbf{x}, t) \xi(\mathbf{x}) d\mathbf{x} = |\dot{\mathbf{c}}(t)| \int_{\mathbf{z}^{-1}(\mathbf{c}(t)) \cap \Omega} \frac{\xi(\mathbf{s})}{|\mathbf{a}(\mathbf{s})|} d\mathcal{H}^2(\mathbf{s}).$$

which is the claimed result. □

4.5 Basic properties of microstructures

According to the results of the previous sections, we can restrict our attention to deformation gradients satisfying for every $t \in (0, T)$

$$\begin{cases} \nabla \mathbf{y}(\mathbf{x}, t) = \mathbf{1} + \mathbf{a}(\mathbf{x}) \otimes \mathbf{n}(\mathbf{x}), & \text{a.e. } \mathbf{x} \in \Omega_M(t), \\ \nabla \mathbf{y}(\mathbf{x}, t) = \mathbf{1}, & \text{a.e. } \mathbf{x} \in \Omega_A(t), \end{cases}$$

for some $\mathbf{a}(\mathbf{x}) \in L^\infty(\Omega, \mathbb{R}^3)$, $\mathbf{n}(\mathbf{x}) \in L^\infty(\Omega, \mathbb{S}^2)$ such that $\mathbf{n}(\mathbf{x})$ is normal to the austenite-martensite interface in $\mathbf{x}$ at a certain time $t^* \in (0, t)$. Here we assumed without loss of generality to have $\mathbf{Q} = \mathbf{1}$ in Definition 4.3.1. In this way the martensitic macroscopic deformation gradient is a function of the moving, possibly curved, austenite-martensite interface during phase transition. In light of the above considerations, we assume that martensitic microstructures arising from austenite to martensite phase transitions are described by deformation gradients $\nabla \mathbf{y}$ of the form

$$\nabla \mathbf{y}(\mathbf{x}) = \mathbf{1} + \mathbf{a}(\mathbf{x}) \otimes \mathbf{n}(\mathbf{x}), \quad \nabla \mathbf{y} \in K^{qc}, \quad \text{a.e. } \mathbf{x} \in \Omega. \quad (\text{H1})$$

As the determinant is constant in K , it follows that

$$\det \nabla \mathbf{y}(\mathbf{x}) = \mathfrak{D} \quad \text{a.e. } \mathbf{x} \in \Omega, \quad (\text{H2})$$

for some constant $\mathfrak{D} > 0$. (H1) and (H2) imply

$$\mathbf{a}(\mathbf{x}) \cdot \mathbf{n}(\mathbf{x}) = \mathfrak{D} - 1, \quad (4.16)$$

for a.e. $\mathbf{x} \in \Omega$, and

$$\nabla \times (a_i(\mathbf{x})\mathbf{n}(\mathbf{x})) = \mathbf{0} \quad \text{for each } i = 1, 2, 3, \quad (4.17)$$

in a weak sense.

In conclusion, in what follows we look at martensite microstructures as a part of the domain where (H1) and (H2) hold. We first begin with an estimate for the norm of $\mathbf{a}(\mathbf{x})$:

Proposition 4.5.1. *Let λ_{max} and λ_{min} be respectively the biggest and the smallest eigenvalues of the martensite deformation matrices $\mathbf{U}_i \in K$. Let also $\mathbf{y} \in W^{1,\infty}(\Omega; \mathbb{R}^3)$ satisfy (H1), (H2). Then, we have*

$$|\mathfrak{D} - 1| \leq |\mathbf{a}| \leq \lambda_{max} - \lambda_{min}, \quad \text{a.e. in } \Omega.$$

Proof. The first inequality follows trivially from Cauchy-Schwarz and the fact that $\mathbf{a} \cdot \mathbf{n} = \mathfrak{D} - 1$. In order to get the other one, we observe that by the polar decomposition theorem we have that

$$\nabla \mathbf{y}(\mathbf{x}) = \mathbf{R}(\mathbf{x})\mathbf{F}(\mathbf{x}),$$

for almost every $\mathbf{x} \in \Omega$, where $\mathbf{R}(\mathbf{x}) \in SO(3)$ and $\mathbf{F}(\mathbf{x})$ is symmetric positive definite. The argument below holds for almost every $\mathbf{x} \in \Omega$. By arguing as in [13] one can deduce that, in order to have a rank-one connection with the identity matrix, the eigenvalues of $\mathbf{F}$, namely $\sigma_{min} \leq \sigma_{mid} \leq \sigma_{max}$, must satisfy

$$\sigma_{mid} = 1, \quad \sigma_{min}\sigma_{max} = \mathfrak{D}.$$

Therefore, we have

$$\det \mathbf{F} = \mathfrak{D} = \sigma_{min}\sigma_{max}, \quad \text{tr}(\mathbf{F}^T \mathbf{F}) = 1 + \sigma_{min}^2 + \sigma_{max}^2. \quad (4.18)$$

On the other hand,

$$\nabla \mathbf{y}^T \nabla \mathbf{y} = \mathbf{F}^T \mathbf{F} = \mathbf{1} + \mathbf{a} \otimes \mathbf{n} + \mathbf{n} \otimes \mathbf{a} + |\mathbf{a}|^2 \mathbf{n} \otimes \mathbf{n},$$

which yields

$$\operatorname{tr}(\mathbf{F}^T \mathbf{F}) = 3 + 2\mathbf{a} \cdot \mathbf{n} + |\mathbf{a}|^2 = 1 + 2\mathfrak{D} + |\mathbf{a}|^2. \quad (4.19)$$

Therefore, by putting together (4.19) and (4.18) we deduce

$$0 \leq |\mathbf{a}|^2 = (\sigma_{\max} - \sigma_{\min})^2 \leq (\lambda_{\max} - \lambda_{\min})^2.$$

Here we also made use of the following relation between eigenvalues proved in [13]:

$$\lambda_{\min} \leq \sigma_{\min} \leq 1 \leq \sigma_{\max} \leq \lambda_{\max}.$$

□

Another interesting property regards the divergence of $\mathbf{n} \otimes \mathbf{a}$:

Proposition 4.5.2. *Let $\mathbf{z} \in W^{1,\infty}(\Omega; \mathbb{R}^3)$ be such that*

$$\nabla \mathbf{z}(\mathbf{x}) = \mathbf{a}(\mathbf{x}) \otimes \mathbf{n}(\mathbf{x}), \quad \mathbf{a} \cdot \mathbf{n} = \mathfrak{D} - 1 \in \mathbb{R}, \quad \text{a.e. in } \Omega.$$

Then, $\nabla \cdot (\mathbf{n} \otimes \mathbf{a}) = \mathbf{0}$ in the sense of distributions. Furthermore, if $\mathbf{n} \in W^{1,1}(\Omega; \mathbb{R}^3)$ with $|\mathbf{n}| = 1$ a.e. in Ω , then $\nabla \cdot \mathbf{a} = 0$ in the sense of distributions.

Proof. On the one hand, we have

$$\nabla(\nabla \cdot \mathbf{z}) = \nabla(\mathfrak{D} - 1) = \mathbf{0}.$$

On the other hand

$$\nabla(\nabla \cdot \mathbf{z}) = \nabla \cdot (\nabla \mathbf{z})^T = \nabla \cdot (\mathbf{n} \otimes \mathbf{a}).$$

Here, both identities should be understood in the sense of distributions. By putting them together we hence get the first statement. As a consequence, if $\mathbf{n} \in W^{1,1}(\Omega; \mathbb{R}^3)$ such that $|\mathbf{n}| = 1$ a.e. in Ω , we can write,

$$\begin{aligned} \int_{\Omega} \mathbf{a} \cdot \nabla \varphi \, d\mathbf{x} &= \int_{\Omega} \mathbf{n}^T \mathbf{n} \otimes \mathbf{a} \cdot \nabla \varphi \, d\mathbf{x} \\ &= \int_{\Omega} \mathbf{n} \otimes \mathbf{a} : \nabla(\mathbf{n} \varphi) \, d\mathbf{x} - \int_{\Omega} \varphi \mathbf{n} \otimes \mathbf{a} : \nabla \mathbf{n} \, d\mathbf{x} \\ &= -\frac{1}{2} \int_{\Omega} \varphi \mathbf{a} \cdot \nabla |\mathbf{n}|^2 \, d\mathbf{x} = 0, \end{aligned}$$

for every $\varphi \in C_c^\infty(\Omega)$.

□

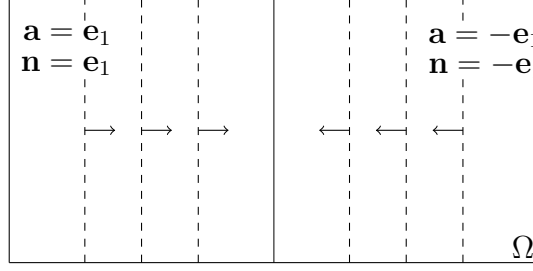


Figure 4.2: Picture of two parallel interfaces moving in opposite directions and meeting at a certain interface where $\mathbf{a}$, $\mathbf{n}$ are discontinuous. In this case, $\nabla \cdot (\mathbf{n} \otimes \mathbf{a}) = \mathbf{0}$, but $\nabla \cdot \mathbf{a} \neq 0$.

In general, it is not true that $\nabla \cdot \mathbf{a} = 0$. Indeed, let us fix $\mathbf{e} \in \mathbb{R}^3$ with $|\mathbf{e}| = 1$, and consider $\mathbf{z}$ to be of the form $\mathbf{z}(\mathbf{x}) = \mathbf{e}(\mathbf{x} \cdot \mathbf{e})$. Let us also fix $c \in \mathbb{R}$ such that $\mathbf{x} \cdot \mathbf{e} = c$ for some $\mathbf{x} \in \Omega$ and define

$$\mathbf{a} = \mathbf{n} = \begin{cases} \mathbf{e}, & \mathbf{x} \cdot \mathbf{e} < c, \\ -\mathbf{e}, & \mathbf{x} \cdot \mathbf{e} > c; \end{cases}$$

In this case, clearly $\mathbf{a} \cdot \mathbf{n} = 1$ a.e. in Ω , and $\nabla \cdot (\mathbf{n} \otimes \mathbf{a}) = 0$ in the distributional sense. However, $\mathbf{a} \in BV(\Omega; \mathbb{R}^3)$ satisfies

$$\nabla \mathbf{a} = -2\mathbf{e} \otimes \mathbf{e} \mathcal{H}^2 \llcorner \{\mathbf{x} : \mathbf{x} \cdot \mathbf{e} = c\},$$

and, as $|\mathbf{e}| = 1$ by hypothesis, $\nabla \cdot \mathbf{a} \neq 0$ in the distributional sense.

Keeping in mind the counterexample above, we extend the validity of the identity $\nabla \cdot \mathbf{a} = 0$ in the following Corollary 4.5.1. In this result we use the space of special functions with bounded variation on Ω , namely $SBV(\Omega)$. If $\varphi \in SBV(\Omega)$, its gradient is the sum of two Radon measures, one absolutely continuous with respect to the Lebesgue measure $\mathcal{L}^3$, and the other concentrated on a 2-rectifiable set S_φ , usually called the jump set. We denote by φ^-, φ^+ the trace of φ on the two sides of the jump set S_φ respectively. We refer the interested reader to [4] and [40] for more details on this space.

Corollary 4.5.1. *Let $\mathbf{z} \in W^{1,\infty}(\Omega; \mathbb{R}^3)$ be such that*

$$\nabla \mathbf{z}(\mathbf{x}) = \mathbf{a}(\mathbf{x}) \otimes \mathbf{n}(\mathbf{x}), \quad \mathbf{a} \cdot \mathbf{n} = \mathfrak{D} - 1 \in \mathbb{R}, \quad \text{a.e. in } \Omega. \quad (4.20)$$

Let $\mathbf{a}, \mathbf{n} \in SBV(\Omega; \mathbb{R}^3) \cap L^\infty(\Omega; \mathbb{R}^3)$, $|\mathbf{n}| = 1$ a.e. in Ω and let

$$\mathbf{n}^-(\mathbf{x}) \neq -\mathbf{n}^+(\mathbf{x}), \quad \text{if } \mathfrak{D} \neq 1, \quad (4.21)$$

$$\mathbf{a}^-(\mathbf{x}) = -\mathbf{a}^+(\mathbf{x}) \Rightarrow \mathbf{n}^-(\mathbf{x}) \neq -\mathbf{n}^+(\mathbf{x}), \quad \text{if } \mathfrak{D} = 1, \quad (4.22)$$

for $\mathcal{H}^2$ -a.e. $\mathbf{x} \in S_{\mathbf{a}} \cap S_{\mathbf{n}}$. Then, $\nabla \cdot \mathbf{a} = 0$ in the sense of distributions.

The following lemma is needed for the proof of Corollary 4.5.1:

Lemma 4.5.1. *Let $\varphi \in H^1(\Omega, \mathbb{R}^3)$, with $\nabla \varphi \in BV(\Omega; \mathbb{R}^{3 \times 3})$. Then, there exists $\mathbf{b}(\mathbf{x}) \in L^1(S_{\nabla \varphi}; \mathbb{R}^3)$ such that*

$$(\nabla \varphi)^+(\mathbf{x}) - (\nabla \varphi)^-(\mathbf{x}) = \mathbf{b}(\mathbf{x}) \otimes \mathbf{m}(\mathbf{x}), \quad \mathcal{H}^2\text{-a.e. } \mathbf{x} \in S_{\nabla \varphi},$$

with $\mathbf{m}(\mathbf{x})$ being the normal to $S_{\nabla \varphi}$ in $\mathbf{x}$.

Proof. The proof of this type of result is usually done via a blow up argument and exploits continuity. This may be possible here, but we use a slightly different proof. Let $\varphi \in H^1(\Omega; \mathbb{R}^3)$ and let A be a Lipschitz open subset of Ω . Since $\nabla \times \nabla \varphi = 0$ in the sense of distribution, we have that the following integration by parts formula holds (see [43, Ch. 2, (2.18)])

$$\int_A \nabla \varphi_i \cdot \nabla \times \psi \, d\mathbf{x} = \int_{\partial A} (\nabla \varphi_i \times \mathbf{m}) \cdot \psi \, d\mathcal{H}^2, \quad \forall \psi \in H^1(\Omega; \mathbb{R}^3), \quad (4.23)$$

with $\mathbf{m}$ being the outpointing normal to A . We remark that, as stated in [43, Ch. 2, Thm 2.5], $(\nabla \varphi_i \times \mathbf{m})$ is a well defined object in $H^{-\frac{1}{2}}(\partial A)$, and the integral on the right hand side should be interpreted as $\langle \nabla \varphi_i \times \mathbf{m}, \psi \rangle_{H^{-\frac{1}{2}}, H^{\frac{1}{2}}}$. Let now S be a Lipschitz 2-graph contained in Ω with normal $\mathbf{m}$ and let U_i be a countable set of open neighbourhoods such that $U_i \subset \Omega$, $U_i \setminus S$ has exactly two connected components, namely U_i^+ and U_i^- , and

$$\mathcal{H}^2(S \setminus \bigcup U_i) = 0.$$

We now define $(\nabla \varphi_i \times \mathbf{m})^\pm$ to be the objects of $H^{-\frac{1}{2}}(S)$ satisfying (4.23) respectively for $A = U_i^\pm$. From (4.23) we deduce

$$\begin{aligned} 0 &= \int_{U_i} \nabla \varphi_i \cdot \nabla \times \psi \, d\mathbf{x} = \int_{U_i^+} \nabla \varphi_i \cdot \nabla \times \psi \, d\mathbf{x} + \int_{U_i^-} \nabla \varphi_i \cdot \nabla \times \psi \, d\mathbf{x} \\ &= \int_{S \cap U_i} ((\nabla \varphi_i \times \mathbf{m})^+ - (\nabla \varphi_i \times \mathbf{m})^-) \cdot \psi \, d\mathcal{H}^2, \end{aligned}$$

for every $\psi \in C_c^\infty(U_i; \mathbb{R}^3)$. By repeating this argument on all U_i , this implies

$$\|(\nabla \varphi_i \times \mathbf{m})^+ - (\nabla \varphi_i \times \mathbf{m})^-\|_{H^{-\frac{1}{2}}(S)} = 0, \quad i = 1, 2, 3, \quad (4.24)$$

which is a weak Hadamard jump condition for $H^1(\Omega; \mathbb{R}^3)$ on Lipschitz 2-graphs. Furthermore, if $\nabla \varphi \in SBV(\Omega)$, we know that the set where it is discontinuous, namely $S_{\nabla \varphi}$ is 2-rectifiable (see e.g., [4]). Therefore, by Proposition (4.2.1) we can cover $S_{\nabla \varphi}$ with countably many Lipschitz graphs where (4.24) holds. On the other hand, from [4, Thm 3.77] we know that the trace of $\nabla \varphi_i$ is well defined for almost every point of $S_{\nabla \varphi}$. Collecting these two facts we thus deduce the desired result. $\square$

Remark 4.5.1. Equation (4.24) is a very weak version of the Hadamard jump condition on Lipschitz surfaces Γ with normal $\mathbf{m}$ for functions $\mathbf{y} \in H^1(\Omega; \mathbb{R}^3)$. Indeed, we can only make sense to the tangential trace $\nabla \mathbf{y} \times \mathbf{m}$ of $\mathbf{y}$ on Γ as an object of $H^{-\frac{1}{2}}(\Gamma)$, and (4.24) states that, across Γ , $\nabla \mathbf{y} \times \mathbf{m}$ must not jump as an object of $H^{-\frac{1}{2}}(\Gamma)$, which is kind of an average sense.

Proof. It follows from the definition of jump points of a BV function (see e.g., [4, 40]) that (4.20) and $|\mathbf{n}| = 1$ hold $\mathcal{H}^2$ -almost everywhere on $S_{\mathbf{a}} \cup S_{\mathbf{n}}$. Therefore, under our hypotheses $\mathbf{a} \otimes \mathbf{n} \in SBV(\Omega) \cap L^\infty(\Omega)$ and $S_{\mathbf{a} \otimes \mathbf{n}} = S_{\mathbf{a}} \cup S_{\mathbf{n}}$ up to an $\mathcal{H}^2$ -negligible set. Here, $\mathbf{a} = \mathbf{n}$ are chosen to be the precise representatives for $\mathbf{a}, \mathbf{n}$. Furthermore, Lemma 4.5.1 implies that a Hadamard jump condition must hold across $S_{\mathbf{a} \otimes \mathbf{n}}$, so that

$$\mathbf{a}^+ \otimes \mathbf{n}^+ - \mathbf{a}^- \otimes \mathbf{n}^- = \mathbf{b} \otimes \mathbf{m}, \quad \mathcal{H}^2\text{-a.e. on } S_{\mathbf{a} \otimes \mathbf{n}}, \quad (4.25)$$

for some $\mathbf{b} \in L^\infty(S_{\mathbf{a} \otimes \mathbf{n}}; \mathbb{R}^3)$ and with $\mathbf{m}$ being the normal to $S_{\mathbf{a} \otimes \mathbf{n}}$. In case $\mathfrak{D} \neq 1$, this, together with (4.21) and $|\mathbf{n}| = 1$, imply that the only possible scenarios are the following on $S_{\mathbf{a}} \cup S_{\mathbf{n}}$:

(I) $\mathbf{n}^+ = \mathbf{n}^- = \mathbf{m}$, in which case $\mathbf{b} = \mathbf{a}^+ - \mathbf{a}^-$;

(II) $\mathbf{a}^+ = \xi \mathbf{a}^-$, in which case $\mathbf{m} \parallel \xi \mathbf{n}^+ - \mathbf{n}^-$ and $\mathbf{b} \parallel \mathbf{a}^+ \parallel \mathbf{a}^-$,

for some $\xi \in L^\infty(S_{\mathbf{n}})$. Taking the trace of (4.25) implies also $\mathbf{b} \cdot \mathbf{m} = 0$. As a consequence, $(\mathbf{a}^+ - \mathbf{a}^-) \cdot \mathbf{m} = 0$ $\mathcal{H}^2$ -a.e. on $S_{\mathbf{a}} \cup S_{\mathbf{n}}$, that is, the divergence of $\mathbf{a}$ has no singular part. In case $\mathfrak{D} = 1$, (4.22) allows also to have $\mathbf{n}^+ = -\mathbf{n}^- = \mathbf{m}$ in $S_{\mathbf{a}}$, when $\mathbf{a}^+ \neq -\mathbf{a}^-$. In which case $(\mathbf{a}^+ - \mathbf{a}^-) \cdot \mathbf{m} = 0$ follows just by the fact that $\mathbf{a} \cdot \mathbf{n} = 0$.

It just remains to check that the part of $\nabla \cdot \mathbf{a}$ which is absolutely continuous with respect to $\mathcal{L}^3$ is null as well. To this aim, we first observe that, by the chain rule for BV functions (see e.g., [4, Example 3.97])

$$\mathbf{0} = \nabla \cdot (\mathbf{n} \otimes \mathbf{a}) = (\mathbf{n}(\nabla^a \cdot \mathbf{a}) + \nabla^a \mathbf{n} \mathbf{a}) \mathcal{L}^3 + (\mathbf{n}^+ \otimes \mathbf{a}^+ - \mathbf{n}^- \otimes \mathbf{a}^-) \mathbf{m} \mathcal{H}^2 \llcorner S_{\mathbf{a} \otimes \mathbf{n}}, \quad (4.26)$$

where we denoted by ∇^a the absolutely continuous part of the gradient. Given (I)–(II) above, under our hypotheses we have $(\mathbf{n}^+ \otimes \mathbf{a}^+ - \mathbf{n}^- \otimes \mathbf{a}^-) \mathbf{m} = \mathbf{0}$ $\mathcal{H}^2$ -a.e. on $S_{\mathbf{a} \otimes \mathbf{n}}$. Furthermore, as $|\mathbf{n}| = 1$ a.e., we have

$$\mathbf{0} = \nabla |\mathbf{n}|^2 = \mathbf{n}^T \nabla^a \mathbf{n} \mathcal{L}^3.$$

Therefore, multiplying (4.26) by $\mathbf{n}$ and exploiting the fact that $|\mathbf{n}| = 1$ a.e., we thus get

$$\nabla^a \cdot \mathbf{a} = 0, \quad \text{a.e. in } \Omega.$$

Therefore, as $\mathbf{a} \in SBV(\Omega)$, for every $\phi \in C_c^1(\Omega)$ we have

$$-\int_{\Omega} \nabla \phi \cdot \mathbf{a} \, d\mathbf{x} = \int_{S_{\mathbf{a}}} \phi(\mathbf{a}^+ - \mathbf{a}^-) \cdot \mathbf{m} \, d\mathcal{H}^2 + \int_{\Omega} \phi \nabla^a \cdot \mathbf{a} \, d\mathbf{x} = 0$$

which concludes the proof. $\square$

We now focus on compound twins. Thanks to Proposition 1.4.2 we know that if two martensite variants $\mathbf{U}_1, \mathbf{U}_2 \in \mathbb{R}_{Sym^+}^{3 \times 3}$ form a compound twin, then there exists $\mu > 0$, $\mathbf{v} \in \mathbb{S}^2$ such that $\mathbf{U}_1 \mathbf{v} = \mathbf{U}_2 \mathbf{v} = \mu \mathbf{v}$. In this case, [36, Theorem 2.5.1] states:

Theorem 4.5.1. *Let $\mathbf{U}_1, \dots, \mathbf{U}_n \in \mathbb{R}_{Sym^+}^{3 \times 3}$, such that $\det \mathbf{U}_i = \mathfrak{D} > 0$, and such that $\mathbf{U}_i \mathbf{v} = \mu \mathbf{v}$ for some $\mu > 0$, $\mathbf{v} \in \mathbb{S}^2$, for each $i = 1, \dots, n$. Let also*

$$H = \bigcup_{i=1}^n SO(3) \mathbf{U}_i.$$

Then, there exists $l \in \mathbb{N}$, $\mathbf{w}_1, \dots, \mathbf{w}_l$ such that

$$H^{qc} = \left\{ \mathbf{F} \in \mathbb{R}^{3 \times 3} \left| \begin{array}{l} \det \mathbf{F} = \mathfrak{D}, \mathbf{F}^T \mathbf{F} \mathbf{v} = \mu^2 \mathbf{v} \text{ and} \\ |\mathbf{F} \mathbf{w}_i|^2 \leq \max_{j=1, \dots, n} |\mathbf{U}_j \mathbf{w}_i|^2 \text{ for each} \\ i = 1, \dots, l. \end{array} \right. \right\}$$

Therefore, in this simple case which includes compound twins, we can actually construct the quasiconvex hull of the set. An interesting result is given by the following lemma:

Lemma 4.5.2. *Let $\mathbf{U}_1, \dots, \mathbf{U}_n \in \mathbb{R}_{Sym^+}^{3 \times 3}$, such that $\det \mathbf{U}_i = \mathfrak{D} > 0$, and such that $\mathbf{U}_i \mathbf{v} = \mu \mathbf{v}$ for some $\mu > 0$, $\mathbf{v} \in \mathbb{S}^2$, for each $i = 1, \dots, n$. Let also*

$$H = \bigcup_{i=1}^n SO(3) \mathbf{U}_i,$$

and $\mu \neq 1$. Then every map $\mathbf{y} \in W^{1, \infty}(\Omega; \mathbb{R}^3)$ satisfying $\nabla \mathbf{y}(\mathbf{x}) \in H^{qc}$ and $\nabla \mathbf{y}(\mathbf{x}) = \mathbf{1} + \mathbf{a}(\mathbf{x}) \otimes \mathbf{n}(\mathbf{x})$ for a.e. $\mathbf{x}$ in Ω is such that $\mathbf{a} \otimes \mathbf{n}$ is constant, and $|\mathbf{a}| = |\mu|^{-1} |\mathfrak{D} - \mu^2|$.

Remark 4.5.2. It comes up in the proof of Lemma 4.5.2 that the following condition

$$1 \geq \frac{\mu^2(1 - \mu^2)}{(\mathfrak{D}^2 - \mu^4)} > 0 \tag{4.27}$$

is necessary in order to have the existence of $\mathbf{a} \in \mathbb{R}^3$, $\mathbf{n} \in \mathbb{S}^2$ such that $\mathbf{1} + \mathbf{a} \otimes \mathbf{n} \in H^{qc}$. We refer the reader to [13] for some stricter necessary conditions for this to hold.

Proof. Thanks to a change of coordinates, we can suppose without loss of generality that $\mathbf{v} = \mathbf{e}_3$. In this case, every $\mathbf{F} \in H^{qc}$ satisfies

$$\mathbf{F}^T \mathbf{F} = \begin{bmatrix} \alpha & \gamma & 0 \\ \gamma & \beta & 0 \\ 0 & 0 & \mu^2 \end{bmatrix} \quad (4.28)$$

and

$$\alpha\beta - \gamma^2 = \mathfrak{D}^2 \mu^{-2}. \quad (4.29)$$

We are first interested in solving $\mathbf{1} + \mathbf{a} \otimes \mathbf{n} \in H^{qc}$ for $\mathbf{a} \in \mathbb{R}^3, \mathbf{n} \in \mathbb{S}^2$. By (4.28), the first step is to solve the following nonlinear system of equations for the components of $\mathbf{a}$ and $\mathbf{n}$:

$$\begin{cases} 1 + 2a_1n_1 + |a|^2n_1^2 & = \alpha \\ 1 + 2a_2n_2 + |a|^2n_2^2 & = \beta \\ a_1n_2 + a_2n_1 + |a|^2n_1n_2 & = \gamma \\ a_3n_1 + a_1n_3 + |a|^2n_1n_3 & = 0 \\ a_3n_2 + a_2n_3 + |a|^2n_2n_3 & = 0 \\ 1 + 2a_3n_3 + |a|^2n_3^2 & = \mu^2 \end{cases} \quad (4.30)$$

subject to the constraint (4.29). As a first step we compute $\alpha\beta - \gamma^2$ using system (4.30). After rearranging terms

$$\alpha\beta - \gamma^2 = 2(a_1n_1 + a_2n_2) + 1 + |\mathbf{a}|^2(n_1^2 + n_2^2) - (a_1n_2 - a_2n_1)^2.$$

Since $|\mathbf{n}| = 1$, using last equation of (4.30) follows that

$$|\mathbf{a}|^2(n_1^2 + n_2^2) = |\mathbf{a}|^2(1 - n_3^2) = |\mathbf{a}|^2 + 1 + 2a_3n_3 - \mu^2,$$

which leads to

$$\alpha\beta - \gamma^2 = 2(\mathbf{a} \cdot \mathbf{n}) + 2 + |\mathbf{a}|^2 - \mu^2 - (a_1n_2 - a_2n_1)^2.$$

Finally, by recalling (4.16) we deduce

$$\alpha\beta - \gamma^2 = 2\mathfrak{D} + |\mathbf{a}|^2 - \mu^2 - (a_1n_2 - a_2n_1)^2. \quad (4.31)$$

Thus, it is immediately seen that (4.31) together with the first of (4.29) implies

$$(a_1n_2 - a_2n_1)^2 = 2\mathfrak{D} + |\mathbf{a}|^2 - \mu^2 - \frac{\mathfrak{D}^2}{\mu^2} = |\mathbf{a}|^2 - \frac{1}{\mu^2}(\mathfrak{D} - \mu^2)^2. \quad (4.32)$$

On the other hand, exploiting the last equation of (4.30) in the fourth and fifth one we have

$$a_2 n_3 - a_3 n_2 = -\frac{n_2}{n_3}(\mu^2 - 1) \quad (4.33)$$

$$a_3 n_1 - a_1 n_3 = \frac{n_1}{n_3}(\mu^2 - 1). \quad (4.34)$$

We note that it is legitimate to divide by n_3 since the last identity in (4.30) together with $\mu \neq 1$ implies $n_3 \neq 0$. By putting together (4.32)-(4.34) we thus get

$$\mathbf{a} \times \mathbf{n} = \begin{bmatrix} -\frac{n_2}{n_3}(\mu^2 - 1) \\ \frac{n_1}{n_3}(\mu^2 - 1) \\ \pm \sqrt{|\mathbf{a}|^2 - \frac{1}{\mu^2}(\mathfrak{D} - \mu^2)^2} \end{bmatrix}. \quad (4.35)$$

Since $(\mathbf{a} \times \mathbf{n}) \cdot \mathbf{n} = 0$ we have

$$n_3 \sqrt{|\mathbf{a}|^2 - \frac{1}{\mu^2}(\mathfrak{D} - \mu^2)^2} = 0, \quad (4.36)$$

which, as $n_3 \neq 0$, leads to $|\mathbf{a}|^2 = \frac{1}{\mu^2}(\mathfrak{D} - \mu^2)^2$. In this case the norm of $\mathbf{a}$ is hence forced to be constant. We restrict ourselves without loss of generality to the case $\mathfrak{D} \neq \mu^2$.

We now want to write $\mathbf{a}$ in terms of the orthogonal vectors $(\mathbf{n}, \mathbf{n}^\perp, \mathbf{n} \times \mathbf{n}^\perp)$, where

$$\mathbf{n}^\perp := \begin{bmatrix} n_2 \\ -n_1 \\ 0 \end{bmatrix}, \quad \mathbf{n} \times \mathbf{n}^\perp = \begin{bmatrix} n_1 n_3 \\ n_2 n_3 \\ n_3^2 - 1 \end{bmatrix}. \quad (4.37)$$

It is important to remark that, if $n_1 = n_2 = 0$ then it is easy to see from (4.30) that $a_1 = a_2 = 0$, so we do not lose any generality with this representation. We have,

$$\mathbf{a} = a^{(1)}\mathbf{n} + a^{(2)}\mathbf{n}^\perp + a^{(3)}\mathbf{n} \times \mathbf{n}^\perp.$$

As a first thing, recalling that $\det(\mathbf{1} + \mathbf{a} \otimes \mathbf{n}) = \mathfrak{D}$, we deduce that $a^{(1)} = \mathbf{a} \cdot \mathbf{n} = \mathfrak{D} - 1$. On the other hand, taking cross product of $\mathbf{a}$ with $\mathbf{n}$ follows that

$$\mathbf{n} \times \mathbf{a} = a^{(2)}\mathbf{n} \times \mathbf{n}^\perp - a^{(3)}\mathbf{n}^\perp. \quad (4.38)$$

A comparison of (4.38) with (4.35) and (4.36) thus leads to $a^{(3)} = \frac{1}{n_3}(1 - \mu^2)$, and $a^{(2)} = 0$, which implies

$$\mathbf{a} = (\mathfrak{D} - 1)\mathbf{n} + \frac{1}{n_3}(1 - \mu^2)\mathbf{n} \times \mathbf{n}^\perp = (\mathfrak{D} - \mu^2)\mathbf{n} - \frac{1}{n_3}(1 - \mu^2)\mathbf{e}_3, \quad (4.39)$$

with $\mathbf{e}_3 = (0, 0, 1)^T$. As a first thing now we want to check whether it is possible to have $|\mathbf{a}|^2 = \frac{1}{\mu^2}(\mathfrak{D} - \mu^2)^2$. Since the orthogonal vectors $\mathbf{n}$ and $\mathbf{n} \times \mathbf{n}^\perp$ satisfy $|\mathbf{n}| = 1$ and $|\mathbf{n} \times \mathbf{n}^\perp|^2 = 1 - n_3^2$, by rearranging the terms it is possible to obtain

$$n_3^2 = \frac{(1 - \mu^2)^2}{(1 - \mu^2)^2 + \frac{1}{\mu^2}(\mu^2 - \mathfrak{D})^2 - (\mathfrak{D} - 1)^2} = \frac{(1 - \mu^2)}{\frac{1}{\mu^2}(\mathfrak{D}^2 - \mu^4)}. \quad (4.40)$$

It is easy to check that a pair of vectors $(\mathbf{n}, \mathbf{a})$ with $\mathbf{a}$ defined in terms of $\mathbf{n}$ as in (4.39), and where n_3 is given by (4.40), satisfies the equations of (4.30). Thus, it turns out that (4.27) is necessary in order not to contradict $|\mathbf{n}| = 1$, $1 - n_3^2 = n_1^2 + n_2^2 \geq 0$ and $n_3^2 > 0$ (see also Remark 4.5.2).

Let us now check when a map $\mathbf{y} \in W^{1,\infty}(\Omega; \mathbb{R}^3)$ such that $\nabla \mathbf{y}(\mathbf{x})$ satisfies (4.28)–(4.29) and $\nabla \mathbf{y}(\mathbf{x}) = \mathbf{1} + \mathbf{a}(\mathbf{x}) \otimes \mathbf{n}(\mathbf{x})$ for a.e. $\mathbf{x}$ in Ω . We have to verify that conditions expressed in (4.17) hold, that is, we have to check when the constructed deformations are actually gradients, by verifying that $\nabla \times (a_i \mathbf{n}) = 0$ for $i = 1, 2, 3$, in the distributional sense. Since $\mathfrak{D}$ and μ are constants, so is $|n_3| = \sqrt{(\mu^2 - \mu^4)(\mathfrak{D}^2 - \mu^4)^{-1}}$. Hence, by (4.39), $|a_3|$ is constant and non zero as long as $\mu^2 \neq \mathfrak{D}$ and $\mu \neq 1$, which is our case. Furthermore, $n_3 = \text{sign}(n_3)|n_3|$ and $a_3 = \pm |a_3| \text{sign}(n_3)$, where the sign depends on $\mu, \mathfrak{D}$ only and is hence fixed. We can hence suppose without loss of generality $a_3 = |a_3| \text{sign}(n_3)$. Therefore,

$$\nabla \times (a_3 \mathbf{n}) = |a_3| \nabla \times (\text{sign}(n_3) \mathbf{n}) = 0,$$

implies

$$\nabla \times (\text{sign}(n_3) \mathbf{n}) = 0,$$

in the sense of distributions. This implies the existence of $\psi \in W^{1,2}(\Omega)$ such that $\text{sign}(n_3) \mathbf{n} = \nabla \psi$ (see e.g., [43]). On the other hand, by Proposition 4.5.2

$$\nabla \cdot (\mathbf{a} n_3) = |n_3| \nabla \cdot (\text{sign}(n_3) \mathbf{a}) = 0,$$

which implies

$$\nabla \cdot (\text{sign}(n_3) \mathbf{a}) = 0,$$

in the sense of distributions. Furthermore, we have from (4.39) that

$$\nabla \cdot (\text{sign}(n_3) \mathbf{a}) = (\mathfrak{D} - \mu^2) \nabla \cdot (\text{sign}(n_3) \mathbf{n})$$

which thus implies that ψ is harmonic in the sense of distributions. By using standard elliptic theory (see e.g., [39]) we can thus deduce that $\psi \in C^\infty(\Omega)$. On the other hand, as $|\nabla \psi|^2 = 1$, we have

$$\nabla \psi^T \nabla^2 \psi = 0, \quad \text{for all } \mathbf{x} \in \Omega,$$

which implies that one eigenvalue of the Hessian matrix is null. Another eigenvalue is 0 as $\mathbf{e}_3$ is a left eigenvector related to 0 for $\nabla^2\psi$, and is not parallel to $\mathbf{n}$ unless $\mathbf{n} = \mathbf{e}_3$, in which case (4.40) forces $\text{sign}(n_3)\mathbf{n} = \mathbf{e}_3$ everywhere. The fact that ψ is harmonic, i.e., $\text{tr } \nabla^2\psi = 0$, thus means that $\nabla\psi$ is constant. Therefore, $\text{sign}(n_3)\mathbf{n}$ is constant, and as a consequence of (4.39), so is $\text{sign}(n_3)\mathbf{a}$ and $\mathbf{a} \otimes \mathbf{n}$. $\square$

Remark 4.5.3. In case $\mu = 1$, it is not possible to deduce the same rigidity as in Lemma 4.5.2. Indeed, it can be deduced from equations (4.30) that $\mu = 1$ implies either $n_3 = a_3 = 0$, or $\mathbf{a} = (\mathfrak{D} - 1)\mathbf{n}$, but in the latter case the last equation in (4.30) implies $\mathbf{a} \otimes \mathbf{n} = 0$. The problem becomes thus two-dimensional, and we can rewrite

$$\mathbf{n} = (n_1, n_2, 0)^T, \quad \mathbf{n}^\perp = (n_2, -n_1, 0)^T, \quad \mathbf{a}(s) = (\mathfrak{D} - 1)\mathbf{n} + s\mathbf{n}^\perp,$$

for some $s \in \mathbb{R}$, $n_1, n_2 \in [-1, 1]$ with $n_1^2 + n_2^2 = 1$. Now, take two martensite variants, for example,

$$\mathbf{U}_1 = \frac{1}{2} \begin{bmatrix} \lambda_m + \lambda_M & \lambda_m - \lambda_M & 0 \\ \lambda_m - \lambda_M & \lambda_m + \lambda_M & 0 \\ 0 & 0 & 2 \end{bmatrix}, \quad \mathbf{U}_2 = \frac{1}{2} \begin{bmatrix} \lambda_m + \lambda_M & \lambda_M - \lambda_m & 0 \\ \lambda_M - \lambda_m & \lambda_m + \lambda_M & 0 \\ 0 & 0 & 2 \end{bmatrix},$$

generating a compound twin, and such that their biggest and smallest eigenvalue, namely λ_M and λ_m , satisfy $\lambda_m < 1 < \lambda_M$. Assume further, for simplicity, $\lambda_m^2 + \lambda_M^2 > 2$ and $\frac{1}{2} \leq \lambda_m \lambda_M < 1$. We can take for example $\lambda_m = 0.9$ and $\lambda_M = 1.1$. It can be computed that $(SO(3)\mathbf{U}_1 \cup SO(3)\mathbf{U}_2)^{qc}$ coincides with the set of matrices $\mathbf{F}$ satisfying (4.28)–(4.29), where $0 \leq \alpha, \beta \leq \frac{1}{2}(\lambda_m^2 + \lambda_M^2)$. Consider now $\mathbf{F} = \mathbf{1} + \mathbf{a}(s) \otimes \mathbf{n}$, then

$$\alpha = 1 + n_1^2(\mathfrak{D}^2 - 1 + s^2) + sn_1n_2, \quad \beta = 1 + n_2^2(\mathfrak{D}^2 - 1 + s^2) - sn_1n_2.$$

Choosing for simplicity $n_2 = 0$ and s small enough, we get

$$0 < \alpha, \beta < \frac{1}{2}(\lambda_m^2 + \lambda_M^2), \quad \alpha\beta > \mathfrak{D}^2.$$

Thus, there exists $\varepsilon > 0$, and an open interval $[\hat{a}_0 - \varepsilon, \hat{a}_0 + \varepsilon]$ such that for every smooth function $f: \mathbb{R} \rightarrow [\hat{a}_0 - \varepsilon, \hat{a}_0 + \varepsilon]$ we have that

$$\mathbf{1} + \mathbf{a}(f(\mathbf{x} \cdot \mathbf{n}^\perp)) \otimes \mathbf{n},$$

is the gradient of a smooth map $\mathbf{y}$, which is not constant, and which satisfies $\nabla\mathbf{y}(\mathbf{x}) \in (SO(3)\mathbf{U}_1 \cup SO(3)\mathbf{U}_2)^{qc}$ and (H1)–(H2) for each $\mathbf{x}$. Therefore, a rigidity result as the one in Lemma 4.5.2 does not hold in general when $\mu = 1$.

By adding the further hypotheses that $\mathbf{y}|_{\partial\Omega}$ is the restriction on $\partial\Omega$ of a 1 – 1 map, we can extend Lemma 4.5.2 to general two well problems

Proposition 4.5.3. *Let $\mathbf{U}, \mathbf{V} \in \mathbb{R}_{Sym^+}^{3 \times 3}$ and $\mathbf{R}_I, \mathbf{R}_{II} \in SO(3)$, $\mathbf{b}_I, \mathbf{b}_{II}, \mathbf{m}_I, \mathbf{m}_{II} \in \mathbb{R}^3 \setminus \{\mathbf{0}\}$ satisfy*

$$\mathbf{R}_i \mathbf{V} = \mathbf{U} + \mathbf{b}_i \otimes \mathbf{m}_i, \quad i = I, II,$$

where $(\mathbf{U}, \mathbf{b}_I, \mathbf{m}_I)$, $(\mathbf{U}, \mathbf{b}_{II}, \mathbf{m}_{II})$ do not fulfil (CC2). Assume $\mathbf{y} \in W^{1,\infty}(\Omega; \mathbb{R}^3)$ is such that $\mathbf{y}|_{\partial\Omega} = \mathbf{y}_0|_{\partial\Omega}$ for some $\mathbf{y}_0 \in C(\overline{\Omega}; \mathbb{R}^3)$ which is 1 – 1 in Ω , and

$$\nabla \mathbf{y}(\mathbf{x}) \in (SO(3)\mathbf{U} \cup SO(3)\mathbf{V})^{qc}, \quad \nabla \mathbf{y}(\mathbf{x}) = \mathbf{1} + \mathbf{a}(\mathbf{x}) \otimes \mathbf{n}(\mathbf{x}), \quad (4.41)$$

a.e. in Ω , for some $\mathbf{a} \in L^\infty(\Omega; \mathbb{R}^3)$, $\mathbf{n} \in L^\infty(\Omega; \mathbb{S}^2)$. Then,

$$\nabla \mathbf{y} = \text{constant}$$

Proof. Following the approach devised in [19], we introduce the orthonormal system of coordinates

$$\mathbf{u}_1^i := \frac{\mathbf{U}_1^{-1} \mathbf{m}_i}{|\mathbf{U}_1^{-1} \mathbf{m}_i|}, \quad \mathbf{u}_3^i := \frac{\mathbf{b}_i}{|\mathbf{b}_i|}, \quad \mathbf{u}_2^i := \mathbf{u}_3^i \times \mathbf{u}_1^i,$$

with $i = I, II$ to be chosen later, and let

$$\mathbf{L}_i := \mathbf{U}_1^{-1} (\mathbf{1} - \delta^i \mathbf{u}_3^i \otimes \mathbf{u}_1^i), \quad \delta^i = \frac{1}{2} |\mathbf{U}_1^{-1} \mathbf{m}_i| |\mathbf{b}_i|.$$

We set $\mathbf{z}^i(\mathbf{x}) := \mathbf{y}(\mathbf{L}_i \mathbf{x})$ and the problem becomes equivalent to find a map $\mathbf{z}^i \in W^{1,\infty}(\Omega; \mathbb{R}^3)$ such that

$$\nabla \mathbf{z}^i(\mathbf{x}) \in (SO(3)\mathbf{S}_i^- \cup SO(3)\mathbf{S}_i^+)^{qc}, \quad \text{a.e. } \mathbf{x} \in \Omega^{L_i} := \{\mathbf{x} : \mathbf{L}_i \mathbf{x} \in \Omega\}, \quad (4.42)$$

with $\mathbf{S}_i^\pm = \mathbf{1} \pm \delta^i \mathbf{u}_3^i \otimes \mathbf{u}_1^i$, and

$$\nabla \mathbf{z}^i(\mathbf{x}) = \mathbf{L}_i + \mathbf{a}(\mathbf{L}_i \mathbf{x}) \otimes \mathbf{L}_i^T \mathbf{n}(\mathbf{L}_i \mathbf{x}), \quad \text{a.e. in } \Omega^{L_i}, \quad (4.43)$$

where $\mathbf{a} \in L^\infty(\Omega; \mathbb{R}^3)$, $\mathbf{n} \in L^\infty(\Omega; \mathbb{S}^2)$ are as in (4.41). By [6], $\mathbf{z}^i$ is a plane strain and satisfies

$$\mathbf{z}^i(\mathbf{x}) = \mathbf{Q} (z_1^i(s_1^i, s_3^i) \mathbf{u}_1^i + s_2^i \mathbf{u}_2^i + z_3^i(s_1^i, s_3^i) \mathbf{u}_3^i), \quad (4.44)$$

for some $\mathbf{Q} \in SO(3)$, some Lipschitz functions z_1^i, z_2^i , and where $s_j^i = \mathbf{x} \cdot \mathbf{u}_j^i$, $j = 1, 2, 3$. As a consequence, from (4.43)–(4.44) we deduce

$$\begin{aligned} \mathbf{u}_2^i &= \mathbf{Q}^T \nabla \mathbf{z}^i(\mathbf{x}) \mathbf{u}_2^i = \mathbf{Q}^T \mathbf{L}_i \mathbf{u}_2^i + \mathbf{Q}^T \mathbf{a}(\mathbf{L}_i \mathbf{x}) (\mathbf{n}(\mathbf{L}_i \mathbf{x}) \cdot \mathbf{L}_i \mathbf{u}_2^i), \\ \mathbf{u}_2^i &= (\nabla \mathbf{z}^i(\mathbf{x}))^T \mathbf{Q} \mathbf{u}_2^i = \mathbf{L}_i^T \mathbf{Q} \mathbf{u}_2^i + \mathbf{L}_i^T \mathbf{n}(\mathbf{L}_i \mathbf{x}) (\mathbf{a}(\mathbf{L}_i \mathbf{x}) \cdot \mathbf{Q} \mathbf{u}_2^i), \end{aligned}$$

a.e. in $\Omega^{\mathbf{L}_i}$. That is,

$$\mathbf{a}(\mathbf{x})(\mathbf{n}(\mathbf{x}) \cdot \mathbf{L}_i \mathbf{u}_2^i) = (\mathbf{Q} - \mathbf{L}_i) \mathbf{u}_2^i, \quad (4.45)$$

$$\mathbf{n}(\mathbf{x})(\mathbf{a}(\mathbf{x}) \cdot \mathbf{Q} \mathbf{u}_2^i) = (\mathbf{L}_i^{-T} - \mathbf{Q}) \mathbf{u}_2^i, \quad (4.46)$$

a.e. in Ω . Let us now consider the function

$$f_i(\mu) = \det((\mathbf{U} + \mu \mathbf{b}_i \otimes \mathbf{m}_i)^T (\mathbf{U} + \mu \mathbf{b}_i \otimes \mathbf{m}_i) - 1), \quad \mu \in [0, 1], \quad i = I, II.$$

Thanks to [18, Prop. 5] we know that f_i is a quadratic polynomial and $f_i(\mu) = f_i(1 - \mu)$. We first notice that

$$\begin{aligned} f_i(\mu) &= (\det \mathbf{U}) \det((\mathbf{U} + \mu \mathbf{b}_i \otimes \mathbf{m}_i) - (\mathbf{U} + \mu \mathbf{b}_i \otimes \mathbf{m}_i)^{-T}) \\ &= (\det \mathbf{U}^2) \det((1 - \mathbf{U}^{-2}) + \mu 2\delta(\mathbf{u}_3^i \otimes \mathbf{u}_1^i + \mathbf{u}_1^i \otimes \mathbf{U}^{-2} \mathbf{u}_3^i)) \\ &= \det((\mathbf{U}^2 - 1) + \mu 2\delta(\mathbf{u}_3^i \otimes \mathbf{U}^2 \mathbf{u}_1^i + \mathbf{u}_1^i \otimes \mathbf{u}_3^i)). \end{aligned}$$

A derivation of f_i with respect to μ leads

$$\begin{aligned} f_i'(\mu) &= 2\delta(\det \mathbf{U}^2) \operatorname{cof}((1 - \mathbf{U}^{-2}) + \mu 2\delta(\mathbf{u}_3^i \otimes \mathbf{u}_1^i + \mathbf{u}_1^i \otimes \mathbf{U}^{-2} \mathbf{u}_3^i)) : (\mathbf{u}_3^i \otimes \mathbf{u}_1^i + \mathbf{u}_1^i \otimes \mathbf{U}^{-2} \mathbf{u}_3^i) \\ &= 2\delta(\det \mathbf{U}^2) (\operatorname{cof}(1 - \mathbf{U}^{-2}) \mathbf{u}_1^i \cdot (\mathbf{u}_3^i + \mathbf{U}^{-2} \mathbf{u}_3^i) + \mu 4\delta(1 - \mathbf{U}^{-2}) \mathbf{u}_2^i \cdot (\mathbf{u}_1^i \times \mathbf{U}^{-2} \mathbf{u}_3^i)) \\ &= 2\delta(\operatorname{cof}(\mathbf{U}^2 - 1) \mathbf{u}_3^i \cdot (\mathbf{u}_1^i + \mathbf{U}^2 \mathbf{u}_1^i) + \mu 4\delta(\mathbf{U}^2 - 1) \mathbf{u}_2^i \cdot (\mathbf{U}^2 \mathbf{u}_1^i \times \mathbf{u}_3^i)). \end{aligned}$$

We now fix $i = I$ and claim that, under our hypotheses, there exist no $\mathbf{Q} \in SO(3)$ such that $(\mathbf{Q} - \mathbf{L}_I^{-T}) \mathbf{u}_2^I = \mathbf{0}$. This, by (4.46) and the fact that $|\mathbf{n}| = 1$ a.e. implies that $\mathbf{n}$ is, up to a change of sign, equal to a constant. That is, $\mathbf{n} \operatorname{sign}(n_j)$ is constant, for some $j = 1, 2, 3$ such that $|n_j| \neq 0$ a.e. in Ω . To prove the claim we argue by contradiction, and notice that the existence of $\mathbf{Q} \in SO(3)$ satisfying $(\mathbf{Q} - \mathbf{L}_I^{-T}) \mathbf{u}_2^I = \mathbf{0}$ implies

$$(\mathbf{U}^2 - 1) \mathbf{u}_2^I \cdot \mathbf{u}_2^I = 0. \quad (4.47)$$

Let us notice now that

$$(\mathbf{U}^2 \mathbf{u}_1^I \times \mathbf{u}_3^I) \times \mathbf{u}_2^I = (\mathbf{U}^2 \mathbf{u}_1^I \cdot (\mathbf{u}_1^I \times \mathbf{u}_3^I)) \mathbf{u}_3^I. \quad (4.48)$$

By making use of (1.12) in (4.48) we show that $\mathbf{U}^2 \mathbf{u}_1^I \cdot (\mathbf{u}_1^I \times \mathbf{u}_3^I) = 0$, which by (4.48) implies that $\mathbf{U}^2 \mathbf{u}_1^I \times \mathbf{u}_3^I$ is parallel to $\mathbf{u}_2^I$. Therefore, by (4.47), we get that f_I' is constant in μ . Furthermore, as f_I is a quadratic polynomial and $f_I(\mu) = f_I(1 - \mu)$ we have that $f_I'(\frac{1}{2}) = 0$ and hence f_I' is identically 0. But, as

$$f_I'(\mu)|_{\mu=0} = 2\mathbf{b}_I \cdot \mathbf{U} \operatorname{cof}(\mathbf{U}^2 - 1) \mathbf{m}_I = 0,$$

we contradict the assumption that $(\mathbf{U}, \mathbf{b}_I, \mathbf{m}_I)$ does not satisfy (CC2) concluding the proof of the claim. We now fix $i = II$ and claim that, under our hypotheses, there exist no $\mathbf{Q} \in SO(3)$ such that $(\mathbf{Q} - \mathbf{L}_{II})\mathbf{u}_2^{II} = \mathbf{0}$. This, by (4.45) and the fact that $\text{sign}(n_j)\mathbf{n}$ is a constant implies that $\text{sign}(n_j)\mathbf{a}$ is also a constant. To prove the claim we argue again by contradiction, and notice that the existence of $\mathbf{Q} \in SO(3)$ satisfying $(\mathbf{Q} - \mathbf{L}_{II})\mathbf{u}_2^{II} = \mathbf{0}$ implies

$$(\mathbf{U}^{-2} - 1)\mathbf{u}_2^{II} \cdot \mathbf{u}_2^{II} = 0. \quad (4.49)$$

By making use of (1.13) we can now show that $\mathbf{u}_1^{II} \times \mathbf{U}^{-2}\mathbf{u}_3^{II}$ is parallel to $\mathbf{u}_2^{II}$, and hence, by (4.49), that f'_{II} is constant in μ and identically 0. But, noticing that

$$f'_{II}(\mu)|_{\mu=0} = 2\mathbf{b}_{II} \cdot \mathbf{U} \text{cof}(\mathbf{U}^2 - 1)\mathbf{m}_{II} = 0,$$

we contradict the assumption that $(\mathbf{U}, \mathbf{b}_{II}, \mathbf{m}_{II})$ does not satisfy (CC2) concluding the proof of the second claim.

In conclusion, we proved that $\text{sign}(n_j)\mathbf{n}$ and $\text{sign}(n_j)\mathbf{a}$ are constants, and therefore so must be $\nabla \mathbf{y}$. $\square$

Remark 4.5.4. *It is clear from the proof of Proposition 4.5.3 that, if the type I solution $(\mathbf{U}, \mathbf{b}_I, \mathbf{m}_I)$ of the twinning equation (1.4.1) does not satisfy (CC2), but the type II solution $(\mathbf{U}, \mathbf{b}_{II}, \mathbf{m}_{II})$ of (1.4.1) does, then we can guarantee that $\mathbf{n}$ in (4.41) is constant up to a change of sign. Similarly, if the type I solution $(\mathbf{U}, \mathbf{b}_I, \mathbf{m}_I)$ of the twinning equation (1.4.1) does satisfy (CC2), but the type II solution $(\mathbf{U}, \mathbf{b}_{II}, \mathbf{m}_{II})$ of (1.4.1) does not, then the direction of $\mathbf{a}$ in (4.41) is constant. That is, there exists $\mathbf{v} \in \mathbb{R}^3$ such that $\mathbf{a} \times \mathbf{v} = 0$ a.e. in Ω . We refer the reader to Proposition 4.6.1 and Proposition 4.6.2 for further examples under these hypotheses.*

4.6 Macroscopic moving interfaces

In this section, we use the theory of the previous sections to prove some results about martensitic microstructures. The results are different for different type of twins. We start with compound twins and we recall that, by Proposition 1.4.2, two martensite variants $\mathbf{U}_1, \mathbf{U}_2 \in \mathbb{R}_{Sym+}^{3 \times 3}$ form a compound twin if and only if there exist $\mu > 0$, $\mathbf{v} \in \mathbb{S}^2$ such that $\mathbf{U}_1\mathbf{v} = \mathbf{U}_2\mathbf{v} = \mu\mathbf{v}$. Thanks to Lemma 4.5.2 we can prove that in this case moving interfaces need to be planar and the related macroscopic gradient constant.

Theorem 4.6.1. *Let $U_1, U_2 \in \mathbb{R}_{Sym+}^{3 \times 3}$ be a compound twin and such that*

$$U_1 \mathbf{w} = U_2 \mathbf{w} = \mu \mathbf{w}$$

for some $\mathbf{w} \in \mathbb{R}^3$, $\mu \neq 1$. Then, every $\mathbf{y}$ satisfying the regular moving mask approximation and such that

$$\nabla \mathbf{y} \in (SO(3)U_1 \cup SO(3)U_2)^{qc}, \quad \text{a.e. in } \Omega$$

is constant, and the related moving interfaces planar.

Proof. As $\mathbf{y}$ satisfies a regular moving mask approximation, by Theorem 4.4.1 we know that $\nabla \mathbf{y} = \mathbf{1} + \mathbf{a} \otimes \mathbf{n}$ for some $\mathbf{a} \in L^\infty(\Omega; \mathbb{R}^3)$, $\mathbf{n} \in L^\infty(\Omega; \mathbb{S}^2)$. Since $\mu \neq 1$, we can apply Lemma 4.5.2 and thus deduce that $\nabla \mathbf{y}$ is constant in Ω . The function $\mathbf{z}(\mathbf{x}) := \mathbf{y}(\mathbf{x}) - \mathbf{x}$ is such that $\mathbf{z} \in W^{1,\infty}(\Omega; \mathbb{R}^3)$ and is constant in every connected component of $\Omega_A(t)$, for each $t \geq 0$. Thus, $\Gamma(t)$ must be a (or at least the union of disconnected subsets of a) level-set for $\mathbf{z}$, and hence a plane (or union of disconnected planes) as $\nabla \mathbf{z}$ is constant and rank-one in Ω . $\square$

By arguing in the same way, Theorem 4.6.1 can be generalised to a wider range of situations as stated in Theorem 4.6.2 below. This is relevant, for example, in the cubic to monoclinic transformation occurring in $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$, where there are 3 sets of four deformation gradients satisfying the hypotheses of Theorem 4.6.2.

Theorem 4.6.2. *Let $U_1, \dots, U_N \in \mathbb{R}_{Sym+}^{3 \times 3}$ be such that $U_i \mathbf{w} = \mu \mathbf{w}$ and $\det U_i = \mathfrak{D}$ for some $\mathbf{w} \in \mathbb{R}^3$, $\mu \neq 1$ and every $i = 1, \dots, N$. Then, every $\mathbf{y}$ satisfying the regular moving mask approximation and such that*

$$\nabla \mathbf{y} \in (\cup_{i=1}^N SO(3)U_i)^{qc}, \quad \text{a.e. in } \Omega$$

is constant, and the related moving interfaces planar.

The second result is related to type I twins satisfying the cofactor conditions. In this case we can prove that simple laminates can form macroscopically curved families of austenite-martensite interfaces with no transition layer. The proof strongly relies on Theorem 1.4.2.

Proposition 4.6.1. *Let U_1 and U_2 be two martensitic variants and $\hat{\mathbf{R}} \in SO(3)$, $\mathbf{b}_I, \mathbf{m}_I \in \mathbb{R}^3$ be a type I solution to (1.10) and satisfying the cofactor conditions. Then, there exist $\mathbf{R}_0, \mathbf{R}_1 \in SO(3)$, $\xi \in \mathbb{R}$, $\mathbf{a}_0 \in \mathbb{R}^3$, $\mathbf{n}_0, \mathbf{n}_1 \in \mathbb{S}^2$ such that for every*

$\lambda \in L^\infty(\Omega; [0, 1])$ satisfying $\nabla \lambda \times (\xi \mathbf{n}_1 - \mathbf{n}_0) = \mathbf{0}$ in the sense of distributions, there exists $\mathbf{y} \in W^{1,\infty}(\Omega; \mathbb{R}^3)$ with

$$\nabla \mathbf{y} = \mathbf{R}_0[(1 - \lambda)\mathbf{U}_1 + \lambda \hat{\mathbf{R}}\mathbf{U}_2] = \mathbf{1} + \mathbf{a}_0 \otimes (\lambda \xi \mathbf{n}_1 + (1 - \lambda)\mathbf{n}_0), \quad a.e. \text{ in } \Omega.$$

Furthermore, $\mathbf{y}$ satisfies a regular moving mask approximation, the related moving interfaces are curved, and $\nabla \cdot (|\lambda \xi \mathbf{n}_1 + (1 - \lambda)\mathbf{n}_0| \mathbf{a}_0) = 0$ in the sense of distributions.

Proof. From Theorem 1.4.2 and in particular from (1.18) we know the existence of $\mathbf{R}_0, \mathbf{R}_1 \in SO(3)$, $\xi \in \mathbb{R}$, $\mathbf{a}_0 \in \mathbb{R}^3$, $\mathbf{n}_0, \mathbf{n}_1 \in \mathbb{S}^2$ such that

$$\mathbf{R}_0[(1 - \lambda)\mathbf{U}_1 + \lambda \hat{\mathbf{R}}\mathbf{U}_2] = \mathbf{1} + \mathbf{a}_0 \otimes (\lambda \xi \mathbf{n}_1 + (1 - \lambda)\mathbf{n}_0), \quad \text{for all } \lambda \in [0, 1].$$

We thus choose $\lambda \in L^\infty(\Omega)$ to be a function such that $\lambda \in [0, 1]$ a.e., and define

$$\mathbf{a}(\mathbf{x}) := \mathbf{a}_0 |\mathbf{n}_0 + \lambda(\mathbf{x})(\xi \mathbf{n}_1 - \mathbf{n}_0)|, \quad \mathbf{n}(\mathbf{x}) := \frac{\mathbf{n}_0 + \lambda(\mathbf{x})(\xi \mathbf{n}_1 - \mathbf{n}_0)}{|\mathbf{n}_0 + \lambda(\mathbf{x})(\xi \mathbf{n}_1 - \mathbf{n}_0)|},$$

so that $\mathbf{n}$ has unitary norm. In the notation of Theorem 1.4.2, we have

$$\det(\mathbf{R}_0[(1 - \lambda)\mathbf{U}_1 + \lambda \hat{\mathbf{R}}\mathbf{U}_2]) = \det(\mathbf{R}_0) \det(\mathbf{U}_1 + \lambda \mathbf{b}_I \otimes \mathbf{m}_I) = \det \mathbf{U}_1$$

where the Sherman-Morrison inversion formula and the fact that $\mathbf{U}_1^{-1} \mathbf{n} \cdot \mathbf{a} = 0$ have been used. That is,

$$\mathbf{a}_0 \cdot (\mathbf{n}_0 + \lambda(\xi \mathbf{n}_1 - \mathbf{n}_0))$$

is constant independently of λ , or, in an equivalent way,

$$\mathbf{a}_0 \cdot (\xi \mathbf{n}_1 - \mathbf{n}_0) = 0. \quad (4.50)$$

We just need to check if it is possible to have

$$\nabla \times (a_i \mathbf{n}) = 0.$$

Exploiting the definition of $\mathbf{a}$ and $\mathbf{n}$ get that this is satisfied if and only if

$$\nabla \lambda \times (\xi \mathbf{n}_1 - \mathbf{n}_0) = \mathbf{0}, \quad (4.51)$$

in a weak sense. Therefore, if Ω is convex λ must satisfy $\lambda(\mathbf{x}) = f(\mathbf{x} \cdot (\xi \mathbf{n}_1 - \mathbf{n}_0))$, for some $f \in L^\infty(\mathbb{R}; [0, 1])$, and thus

$$\mathbf{y} = \mathbf{x} + \mathbf{a}_0(\mathbf{n}_0 \cdot \mathbf{x} + F(\mathbf{x} \cdot (\xi \mathbf{n}_1 - \mathbf{n}_0))) + \mathbf{c},$$

for some constant $\mathbf{c} \in \mathbb{R}^3$, and where $F(s) = \int_0^s f(s) ds$. Therefore, after choosing

$$I_T := \left(\inf_{\mathbf{x} \in \Omega} G(\mathbf{x}), \sup_{\mathbf{x} \in \Omega} G(\mathbf{x}) \right),$$

where $G(\mathbf{x}) = \mathbf{n}_0 \cdot \mathbf{x} + F(\mathbf{x} \cdot (\xi \mathbf{n}_1 - \mathbf{n}_0))$, we deduce that the image of $\mathbf{z}(\mathbf{x}) = \mathbf{y}(\mathbf{x}) - \mathbf{x}$ is $\mathbf{c} + t\mathbf{a}_0$, for $t \in I_T$. If Ω is connected but not convex, then λ might not be of the form $f(\mathbf{x} \cdot (\xi \mathbf{n}_1 - \mathbf{n}_0))$, but the image of $\mathbf{z}(\mathbf{x}) = \mathbf{y}(\mathbf{x}) - \mathbf{x}$ is still contained in the one-dimensional line $\mathbf{c} + t\mathbf{a}_0$, for some $\mathbf{c} \in \mathbb{R}^3$ and t in some bounded interval $\hat{I}_T$. We can thus use Corollary 4.4.1 and deduce the existence of a family of moving interfaces, which are also level sets for $\mathbf{z}(\mathbf{x})$.

In order to prove that $\nabla \cdot \mathbf{a}$, we first mollify λ and, defined $\mathbf{m}$ as $\mathbf{m} := \xi \mathbf{n}_1 - \mathbf{n}_0$, notice that thanks to Fubini's theorem for distributions we can write

$$\langle \nabla \lambda_\varepsilon \times \mathbf{m}, \boldsymbol{\psi} \rangle_{\mathcal{D}', \mathcal{D}} = \langle \nabla \lambda, (\mathbf{m} \times \boldsymbol{\psi}_\varepsilon) \rangle_{\mathcal{D}', \mathcal{D}} = \langle \nabla \lambda \times \mathbf{m}, \boldsymbol{\psi}_\varepsilon \rangle_{\mathcal{D}', \mathcal{D}} = 0,$$

thanks to (4.51), for all $\boldsymbol{\psi} \in C_c^\infty(\Omega, \mathbb{R}^3)$. Therefore, $\nabla \lambda_\varepsilon \parallel \mathbf{m}$, and, by (4.50),

$$\nabla \lambda_\varepsilon \cdot \mathbf{a}_0 = 0. \quad (4.52)$$

On the other hand, exploiting the smooth dependence on λ of $\mathbf{a}$ and the fact that $\lambda_\varepsilon \rightarrow \lambda$ in $L^p(\Omega)$ for every $p \in [1, \infty)$, we have

$$\int_{\Omega} \mathbf{a}(\lambda) \cdot \nabla \psi \, d\mathbf{x} = \lim_{\varepsilon \rightarrow 0} \int_{\Omega} \mathbf{a}(\lambda_\varepsilon) \cdot \nabla \psi \, d\mathbf{x} = - \lim_{\varepsilon \rightarrow 0} \int_{\Omega} g'(\lambda_\varepsilon) \mathbf{a}_0 \cdot \nabla \lambda_\varepsilon \psi \, d\mathbf{x}$$

for every $\psi \in C_c^\infty(\Omega)$ and where $g(s) = |\mathbf{n}_0 + s(\xi \mathbf{n}_1 - \mathbf{n}_0)|$. The last term in the chain of identities above is null due to (4.52), and therefore the proof is concluded. $\square$

Finally, a result related to type II twins satisfying the cofactor conditions:

Proposition 4.6.2. *Let U_1 and U_2 be two martensitic variants and $\hat{R} \in SO(3)$, $\mathbf{b}_{II}, \mathbf{m}_{II} \in \mathbb{R}^3$ be a type II solution to (1.10) and satisfying the cofactor conditions. Then, there exist $R_0 \in SO(3)$, $\xi \in \mathbb{R}$, $\mathbf{a}_0, \mathbf{a}_1 \in \mathbb{R}^3$, $\mathbf{n}_0 \in \mathbb{S}^2$ such that for every $\lambda \in L^\infty(\Omega; [0, 1])$ satisfying $\nabla \lambda \times \mathbf{n}_0 = \mathbf{0}$ in the sense of distributions, there exists $\mathbf{y} \in W^{1, \infty}(\Omega; \mathbb{R}^3)$ with*

$$\nabla \mathbf{y} = R_0[(1 - \lambda)U_1 + \lambda \hat{R}U_2] = \mathbf{1} + (\lambda \xi \mathbf{a}_1 + (1 - \lambda)\mathbf{a}_0) \otimes \mathbf{n}_0, \quad \text{a.e. in } \Omega.$$

Furthermore, $\mathbf{y}$ satisfies a regular moving mask approximation, and

$$\nabla \cdot (\lambda \xi \mathbf{a}_1 + (1 - \lambda)\mathbf{a}_0) = 0$$

in the sense of distributions.

Proof. From Theorem (1.4.3) and in particular from (1.20) we know the existence of $R_0, R_1 \in SO(3)$, $\xi \in \mathbb{R}$, $\mathbf{a}_0, \mathbf{a}_1 \in \mathbb{R}^3$, $\mathbf{n}_0 \in \mathbb{S}^2$ satisfying

$$R_0[(1 - \lambda)U_1 + \lambda \hat{R}U_2] = 1 + (\lambda \xi \mathbf{a}_1 + (1 - \lambda)\mathbf{a}_0) \otimes \mathbf{n}_0, \quad \text{for all } \lambda \in [0, 1].$$

Thus choose $\lambda \in L^\infty(\Omega)$ to be a function such that $\lambda \in [0, 1]$ a.e., and define

$$\mathbf{a}(\mathbf{x}) := (\mathbf{a}_0 + \lambda(\mathbf{x})(\xi \mathbf{a}_1 - \mathbf{a}_0))$$

It is trivial to check, by arguing as in the proof of Proposition 4.6.1, that also (H2) holds.

Now, by taking the curl of $\nabla \mathbf{y}_i$ we deduce that

$$\nabla \times (\lambda \mathbf{n}_0) = \nabla \lambda \times \mathbf{n}_0 = 0,$$

in the sense of distributions. Therefore taking λ such that $\nabla \lambda \parallel \mathbf{n}$ in a weak sense, by Proposition 4.4.1 and Remark 4.4.1 follows the existence of a family of moving planar interfaces. The fact that $\nabla \cdot \mathbf{a} = 0$ in the sense of distributions trivially follows from Proposition 4.5.2. $\square$

4.7 Experimental evidence

The physical assumptions which were made in this work and some of the properties that have been deduced here are currently being investigated from an experimental perspective. Indeed, the authors of [27] have used X-ray Laue microdiffraction to measure the orientations and structural parameters of variants and phases in $\text{Zn}_{45}\text{Au}_{30}\text{Cu}_{25}$. With this modern technique, the scanned area is meshed with small rectangles (e.g., in [27] authors use $2\mu\text{m}$ wide squares), and one can identify the phase and variant in each cuboid which has as a basis a rectangle of the mesh and depth of approximate $2\mu\text{m}$ from the sample surface. In cubes where a single phase or variant has been recognised, one can also measure the lattice parameters necessary to compute the average deformation gradient. In this way one can investigate what is happening at the phase interface by studying the lattice parameters in mesh rectangles where a martensite variant has been recognised and which have at least one neighbouring rectangle where the Laue microdiffraction was able to identify austenite.

In this way, the authors of [27] compute in some of the mesh cubes lying on the

interface the number $\|\text{cof}(\nabla \mathbf{y} - \mathbf{1})\|$, which, as it is easy to verify, is zero if and only if $\nabla \mathbf{y} - \mathbf{1}$ is rank-one. Experimental results give $\|\text{cof}(\nabla \mathbf{y} - \mathbf{1})\|$ to be of the order of 10^{-4} , which seems small enough to be considered zero, and hence to justify (H1).

Further investigations are ongoing to verify that $\nabla \mathbf{y}(\mathbf{x})$ remains constant in time when $\mathbf{x}$ is not on the interface. This seems a reasonable assumption, as long as no external force acts on the sample and as long as one neglects other internal stresses giving rise to elastic deformations which, anyway, seem to be small compared to the deformations induced by the phase transition.

In conclusion, the data collected up till now seem to confirm the validity of the assumptions that we made in the present work. However, in the images in [27] there are many mesh cubes close to some of the phase interfaces where the X-ray Laue microdiffraction is not able to recognize any single variant or phase, and hence where the validity of the assumptions to get (H1) could be questioned or should be verified in some other way.

Chapter 5

Vanishing interfacial energy as a selection mechanism for minimising gradient Young measures

A common problem that arises when studying martensitic transformations in the context of nonlinear elasticity (see e.g., [18, 19, 23, 62]) is to minimize an energy functional

$$E(y) = \int_{\Omega} \phi(\nabla y(x)) \, dx,$$

where Ω is an open and bounded Lipschitz domain, and $y: \Omega \rightarrow \mathbb{R}^3$ is a map in a suitable Sobolev space satisfying $y = \bar{y}$ on $\partial\Omega$, for some smooth enough mapping $\bar{y}$. In this context, the continuous function $\phi: \mathbb{R}^{3 \times 3} \rightarrow [0, +\infty]$ is generally such that

$$\phi(F) = 0 \quad \Longleftrightarrow \quad F \in \mathcal{K} := \sum_{i=1}^n SO(3)U_i,$$

where $n \geq 1$ and U_i are positive definite symmetric matrices representing the different variants of martensite. As in general E is not quasiconvex, minimizers for this energy might not exist. Therefore, following the idea of [18] one can study the behaviour of minimizing sequences, having a gradient that tends in measure to $\mathcal{K}$, and characterised by interesting microstructures. In order to capture the limiting behaviour of the minimising sequences, one can study the relaxed functional

$$\bar{E}(\nu) = \int_{\Omega} \int_{\mathbb{R}^{3 \times 3}} \phi(F) \, d\nu_x(F) \, dx,$$

where ν is a gradient Young measure containing the information about microstructures in the crystal (see e.g., [19, 62, 69]). Defining $\mathcal{M}_1(\mathbb{R}^{3 \times 3})$ as the set of probability

measures on $\mathbb{R}^{3 \times 3}$, let us consider

$$\mathcal{A} := \left\{ \nu \in L_{w^*}^\infty(\Omega; \mathcal{M}_1(\mathbb{R}^{3 \times 3})) \left| \begin{array}{l} \text{supp } \nu_x \subset \mathcal{K}, \exists y \in W^{1,\infty}(\Omega; \mathbb{R}^3) \text{ s.t. } y = \bar{y} \text{ on } \partial\Omega, \\ \text{and } \int_{\mathbb{R}^{3 \times 3}} F d\nu_x(F) = \nabla y(x) \text{ a.e. in } \Omega \end{array} \right. \right\},$$

and notice that this set is the set of minimizers of $\bar{E}$ whenever $\min \bar{E} = 0$. Here, we denoted by $L_{w^*}^\infty(\Omega; \mathcal{M}_1(\mathbb{R}^{3 \times 3}))$ the space $L^\infty(\Omega; \mathcal{M}_1(\mathbb{R}^{3 \times 3}))$ endowed with the weak* topology. The solutions constructed in [65] with the technique of convex integration, show that the set $\mathcal{A}$ might contain infinitely many minimizers for $\bar{E}$, and its elements might sometimes appear non-physical. In agreement with the physics, many authors in the literature (see e.g., [16, 18, 37, 60, 32, 33]) have considered a regularization of E that penalizes the second derivatives of y such as

$$E^\varepsilon(y) = \int_{\Omega} (\varepsilon^2 |\nabla^2 y|^2 + \phi(\nabla y(x))) dx, \quad \text{or} \quad \tilde{E}^\varepsilon(y) = \varepsilon |\nabla^2 y|(\Omega) + \int_{\Omega} \phi(\nabla y(x)) dx. \quad (5.1)$$

Here, $\varepsilon > 0$ is small and $|\nabla^2 y|(\Omega)$ is the norm of $\nabla^2 y$ as a measure on Ω . Many results have been proved in the case $n = 2$ and without boundary conditions. For example, it is proved in [37] that the requirement $\nabla y \in BV(\Omega; \mathcal{K})$ forces the gradient discontinuities to be just on planes that never intersect in Ω . In [32] the limit solutions for E^ε as $\varepsilon \rightarrow 0$ when $\mathcal{K} = \{A, B\}$ are characterized via a Γ -limit argument. In the two-dimensional setting with $\mathcal{K} = \{SO(2)A, SO(2)B\}$ the generalised Γ -limit has been analysed in [33], and strongly exploits the above mentioned result of [37].

More generally, we could argue that the physically relevant minimizers of $\bar{E}$ are not those in $\mathcal{A}$, but those belonging to the subset

$$\mathcal{B} := \left\{ \nu \in \mathcal{A} \left| \begin{array}{l} \exists \text{ minimizers } u^{\varepsilon_j} \text{ of } E^{\varepsilon_j}, \text{ with } \varepsilon_j \downarrow 0, \text{ such that} \\ \delta_{\nabla u^{\varepsilon_j}} \rightarrow \nu \text{ in } L_{w^*}^\infty(\Omega; \mathcal{M}_1(\mathbb{R}^{3 \times 3})) \end{array} \right. \right\},$$

or equivalently $\tilde{\mathcal{B}}$ where E^{ε_j} is replaced by $\tilde{E}^{\varepsilon_j}$.

Finding an explicit characterization for $\mathcal{B}$ seems however out of reach for the general three-dimensional problem. For this reason, in this work we focus on the one-dimensional energy functional

$$\mathcal{E}(u) = \int_0^1 (W(u_x) + u^2) dx, \quad (5.2)$$

which has been often considered in the literature (see e.g., [17, 60, 61, 68]) as a one-dimensional prototype for E . Indeed, the role of the boundary conditions in more

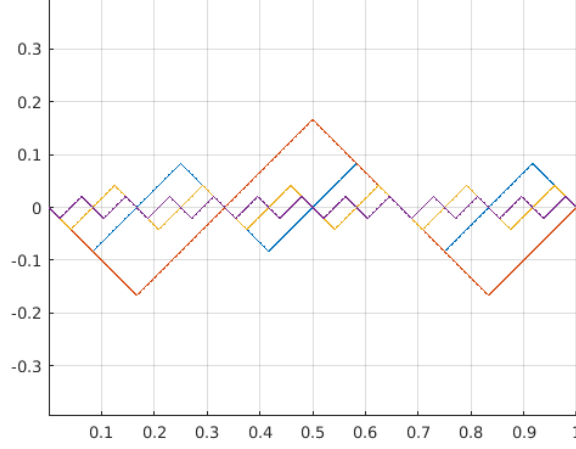


Figure 5.1: Minimizing sequences for $\mathcal{E}$ when $0 \notin \mathcal{Z}$.

dimensions is played here by the term u^2 in the energy, which forces the L^2 -norm of the minimisers (or of the minimising sequences) to be close to a prescribed value, which is chosen to be null for simplicity, and whose gradient does not sit on the wells. Suppose W satisfies

(H1) $W : \mathbb{R} \mapsto \mathbb{R}_+$ is a continuous non-negative function;

(H2) there exist $c_1, c_2, c_3 > 0$ and $p \in (1, \infty)$ such that

$$c_1|s|^p - c_2 \leq W(s) \leq c_3(|s|^p + 1), \quad \forall s \in \mathbb{R};$$

(H3) $W(s) = 0, \quad \forall s \in \mathcal{Z}, \quad W(s) > 0, \quad \text{otherwise, where}$

$$\mathcal{Z} := \{s \in \mathbb{R} : s \in \operatorname{argmin}(W)\}.$$

If $\mathcal{Z}$ has a finite number of elements, if there exist $z_1, z_2 \in \mathcal{Z}$ with $z_1 < 0 < z_2$, and if $0 \notin \mathcal{Z}$, then W is not convex and $\mathcal{E}$ does not have minimizers in $W_0^{1,p}(0, 1)$. Indeed, by constructing arbitrarily small saw-tooth functions (cf. Figure 5.1) with gradient in $\mathcal{Z}$ one can show that $\inf \mathcal{E} = 0$. Therefore, the existence of a minimizer $u \in W_0^{1,p}$ would imply $u = 0$, and hence $u_x = 0$ a.e. in $(0, 1)$, which is in contradiction with the fact that, by (H3), $W(0) \neq 0$. For this reason, we consider the regularized problem

$$\mathcal{E}^\varepsilon(u) = \begin{cases} \int_0^1 (\varepsilon^6 |u_{xx}|^2 + W(u_x) + u^2) \, dx, & \text{if } u \in W_0^{1,p}(0, 1) \cap H^2(0, 1), \\ +\infty, & \text{otherwise,} \end{cases} \quad (5.3)$$

which is the one-dimensional analogue of (5.1). $\mathcal{E}^\varepsilon$ can also be rewritten by using gradient Young measures (see e.g., [62, 69]) as

$$\bar{\mathcal{E}}^\varepsilon(u, \nu) = \begin{cases} \mathcal{E}^\varepsilon(u), & \text{if } u \in W_0^{1,p}(0, 1) \text{ and } \nu_x = \delta_{u_x(x)} \text{ a.e. in } (0, 1), \\ +\infty, & \text{otherwise,} \end{cases} \quad (5.4)$$

with δ_s denoting the Dirac mass at s . In this case the problem admits a solution in $W_0^{1,p}(0, 1) \cap H^2(0, 1)$ and the question arises as to what happens to the limit of the minimizers u^ε as $\varepsilon \downarrow 0$. For every $u \in W_0^{1,p}(0, 1)$ let us define the set of its gradient Young measures

$$\begin{aligned} \text{GYM}^p(u) &:= \left\{ \nu \in L_{w^*}^\infty(0, 1; \mathcal{M}_1(\mathbb{R})) \left| \begin{array}{l} \int_{\mathbb{R}} s \, d\nu_x(s) = u_x(x) \text{ a.e. } x \in (0, 1), \\ \int_0^1 \int_{\mathbb{R}} |s|^p \, d\nu_x(s) \, dx < \infty \end{array} \right. \right\}, \\ \text{GYM}^\infty(u) &:= \left\{ \nu \in L_{w^*}^\infty(0, 1; \mathcal{M}_1(\mathbb{R})) \left| \begin{array}{l} \int_{\mathbb{R}} s \, d\nu_x(s) = u_x, \text{ supp } \nu_x \subset K \text{ a.e. in } \\ (0, 1), K \subset \mathbb{R} \text{ compact} \end{array} \right. \right\}. \end{aligned}$$

Here, $\mathcal{M}(\mathbb{R})$ and $\mathcal{M}_1(\mathbb{R})$, often abbreviated below by $\mathcal{M}$ and $\mathcal{M}_1$, are respectively the space of bounded Radon measures μ on $\mathbb{R}$, and its subset of probability measures. A preliminary result that is proved later in Section 5.1 is the following

Proposition 5.0.1. *Let W satisfy (H1)–(H3). Then, $\bar{\mathcal{E}}^\varepsilon$ Γ -converges to*

$$\mathcal{E}^0(u, \nu) = \begin{cases} \int_0^1 (\langle \nu_x, W \rangle + u^2) dx, & \text{if } u \in W_0^{1,p}(0, 1), \nu \in \text{GYM}^p(u), \\ +\infty, & \text{otherwise,} \end{cases}$$

in the $L^2(0, 1) \times L_{w^*}^\infty(0, 1; \mathcal{M})$ topology as ε tends to 0.

If $\mathcal{Z} = \{z_1, z_2\}$ with $z_1 < 0 < z_2$, then under (H1)–(H3) minimizers (u, ν) of $\mathcal{E}^0$ must satisfy

$$u(x) = 0, \quad \nu \in \text{GYM}^p(0), \quad \text{supp } \nu_x \in \mathcal{Z}, \quad \text{a.e. in } (0, 1). \quad (5.5)$$

These conditions determine a unique minimizer (u, ν) to $\mathcal{E}^0$, namely

$$u(x) = 0, \quad \nu_x = \frac{z_2}{z_2 - z_1} \delta_{z_1} - \frac{z_1}{z_2 - z_1} \delta_{z_2}, \quad \text{a.e. } x \text{ in } (0, 1).$$

Let us assume

(H4) $\mathcal{Z} = \{z_1, z_2, z_3\}$, and, without loss of generality, that $z_1 < 0 < z_2 < z_3$.

In this case, given any arbitrary measurable

$$\lambda: (0, 1) \rightarrow \left[0, \left(1 - \frac{z_2}{z_1}\right)^{-1}\right],$$

the pair (u, ν) defined for almost every $x \in (0, 1)$ as $u(x) = 0$ and

$$\nu_x = -\frac{z_3 + \lambda(x)(z_2 - z_3)}{z_1 - z_3} \delta_{z_1} + \lambda(x) \delta_{z_2} + \frac{z_1 + \lambda(x)(z_2 - z_1)}{z_1 - z_3} \delta_{z_3},$$

minimises $\mathcal{E}^0$. As a consequence, by assuming (H1)–(H4), uniqueness of minimizers for $\mathcal{E}^0$ is lost, that is, the gradient of the minimizing sequences for $\mathcal{E}$ oscillate, and converge in measure to $\{z_1, z_2, z_3\}$ without any particular preference. The aim of this work is to prove that minimizers of $\mathcal{E}^\varepsilon$ generate gradient Young measures supported in $\{z_1, z_2\}$, but not in z_3 . Therefore, by choosing minimisers of $\mathcal{E}^\varepsilon$ with $\varepsilon \downarrow 0$ as minimizing sequences for $\mathcal{E}$ we can select a unique minimising gradient Young Measure, out of the infinitely many given above.

Let

$$V := H^2(0, 1) \cap W_0^{1,p}(0, 1).$$

Then we define I^ε by

$$I^\varepsilon(u) = I^\varepsilon(u, \nu) := \begin{cases} \varepsilon^{-2} \int_0^1 (\varepsilon^6 u_{xx}^2 + W(u_x) + u^2) dx, & \text{if } u \in V, \nu_x = \delta_{u_x(x)}, \\ +\infty, & \text{otherwise.} \end{cases}$$

We remark that this problem was thoroughly studied in [60, 61], under the assumption that W is a double-well potential, and where quasi-periodicity of the minimizers was also proved. As shown below, however, generalization to a three well problem is non-trivial and requires a good understanding on the possible shape of the minimizing sequences. We also point out that the behaviour of I^ε is different from the one of Modica-Mortola type functionals (see e.g., [29, 57]) as $\varepsilon \downarrow 0$. Indeed, in our case the term in u^2 forces minimizers of I^ε to oscillate faster and faster as $\varepsilon \downarrow 0$, making the number of oscillations in the gradient tend to infinity. In what follows we define

$$E_0 := 2 \int_{z_1}^{z_2} |W(s)|^{\frac{1}{2}} ds, \quad E_1 := 2 \int_{z_2}^{z_3} |W(s)|^{\frac{1}{2}} ds,$$

and

$$A_0 := \inf_d (3^{-1} z_2^2 z_{21} d^2 + E_0 d^{-1}) = (2^{-1} 3)^{\frac{2}{3}} E_0^{\frac{2}{3}} (z_2^2 z_{21})^{\frac{1}{3}},$$

$$B_0 := \inf_d (3^{-1} z_3^2 z_{31} d^2 + (E_0 + E_1) d^{-1}) = (2^{-1} 3)^{\frac{2}{3}} (E_0 + E_1)^{\frac{2}{3}} (z_3^2 z_{31})^{\frac{1}{3}}.$$

where $z_{i1} := (1 - \frac{z_i}{z_1})$ for $i = 2, 3$. Further assumptions on W are:

(H5) (Coercivity) There exist $\eta_0 \in (0, \min\{1, -z_1, z_2, \frac{z_3 - z_2}{2}\})$, $c_0 > 0$, $q > 0$ such that

$$W(s) \geq c_0 \min\{\min_i |s - z_i|^q, |\eta_0|^q\}, \quad \forall s \in \mathbb{R};$$

(H6) Let $f_6(y) := 9(E_0 + E_1)^2(z_2^2 + y^3 z_3^2 + 3yz_2(yz_3 + z_2))$, then

$$f_6(y) - (A_0 + B_0 y)^3 \geq 0, \quad \text{for every } y \geq 0;$$

(H7) Let $f_7(y) := \frac{9}{4}(E_0 + 2E_1)^2(z_2^2 z_{21} + y^3 z_3^2 z_{31} + 3yz_2 z_{31}(yz_3 + z_2))$, then

$$f_7(y) - (A_0 + B_0 y)^3 \geq 0, \quad \text{for every } y \geq 0;$$

(H8) Let

$$f_8(y) := 9(E_0 + E_1)^2 \left(z_2^2 z_{21} + y^3 z_3^2 z_{31} - 3 \frac{(y^2 z_{31} z_3 - z_2 z_{21})^2}{4(z_{21} + y z_{31})} \right),$$

then,

$$f_8(y) - (A_0 + B_0 y)^3 \geq 0, \quad \text{for every } y \geq 0;$$

These technical assumptions are used to guarantee that the microstructures constructed in Section 5.2 are energetically preferable to those constructed in Proposition 5.6.1 and Proposition 5.6.2 (see also Figure 5.7). Here by microstructure we mean the shape of a building block which is repeated quasi-periodically in configurations of low energy for I^ε . The period gets smaller with ε . The preferred microstructure clearly depends on the position of the wells, that is on z_1, z_2, z_3 , and on the cost of passing from one well to the other, that is on E_0, E_1 . (H6) and (H7) reduce to checking that two cubic polynomials are non-negative on $\mathbb{R}_+$. (H6)–(H8) can be verified easily with a computer and hold in a wide range of cases. We refer the reader to Section 5.6.1 for more details and for a couple of examples.

The first result that we prove is a second Γ –limit for $\bar{\mathcal{E}}$

Theorem 5.0.1. *Assume (H1)–(H8). Then $I^\varepsilon(u, \nu)$ Γ –converges in the $L^2(0, 1) \times L_{w*}^\infty(0, 1; \mathcal{M})$ topology to*

$$I^0(u, \nu) = \begin{cases} A_0 \int_0^1 \nu_x(z_2) dx + B_0 \int_0^1 \nu_x(z_3) dx, & \text{if } u = 0, \nu \in GYM^\infty(0), \\ & \text{supp } \nu \subset \mathcal{Z} \text{ a.e.,} \\ +\infty, & \text{otherwise.} \end{cases}$$

We remark that, as $\nu \in \text{GYM}^\infty(0)$, and $\text{supp } \nu \subset \mathcal{Z}$ a.e., we must have

$$\int_0^1 \nu_x(z_3) dx = z_{31}^{-1} - \frac{z_{21}}{z_{31}} \int_0^1 \nu_x(z_2) dx.$$

On the other hand $A_0 < \frac{z_{21}}{z_{31}} B_0$, so that $I^0(0, \nu)$ is a linearly decreasing function of $\int_0^1 \nu_x(z_2) dx$. Therefore, the minimum for I^0 is attained for

$$\int_0^1 \nu_x(z_3) dx = 0, \quad \int_0^1 \nu_x(z_2) dx = z_{21}^{-1}.$$

Thus minimizing sequences for $\mathcal{E}^\varepsilon$ have gradients tending in measure to $\{z_1, z_2\}$, and z_3 is not seen in the limit. That is, the vanishing interfacial energy limit selects a unique minimizer out of the infinitely many minimizers of $\mathcal{E}^0$.

As shown in Section 5.6, (H7) and (H8) are necessary conditions to prove the above Γ -limit result. Nonetheless, it turns out that we can characterize the set of gradient Young measures generated by minimizing sequences for I^ε , even without the second Γ -limit for $\bar{\mathcal{E}}$. This is the result of the following theorem, where also (H6) is relaxed:

Theorem 5.0.2. *Assume (H1)–(H5) and $z_3 \leq 3|z_1|$. Then every gradient Young measure $\nu \in \text{GYM}^\infty(0)$ generated by a sequence u^{ε_j} of minimizers for $\mathcal{E}^{\varepsilon_j}$, satisfies*

$$\text{supp } \nu_x \in \{z_1, z_2\}, \quad \nu_x = \frac{z_2}{z_2 - z_1} \delta_{z_1} - \frac{z_1}{z_2 - z_1} \delta_{z_2}, \quad \text{a.e. } x \in (0, 1).$$

In this way we have shown that, in our case, even if the set of gradient Young measures minimizing $\mathcal{E}^0$ has infinitely many elements, its subset generated by minimizers for the regularized and rescaled problem I^ε , which are also minimizers for $\mathcal{E}^\varepsilon$, contains just one element.

Therefore, the one-dimensional model problem studied in this chapter confirms that vanishing interface energy can be used as a tool to select minimizing gradient Young measures. This suggests that for the three-dimensional problem E the set $\mathcal{B}$ is actually much smaller than $\mathcal{A}$. Furthermore, our results show that the shape of the second Γ -limit for E^ε might change with the shape of ϕ . Nonetheless, as in our model problem, it might be possible to characterize $\mathcal{B}$ independently of the second Γ -limit for E^ε .

The plan for the chapter is the following: in Section 5.1 we prove Proposition 5.0.1, in Section 5.2 and 5.3 we compute some upper and lower bounds for I^ε . Section 5.4 is devoted to prove Theorem 5.0.1, while Section 5.5 is devoted to prove Theorem

5.0.2. Finally, in Section 5.6 we sketch necessity of (H7)–(H8) and give an example where (H7)–(H8) hold, and one where they don't.

In the following sections we will denote by c a generic positive constant depending only on the parameters of the problem, and not on the quantities $N, M, \eta, \varepsilon, \mu, j$ appearing below. Its value may change from line to line or even within the same line.

5.1 Proof of the first Γ –limit

In this section we prove Proposition 5.0.1.

We first observe that, as $\bar{\mathcal{E}}^\varepsilon(u, \nu)$ is a monotone sequence in ε , the Γ –limit exists and is given by the lower semicontinuous envelope of the pointwise limit of the sequence (cf. [25, Remark 1.40]), that is, the Γ –limit is given by

$$sc \begin{cases} \int_0^1 (\langle \nu_x, W \rangle + u^2) dx, & \text{if } u \in V, \nu_x = \delta_{u_x(x)} \text{ a.e. in } (0, 1), \\ +\infty, & \text{otherwise,} \end{cases} \quad (5.6)$$

where sc denotes the lower semicontinuous envelope with respect to the topology $L^2(0, 1) \times L_{w^*}^\infty(0, 1; \mathcal{M})$. We first claim that (5.6) is equal to

$$sc \begin{cases} \int_0^1 (\langle \nu_x, W \rangle + u^2) dx, & \text{if } u \in W^{1,p}(0, 1), \nu_x \in GYM^p(u), \\ +\infty, & \text{otherwise,} \end{cases} \quad (5.7)$$

that is we can relax the requirement $u \in V, \nu_x = \delta_{u_x(x)}$ a.e. in $(0, 1)$. Indeed, given an $u \in W_0^{1,p}(0, 1)$, we can approximate it by $u^j \in H^2(0, 1)$ such that $u^j \rightarrow u$ strongly in $W_0^{1,p}(0, 1)$. Therefore, by passing into the limit as j tends to ∞ we can drop the requirement $u \in H^2(0, 1)$ in (5.6). Now, let $\nu \in GYM^p(u)$ for some $u \in W_0^{1,p}(0, 1)$. Then by [69, Thm. 8.7] we know the existence of a sequence $u^j \in W^{1,p}(0, 1)$ converging weakly to u in $W^{1,p}(0, 1)$, strongly in $L^2(0, 1)$, such that $\delta_{u_x^j}$ converges to ν in $L_{w^*}^\infty(0, 1; \mathcal{M})$. Thanks to [69, Lemma 8.3] the sequence can actually be chosen in $W_0^{1,p}(0, 1)$. Therefore, the fact that (cf. [69, Thm. 6.11])

$$\liminf_j \int_0^1 \langle \delta_{u_x^j}, W \rangle dx \geq \int_0^1 \langle \nu_x, W \rangle dx,$$

allows us to drop also the requirement on ν that $\nu_x = \delta_{u_x(x)}$ in (5.6), concluding the proof that (5.6) is equal to (5.7). We now claim that we can drop sc from (5.7), that means, that

$$\begin{cases} \int_0^1 (\langle \nu_x, W \rangle + u^2) dx, & \text{if } u \in W^{1,p}(0, 1), \nu_x \in GYM^p(u), \\ +\infty, & \text{otherwise,} \end{cases} \quad (5.8)$$

is already lower semicontinuous in the $L^2(0, 1) \times L_{w^*}^\infty(0, 1; \mathcal{M})$ topology. To prove this claim, it is sufficient to show that for every sequence $(u_j, \nu^j) \in L^2(0, 1) \times \text{GYM}(u_j)$ converging to (u, ν) in $L^2(0, 1) \times L_{w^*}^\infty(0, 1; \mathcal{M})$, we have $\liminf_j \mathcal{E}^0(u_j, \nu^j) \geq \mathcal{E}^0(u, \nu)$. We will follow the approach devised in [20]. If $\liminf_j \mathcal{E}^0(u_j, \nu^j) = \infty$, the thesis follows trivially. Therefore, by passing without loss of generality to a subsequence, we can assume $\mathcal{E}^0(u_j, \nu^j) \leq C$. By (H2), this implies that

$$\int_0^1 \langle \nu_x^j, |\cdot|^p \rangle dx \leq C,$$

and, by [77, Thm. 3.6], we deduce that ν_x is a probability measure for almost every $x \in (0, 1)$. Jensen's inequality and the fact that $|\cdot|^p$ is convex yield

$$\int_0^1 |\bar{\nu}_x^j|^p dx \leq \int_0^1 \langle \nu_x^j, |\cdot|^p \rangle dx \leq C, \quad \text{where} \quad \bar{\nu}^j := \int_{\mathbb{R}} s d\nu^j(s).$$

It follows that $\bar{\nu}^j \rightharpoonup \bar{\nu}$ in $L^p(0, 1)$ and, therefore, that $u_j \rightharpoonup u$ in $W^{1,p}(0, 1)$, where $u_x = \bar{\nu}$. A result like the one in [77, Prop. 4.5] finally gives us that $\nu \in \text{GYM}^p(u)$. At this point, an application of [77, Prop. 3.7] allows us to deduce that $\liminf_j \mathcal{E}^0(u_j, \nu^j) \geq \mathcal{E}^0(u, \nu)$, thus concluding the proof.

Remark 5.1.1. Following the same strategy it is actually possible to prove that E^ε and $\tilde{E}^\varepsilon$ Γ -converge in the $L^1(\Omega) \times L_{w^*}^\infty(\Omega; \mathcal{M}_1(\mathbb{R}^{3 \times 3}))$ topology to $\bar{E}$ as $\varepsilon \rightarrow 0$.

5.2 Construction of an upper bound

In this section we prove the following proposition:

Proposition 5.2.1. *Assume (H1)–(H5). There exist $\zeta > 0$ and $\varepsilon_0 > 0$, such that for every $\lambda_1, \lambda_2, \lambda_3 \in [0, 1]$ satisfying*

$$\lambda_1 + \lambda_2 + \lambda_3 = 1, \quad z_1 \lambda_1 + z_2 \lambda_2 + z_3 \lambda_3 = 0, \quad (5.9)$$

and every $\varepsilon \leq \varepsilon_0$ we can find $u \in V$ with

$$I^\varepsilon(u) \leq (A_0 \lambda_2 + B_0 \lambda_3) + c\varepsilon^\zeta. \quad (5.10)$$

Furthermore, for every $\sigma \in (\varepsilon^{\frac{1}{\max\{3, q\}}}, \eta_0)$

$$|\mathcal{L}((0, 1) \cap \{|u_x - z_2| \leq \sigma\}) - \lambda_2| + |\mathcal{L}((0, 1) \cap \{|u_x - z_3| \leq \sigma\}) - \lambda_3| \leq c\varepsilon^\zeta. \quad (5.11)$$

Proof. Here we generalise the approach devised in [60]. Let us divide the interval $(0, 1)$ into the subintervals $(0, l_0)$, with $l_0 := (1 - \frac{z_2}{z_1})\lambda_2 = z_{21}\lambda_2$ and $(l_0, 1)$, noticing that, by (5.9), $1 - l_0 = (1 - \frac{z_3}{z_1})\lambda_3 = z_{31}\lambda_3$. We first construct the bit of u with energy $A_0\lambda_2$ in $(0, l_0)$, and then use the same argument to construct on $(l_0, 1)$ the bit of u which has energy $B_0\lambda_3$.

We start by splitting the interval $(0, l_0)$ into N pieces of length $l_N := \frac{z_{21}\lambda_2}{N} = \frac{l_0}{N}$. Let us also consider $\hat{w}(x)$, solution of

$$\varepsilon^3 \hat{w}_x = \sqrt{W(\hat{w})}, \quad \hat{w}(0) = 0. \quad (5.12)$$

Standard ODE theory tells us that $\hat{w}$ exists, and that $\hat{w}$ is strictly increasing with x when $\hat{w}(x) \in (z_1, z_2)$. We point out that, in case $q < 2$, the solution might not be unique. In this case, when solutions encounter z_1 or z_2 we choose the one that stays bounded in $[z_1, z_2]$ and does not decrease/increase further. As $\hat{w}(x - \omega)$ still satisfies the equation in (5.12) for every $\omega \in \mathbb{R}$, we will choose $\omega = \omega^*$ so that

$$F(\omega^*) := \int_0^{l_N} \hat{w}(s - \omega^*) ds = 0. \quad (5.13)$$

Indeed, this is possible as F is negative for $\omega \rightarrow \infty$, positive when $\omega \rightarrow -\infty$, continuous and decreasing. Now we define w as

$$w(x) = \begin{cases} \hat{w}(x - \omega^* - x_i), & \text{if } x \in (x_i, x_{i+1}) \text{ when } i \text{ even,} \\ \hat{w}(x_{i+1} - \omega^* - x), & \text{if } x \in (x_i, x_{i+1}) \text{ when } i \text{ odd,} \end{cases} \quad (5.14)$$

where $x_i := il_N$ for $i = 0, \dots, N$. We are now ready to construct u as

$$u(x) := \int_0^x w(s) ds, \quad (5.15)$$

and to notice that $u(x_i) = 0$ for each $i = 0, \dots, N$ by (5.13). By (5.12) we have

$$\int_{x_i}^{x_{i+1}} (\varepsilon^4 u_{xx}^2 + \varepsilon^{-2} W(u_x)) dx = 2\varepsilon \int_{x_i}^{x_{i+1}} |W(u_x)|^{\frac{1}{2}} |u_{xx}| dx \leq 2\varepsilon \int_{z_1}^{z_2} |W(s)|^{\frac{1}{2}} ds = \varepsilon E_0.$$

On the other hand, called x_i^* the point in (x_i, x_{i+1}) such that $u_x(x_i^*) = 0$, and assuming without loss of generality that $u_x > 0$ in (x_i, x_i^*) (the case $u_x < 0$ is similar), we have

$$\int_{x_i}^{x_{i+1}} u^2 dx \leq \int_{x_i}^{x_i^*} (z_2(x - x_i))^2 dx + \int_{x_i^*}^{x_{i+1}} (z_1(x_{i+1} - x))^2 dx = 3^{-1} (z_2^2 \alpha^3 + z_1^2 \gamma^3) l_N^3, \quad (5.16)$$

where

$$\alpha := \mathcal{L}((x_0, x_1) \cap \{w \geq 0\})l_N^{-1}, \quad \gamma := \mathcal{L}((x_0, x_1) \cap \{w < 0\})l_N^{-1}.$$

Therefore,

$$\begin{aligned} \int_0^{l_0} (\varepsilon^4 u_{xx}^2 + \varepsilon^{-2} W(u_x) + \varepsilon^{-2} u^2) dx &\leq N \left(3^{-1} \varepsilon^{-2} (z_2^2 \alpha^3 + z_1^2 \gamma^3) l_N^3 + \varepsilon E_0 \right) \\ &= l_0 \left(3^{-1} (z_2^2 \alpha^3 + z_1^2 \gamma^3) d_\varepsilon^2 + E_0 d_\varepsilon^{-1} \right), \end{aligned} \quad (5.17)$$

where $d_\varepsilon = \frac{l_N}{\varepsilon}$. Now, chosen $\eta \in (0, \eta_0)$ we notice that

$$0 = \int_0^{l_N} \hat{w}(s - \omega^*) ds \leq (z_2 \alpha + (z_1 + \eta) \gamma) l_N + r, \quad (5.18)$$

with $r := -(z_1 + \eta) \mathcal{L}(\{s: \hat{w}(s - \omega^*) \in (z_1 + \eta, 0)\})$. But r can be estimated as follows: we can rewrite (5.12) in terms of $\hat{v} := \hat{w} - z_1$ as

$$\varepsilon^3 \hat{v}_y = -\sqrt{W(\hat{v} + z_1)}, \quad \hat{v}(0) = -z_1,$$

where we also made the change of variable $y = -x$. Now, called y^* the point in $\mathbb{R}_+$ where $\hat{v}(y^*) = \eta_0$, by (H5) we have

$$\varepsilon^3 \hat{v}_y(y) \leq -\hat{c}, \quad \text{for all } y \in (0, y^*],$$

for some $\hat{c} > 0$. After an integration in y between 0 and y^* , this leads to

$$y^* \leq \hat{c}^{-1} \varepsilon^3 (|z_1| - \eta_0) \leq c \varepsilon^3.$$

In the same way, when $\eta \leq \hat{v} < \eta_0$

$$\varepsilon^3 \hat{v}_y \leq -c_0 |\hat{v}|^{\frac{q}{2}} \leq -c_0 |\hat{v}|^{\frac{\max\{3, q\}}{2}},$$

where we made also use of (H5). Let us now denote by $\tilde{y}$ the point in $\mathbb{R}^+$ such that $\hat{v}(\tilde{y}) = \eta$. An integration between y^* and $\tilde{y}$ yields

$$\tilde{y} - y^* \leq c \varepsilon^3 \eta^{1 - \frac{\max\{3, q\}}{2}}.$$

Thus, as $\mathcal{L}(\{s: \hat{w}(s - \omega^*) \in (z_1 + \eta, 0)\}) = \tilde{y}$, we have obtained

$$|r| \leq |z_1| \tilde{y} \leq c \varepsilon^3 \eta^{1 - \frac{\max\{3, q\}}{2}}. \quad (5.19)$$

This together with (5.18) thus imply

$$\gamma \leq \frac{z_2}{|z_1|} \alpha + c \bar{r}, \quad (5.20)$$

where $\bar{r} := \eta + N\varepsilon^3\eta^{1-\frac{\max\{3,q\}}{2}}$. On the other hand,

$$0 = \int_0^l \hat{w}(s - \omega^*) \, ds \geq ((z_2 - \eta)\alpha + z_1\gamma)l_N + r_1,$$

where now $r_1 := (z_2 - \eta)\mathcal{L}(\{s: \hat{w}(s - \omega^*) \in (0, z_2 - \eta)\})$. By arguing as to get (5.20), we have

$$\frac{z_2}{|z_1|}\alpha \leq \gamma + c\bar{r},$$

and, as $\alpha + \gamma = 1$, by (5.20) we thus deduce

$$\frac{z_1}{z_1 - z_2} - c\bar{r} \leq \alpha \leq \frac{z_1}{z_1 - z_2} + c\bar{r}.$$

The fact that, by construction, $l_0 \frac{z_1}{z_1 - z_2} = \lambda_2$ also implies

$$\lambda_2 - c\bar{r} \leq \alpha l_0 \leq \lambda_2 + c\bar{r}. \quad (5.21)$$

We now choose N as the smallest even integer larger than $l_0(\varepsilon d^*)^{-1}$, where

$$d^* = (3E_0)^{\frac{1}{3}} \left(2\lambda_2^3 l_0^{-3} z_2^2 \left(1 - \frac{z_2}{z_1} \right) \right)^{-\frac{1}{3}}.$$

In this way, $d_\varepsilon \leq d^*$, $d_\varepsilon^{-1} \leq c\varepsilon + (d^*)^{-1}$ and $N \leq c\varepsilon^{-1}$. Let us also choose $\eta = \varepsilon^{\frac{4}{\max\{3,q\}}}$, so that $\bar{r} \leq c\varepsilon^{\frac{4}{\max\{3,q\}}}$ and $\varepsilon^{\frac{4}{\max\{3,q\}}} < 1$ for each $\varepsilon \leq \varepsilon_0 < 1$. By exploiting (5.20) in (5.17) we thus get

$$\begin{aligned} \int_0^{l_0} (\varepsilon^4 u_{xx}^2 + \varepsilon^{-2} W(u_x) + \varepsilon^{-2} u^2) \, dx &\leq l_0 \left(3^{-1} \alpha^3 z_2^2 \left(1 - \frac{z_2}{z_1} \right) d_\varepsilon^2 + E_0 d_\varepsilon^{-1} \right) + c\tilde{r} \\ &\leq l_0 \left(3^{-1} \lambda_2^3 l_0^{-3} z_2^2 \left(1 - \frac{z_2}{z_1} \right) d_\varepsilon^2 + E_0 d_\varepsilon^{-1} \right) + c\tilde{r} \quad (5.22) \\ &\leq \lambda_2 A_0 + c\tilde{r}. \end{aligned}$$

Here and below $\tilde{r} := \varepsilon^{\frac{4}{\max\{3,q\}}} + \varepsilon$. We remark that α depends on ε , but for every $\sigma \in (\varepsilon^{\frac{1}{\max\{3,q\}}}, \eta_0)$, we have that

$$\mathcal{L}((0, l_0) \cap \{|u_x - z_2| \leq \sigma\}) = \alpha l_0 - R, \quad (5.23)$$

where $R = N\mathcal{L}((x_0, x_1) \cap \{s: 0 < w(s) < z_2 - \sigma\})$. By arguing as in the proof of (5.19) with η replaced by σ , we have that $|R| \leq cN\sigma^{1-\max\{3,q\}}\varepsilon^3 \leq c\sigma^{-\max\{3,q\}}\varepsilon^2$. Thus, since we assumed $\sigma \geq \varepsilon^{\frac{1}{\max\{3,q\}}}$, we deduce $|R| \leq c\varepsilon$. Therefore, by recalling (5.21) with $\eta = \varepsilon^{\frac{4}{\max\{3,q\}}}$, from (5.23) we finally obtain

$$|\mathcal{L}((0, l_0) \cap \{|u_x - z_2| \leq \sigma\}) - \lambda_2| \leq c\varepsilon^{\frac{4}{\max\{3,q\}}} + c\varepsilon. \quad (5.24)$$

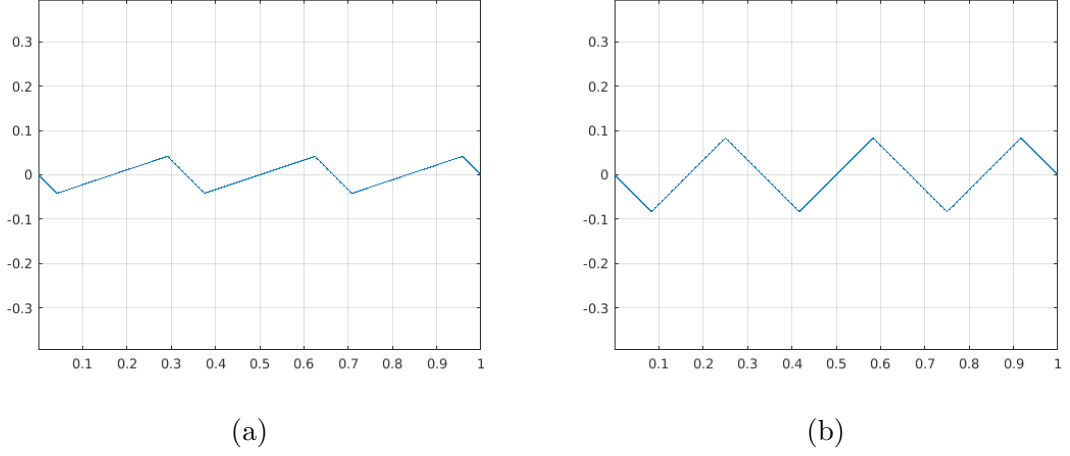


Figure 5.2: Piecewise approximation of the constructed function: in Figure 5.2a we show the function constructed in $(0, l_0)$, whose gradient oscillates between z_1 and z_2 . In Figure 5.2b we show the function constructed in $(l_0, 1)$, whose gradient oscillates between z_1 and z_3 .

Let us now focus on the interval $(l_0, 1)$, where we want to construct the part of u related to the B_0 -term of the energy in (5.10). This part of the argument is very similar to the one above, but, as there might be no solution to (5.12) connecting z_1 to z_3 , this time we need to construct an u whose gradient is slightly more complicated. Below, we try to highlight the differences from the case above without incurring into many repetitions. Let us consider $\tilde{w}$ to be the solution to

$$\varepsilon^3 \tilde{w}_x = \sqrt{W(\tilde{w})}, \quad \tilde{w}(s_0 + 2\mu^{\theta+1}) = z_2 + \mu, \quad (5.25)$$

where $s_0 > 0$ is such that $\hat{w}(s_0) = z_2 - \mu$, $\hat{w}$ is as in (5.12), and $\theta = \frac{3}{2}(\max\{q, 3\} - 2)$. Here and below $\mu = \varepsilon^{\frac{2}{\max\{3, q\} - 2}}$, so that $\mu^\theta = \varepsilon^3$. We remark that an argument as the one to prove (5.19) yields

$$s_0 \leq c\varepsilon^3 \mu^{1 - \frac{\max\{3, q\}}{2}} \leq c\varepsilon^2,$$

so that s_0 does not explode but actually goes to zero faster than ε . Again, if $q < 2$ $\hat{w}$ might not be unique, but we choose the one which stays bounded in $[z_2, z_3]$. Let us define v as

$$v(s) = \begin{cases} \hat{w}(s), & \text{if } s \leq s_0, \\ \mu^{-\theta}(s - s_0) + (z_2 - \mu), & \text{if } s_0 < s \leq s_0 + 2\mu^{\theta+1}, \\ \tilde{w}(s), & \text{if } s_0 + 2\mu^{\theta+1} < s. \end{cases}$$

Again, we divide $(l_0, 1)$ into M subintervals of equal length $l_M := M^{-1}(1 - l_0)$, and notice that, as v is strictly increasing, we can find ω_* such that

$$\int_0^{l_M} v(s - \omega_*) \, ds = 0.$$

As in the previous part of the proof, we construct

$$w(x) = \begin{cases} v(x - \omega_* - y_i), & \text{if } x \in (y_i, y_{i+1}) \text{ when } i \text{ even and } i \neq 0, \\ v(y_{i+1} - \omega_* - x), & \text{if } x \in (y_i, y_{i+1}) \text{ when } i \text{ odd,} \end{cases} \quad (5.26)$$

with $y_i := l_0 + il_M$, for $i = 0, \dots, M$, and u as in (5.15). We remark that, as in general

$$\lim_{s \uparrow l_0} w(s) = \hat{w}(-\omega^*) \neq \hat{w}(-\omega_*),$$

w needs to be defined differently in (y_0, y_1) in order to be continuous and to have $u \in H^2(0, 1)$. For this reason, we construct w as follows in (y_0, y_1) :

$$w(x) = \begin{cases} \hat{w}(-\omega^*) + \frac{s-y_0}{\rho}(z_1 - \hat{w}(-\omega^*)), & \text{if } x \in (y_0, y_0 + \rho), \\ z_1, & \text{if } x \in (y_0 + \rho, y_0 + \rho + a), \\ z_1 + \frac{s-y_0-\rho-a}{\rho}(z_3 - z_1), & \text{if } x \in (y_0 + \rho + a, y_0 + 2\rho + a), \\ z_3, & \text{if } x \in (y_0 + 2\rho + a, y_1 - \rho), \\ \hat{w}(-\omega_*) + \frac{y_1-s}{\rho}(z_3 - \hat{w}(-\omega_*)), & \text{if } x \in (y_1 - \rho, y_1), \end{cases}$$

where $\rho = \varepsilon^3$, and a is such that $\int_{y_0}^{y_1} w(s) \, ds = 0$. We point out that such a exists for each $\varepsilon \leq \varepsilon_0$, for some $\varepsilon_0 < 1$ depending on z_1 and z_3 only. After defining u in (y_0, y_1) as in (5.15), we have $u(y_0) = u(y_1) = 0$ and

$$\begin{aligned} \int_{y_0}^{y_1} u^2(s) \, ds &\leq (\max\{|z_1|, z_3\})^2 \int_{y_0}^{y_1} (s - y_0)^2 \, ds \leq cl_M^3, \\ \int_{y_0}^{y_1} W(u_x(s)) \, ds &\leq c\rho, \\ \int_{y_0}^{y_1} |u_{xx}(s)|^2 \, ds &\leq c\rho^{-1}. \end{aligned}$$

Thus,

$$\int_{y_0}^{y_1} (\varepsilon^4 u_{xx}^2 + \varepsilon^{-2} W(u_x) + \varepsilon^{-2} u^2) \, ds \leq c(\varepsilon + \varepsilon^{-2} M^{-3}).$$

On the other hand, if $i > 0$, by the definition of E_0, E_1 and by the way we constructed u we have

$$\begin{aligned} \int_{y_i}^{y_{i+1}} (\varepsilon^4 u_{xx}^2 + \varepsilon^{-2} W(u_x)) \, dx \\ \leq 2\varepsilon \left(\int_{z_1}^{z_2} |W(s)|^{\frac{1}{2}} \, ds + \int_{z_2}^{z_3} |W(s)|^{\frac{1}{2}} \, ds \right) + c\mu^{\theta+1} (\varepsilon^4 \mu^{-2\theta} + \varepsilon^{-2}) \\ \leq \varepsilon(E_0 + E_1) + c\mu\varepsilon. \end{aligned}$$

Furthermore, once defined

$$\beta := \mathcal{L}((y_1, y_2) \cap \{w \geq 0\}) l_M^{-1}, \quad \gamma := \mathcal{L}((y_1, y_2) \cap \{w < 0\}) l_M^{-1},$$

by arguing as in the proof of (5.16) we deduce

$$\int_{y_i}^{y_{i+1}} u^2 \leq 3^{-1} (z_3^2 \beta^3 + z_1^2 \gamma^3) l_M^3.$$

Therefore, collecting the inequalities above

$$\begin{aligned} \int_{l_0}^1 (\varepsilon^4 u_{xx}^2 + \varepsilon^{-2} W(u_x) + \varepsilon^{-2} u^2) \, dx \leq c(\mu\varepsilon M + \varepsilon + \varepsilon^{-2} M^{-3}) \\ + (1 - l_0) \left((E_0 + E_1) h_\varepsilon^{-1} + 3^{-1} (z_3^2 \beta^3 + z_1^2 \gamma^3) h_\varepsilon^2 \right), \end{aligned} \quad (5.27)$$

where now $h_\varepsilon = \frac{l_M}{\varepsilon}$. As in (5.18) we have

$$0 = \int_0^{l_M} v(s - \omega^*) \, ds \geq (z_1 \gamma + (z_3 - \mu) \beta) l_M - r_2,$$

with $r_2 := (z_3 - \mu) \mathcal{L}(\{s : v(s - \omega^*) \in (0, z_3 - \mu)\})$. We first notice that

$$r_2 = (z_3 - \mu) \left(\mu^{\theta+1} + s_0 + \mathcal{L}(\{s : \tilde{w}(s) \in (z_2 + \mu, z_3 - \mu)\}) \right).$$

Thus, by arguing as in the proof of (5.19) we first deduce

$$\mathcal{L}(\{s : \tilde{w}(s) \in (z_2 + \mu, z_3 - \mu)\}) \leq c\varepsilon^3 \mu^{1 - \frac{\max\{3, q\}}{2}},$$

and therefore

$$|r_2| \leq c(\varepsilon^3 \mu^{1 - \frac{\max\{3, q\}}{2}} + \mu^{\theta+1}) \leq c\varepsilon^2. \quad (5.28)$$

Define $\bar{r}_M := M\varepsilon^2 + \mu$, then (5.28) implies

$$\frac{z_3}{|z_1|} \beta \leq \gamma + c\bar{r}_M. \quad (5.29)$$

In the same way, we can prove that

$$\gamma \leq \frac{z_3}{|z_1|} \beta + c\bar{r}_M, \quad (5.30)$$

and, recalling that $\beta + \gamma = 1$, by (5.29) we obtain

$$\frac{z_1}{z_1 - z_3} - c\bar{r}_M \leq \beta \leq \frac{z_1}{z_1 - z_3} + c\bar{r}_M.$$

Since, by construction, $(1 - l_0) \frac{z_1}{z_1 - z_3} = \lambda_3$, we also have

$$\lambda_3 - c\bar{r}_M \leq \beta(1 - l_0) \leq \lambda_3 + c\bar{r}_M. \quad (5.31)$$

Then, after choosing M to be the smallest integer larger than $(1 - l_0)(\varepsilon d_2^*)^{-1}$, with

$$d_2^* := (3(E_0 + E_1))^{\frac{1}{3}} \left(2\lambda_3^3(1 - l_0)^{-3} z_3^2 \left(1 - \frac{z_3}{z_1} \right) \right)^{-\frac{1}{3}},$$

and exploiting (5.30)–(5.31), (5.27) becomes

$$\begin{aligned} & \int_{l_0}^1 (\varepsilon^4 u_{xx}^2 + \varepsilon^{-2} W(u_x) + \varepsilon^{-2} u^2) \, dx \\ & \leq (1 - l_0) \left((E_0 + E_1) h_\varepsilon^{-1} + 3^{-1} z_3^2 \lambda_3^3 (1 - l_0)^{-3} \left(1 - \frac{z_3}{z_1} \right) h_\varepsilon^2 \right) + c\hat{r} \\ & \leq \beta B_0 (1 - l_0) + c\hat{r} \\ & \leq B_0 \lambda_3 + c\hat{r}, \end{aligned}$$

where $\hat{r} = \varepsilon + \mu$. Here, we repeatedly used the fact that $M \leq c\varepsilon^{-1}$, $M^{-1} \leq c\varepsilon$ and that $\mu, \varepsilon < 1$ in order to estimate the above error. This together with (5.22) proves (5.10).

Now, since $\sigma > \varepsilon^{\frac{1}{\max\{3, q\}}} > \mu$,

$$\begin{aligned} & \mathcal{L}((l_0, 1) \cap \{|u_x - z_2| \leq \sigma\}) \\ & \leq \mathcal{L}((l_0, 1) \cap \{|u_x - z_2| \leq \mu\}) + \mathcal{L}((l_0, 1) \cap \{\mu \leq |u_x - z_2| \leq \sigma\}) \\ & \leq cM(\mu^{\theta+1} + \varepsilon^3 \sigma^{1 - \frac{\max\{3, q\}}{2}}) + c\rho \leq c(\varepsilon^2 \sigma^{-\max\{3, q\}} + \varepsilon^{-1} \mu^{\theta+1}) + c\rho \leq c\varepsilon, \end{aligned}$$

where we argued as to get (5.19) in order to bound $\mathcal{L}((l_0, 1) \cap \{\mu \leq |u_x - z_2| \leq \sigma\})$. Furthermore,

$$\mathcal{L}((l_0, 1) \cap \{|u_x - z_3| \leq \sigma\}) = \beta(1 - l_0) + R_2 - \beta l_M + \mathcal{L}((y_0, y_1) \cap \{|u_x - z_3| \leq \sigma\}),$$

where $R_2 := (M - 1) \mathcal{L}((y_1, y_2) \cap \{v(s - \omega^*) \in (0, z_3 - \sigma)\})$. As $R_2 \leq cM|r_2| \leq c\varepsilon$ and $l_M \leq c\varepsilon$ we have

$$|\mathcal{L}((l_0, 1) \cap \{|u_x - z_3| \leq \sigma\}) - \beta(1 - l_0)| \leq c\varepsilon.$$

Now, (5.31) implies

$$|\lambda_3 - \beta(1 - l_0)| \leq c\bar{r}_M \leq c\hat{r},$$

and hence, by the triangular inequality,

$$|\mathcal{L}((l_0, 1) \cap \{|u_x - z_3| \leq \sigma\}) - \lambda_3| \leq c\hat{r}.$$

This together with (5.24) lead to the second statement of the result. □

5.3 Construction of a lower bound

This section is the core of this chapter, and is where we prove a lower bound for the energy depending on the volume fractions $\lambda_k^\eta(v)$ defined in (5.32) below. We point out that the presence of a third well gives the possibility of many different microstructures (see e.g., figure 5.2b and figure 5.7), and makes the estimates below long and technical.

Let $\eta_0 > 0$ be as in (H4). Given a generic $v \in H^2(0, 1)$, $\eta \in (0, \eta_0)$ let us define the k -th volume fraction for v as

$$\lambda_k^\eta(v) := \mathcal{L}(\{x \in (0, 1) : |v_x(x) - z_k| \leq \eta\}), \quad k = 1, 2, 3, \quad (5.32)$$

and let us also generalize the definition of transition layers given in [60]

Definition 5.3.1. *Let $v \in H^2(a, b)$ and $\eta \in (0, \eta_0)$. An interval (x^-, x^+) is called A_+^η -transition (resp. A_-^η -transition) layer for v if*

$$\begin{aligned} v_x(x) &\in (z_1 + \eta, z_2 - \eta), & \forall x \in (x^-, x^+), \\ v_x(x^-) &= z_1 + \eta, & (\text{resp. } v_x(x^+) = z_1 + \eta), \\ v_x(x^+) &= z_2 - \eta, & (\text{resp. } v_x(x^-) = z_2 - \eta). \end{aligned}$$

An interval (x^-, x^+) is called B_+^η -transition (resp. B_-^η -transition) layer for v if

$$\begin{aligned} v_x(x) &\in (z_2 + \eta, z_3 - \eta), & \forall x \in (x^-, x^+), \\ v_x(x^-) &= z_2 + \eta, & (\text{resp. } v_x(x^+) = z_2 + \eta), \\ v_x(x^+) &= z_3 - \eta, & (\text{resp. } v_x(x^-) = z_3 - \eta). \end{aligned}$$

Given a function $v \in H^2(0, 1)$ and $\eta \in (0, \eta_0)$ we denote by $\#A_+^\eta$ (or by $\#A_-^\eta$, $\#B_+^\eta$, $\#B_-^\eta$) the number of A_+^η –transition layers for v (resp. $\#A_-^\eta$, $\#B_+^\eta$, $\#B_-^\eta$ transition layers for v) in the interval $(0, 1)$. The number of transition layers of a function v with bounded energy can be controlled nicely:

Lemma 5.3.1. *Assume (H1)–(H5), $\eta < \eta_0$, and let $\varepsilon > 0$ and $v \in H^2(0, 1)$ be such that $I^\varepsilon(v) \leq C$. Then, there exists $c = c(C) > 0$ such that*

$$\max\{\#A_+^\eta, \#A_-^\eta, \#B_+^\eta, \#B_-^\eta\} \leq c\varepsilon^{-1}.$$

Proof. Let us first recall that, given $0 \leq a \leq b \leq 1$ we have

$$\int_a^b (\varepsilon^4 v_{xx}^2 + \varepsilon^{-2} W(v_x)) dx \geq 2\varepsilon \int_a^b |\sqrt{W(v_x)} v_{xx}| dx \geq 2\varepsilon |H(v_x(b)) - H(v_x(a))|, \quad (5.33)$$

where $H(s) = \int_0^s \sqrt{W(r)} dr$.

Now, let us restrict ourselves to the case of the $A_\pm^\eta$ –transition layers, as the proof for the $B_\pm^\eta$ –transition layers follows the same strategy. Let (x^-, x^+) be an $A_\pm^\eta$ –transition layer, then by (5.33) and the fact that $\eta < \eta_0$ we have

$$\int_{x^-}^{x^+} (\varepsilon^4 v_{xx}^2 + \varepsilon^{-2} W(v_x)) dx \geq 2\varepsilon (H(z_2 - \eta_0) - H(z_1 + \eta_0)) > \varepsilon \tilde{c},$$

for some positive constant $\tilde{c}$. Summing all the $A_\pm^\eta$ –transition layers we thus get

$$C \geq I^\varepsilon(v) \geq \varepsilon \tilde{c} (\#A_+^\eta + \#A_-^\eta),$$

which concludes the proof. $\square$

We can now introduce also the D –intervals, which are the intervals between an A_+^η –transition layer (y^-, x^-) and the first A_-^η –transition layer (y^+, x^+) in order of appearance in $(0, 1)$ after (y^-, x^-) (see Figure 5.3):

Definition 5.3.2. *Let $v \in H^2(a, b)$ and $\eta \in (0, \eta_0)$. Let (y^-, x^-) be an A_+^η –transition layer for v and (y^+, x^+) be an A_-^η –transition layer for v , with $x^- \leq x^+$. We say that (y^-, y^+) is a D –interval for v , if*

$$v_x(x) > z_1 + \eta, \quad \text{for each } x \in (y^-, y^+), \quad \text{and} \quad v_x(y^+) = v_x(y^-) = z_1 + \eta. \quad (5.34)$$

If (5.34) holds, the interval (x^-, x^+) is called a D –interval for v .

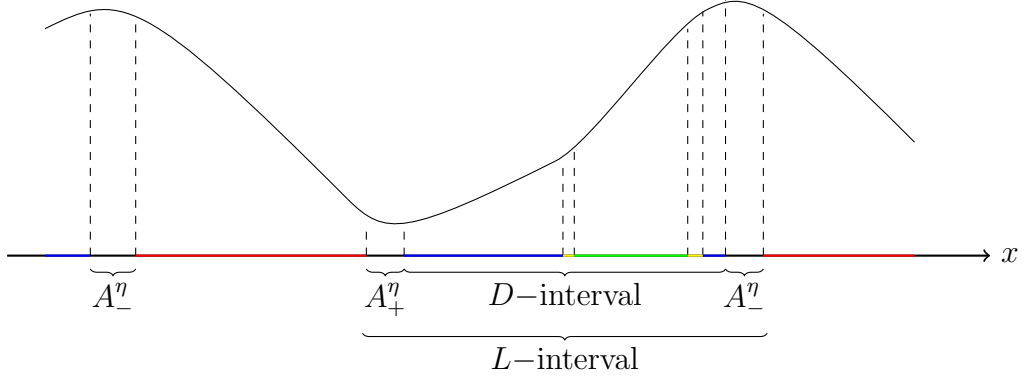


Figure 5.3: Example of L -interval and D -interval defined in Definition 5.3.2. In this picture, the red, the blue and the green intervals are respectively the sets of points where $|v_x - z_1| \leq \eta$, $|v_x - z_2| \leq \eta$, and $|v_x - z_3| \leq \eta$. The $B_{\pm}^{\eta}$ -transition layers are coloured in yellow.

It is important to notice that v_x might take negative values in a D -interval. For every $v \in V$, the number

$$N_v = \text{number of } D\text{-intervals for } v \text{ in } (0, 1),$$

is finite. Indeed, for every $v \in V$, v_x is continuous and N_v is equal to the number of A_+^{η} -transition layers, which is finite by Lemma 5.3.1. We denote by D_i the i -th D -interval in order of appearance in the interval $(0, 1)$, where i goes from 1 to N_v . This means that, given two D -intervals $D_i = (x_i^-, x_i^+)$ and $D_j = (x_j^-, x_j^+)$, we have that $x_i^+ < x_j^-$ if $i < j$. The same can be done for L -intervals. Given $v \in H^2(0, 1)$ we define also the following quantities

$$\begin{aligned} \alpha_i^{\eta}(v) &:= \mathcal{L}(\{x \in D_i : |v_x(x) - z_2| \leq \eta\}), \\ \beta_i^{\eta}(v) &:= \mathcal{L}(\{x \in D_i : |v_x(x) - z_3| \leq \eta\}), \end{aligned}$$

measuring the subset of D_i where v_x is respectively close to z_2 and z_3 . For ease of notation, we omit the dependence on v of $\alpha_i^{\eta}, \beta_i^{\eta}$ and λ_k^{η} . In what follows we will also drop the η from $\alpha_i^{\eta}, \beta_i^{\eta}$, keeping their dependence from this variable implicit. We remark that, denoting $D_1 = (x_1^-, x_1^+)$, $D_{N_v} = (x_{N_v}^-, x_{N_v}^+)$, we have

$$\begin{aligned} \sum_{i=1}^{N_v} \alpha_i + \mathcal{L}(\{x \in (0, x_1^-) \cup (x_{N_v}^+, 1) : |v_x(x) - z_2| \leq \eta\}) &= \lambda_2^{\eta}, \\ \sum_{i=1}^{N_v} \beta_i + \mathcal{L}(\{x \in (0, x_1^-) \cup (x_{N_v}^+, 1) : |v_x(x) - z_3| \leq \eta\}) &= \lambda_3^{\eta}, \end{aligned}$$

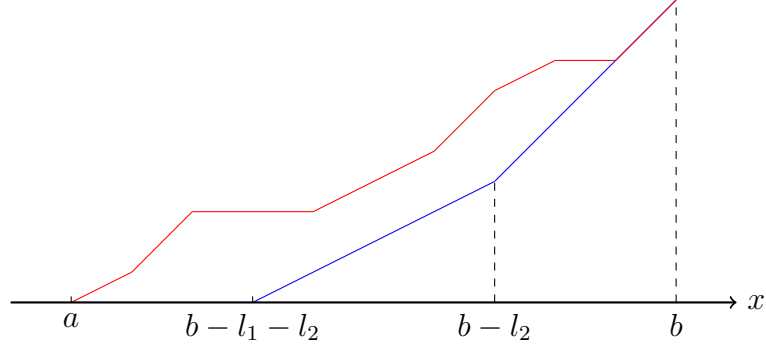


Figure 5.4: The aim of Lemma 5.3.2 is to show that, on the interval (a, b) , the L^2 -norm of the red function in this picture is bigger or equal than the L^2 -norm of the blue one. Here both, the red and the blue functions, are piecewise linear, with derivative in $\{0, \tau_1, \tau_2\}$ for almost every $x \in (a, b)$. Furthermore, $\mathcal{C}_1, \mathcal{C}_2$ are the sets where the derivative of the red function is respectively τ_1 and τ_2 . We have $l_1 = \mathcal{L}(\mathcal{C}_1)$, $l_2 = \mathcal{L}(\mathcal{C}_2)$, that is, the measure of the set where the blue curve has gradient τ_1 (resp. τ_2) is equal to the measure of the set where the blue curve has gradient τ_1 (resp. τ_2).

and that, in general,

$$\alpha_i + \beta_i < \mathcal{L}(D_i).$$

Below, we estimate the energy of a generic $v \in V$ on every D -interval in terms of the quantities α_i, β_i . In order to do that, we first need to prove the following lemma whose meaning is given in Figure 5.4

Lemma 5.3.2. *Let $0 \leq a < b \leq 1$, and let $\mathcal{C}_1, \mathcal{C}_2 \subset (a, b)$ be two Borel sets such that $\mathcal{C}_1 \cap \mathcal{C}_2 = \emptyset$. Let also $\tau_2 > \tau_1 \geq 0$. Then,*

$$\int_a^b \left(\sum_{i=1,2} \tau_i \mathcal{L}((a, x) \cap \mathcal{C}_i) \right)^2 dx \geq \int_0^{\mathcal{L}(\mathcal{C}_1)} (\tau_1 x)^2 dx + \int_0^{\mathcal{L}(\mathcal{C}_2)} (\tau_1 \mathcal{L}(\mathcal{C}_1) + \tau_2 x)^2 dx. \quad (5.35)$$

Proof. Let us define $l_i := \mathcal{L}(\mathcal{C}_i)$, $i = 1, 2$. The right hand side of (5.35) can be rewritten as

$$\int_0^{l_1+l_2} g^2(x) dx = \int_a^b \chi_{(b-l_1-l_2, b)}(x) g^2(x - b + l_1 + l_2) dx,$$

where $g(x) := \tau_1 x \chi_{(0, l_1]}(x) + ((\tau_1 - \tau_2)l_1 + \tau_2 x) \chi_{(l_1, l_1+l_2)}(x)$ and $\chi_{\mathcal{B}}$ is the indicator function on the Borel set $\mathcal{B}$. Therefore, in order to prove (5.35) it is sufficient to show that

$$\hat{g}(x) := \sum_{i=1,2} \tau_i \mathcal{L}((a, x) \cap \mathcal{C}_i) \geq \chi_{(b-l_1-l_2, b)}(x) g(x - b + l_1 + l_2),$$

for every $x \in (a, b)$. Suppose not, then there exists $\hat{x} \in (a, b)$ such that $\hat{g}(\hat{x}) < \chi_{(b-l_1-l_2, b)}(\hat{x})g(\hat{x}-b+l_1+l_2)$. As $g, \hat{g} \geq 0$ in (a, b) , $\hat{x}$ cannot be in $(a, b-l_1-l_2]$. Suppose without loss of generality that $\hat{x} \in (b-l_1-l_2, b-l_2]$. The same argument can be used to get a contradiction in the case where $\hat{x} \in (b-l_2, b)$. Then,

$$\hat{g}(\hat{x}) = \tau_1 l_1 + \tau_2 l_2 - \sum_{i=1,2} \tau_i \mathcal{L}((\hat{x}, b) \cap \mathcal{C}_i) < \tau_1(\hat{x} - b + l_1 + l_2) = g(\hat{x}),$$

that is,

$$\tau_2 l_2 + \tau_1(b-l_2-\hat{x}) < \sum_{i=1,2} \tau_i \mathcal{L}((\hat{x}, b) \cap \mathcal{C}_i).$$

On the other hand, as $\mathcal{C}_1, \mathcal{C}_2$ are disjoint, we have

$$\begin{aligned} \sum_{i=1,2} \tau_i \mathcal{L}((\hat{x}, b) \cap \mathcal{C}_i) &= \tau_2 l_2 + \tau_2(\mathcal{L}((\hat{x}, b) \cap \mathcal{C}_2) - l_2) + \tau_1 \mathcal{L}((\hat{x}, b) \cap \mathcal{C}_1) \\ &\leq \tau_2 l_2 + \tau_1(\mathcal{L}((\hat{x}, b) \cap (\mathcal{C}_1 \cup \mathcal{C}_2)) - l_2) \leq \tau_2 l_2 + \tau_1(b - \hat{x} - l_2). \end{aligned}$$

Combining the last two inequalities we thus get a contradiction. $\square$

Remark 5.3.1. By arguing as in the proof of Lemma 5.3.2, and under the same hypotheses, (5.35) can be easily generalised to

$$\int_a^b \left(\tau_0 + \sum_{i=1,2} \tau_i \mathcal{L}((a, x) \cap \mathcal{C}_i) \right)^2 dx \geq \int_0^{\mathcal{L}(\mathcal{C}_1)} (\tau_0 + \tau_1 x)^2 dx + \int_0^{\mathcal{L}(\mathcal{C}_2)} (\tau_0 + \tau_1 \mathcal{L}(\mathcal{C}_1) + \tau_2 x)^2 dx$$

for any $\tau_0 \geq 0$.

We can now start to estimate the energy in the D -intervals. Thanks to the following lemma, we can restrict our estimates to those D_i 's where α_i and β_i are neither too small nor too big.

Lemma 5.3.3. Assume (H1)–(H5), and let $v \in H^2(0, 1)$, $\eta \in (0, \eta_0)$ and $\varepsilon \leq \eta^q$ be such that $I^\varepsilon(v) \leq C$, with $C > 0$. Then, there exist $R_*, R^*(C) > 0$ such that for any $L_i = (y_i^-, y_i^+)$, with $\max\{\alpha_i, \beta_i\} \geq R^* \varepsilon$ or $\max\{\alpha_i, \beta_i\} \leq R_* \varepsilon$,

$$\int_{y_i^-}^{y_i^+} (\varepsilon^4 v_{xx}^2 + \varepsilon^{-2} W(v_x) + \varepsilon^{-2} v^2) dx \geq A_0 \alpha_i + B_0 \beta_i. \quad (5.36)$$

Proof. First we want to prove that if α_i or β_i is too big, then v is growing too much, and hence its L^2 -norm on D_i is big enough to show (5.36). In order to do this, let $\max\{\alpha_i, \beta_i\} = R\varepsilon$ for some $R > 0$. As

$$\int_{y_i^-}^{y_i^+} (\varepsilon^4 v_{xx}^2 + \varepsilon^{-2} W(v_x) + \varepsilon^{-2} v^2) dx \geq \varepsilon^{-2} \int_{x_i^-}^{x_i^+} v^2 dx,$$

we are interested in an estimate of the right hand side. We assume the existence of $x_i^* \in (x_i^-, x_i^+)$ such that $v(x_i^*) = 0$, but the following estimates hold in the case $v > 0$ (or $v < 0$) in (x_i^-, x_i^+) by taking $x_i^* = x_i^-$ (resp. $x_i^* = x_i^+$). Assume also without loss of generality that

$$\max\{\mathcal{L}((x_i^*, x_i^+) \cap \{|v_x - z_2| \leq \eta\}), \mathcal{L}((x_i^*, x_i^+) \cap \{|v_x - z_3| \leq \eta\})\} \geq \varepsilon \frac{R}{2}, \quad (5.37)$$

the alternative case where

$$\max\{\mathcal{L}((x_i^-, x_i^*) \cap \{|v_x - z_2| \leq \eta\}), \mathcal{L}((x_i^-, x_i^*) \cap \{|v_x - z_3| \leq \eta\})\} \geq \varepsilon \frac{R}{2},$$

can be proven similarly. We remark that, as (y_i^-, y_i^+) is a L -interval for v , $v_x(x) > z_1 + \eta$ for each $x \in (y_i^-, y_i^+)$. Therefore,

$$\begin{aligned} v(x) &\geq \int_{x_i^*}^x v_x \, dx \geq (z_2 - \eta) \mathcal{L}((x_i^*, x) \cap \{v_x - z_2 \geq -\eta\}) \\ &\quad + (z_1 + \eta) \mathcal{L}((x_i^*, x) \cap \{z_1 + \eta < v_x \leq 0\}). \end{aligned} \quad (5.38)$$

Now, we notice that

$$\mathcal{L}((x_i^*, x) \cap \{z_1 + \eta < v_x \leq 0\}) \leq \mathcal{L}(\Sigma^\eta),$$

where

$$\Sigma^\eta := \{x \in (0, 1) : |v_x - z_k| > \eta, \forall k = 1, 2, 3\}. \quad (5.39)$$

Thus, by the boundedness of $I^\varepsilon(v)$ and (H5), we can write

$$c_0 \mathcal{L}(\Sigma^\eta) \eta^q \leq \int_{\Sigma^\eta} W(v_x) \, dx \leq \int_0^1 W(v_x) \, dx \leq C \varepsilon^2,$$

which implies

$$\mathcal{L}(\Sigma^\eta) \leq c \varepsilon^2 \eta^{-q}. \quad (5.40)$$

It follows then from (5.38), (5.40) and $\varepsilon \leq \eta^q$ that

$$v(x) \geq \hat{c} \mathcal{L}((x_i^*, x) \cap \{v_x - z_2 \geq -\eta\}) - c \varepsilon, \quad (5.41)$$

for every $x \in (x_i^*, x_i^+)$ and some positive constant $\hat{c}$. Therefore, thanks to (5.37), together with Lemma 5.3.2 with $\tau_1 = 0$,

$$\begin{aligned} 2 \int_{x_i^*}^{x_i^+} v^2 \, dx + c \varepsilon^2 (x_i^+ - x_i^*) &\geq \hat{c}^2 \int_{x_i^*}^{x_i^+} \left(\mathcal{L}((x_i^*, x) \cap \{v_x - z_2 \geq -\eta\}) \right)^2 dx \\ &\geq \hat{c}^2 \int_0^{\varepsilon \frac{R}{2}} x^2 \, dx \geq \frac{\hat{c}^2}{8} \varepsilon^3 R^3, \end{aligned}$$

which, by using the fact that

$$(x_i^+ - x_i^*) \leq 2R\varepsilon + \mathcal{L}(\Sigma^\eta) \leq c\varepsilon(1 + R), \quad (5.42)$$

yields

$$\int_{y_i^-}^{y_i^+} (\varepsilon^4 v_{xx}^2 + \varepsilon^{-2} W(v_x) + \varepsilon^{-2} v^2) dx \geq \bar{c}\varepsilon R^3 - c\varepsilon(1 + R), \quad (5.43)$$

for some $\bar{c} > 0$. On the other hand, there exists $c^* > 0$ such that

$$A_0\alpha_i + B_0\beta_i \leq c^* R\varepsilon. \quad (5.44)$$

Therefore, setting R^* as the biggest root of $\bar{c}R^3 = (c + c^*)(R + 1)$, we deduce that, if $R \geq R^*$, (5.43)–(5.44) imply (5.36). On the other hand, we recall that in every L -interval there are exactly two $A_\pm^\eta$ -transition layers. By (5.33) we thus have

$$\begin{aligned} \int_{y_i^-}^{y_i^+} (\varepsilon^4 v_{xx}^2 + \varepsilon^{-2} W(v_x) + \varepsilon^{-2} v^2) dx &\geq \int_{y_i^-}^{y_i^+} (\varepsilon^4 v_{xx}^2 + \varepsilon^{-2} W(v_x)) dx \\ &\geq 4\varepsilon(H(z_2 - \eta_0) - H(z_1 + \eta_0)) \geq \varepsilon\tilde{c}, \end{aligned}$$

for some $\tilde{c} > 0$. Hence, if we set $R_* = \frac{\tilde{c}}{c^*}$, by (5.44) we deduce that (5.36) holds for every $R < R_*$. We remark that R^*, R_* do not depend on i, v, ε, η in any way. R_* does not depend on C either. $\square$

Let n_i be the even number of $B_\pm^\eta$ -transition layers for v in D_i . If $n_i = 2$, we denote the B_+^η and the B_-^η -transition layers for v in D_i respectively by $(z_{i,1}^-, z_{i,1}^+)$ and $(z_{i,2}^-, z_{i,2}^+)$, and we define E_i as $E_i := (z_{i,1}^+, z_{i,2}^-)$. We remark that, in general, $\beta_i < \mathcal{L}(E_i)$, but that $v_x(x) > z_2 + \eta$ for every $x \in E_i$. We now prove an estimate for the L^2 -norm of v in the D_i 's. The estimates are different for the different types of D -intervals given in the following definition (see Figure 5.5)

Definition 5.3.3. Let $v \in V$, $\eta \in (0, \eta_0)$ and $\varepsilon \leq \eta^q$ be such that $I^\varepsilon(v) \leq C$, with $C > 0$. Let $D_i = (x_i^-, x_i^+)$ be the i -th D -interval for v , and $n_i \in 2\mathbb{N}$ be the number of $B_\pm^\eta$ -transition layers for v in D_i . Then we say that D_i is of

- **type 0:** if $\max\{\alpha_i, \beta_i\} \notin (\varepsilon R_*, \varepsilon R^*)$;
- **type I:** if $\max\{\alpha_i, \beta_i\} \in (\varepsilon R_*, \varepsilon R^*)$ and $v(x_i^-)v(x_i^+) \geq 0$;
- **type II:** if $\max\{\alpha_i, \beta_i\} \in (\varepsilon R_*, \varepsilon R^*)$, $v(x_i^-)v(x_i^+) < 0$ and either $n_i = 0$ or $n_i \geq 4$;

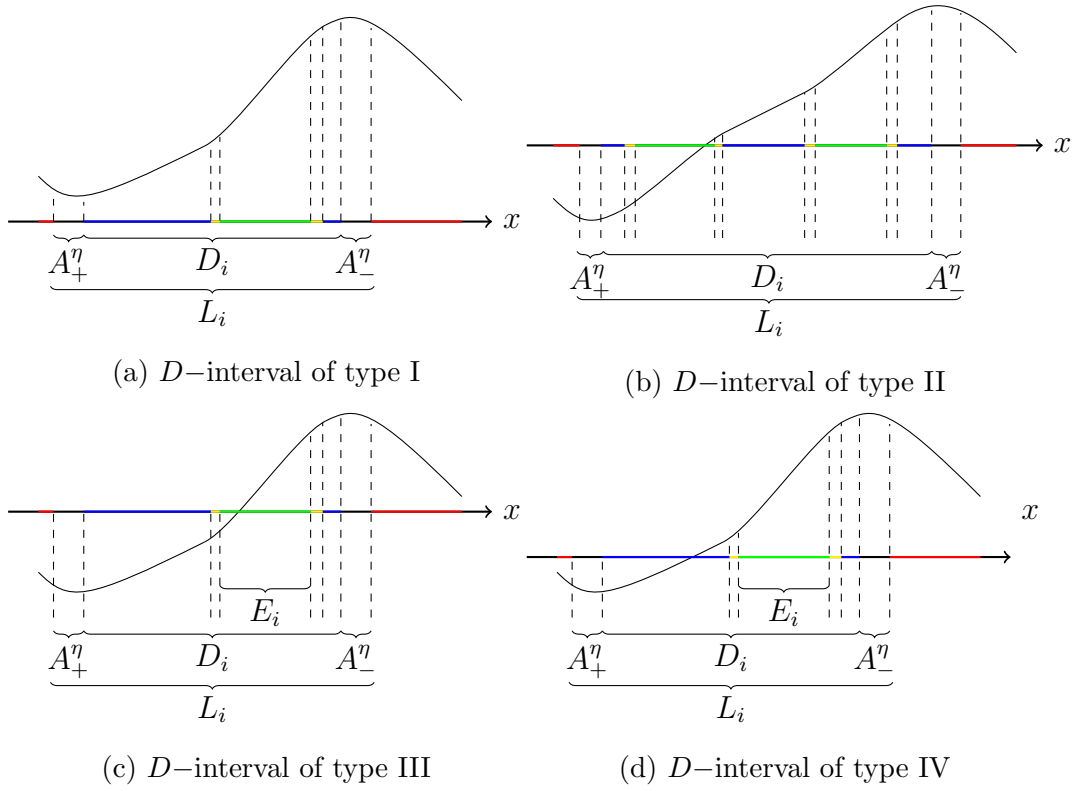


Figure 5.5: Classification of D -intervals. In this picture, the red, the blue and the green intervals are respectively the sets of points where $|v_x - z_1| \leq \eta$, $|v_x - z_2| \leq \eta$, and $|v_x - z_3| \leq \eta$. The $B_{\pm}^{\eta}$ -transition layers are coloured in yellow.

- **type III:** if $\max\{\alpha_i, \beta_i\} \in (\varepsilon R_*, \varepsilon R^*)$, $v(x_i^-)v(x_i^+) < 0$, $n_i = 2$ and there exists $x_i^* \in E_i$ such that $v(x_i^*) = 0$;
- **type IV:** if $\max\{\alpha_i, \beta_i\} \in (\varepsilon R_*, \varepsilon R^*)$, $v(x_i^-)v(x_i^+) < 0$, $n_i = 2$ and there exists no $x_i^* \in E_i$ such that $v(x_i^*) = 0$.

In this definition R_*, R^* are as in the statement of Lemma 5.3.3.

We are now ready to prove the following lemma:

Lemma 5.3.4. Assume (H1)–(H5), and let $\eta \in (0, \eta_0)$, $\varepsilon \leq \eta^{q+1}$, and $v \in H^2(0, 1)$ be such that $I^\varepsilon(v) \leq C$. Then, for every $D_i = (x_i^-, x_i^+)$, such that $\max\{\alpha_i, \beta_i\} \in (R_*\varepsilon, R^*\varepsilon)$, with R_*, R^* as in Lemma 5.3.3, we have

$$\int_{x_i^-}^{x_i^+} v^2 dx \geq 3^{-1}h(\alpha_i, \beta_i, 1) - c\varepsilon^3\eta, \quad \text{if } D_i \text{ is of type I,} \quad (5.45)$$

$$\int_{x_i^-}^{x_i^+} v^2 dx \geq 12^{-1}h(\alpha_i, \beta_i, 1) - c\varepsilon^3\eta, \quad \text{if } D_i \text{ is of type II,} \quad (5.46)$$

$$\int_{x_i^-}^{x_i^+} v^2 dx \geq 3^{-1}\hat{h}(\alpha_i, \beta_i, \omega_i^a, \omega_i^b) - c\varepsilon^3\eta, \quad D_i \text{ is of type III,} \quad (5.47)$$

$$\int_{x_i^-}^{x_i^+} v^2 dx \geq 3^{-1}h(\alpha_i, \beta_i, \omega_i^a) - c\varepsilon^3\eta, \quad D_i \text{ is of type IV,} \quad (5.48)$$

where

$$\begin{aligned} h(a, b, \omega_i^a) &:= z_2^2 a^3 + z_3^2 b^3 + 3(z_2^2 \omega_i^a (\omega_i^a - 1) a^3 + ab z_2 \omega_i^a (z_2 \omega_i^a \alpha_i + z_3 \beta_i)), \\ \hat{h}(a, b, \omega_i^a, \omega_i^b) &:= 3ab z_3 (\omega_i^a \omega_i^b (z_2 \omega_i^a \alpha_i + z_3 \omega_i^b \beta_i)) + 3(z_2^2 \omega_i^a (\omega_i^a - 1) a^3 + z_3^2 \omega_i^b (\omega_i^b - 1) b^3) \\ &\quad + z_2^2 a^3 + z_3^2 b^3 + 3ab z_3 ((1 - \omega_i^a)(1 - \omega_i^b)(z_2(1 - \omega_i^a)\alpha_i + z_3(1 - \omega_i^b)\beta_i)), \end{aligned}$$

and $\omega_i^a, \omega_i^b \in [0, 1]$.

Proof. Let us first focus on the case $v(x_i^-)v(x_i^+) < 0$. We notice that the continuity of v , implies the existence of $x_i^* \in (x_i^-, x_i^+)$ such that $v(x_i^*) = 0$. We estimate separately $\int_{x_i^*}^{x_i^+} v^2 ds$ and $\int_{x_i^-}^{x_i^*} v^2 ds$. Let us start with the first. As $v_x > z_1 + \eta$ in (x_i^-, x_i^+) , we have

$$v(x) \geq \int_{x_i^*}^x v_x dx \geq \sum_{k=2,3} (z_k - \eta) \mathcal{L}((x_i^*, x) \cap \{|v_x - z_k| \leq \eta\}) + (z_1 + \eta) \mathcal{L}(\Sigma^\eta) \quad (5.49)$$

where Σ^η is as in (5.39) and satisfies (5.40). It follows then from (5.49) and the assumption $\max\{\alpha_i, \beta_i\} \leq R^*\varepsilon$ that

$$v(x) \geq \sum_{k=2,3} z_k \mathcal{L}((x_i^*, x) \cap \{|v_x - z_k| \leq \eta\}) - C \frac{\varepsilon^2}{\eta^q} - 2\varepsilon\eta R^*, \quad (5.50)$$

for every $x \in (x_i^*, x_i^+)$. We now use the fact that, as $\eta \in (0, 1)$, $(1 - \eta)(a + b)^2 \leq a^2 + 2\eta^{-1}b^2$. This yields

$$(1 - \eta) \int_{x_i^*}^{x_i^+} (v(x) + c(\varepsilon^2 \eta^{-q} + \varepsilon \eta))^2 dx \leq \int_{x_i^*}^{x_i^+} v^2 dx + c\hat{r}(x_i^+ - x_i^*), \quad (5.51)$$

with $\hat{r} := \eta^{-1}(\varepsilon^2 \eta^{-q} + \varepsilon \eta)^2$. On the other hand, by Lemma 5.3.2,

$$\begin{aligned} \int_{x_i^*}^{x_i^+} (v(x) + c(\varepsilon^2 \eta^{-q} + \varepsilon \eta))^2 dx &\geq \int_{x_i^*}^{x_i^+} \left(\sum_{k=2,3} z_k \mathcal{L}((x_i^*, x) \cap \{|v_x - z_k| \leq \eta\}) \right)^2 dx \\ &\geq z_2^2 \int_0^{\omega_i^a \alpha_i} x^2 dx + \int_0^{\omega_i^b \beta_i} (z_3 x + z_2 \omega_i^a \alpha_i)^2 dx \\ &\geq 3^{-1} (z_2^2 (\omega_i^a \alpha_i)^3 + z_3^2 (\omega_i^b \beta_i)^3 + g_0(\alpha_i, \beta_i, \omega_i^a, \omega_i^b)) \end{aligned} \quad (5.52)$$

where, $\omega_i^{a,b} \in [0, 1]$ are such that

$$\mathcal{L}((x_i^*, x_i^+) \cap \{|v_x - z_2| \leq \eta\}) = \omega_i^a \alpha_i, \quad \mathcal{L}((x_i^*, x_i^+) \cap \{|v_x - z_3| \leq \eta\}) = \omega_i^b \beta_i, \quad (5.53)$$

and $g_0(a, b, \omega^a, \omega^b) = 3z_2 ab \omega^a \omega^b (z_2 a \omega^a + z_3 b \omega^b)$. Therefore, putting together (5.51)–(5.52), by (5.42) we obtain

$$\int_{x_i^*}^{x_i^+} v^2 dx + c\varepsilon(1 + R^*)\hat{r} \geq (1 - \eta)3^{-1} (z_2^2 (\omega_i^a \alpha_i)^3 + z_3^2 (\omega_i^b \beta_i)^3 + g_0(\alpha_i, \beta_i, \omega_i^a, \omega_i^b)).$$

Recalling $\max\{\alpha_i, \beta_i\} \leq R^* \varepsilon$ and noticing that $\varepsilon \leq \eta^{q+1}$ implies $\hat{r} \leq 4\varepsilon^2 \eta$ we hence have

$$\int_{x_i^*}^{x_i^+} v^2 dx + c\varepsilon^3 \eta \geq 3^{-1} (z_2^2 (\omega_i^a \alpha_i)^3 + z_3^2 (\omega_i^b \beta_i)^3 + g_0(\alpha_i, \beta_i, \omega_i^a, \omega_i^b)). \quad (5.54)$$

In the same way, we prove

$$\int_{x_i^-}^{x_i^*} v^2 dx + c\varepsilon^3 \eta \geq 3^{-1} (z_2^2 ((1 - \omega_i^a) \alpha_i)^3 + z_3^2 ((1 - \omega_i^b) \beta_i)^3 + g_0(\alpha_i, \beta_i, 1 - \omega_i^a, 1 - \omega_i^b)). \quad (5.55)$$

It turns out that if we sum (5.54) to (5.55) the right hand side is a convex quadratic polynomial in ω_i^a, ω_i^b with minimum in $\omega_i^a = \omega_i^b = \frac{1}{2}$. Therefore, optimizing over $\omega_i^{a,b}$, from (5.54) and (5.55) we finally get (5.46). The estimate (5.45) follows again from (5.54) (or (5.55)) if $v(x_i^-) > 0$ (resp. $v(x_i^+) < 0$) which can be deduced via the same argument by setting $x_i^* = x_i^-$ (resp. $x_i^* = x_i^+$). In this case, the above estimates hold with $\omega_i^a, \omega_i^b = 1$ (resp. $\omega_i^a, \omega_i^b = 0$).

Let us suppose now that $v(x_i^-)v(x_i^+) < 0$ and that $n_i = 2$. We notice that if $x_i^* \notin E_i$, it must hold $\omega_i^b \in \{0, 1\}$ and we get, after assuming without loss of generality that $\omega_i^b = 1$, (5.48). On the other hand, if $x_i^* \in E_i$, by using Lemma 5.3.2 first with $\tau_2 = z_3, \tau_1 = 0, \mathcal{A}_2 = \{x \in (x_i^*, x_i^+): |v_x(x) - z_3| \leq \eta\}$, $(a, b) = E_i$ and then with $\tau_2 = z_2, \tau_1 = 0, \mathcal{A}_2 = \{x \in (x_i^*, x_i^+): |v_x(x) - z_2| \leq \eta\}$ (cf. Remark 5.3.1), we can modify (5.52) as follows

$$\begin{aligned} & \int_{x_i^*}^{x_i^+} \left(\sum_{k=2,3} z_k \mathcal{L}((x_i^*, x) \cap \{|v_x - z_k| \leq \eta\}) \right)^2 dx \\ & \geq z_3^2 \int_0^{\omega_i^b \beta_i} x^2 dx + \int_0^{\omega_i^a \alpha_i} (z_2 x + z_3 \omega_i^b \beta_i)^2 dx \\ & \geq 3^{-1} (z_2^2 (\omega_i^a \alpha_i)^3 + z_3^2 (\omega_i^b \beta_i)^3 + \frac{z_3}{z_2} g_0(\alpha_i, \beta_i, \omega_i^a, \omega_i^b)), \end{aligned}$$

which by (5.51) becomes

$$\int_{x_i^*}^{x_i^+} v^2 dx + c\varepsilon^3 \eta \geq 3^{-1} (z_2^2 (\omega_i^a \alpha_i)^3 + z_3^2 (\omega_i^b \beta_i)^3 + \frac{z_3}{z_2} g_0(\alpha_i, \beta_i, \omega_i^a, \omega_i^b)).$$

In the same way, we prove

$$\begin{aligned} & \int_{x_i^-}^{x_i^*} v^2 dx + c\varepsilon^3 \eta \\ & \geq 3^{-1} (z_2^2 ((1 - \omega_i^a) \alpha_i)^3 + z_3^2 ((1 - \omega_i^b) \beta_i)^3 + \frac{z_3}{z_2} g_0(\alpha_i, \beta_i, 1 - \omega_i^a, 1 - \omega_i^b)) \end{aligned} \quad (5.56)$$

By summing up the last two inequalities, we hence get (5.45). We remark that, in this case, the determinant of the Hessian matrix of $\hat{h}$ with respect to ω_i^a, ω_i^b is negative, and hence $\hat{h}$ cannot be bounded from below by choosing $\omega_i^a = \omega_i^b = \frac{1}{2}$. $\square$

The lower bound for $\int_0^1 v^2(s) ds$ on the D -intervals is sufficient just in the case where a D -interval is of type I; in the other cases we need also a lower bound of the L^2 -norm of v in the region where $v_x \leq z_1 + \eta$ (see Figure 5.6). This is done in Proposition 5.3.1 below.

Proposition 5.3.1. *Assume (H1)–(H5) and let $C > 0$. Then, there exists $\eta_1 = \eta_1(C) \in (0, \eta_0)$ such that, for every $\eta \in (0, \eta_1)$, $\varepsilon \leq \eta^{q+1}$ and $v \in V$ satisfying $I^\varepsilon(v) \leq C$, there exists a collection of (possibly empty) Borel sets Σ_i , $i = 1, \dots, N_v$ such that*

$$\Sigma_i \cap \Sigma_j = \Sigma_i \cap L_j = \Sigma_i \cap L_i = \emptyset, \quad \text{for every } j \neq i \in \{1, \dots, N_v\},$$

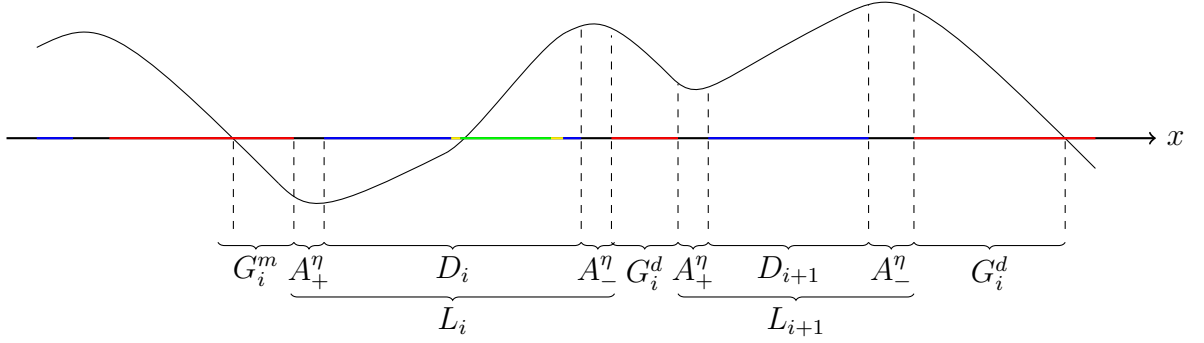


Figure 5.6: The set Σ_i constructed in the proof of Proposition 5.3.1: $\Sigma_{i+1} = \emptyset$, while $\Sigma_i = \Sigma_i^d \cup \Sigma_i^m$, where $\Sigma_i^d \subset G_i^d$, $\Sigma_i^m \subset G_i^m$. In this picture, the red, the blue and the green intervals are respectively the sets of points where $|v_x - z_1| \leq \eta$, $|v_x - z_2| \leq \eta$, and $|v_x - z_3| \leq \eta$. The $B_\pm^\eta$ -transition layers are coloured in yellow.

and, for every $i = 1, \dots, N_v$,

$$\begin{aligned}
\int_{L_i \cup \Sigma_i} (\varepsilon^4 v_{xx}^2 + \varepsilon^{-2} W(v_x) + \varepsilon^{-2} v^2) dx &\geq A_0 \alpha_i + B_0 \beta_i, & \text{if } D_i \text{ is type 0,} \\
\int_{L_i \cup \Sigma_i} (\varepsilon^4 v_{xx}^2 + \varepsilon^{-2} W(v_x) + \varepsilon^{-2} v^2) dx &\geq A_0 \alpha_i - c \varepsilon \eta, & \text{if } D_i \text{ is type I/II, } n_i = 0, \\
\int_{L_i \cup \Sigma_i} (\varepsilon^4 v_{xx}^2 + \varepsilon^{-2} W(v_x) + \varepsilon^{-2} v^2) dx &\geq \alpha_i f_6^{\frac{1}{3}} \left(\frac{\beta_i}{\alpha_i} \right) - c \varepsilon \eta, & \text{if } D_i \text{ is type I, } n_i > 0, \\
\int_{L_i \cup \Sigma_i} (\varepsilon^4 v_{xx}^2 + \varepsilon^{-2} W(v_x) + \varepsilon^{-2} v^2) dx &\geq \alpha_i f_7^{\frac{1}{3}} \left(\frac{\beta_i}{\alpha_i} \right) - c \varepsilon \eta, & \text{if } D_i \text{ is type II, } n_i > 0, \\
\int_{L_i \cup \Sigma_i} (\varepsilon^4 v_{xx}^2 + \varepsilon^{-2} W(v_x) + \varepsilon^{-2} v^2) dx &\geq \alpha_i f_8^{\frac{1}{3}} \left(\frac{\beta_i}{\alpha_i} \right) - c \varepsilon \eta, & \text{if } D_i \text{ is type III/IV}
\end{aligned} \tag{5.57}$$

where f_6, f_7, f_8 are as in (H6)–(H8).

Remark 5.3.2. The first, the third and the fourth lower bounds in Proposition 5.3.1 are sharp up to an error proportional to some positive power of ε . Sharpness of the first bound is given by Proposition 5.2.1. For the third and the fourth bound we refer the reader to the proofs of Proposition 5.6.1 and Proposition 5.6.2 respectively.

Proof. Let $C > 0$, $v \in V$, $\varepsilon \leq \eta^{q+1}$ with $I^\varepsilon(v) \leq C$ and $\eta \in (0, \eta_1)$ with η_1 to be determined later. For every $i = 1, \dots, N_v$, we now construct Σ_i such that (5.57) hold.

Thanks to Lemma 5.3.3 we can set $\Sigma_i = \emptyset$ for all the i 's such that $\max\{\alpha_i, \beta_i\} \notin (R_* \varepsilon, R^* \varepsilon)$, and restrict ourselves to the case $\max\{\alpha_i, \beta_i\} \in (R_* \varepsilon, R^* \varepsilon)$. Let us notice

that, by (5.33)

$$\begin{aligned}
& \int_{L_i \cup \Sigma_i} (\varepsilon^4 v_{xx}^2 + \varepsilon^{-2} W(v_x)) \, dx \\
& \geq 2\varepsilon (H(z_2 - \eta) - H(z_1 + \eta)) + n_i \varepsilon (H(z_3 - \eta) - H(z_2 + \eta)) \\
& \geq \varepsilon (2E_0 + n_i E_1) - c\eta \varepsilon.
\end{aligned} \tag{5.58}$$

We first deal with the i 's where D_i is of type I, and set $\Sigma_i = \emptyset$. Recalling (5.45) we thus get

$$\begin{aligned}
\int_{y_i^-}^{y_i^+} (\varepsilon^4 v_{xx}^2 + \varepsilon^{-2} W(v_x) + \varepsilon^{-2} v^2) \, dx + c\varepsilon \eta & \geq (\varepsilon (2E_0 + n_i E_1) + 3^{-1} \varepsilon^{-2} h(\alpha_i, \beta_i, 1)) \\
& \geq \min_d (d(2E_0 + n_i E_1) + 3^{-1} d^{-2} h(\alpha_i, \beta_i, 1)) \\
& \geq 3^{\frac{2}{3}} (E_0 + \frac{n_i}{2} E_1)^{\frac{2}{3}} h^{\frac{1}{3}}(\alpha_i, \beta_i, 1),
\end{aligned} \tag{5.59}$$

which leads to (5.57) both in the case where $n_i = 0$ and in the case where $n_i > 0$. Here h is as in the statement of Lemma 5.3.4.

We can hence focus on the i 's where $D_i = (x_i^-, x_i^+)$ is such that $v(x_i^+)v(x_i^-) < 0$. As in the proof of Lemma 5.3.4 we will denote by x_i^* a point in D_i such that $v(x_i^*) = 0$. From (5.50) we have

$$v(x_i^+) \geq \sum_{k=2,3} z_k \mathcal{L}((x_i^*, x_i^+) \cap \{|v_x - z_k| \leq \eta\}) - c\varepsilon^2 \eta^{-q} - \eta \varepsilon R^* \geq z_2 \omega_i^a \alpha_i + z_3 \omega_i^b \beta_i - c\eta \varepsilon, \tag{5.60}$$

and, in a similar way, we can prove

$$\begin{aligned}
v(x_i^-) & \leq - \sum_{k=2,3} z_k \mathcal{L}((x_i^-, x_i^*) \cap \{|v_x - z_k| \leq \eta\}) + c\varepsilon^2 \eta^{-q} + \eta \varepsilon R^* \\
& \leq -(z_2(1 - \omega_i^a) \alpha_i + z_3(1 - \omega_i^b) \beta_i) + c\eta \varepsilon,
\end{aligned} \tag{5.61}$$

where ω_i^a, ω_i^b are as in (5.53). We now claim that there exist $\eta_1 \in (0, \eta_0]$ depending just on R_*, R^* and C , such that, for every $\eta < \eta_1$, $v(x_i^+) > 0$ and $v(x_i^-) < 0$. Indeed, we recall that we are working under the assumption $\max\{\alpha_i, \beta_i\} \geq R_* \varepsilon$, and we suppose without loss of generality that $\alpha_i \geq R_* \varepsilon$; the case $\beta_i \geq R_* \varepsilon$ can be treated similarly. Suppose first that $\omega_i \geq \frac{1}{2}$, then (5.60) implies

$$v(x_i^+) \geq \frac{1}{2} z_2 \varepsilon R_* - c\eta \varepsilon.$$

Thus, choosing η small enough, we get $v(x_i^+) > 0$, and, by the fact that $v(x_i^+)v(x_i^-) < 0$, also that $v(x_i^-) < 0$. If $\omega_i^a \leq \frac{1}{2}$, then $(1 - \omega_i^a) \geq \frac{1}{2}$, and (5.61) yields

$$v(x_i^-) \leq -\frac{1}{2}z_2\varepsilon R_* + c\eta\varepsilon,$$

so that for every η small enough $v(x_i^-) < 0$, and hence $v(x_i^+) > 0$.

In the same spirit of (5.60) we have

$$v(x) - v(x_i^+) \geq \int_{x_i^+}^x v_x \, dx \geq (z_1 - \eta)\mathcal{L}((x_i^+, x) \cap \{|v_x - z_1| \leq \eta\}) - \tilde{r}, \quad x \geq x_i^+, \quad (5.62)$$

$$v(x_i^+) \leq \int_{x_i^*}^{x_i^+} v_x \, dx \leq \sum_{k=2,3} z_k \mathcal{L}((x_i^*, x_i^+) \cap \{|v_x - z_k| \leq \eta\}) + \tilde{r} + \eta\varepsilon R^*, \quad (5.63)$$

where

$$\tilde{r} = \int_{\Sigma^\eta} |v_x| \, dx, \quad \Sigma^\eta = \{x \in (0, 1) : |v_x(x) - z_k| > \eta, \, k = 1, 2, 3\}. \quad (5.64)$$

We now give an estimate for $\tilde{r}$. To this aim, we split Σ^η into

$$\Sigma_1^\eta := \{x \in \Sigma^\eta : |v_x(x)| \leq t_0\}, \quad \Sigma_2^\eta := \{x \in \Sigma^\eta : t_0 < |v_x(x)|\},$$

where t_0 is such that $W(s) \geq \frac{c_1}{2}|s|^p$ for each s satisfying $|s| > t_0$, and its existence is guaranteed by (H2). By (5.40), we have

$$\int_{\Sigma_1^\eta} |v_x| \, dx \leq t_0 \mathcal{L}(\Sigma^\eta) \leq c \frac{\varepsilon^2}{\eta^q}. \quad (5.65)$$

On the other hand,

$$\frac{c_1}{2} \mathcal{L}(\Sigma_2^\eta) |t_0|^p \leq \int_{\Sigma_2^\eta} W(v_x) \, dx \leq C\varepsilon^2,$$

implies

$$\mathcal{L}(\Sigma_2^\eta) \leq c\varepsilon^2.$$

Therefore,

$$\int_{\Sigma_2^\eta} |v_x(x)| \, dx \leq c \int_{\Sigma_2^\eta} W^{\frac{1}{p}}(v_x(x)) \, dx \leq c(\mathcal{L}(\Sigma_2^\eta))^{\frac{p-1}{p}} \left(\int_{\Sigma_2^\eta} W(v_x(x)) \, dx \right)^{\frac{1}{p}} \leq c\varepsilon^2, \quad (5.66)$$

where we also made use of the fact that $I^\varepsilon(v) \leq C$. Collecting (5.65)–(5.66) we thus get

$$\tilde{r} \leq \int_{\cup_k \Sigma_k^\eta} |v_x| dx \leq c(\varepsilon^2 \eta^{-q} + \varepsilon^2) \leq cr^*, \quad (5.67)$$

with $r^* := \varepsilon \eta$. By (5.62)–(5.63), (5.67) together with $\max\{\alpha_i, \beta_i\} \leq R^* \varepsilon$, we obtain

$$v(x) - v(x_i^+) \geq (z_1 - \eta) \mathcal{L}((x_i^+, x) \cap \{|v_x - z_1| \leq \eta\}) - cr^*, \quad x \geq x_i^+, \quad (5.68)$$

and

$$0 \leq v(x_i^+) \leq c(\varepsilon \eta + r^* + \varepsilon) \leq c\varepsilon. \quad (5.69)$$

Now let

$$g_b(x) = v(x_i^+) - cr^* - (|z_1| + \eta) \mathcal{L}((x_i^+, x) \cap \{|v_x - z_1| \leq \eta\}).$$

Let $\bar{x}_i \in [x_i^+, 1]$ be the smallest x in $[x_i^+, 1]$ such that $g_b(x) \leq 0$, and let us set

$$\Sigma_i^d := (x_i^+, \bar{x}_i) \cap \{|v_x - z_1| \leq \eta\}.$$

The existence of $\bar{x}_i$ is guaranteed by the continuity of v and the fact that $v(1) = 0$ together with (5.68) imply $g_b(1) \leq 0$. Thus, $\Sigma_i^d = \emptyset$ if and only if $g_b(\bar{x}_i) < 0$. Now, if $\Sigma_i^d \neq \emptyset$, that is $v(x_i^+) > cr^*$, by (5.68) we have that

$$v(x) \geq g_b(x) > 0, \quad \text{for ever } x \in (x_i^+, \bar{x}_i). \quad (5.70)$$

Now, from (5.70) we deduce that

$$\begin{aligned} \int_{\Sigma_i^d} v^2 dx &\geq \int_{\Sigma_i^d} g_b^2(x) dx \\ &\geq (|z_1| + \eta)^{-1} \int_0^{v(x_i^+) - r^*} x^2 dx \geq (|z_1| + \eta)^{-1} v^3(x_i^+) - c\varepsilon^3 \eta. \end{aligned}$$

Here we have used the change of variable $y = g_b(x)$, and, in the last inequality, we exploited (5.69). The same conclusion holds if $\Sigma_i^d = \emptyset$, and hence $v(x_i^+) \leq cr^*$. Indeed,

$$\int_{\Sigma_i^d} v^2 dx = 0 \geq (|z_1| + \eta)^{-1} v^3(x_i^+) - c\varepsilon^3 \eta.$$

Therefore, as $(|z_1| + \eta)^{-1} \geq |z_1|^{-1} - |z_1|^{-2} \eta$, by (5.60) we obtain

$$\int_{\Sigma_i^d} v^2 dx \geq 3^{-1} |z_1|^{-1} (z_2 \omega_i^a \alpha_i + \omega_i^b \beta_i z_3)^3 - cr_b, \quad (5.71)$$

where $r_b := \varepsilon^3 \eta$ and we made use of (5.69) and the fact that $\max\{\alpha_i, \beta_i\} \leq \varepsilon R^*$. In the same way, letting

$$g_m(x) = v(x_i^-) + cr^* + (|z_1| + \eta) \mathcal{L}((x, x_i^-) \cap \{|v_x - z_1| \leq \eta\}),$$

and $\tilde{x}_i \in [0, x_i^-]$ be the largest $x \in [0, x_i^-]$ such that $g_m(x) \geq 0$, we can define

$$\Sigma_i^m := (\tilde{x}_i, x_i^-) \cap \{|v_x - z_1| \leq \eta\}$$

and deduce

$$\int_{\Sigma_i^m} v^2 dx \geq 3^{-1} |z_1|^{-1} (z_2(1 - \omega_i^a) \alpha_i + (1 - \omega_i^b) \beta_i z_3)^3 - cr_b. \quad (5.72)$$

Define now $\Sigma_i := \Sigma_i^m \cup \Sigma_i^d$. We remark that $v_x(x) \leq z_1 + \eta$ for each $x \in \Sigma_i$, while $v_x(x) > z_1 + \eta$ in every L_j , $j = 1, \dots, N_v$. In this way $L_j \cap \Sigma_i = \emptyset$ for every $i, j = 1, \dots, N_v$. We now claim that $\Sigma_j \cap \Sigma_i = \emptyset$ for every $j = 1, \dots, N_v$ with $i \neq j$. Indeed, let x_j^-, x_j^+ be such that $D_j = (x_j^-, x_j^+)$, then in case $v(x_j^-)v(x_j^+) \geq 0$ we defined $\Sigma_j = \emptyset$ and the conclusion follows trivially. We can hence focus on the case $v(x_j^-)v(x_j^+) < 0$, and recall that $\eta < \eta_1$, implies $v(x_j^+) > 0$ and $v(x_j^-) < 0$. Now, the construction of the Σ_j^d, Σ_j^m is such that $\Sigma_j^d \subset (x_j^+, \bar{x}_j)$ (resp. $\Sigma_j^m \subset (\tilde{x}_j, x_j^-)$). Furthermore, as stated in (5.70), v is strictly positive in $(x_i^+, \bar{x}_i)$ (resp. strictly negative in $(\tilde{x}_i, x_i^-)$). Therefore, $\Sigma_i^d \cap \Sigma_j^m = \emptyset$ for every $i, j = 1, \dots, N_v$. Finally, assuming without loss of generality $i < j$, we have that $\Sigma_i^d \subset (x_i^+, \bar{x}_i)$ and $\Sigma_j^d \subset (x_j^+, \bar{x}_j)$ with $v(x) > 0$ for every $x \in (x_i^+, \bar{x}_i) \cup (x_j^+, \bar{x}_j)$. But as $v(x_j^-) < 0$ and $x_j^- \in (x_i^+, x_j^+)$, $(x_i^+, \bar{x}_i)$ and $(x_j^+, \bar{x}_j)$ have to be disjoint, and so must be Σ_i^d and Σ_j^d . In the same way we prove $\Sigma_i^m \cap \Sigma_j^m = \emptyset$ for every $j = 1, \dots, N_v$, $i \neq j$, concluding the proof of the claim.

Now, if $n_i = 0$ or $n_i \geq 4$, we sum (5.71)–(5.72) and minimise the right hand side over $\omega_i^a, \omega_i^b \in [0, 1]$. This is a convex quadratic function attaining its minimum at $\omega_i^a = \omega_i^b = \frac{1}{2}$. We thus have

$$\int_{\Sigma_i} v^2 dx \geq 12^{-1} |z_1|^{-1} (z_2 \alpha_i + z_3 \beta_i)^3 - cr_b. \quad (5.73)$$

Hence, by (5.58) and Lemma 5.3.4,

$$\begin{aligned}
& \int_{L_i \cup \Sigma_i} (\varepsilon^2 v_{xx}^2 + \varepsilon^{-2} W(v_x) + \varepsilon^{-2} v^2) dx + cr_b \\
& \geq \varepsilon(2E_0 + n_i E_1) + \frac{1}{12\varepsilon^2} \left(z_2^2 z_{21} \alpha_i^3 + z_3^2 z_{31} \beta_i^3 + 3\alpha_i \beta_i z_2 z_{31} (\beta_i z_3 + \alpha_i z_2) \right) \\
& \geq 3 \frac{(E_0 + \frac{n_i}{2} E_1)^{\frac{2}{3}}}{12^{\frac{1}{3}}} \left(z_2^2 z_{21} \alpha_i^3 + z_3^2 z_{31} \beta_i^3 + 3\alpha_i \beta_i z_2 z_{31} (\beta_i z_3 + \alpha_i z_2) \right)^{\frac{1}{3}} \\
& = \alpha_i f_7^{\frac{1}{3}} \left(\frac{\alpha_i}{\beta_i} \right).
\end{aligned} \tag{5.74}$$

In case of type IV D -intervals, we recall that $\omega_i^b \in \{0, 1\}$. Therefore, coherently with the proof of Lemma 5.3.4 we assume without loss of generality that $\omega_i^b = 1$ and deduce

$$\begin{aligned}
& cr_b + \int_{\Sigma_i \cup D_i} v^2 dx \\
& \geq 3^{-1} \left(h(\alpha_i, \beta_i, \omega_i^a) + |z_1|^{-1} (z_2 \alpha_i \omega_i^a + z_3 \beta_i)^3 + |z_1|^{-1} (z_2 \alpha_i (1 - \omega_i^a))^3 \right) \\
& \geq \min_{\omega^a \in [0, 1]} 3^{-1} \left(h(\alpha_i, \beta_i, \omega^a) + |z_1|^{-1} (z_2 \alpha_i \omega^a + z_3 \beta_i)^3 + |z_1|^{-1} (z_2 \alpha_i (1 - \omega^a))^3 \right) \\
& \geq 3^{-1} \left(z_2^2 z_{21} \alpha_i^3 + z_3^2 z_{31} \beta_i^3 - 3\alpha_i^3 f_0 \left(\frac{\beta_i}{\alpha_i} \right) \right),
\end{aligned} \tag{5.75}$$

where, f_0 is defined by

$$f_0(y) = \frac{(y^2 z_{31} z_3 - z_2 z_{21})^2}{4(z_{21} + y z_{31})}. \tag{5.76}$$

We remark that the last lower bound in (5.75) is sharp if and only if $\frac{\beta_i}{\alpha_i} \leq \sqrt{\frac{z_2 z_{21}}{z_3 z_{31}}}$. For type III D -intervals, (5.71) and (5.72), together with (5.47) yield

$$cr_b + \int_{\Sigma_i \cup D_i} v^2 dx \geq h^*(\alpha_i, \beta_i, \omega_i^a, \omega_i^b) \tag{5.77}$$

where,

$$\begin{aligned}
h^*(\alpha_i, \beta_i, \omega_i^a, \omega_i^b) &:= 3^{-1} |z_1|^{-1} (z_2 \alpha_i (1 - \omega_i^a) + z_3 \beta_i (1 - \omega_i^b))^3 \\
&\quad + 3^{-1} \left(\hat{h}(\alpha_i, \beta_i, \omega_i^a, \omega_i^b) + |z_1|^{-1} (z_2 \alpha_i \omega_i^a + z_3 \beta_i \omega_i^b)^3 \right).
\end{aligned}$$

We claim, that

$$\min_{(\omega^a, \omega^b) \in [0, 1]^2} h^*(\alpha_i, \beta_i, \omega^a, \omega^b) \geq 3^{-1} \left(z_2^2 z_{21} \alpha_i^3 + z_3^2 z_{31} \beta_i^3 - 3\alpha_i^3 f_0 \left(\frac{\beta_i}{\alpha_i} \right) \right),$$

with f_0 is as in (5.76). Indeed, h^* is a second order polynomial in ω_i^a, ω_i^b with negative Hessian determinant. Therefore, the minimum among the $\omega_i^a, \omega_i^b \in [0, 1]$ is attained at $\omega_i^a \in \{0, 1\}$, or $\omega_i^b \in \{0, 1\}$. More precisely, if $\frac{\beta_i}{\alpha_i} \geq \sqrt{\frac{z_2 z_{21}}{z_3 z_{31}}}$, the minimum is attained at $\omega_i^a \in \{0, 1\}$ and a minimization over $\omega_i^b \in [0, 1]$ gives

$$h^*(\alpha_i, \beta_i, \omega_i^a, \omega_i^b) \geq 3^{-1} \left(z_2^2 z_{21} \alpha_i^3 + z_3^2 z_{31} \beta_i^3 - 3 \alpha_i^3 f\left(\frac{\beta_i}{\alpha_i}\right) \right). \quad (5.78)$$

The same lower bound can be achieved when $\frac{\beta_i}{\alpha_i} < \sqrt{\frac{z_2 z_{21}}{z_3 z_{31}}}$. Indeed, in this case, the minimum is attained at $\omega_i^b \in \{0, 1\}$, and, by using the fact that $z_2 < z_3$, we can bound from below $h^*(\alpha_i, \beta_i, \omega_i^a, \omega_i^b)$ with

$$3^{-1} \left(h(\alpha_i, \beta_i, \omega_i^a) + |z_1|^{-1} (z_2 \alpha_i \omega_i^a + z_3 \beta_i)^3 + |z_1|^{-1} (z_2 \alpha_i (1 - \omega_i^a))^3 \right).$$

The minimization of (5.75) over $\omega_i^a \in [0, 1]$ thus yields again to (5.78). Therefore, for type III/IV D -intervals (5.75)-(5.78) yield

$$cr_b + \int_{\Sigma_i \cup D_i} v^2 dx \geq 3^{-1} \left(z_2^2 z_{21} \alpha_i^3 + z_3^2 z_{31} \beta_i^3 - 3 \alpha_i^3 f\left(\frac{\beta_i}{\alpha_i}\right) \right).$$

Finally, by (5.58) we get

$$\begin{aligned} & \int_{L_i \cup \Sigma_i} (\varepsilon^2 v_{xx}^2 + \varepsilon^{-2} W(v_x) + v^2) dx + c\varepsilon\eta \\ & \geq \min_{\varepsilon > 0} \left(\varepsilon(2E_0 + 2E_1) + \frac{1}{3\varepsilon^2} \left(z_2^2 z_{21} \alpha_i^3 + z_3^2 z_{31} \beta_i^3 - 3 \alpha_i^3 f\left(\frac{\beta_i}{\alpha_i}\right) \right) \right) \\ & \geq 3^{\frac{2}{3}} (E_0 + E_1)^{\frac{2}{3}} \left(z_2^2 z_{21} \alpha_i^3 + z_3^2 z_{31} \beta_i^3 - 3 \alpha_i^3 f\left(\frac{\beta_i}{\alpha_i}\right) \right)^{\frac{1}{3}} \\ & = \alpha_i f_8 \left(\frac{\alpha_i}{\beta_i} \right), \end{aligned}$$

which, together with (5.59) and (5.74) conclude the proof. $\square$

As a corollary of the previous results, we can prove

Theorem 5.3.1. *Assume (H1)–(H8), and let $C > 0$. Then there exists $\eta_1 = \eta_1(C) \in (0, \eta_0)$ such that, if $\eta \in (0, \eta_1)$, $\varepsilon \leq \eta^{q+1}$ and $v \in V$ satisfies $I^\varepsilon(v) \leq C$, it holds*

$$\int_0^1 (\varepsilon^4 v_{xx}^2 + \varepsilon^{-2} W(v_x) + \varepsilon^{-2} v^2) dx \geq (A_0 \lambda_2^\eta + B_0 \lambda_3^\eta) - c\eta. \quad (5.79)$$

Proof. Thanks to Proposition 5.3.1 and (H6)–(H8) we have

$$\int_{L_i \cup \Sigma_i} (\varepsilon^4 v_{xx}^2 + \varepsilon^{-2} W(v_x) + \varepsilon^{-2} v^2) dx \geq A_0 \alpha_i + B_0 \beta_i - c\eta\varepsilon,$$

for every $i = 1, \dots, N_v$. Thus, we can write

$$\begin{aligned} \int_{\sum_i^{N_v} L_i \cup \Sigma_i} (\varepsilon^4 v_{xx}^2 + \varepsilon^{-2} W(v_x) + \varepsilon^{-2} v^2) dx &\geq A_0 \sum_{i=1}^{N_v} \alpha_i + B_0 \sum_{i=1}^{N_v} \beta_i - c N_v \eta \varepsilon \\ &\geq A_0 \sum_{i=1}^{N_v} \alpha_i + B_0 \sum_{i=1}^{N_v} \beta_i - c \eta, \end{aligned}$$

where we made use of Lemma 5.3.1 to bound N_v in the last inequality. It just remains to provide an estimate for the intervals $D_0 := (0, y_1^-)$ and $D_{N_v+1} := (y_{N_v}^+, 1)$. We deal with the first case, as the second can be treated similarly. If $\mathcal{L}(\{x \in D_0 : |v_x(x) - z_k| \leq \eta\}) > R^* \varepsilon$ for some $k = 2, 3$, then, by arguing as in the proof of Lemma 5.3.3 we deduce

$$\int_{D_0} (\varepsilon^4 v_{xx}^2 + \varepsilon^{-2} W(v_x) + \varepsilon^{-2} v^2) dx \geq A_0 \alpha_0 + B_0 \beta_0.$$

On the other hand, if $\mathcal{L}(\{x \in D_0 : |v_x(x) - z_k| \leq \eta\}) \leq R^* \varepsilon$ for $k = 1, 2$, then

$$\int_{D_0} (\varepsilon^4 v_{xx}^2 + \varepsilon^{-2} W(v_x) + \varepsilon^{-2} v^2) dx \geq A_0 \alpha_0 + B_0 \beta_0 - (A_0 + B_0) R^* \varepsilon. \quad (5.80)$$

Therefore,

$$\int_0^1 (\varepsilon^4 v_{xx}^2 + \varepsilon^{-2} W(v_x) + \varepsilon^{-2} v^2) dx \geq A_0 \sum_{i=0}^{N_v+1} \alpha_i + B_0 \sum_{i=0}^{N_v+1} \beta_i - c \eta,$$

which, by the definition of the α_i 's, β_i 's, and of $\lambda_2^\eta, \lambda_3^\eta$ coincides with (5.79). $\square$

5.4 The second Γ -limit

In this section we prove the Γ -limit for I^ε , that is a second Γ -limit for $\mathcal{E}^\varepsilon$, as stated in Theorem 5.0.1. The first step is to prove compactness for the family of energy functionals I^ε .

Proposition 5.4.1. *Assume (H1)–(H5). Let $C > 0$, $\varepsilon_j \downarrow 0$ and (u_j, ν_j) be such that*

$$I^{\varepsilon_j}(u_j, \nu_j) \leq C, \quad \forall j \in \mathbb{N}. \quad (5.81)$$

Then, up to a subsequence, (u_j, ν_j) converges to (u, ν) in $L^2(0, 1) \times L_{w^}^\infty(0, 1; \mathcal{M})$. Furthermore, $u = 0$ and $\nu \in GYM^\infty(u)$ satisfies*

$$\nu_x = \lambda_1(x) \delta_{z_1} + \lambda_2(x) \delta_{z_2} + \lambda_3(x) \delta_{z_3}, \quad a.e. \ x \in (0, 1), \quad (5.82)$$

and $\lambda_1, \lambda_2, \lambda_3 \in L^\infty(0, 1; [0, 1])$ are such that

$$\lambda_1(x) + \lambda_2(x) + \lambda_3(x) = 1, \quad \text{and} \quad \sum_{k=1}^3 z_k \lambda_k(x) = 0, \quad (5.83)$$

for a.e. $x \in (0, 1)$.

Proof. We first notice that (5.81) implies strong convergence of u_j to $u = 0$ in $L^2(0, 1)$. Furthermore, as $W(s) \geq c_1|s|^p - c_2$, we also have

$$\|u_{j,x}\|_{L^p} \leq c, \quad (5.84)$$

and, therefore, up to a subsequence $u_j \rightarrow 0$, weakly in $W_0^{1,p}(0, 1)$. In fact, (5.84) also implies that, up to a further non-relabelled subsequence, $u_{j,x}$ generates a gradient Young measure ν_x , weak* limit of ν_j in $L_w^\infty(0, 1; \mathcal{M})$. Defined Σ_j^η as

$$\Sigma_j^\eta := \{x \in (0, 1) : |u_{j,x}(x) - z_k| > \eta, k = 1, 2, 3\}, \quad (5.85)$$

for some $\eta \in (0, \eta_0)$, by (H5) we have

$$C\varepsilon_j^2 \geq \int_0^1 W(u_{j,x}) dx \geq \int_{\Sigma_j^\eta} W(u_{j,x}) dx \geq c_0 \eta^q \mathcal{L}(\Sigma_j^\eta).$$

This implies

$$\mathcal{L}(\Sigma_j^\eta) \leq c \frac{\varepsilon_j^2}{\eta^q}, \quad (5.86)$$

which is convergence in measure of $u_{j,x}$ to $\mathcal{Z}$. Therefore, ν_x is a probability measure supported on $\mathcal{Z}$ (see e.g., Theorem 1.2.3), and hence $\nu_x = \lambda_1(x)\delta_{z_1} + \lambda_2(x)\delta_{z_2} + \lambda_3(x)\delta_{z_3}$ for a.e. $x \in (0, 1)$, as claimed. The fact that ν is a probability measure implies the first identity in (5.83). By [69, Thm. 8.7] we also know that ν is the gradient Young measure related to $u = 0$, and therefore the average of ν must be 0, that is $\sum_{k=1}^3 \lambda_k(x)z_k = 0$ for a.e. $x \in (0, 1)$, which is the last identity in (5.83). $\square$

Given a sequence $u_j \in W_0^{1,p}(0, 1)$ and $\eta > 0$, let us define

$$\lambda_{k,j}^\eta := \mathcal{L}(\{x \in (0, 1) : |u_{j,x}(x) - z_k| \leq \eta\}), \quad k = 1, 2, 3.$$

The following result is used below:

Lemma 5.4.1. *Assume (H1)–(H5). Let $C > 0$, $\eta \in (0, \eta_0)$, $\varepsilon_j \downarrow 0$ and (u_j, ν_j) be a sequence converging to $(0, \nu)$ in $L^2(0, 1) \times L_w^\infty(0, 1; \mathcal{M})$ such that $I^{\varepsilon_j}(u_j) \leq C$ for each j . Then ν satisfies (5.82)–(5.83). and*

$$\lim_j \lambda_{k,j}^\eta = \int_0^1 \lambda_k dx. \quad (5.87)$$

Proof. The fact that ν satisfies (5.82)–(5.83) follows directly from Proposition 5.4.1. We just need to prove (5.87). Let us consider a continuous function $f_k: \mathbb{R} \rightarrow [0, 1]$, which is equal to 1 for those s such that $|s - z_k| \leq \eta$, and equal to 0 for $|s - z_k| \geq \eta_0$. We have

$$\int_0^1 \lambda_k dx = \int_0^1 \langle \nu, f_k \rangle dx = \lim_j \int_0^1 \langle \nu_j, f_k \rangle dx = \lim_j \int_0^1 f_k(u_{j,x}) dx.$$

Now, we notice that, as $\eta_0 < \frac{|z_k - z_h|}{2}$ for each $h, k = 1, 2, 3$, $h \neq k$ (cf. (H5)),

$$\int_0^1 f_k(u_{j,x}) dx = \lambda_k^\eta(u_j) + r,$$

where

$$0 < r \leq \mathcal{L}(\{x \in (0, 1) : \eta < |u_{j,x}(x) - z_k| \leq \eta_0\}) \leq \mathcal{L}(\Sigma_j^\eta) \leq c\varepsilon_j^2 \eta^{-q}.$$

Here, Σ_j^η is as in (5.85), and we also made use of (5.86). Therefore, collecting all previous identities we finally get

$$\lim_j \lambda_k^\eta(u_j) = \int_0^1 \lambda_k dx,$$

which concludes the proof. $\square$

We are now ready to derive the Γ –limit of I^ε

Theorem 5.4.1. *Assume (H1)–(H8). Then I^ε Γ –converges to*

$$I^0(u, \nu) = \begin{cases} A_0 \int_0^1 \nu_x(z_2) dx + B_0 \int_0^1 \nu_x(z_3) dx, & \text{if } u = 0, \nu \in GYM^\infty(0), \\ & \text{supp } \nu \subset \mathcal{Z} \text{ a.e.}, \\ +\infty, & \text{otherwise.} \end{cases}$$

in the $L^2(0, 1) \times L_{w*}^\infty(0, 1; \mathcal{M})$ topology as $\varepsilon \downarrow 0$.

Proof. The Γ – lim sup inequality is a direct consequence of Proposition 5.2.1 and Lemma 5.4.1. Therefore we just need to prove the Γ – lim inf inequality. In order to do that, we need to consider a generic sequence ε_j converging to 0, a sequence $u_j \in V$ converging strongly in $L^2(0, 1)$ to $u \in L^2(0, 1)$ and a sequence of parametrized measures ν_j converging weakly* to ν in $L_{w*}^\infty(0, 1; \mathcal{M})$. If $\liminf_j I^{\varepsilon_j}(u_j) = \infty$, the liminf inequality is trivial. Otherwise, up to a subsequence we can assume the existence of $C > 0$, independent of ε_j , such that

$$I^{\varepsilon_j}(u_j) \leq C, \quad \forall j \in \mathbb{N}.$$

As a consequence of Proposition 5.4.1, $u = 0$ and the related gradient Young measure ν must satisfy (5.82)–(5.83). Now, Theorem 5.3.1 guarantees the existence of $\eta_1 > 0$ such that, fixed $\eta \in (0, \eta_1)$,

$$\int_0^1 (\varepsilon_j^4 u_{j,xx}^2 + \varepsilon_j^{-2} W(u_{j,x}) + \varepsilon_j^{-2} u_j^2) dx \geq (A_0 \lambda_{2,j}^\eta + B_0 \lambda_{3,j}^\eta) - c\eta,$$

for all $\varepsilon_j < \eta^{q+1}$. Now, by taking the $\liminf$ on both sides, and recalling Lemma 5.4.1, we deduce

$$\liminf_j \int_0^1 (\varepsilon_j^4 u_{j,xx}^2 + \varepsilon_j^{-2} W(u_{j,x}) + \varepsilon_j^{-2} u_j^2) dx \geq A_0 \int_0^1 \lambda_2 dx + B_0 \int_0^1 \lambda_3 dx - c\eta.$$

The arbitrariness of η yields to the desired $\Gamma - \liminf$ inequality. $\square$

5.5 Selecting Minimizing sequences without Γ –convergence

In this section we prove Theorem 5.0.2. In order to do this we strongly rely on the estimates of Section 5.3, to which we refer the reader for the notation. We start with the following theorem:

Theorem 5.5.1. *Assume (H1)–(H5) and $z_3 \leq 3|z_1|$. Then, there exist $\varepsilon_1 > 0$ and $\xi > 0$ such that, if $\varepsilon < \varepsilon_1$*

$$\inf_{v \in V} \int_0^1 (\varepsilon^4 v_{xx}^2 + \varepsilon^{-2} W(v_x) + \varepsilon^{-2} v^2) dx = A_0 z_{21}^{-1} + o(\varepsilon^\xi). \quad (5.88)$$

Furthermore, every minimizer u of I^ε satisfies

$$\mathcal{L}(\{x \in (0, 1) : |u_x(x) - z_3| \leq \varepsilon^{\frac{1}{q+1}}\}) \leq c\varepsilon^\xi. \quad (5.89)$$

Proof. We first notice that Proposition 5.2.1 implies

$$\inf_{v \in V} \int_0^1 (\varepsilon^4 v_{xx}^2 + \varepsilon^{-2} W(v_x) + \varepsilon^{-2} v^2) dx \leq A_0 z_{21}^{-1} + c\varepsilon^\xi. \quad (5.90)$$

Let us assume $\eta \in (0, \eta_1)$, with η_1 as in the statement of Proposition 5.3.1, $\varepsilon \leq \eta^{q+1}$, and define

$$K := \frac{z_{31}}{z_{21}} \sqrt[3]{\frac{(E_0 + E_1)^2}{E_0^2}} > 1.$$

We notice that $B_0 > KA_0$. We now look for a lower bound for the energy I^ε of the type (5.79), but with B_0 replaced by KA_0 . This new bound does not rely on

(H6)–(H8), but is deduced by strongly exploiting the estimates in Section 5.3.

It can be checked by using

$$z_2 < z_3, \quad z_{21} < z_{31}, \quad z_2 z_{31} < z_3 z_{21},$$

that f_7, f_8 given in (H7)–(H8) satisfy

$$f_7(y) \geq A_0 + K A_0 y, \quad f_8(y) \geq A_0 + K A_0 y, \quad \text{for every } y \geq 0. \quad (5.91)$$

Furthermore, as we assumed $z_3 \leq 3|z_1|$, we also have

$$f_6(y) \geq A_0 + K A_0 y, \quad \text{for every } y \geq 0. \quad (5.92)$$

Therefore, collecting the estimates for every L –interval, from Proposition 5.3.1 together with (5.91)–(5.92) and the fact that $B_0 > K A_0$ imply

$$\int_{\sum_i (L_i \cup \Sigma_i)} (\varepsilon^4 v_{xx}^2 + \varepsilon^{-2} W(v_x) + \varepsilon^{-2} v^2) dx + c\eta \geq A_0 \sum_i (\alpha_i + K \beta_i).$$

Here we also made use of Lemma (5.3.1) to bound N_v with $c\varepsilon^{-1}$. Finally, recalling (5.80) we deduce

$$\int_0^1 (\varepsilon^4 v_{xx}^2 + \varepsilon^{-2} W(v_x) + \varepsilon^{-2} v^2) dx + c\eta \geq A_0 (\lambda_2^\eta + K \lambda_3^\eta). \quad (5.93)$$

Now, thanks to (5.39)–(5.40), and the fact that $\varepsilon \leq \eta^{q+1}$, we can write

$$1 \geq \lambda_1^\eta + \lambda_2^\eta + \lambda_3^\eta = 1 - \mathcal{L}(\Sigma^\eta) \geq 1 - c\varepsilon\eta, \quad (5.94)$$

while, on the other hand,

$$0 = \int_0^1 v_x dx \leq \sum_{k=1}^3 (z_k + \eta) \lambda_k^\eta + \tilde{r} \leq \sum_{k=1}^3 z_k \lambda_k^\eta + c\eta, \quad 0 = \int_0^1 v_x dx \geq \sum_{k=1}^3 z_k \lambda_k^\eta - c\eta \quad (5.95)$$

where $\tilde{r}$ is defined as in (5.64), and has been bounded according to the estimate in (5.67). By combining (5.94)–(5.95) we are led to

$$z_{31}^{-1} (1 - z_{21} \lambda_2^\eta) + c\eta \geq \lambda_3^\eta \geq z_{31}^{-1} (1 - z_{21} \lambda_2^\eta) - c\eta, \quad (5.96)$$

so that, by (5.93),

$$\int_0^1 (\varepsilon^4 v_{xx}^2 + \varepsilon^{-2} W(v_x) + \varepsilon^{-2} v^2) dx + c\eta \geq A_0 \left(\lambda_2^\eta + \frac{K}{z_{31}} (1 - z_{21} \lambda_2^\eta) \right). \quad (5.97)$$

The right hand side of the above inequality is a decreasing function of λ_2^η so minimised by the biggest admissible λ_2^η . But as $\lambda_3^\eta \geq 0$, (5.96) entails $\lambda_2^\eta \leq z_{21}^{-1} + c\eta$, thus implying

$$\int_0^1 (\varepsilon^4 v_{xx}^2 + \varepsilon^{-2} W(v_x) + \varepsilon^{-2} v^2) dx + c\eta \geq A_0 z_{21}^{-1}. \quad (5.98)$$

Choosing $\eta = \varepsilon^{\frac{1}{q+1}}$ and $\varepsilon_1 = \min\{\eta_1^{q+1}, \varepsilon_0\}$, where ε_0 is as in Proposition 5.2.1 we complete the proof of (5.88). Finally, combining (5.90) and (5.97) we get

$$c\varepsilon^\xi + A_0 z_{21}^{-1} \geq A_0 \lambda_2^\eta + z_{31}^{-1} K A_0 (1 - z_{21} \lambda_2^\eta) \geq A_0 z_{21}^{-1} - c\varepsilon^\xi,$$

which is

$$|(\lambda_2^\eta - z_{21}^{-1})(1 - K \frac{z_{21}}{z_{31}})| \leq c\varepsilon^\xi.$$

This implies $|\lambda_2^\eta - z_{21}^{-1}| \leq c\varepsilon^\xi$ which, by (5.96), concludes the proof. $\square$

This finally allows us to prove the following

Corollary 5.5.1. *Assume (H1)–(H5) and $z_3 \leq 3|z_1|$. Then, every gradient Young measure $\nu \in GYM^\infty(0)$ generated by a sequence u^ε of minimizers for I^ε as $\varepsilon \downarrow 0$, satisfies*

$$\text{supp } \nu_x \in \{z_1, z_2\}, \quad \nu_x = \frac{z_2}{z_2 - z_1} \delta_{z_1} - \frac{z_1}{z_2 - z_1} \delta_{z_2}, \quad \text{a.e. } x \in (0, 1).$$

Proof. By Proposition 5.0.1 we know that $\text{supp } \nu_x \subset \mathcal{Z}$ almost everywhere in $(0, 1)$, and so ν_x is of the form

$$\nu_x = \lambda_1(x) \delta_{z_1} + \lambda_2(x) \delta_{z_2} + \lambda_3(x) \delta_{z_3},$$

with the λ_i 's satisfying (5.9) for a.e. $x \in (0, 1)$. Therefore we just need to show that $\lambda_3 = 0$. Let us consider a continuous function $f_3: \mathbb{R} \rightarrow [0, 1]$ which is equal to 1 for all $s \in \mathbb{R}$ such that $|s - z_3| \leq \frac{\eta_0}{2}$, and 0 if $|s - z_3| \geq \eta_0$. Take a sequence u_{ε_j} of minimizers of I^{ε_j} generating ν . By arguing as in the proof of Lemma 5.4.1 we get

$$0 \leq \int_0^1 \lambda_3 dx = \lim_j \int_0^1 f_3(u_{\varepsilon_j, x}) dx \leq \lim (\lambda_3^\eta(u_{\varepsilon_j}) + c\varepsilon_j^2 \eta^{-q}).$$

After choosing $\eta = \varepsilon_j^{\frac{1}{q+1}}$, (5.89) gives the sought result. $\square$

5.6 Some remarks on the assumptions

It is worth spending some words on assumptions (H6)–(H8). Hypothesis (H6) is needed in our construction of a lower bound, but it might be possible to remove it by making the arguments of Section 5.3 more involved. It is easy to check that it fails whenever $z_3 > 3|z_1|$. On the other hand, as mentioned in the introduction, it turns out that (H7)–(H8) are necessary conditions in order to prove Theorem 5.0.1, and the second Γ –limit would have a different form without these assumptions. Indeed, as explained in the introduction, these hypotheses guarantee that the microstructures used in the construction of Proposition 5.2.1 are energetically preferable to those constructed in the following Propositions, and shown in Figure 5.7.

Proposition 5.6.1. *Assume (H7) is not satisfied. Then, there exist $\lambda_1, \lambda_2, \lambda_3 \in (0, 1)$ satisfying (5.83), $u_\varepsilon \in V$ such that $(u_\varepsilon, \delta_{u_\varepsilon, x}) \rightarrow (0, \nu)$ in $L^2(0, 1) \times L_{w^*}^\infty(0, 1; \mathcal{M})$, where $\nu_x = \lambda_1 \delta_{z_1} + \lambda_2 \delta_{z_2} + \lambda_3 \delta_{z_3}$ a.e. in $(0, 1)$, and*

$$\limsup_{\varepsilon \downarrow 0} I^\varepsilon(u_\varepsilon) < A_0 \lambda_2 + B_0 \lambda_3.$$

Proof. The proof of the Proposition is very similar to the one of Proposition 5.2.1 in many details. For this reason we skip some long computation and just give the idea of the proof.

Let $\hat{y} \geq 0$ such that the inequality in (H7) does not hold. Let us choose

$$\lambda_2 = (\hat{y} z_{31} + z_{21})^{-1}, \quad \lambda_3 = z_{31}^{-1} - \frac{z_{21}}{z_{31}} \lambda_2, \quad \lambda_1 = 1 - \lambda_2 - \lambda_3, \quad (5.99)$$

so that $\frac{\lambda_3}{\lambda_2} = \hat{y}$. It is easy to check that $\lambda_1, \lambda_2, \lambda_3 \in (0, 1)$ and satisfy (5.9). We divide $(0, 1)$ in N_ε subintervals (x_i, x_{i+1}) of length N_ε^{-1} , where $x_i = i N_\varepsilon^{-1}$ for $i = 0, \dots, N_\varepsilon$. On every interval we construct v_ε as a suitable continuous approximation (see the proof of Proposition 5.2.1) of the function (see the derivative of the function in Figure 5.7a)

$$\hat{v}_\varepsilon(s) = \begin{cases} z_2, & \text{if } 0 \leq |s - (2N_\varepsilon)^{-1}| \leq \lambda_2(2N_\varepsilon)^{-1}, \\ z_3, & \text{if } \lambda_2 N_\varepsilon^{-1} < |s - (2N_\varepsilon)^{-1}| \leq \lambda_2(2N_\varepsilon)^{-1} + \lambda_3(2N_\varepsilon)^{-1} \\ z_1, & \text{if } \lambda_2(2N_\varepsilon)^{-1} + \lambda_3(2N_\varepsilon)^{-1} \leq |s - (2N_\varepsilon)^{-1}| \leq (2N_\varepsilon)^{-1}. \end{cases}$$

Therefore, after defining w_ε as the N_ε^{-1} –periodic extension of v_ε , we construct u_ε as in (5.15). Now, an argument as the one in the proof of Proposition 5.2.1, allows us

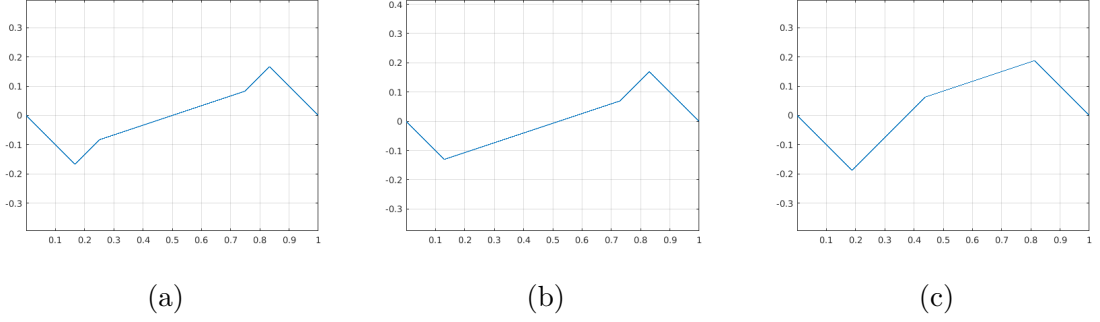


Figure 5.7: 5.7a. Microstructure of lower energy in case (H7) does not hold. Microstructures of low energy in case (H8) does not hold are represented in Figures 5.7b and 5.7c, respectively in the case where $\hat{y} \leq \sqrt{\frac{z_2 z_{21}}{z_3 z_{31}}}$ and $\hat{y} > \sqrt{\frac{z_2 z_{21}}{z_3 z_{31}}}$.

to prove that

$$\begin{aligned} I^\varepsilon(u_\varepsilon) &\leq 3^{\frac{2}{3}} 2^{-\frac{2}{3}} (E_0 + 2E_1)^{\frac{2}{3}} \left(z_2^2 z_{21} \lambda_2^3 + z_3^2 z_{31} \lambda_3^3 + 3\lambda_2 \lambda_3 z_2 z_{31} (\lambda_3 z_3 + \lambda_2 z_2) \right)^{\frac{1}{3}} + c\varepsilon^\xi \\ &\leq \lambda_2 3^{\frac{2}{3}} 2^{-\frac{2}{3}} (E_0 + 2E_1)^{\frac{2}{3}} \left(z_2^2 z_{21} + z_3^2 z_{31} \hat{y}^3 + 3\hat{y} z_2 z_{31} (\hat{y} z_3 + z_2) \right)^{\frac{1}{3}} + c\varepsilon^\xi, \end{aligned}$$

for some $\xi > 0$, and that (5.11) holds. Here we have used that, by construction, $\frac{\lambda_3}{\lambda_2} = \hat{y}$. As $\hat{y}$ contradicts (H7), we have

$$I^\varepsilon(u_\varepsilon) < A_0 \lambda_2 + B_0 \lambda_3 + c\varepsilon^\xi. \quad (5.100)$$

Furthermore, we can deduce that estimates as the one in (5.11) hold. Thus, taking the lim sup in (5.100), by Lemma 5.4.1 we obtain the sought result. $\square$

Proposition 5.6.2. *Assume (H8) is not satisfied. Then, there exist $\lambda_1, \lambda_2, \lambda_3 \in (0, 1)$ satisfying (5.83), $u_\varepsilon \in V$ such that $(u_\varepsilon, \delta_{u_\varepsilon, x}) \rightarrow (0, \nu)$ in $L^2(0, 1) \times L_{w^*}^\infty(0, 1; \mathcal{M})$, where $\nu_x = \lambda_1 \delta_{z_1} + \lambda_2 \delta_{z_2} + \lambda_3 \delta_{z_3}$ a.e. in $(0, 1)$, and*

$$\limsup_{\varepsilon \downarrow 0} I^\varepsilon(u_\varepsilon) < A_0 \lambda_2 + B_0 \lambda_3.$$

Proof. Again, the proof of this Proposition is very similar to the one of Proposition 5.2.1 and of Proposition 5.6.1 in many details. For this reason we skip some long computation and just give the idea of the proof.

Let $\hat{y} \geq 0$ such that (H8) does not hold. Let us choose $\lambda_1, \lambda_2, \lambda_3$ that satisfy (5.9) as in (5.99). Again, we divide $(0, 1)$ in M_ε subintervals (x_i, x_{i+1}) of length M_ε^{-1} , where $x_i = iM_\varepsilon^{-1}$ for $i = 0, \dots, M_\varepsilon$. We have two cases:

$$0 \leq \hat{y} \leq \sqrt{\frac{z_2 z_{21}}{z_3 z_{31}}}, \quad \text{and} \quad \sqrt{\frac{z_2 z_{21}}{z_3 z_{31}}} < \hat{y}.$$

In the first case (see Figure 5.7b), we define $\omega^a := (z_2 z_{21} - z_3 z_{31} \hat{y}^2)(2z_2(z_{21} + \hat{y} z_{31}))^{-1}$ and

$$\hat{v}_a^\varepsilon(s) = \begin{cases} z_1, & \text{if } 0 \leq s < \frac{z_2}{z_1}(1 - \omega^a)\lambda_2 M_\varepsilon^{-1}, \\ z_2, & \text{if } \frac{z_2}{z_1}(1 - \omega^a)\lambda_2 M_\varepsilon^{-1} \leq s < \left(\frac{z_2}{z_1}(1 - \omega^a) + 1\right)\lambda_2 M_\varepsilon^{-1}, \\ z_3, & \text{if } \left(\frac{z_2}{z_1}(1 - \omega^a) + 1\right)\lambda_2 M_\varepsilon^{-1} \leq s < \left(\frac{z_2}{z_1}(1 - \omega^a) + 1\right)\lambda_2 M_\varepsilon^{-1} + \lambda_3 M_\varepsilon^{-1}, \\ z_1, & \text{if } \left(\frac{z_2}{z_1}(1 - \omega^a) + 1\right)\lambda_2 M_\varepsilon^{-1} + \lambda_3 M_\varepsilon^{-1} \leq s. \end{cases}$$

In the second (see Figure 5.7c), we define $\omega^b := -\omega^a$ and

$$\hat{v}_b^\varepsilon(s) = \begin{cases} z_1, & \text{if } 0 \leq s < \frac{z_3}{z_1}(1 - \omega^b)\lambda_3 M_\varepsilon^{-1}, \\ z_3, & \text{if } \frac{z_3}{z_1}(1 - \omega^b)\lambda_3 M_\varepsilon^{-1} \leq s < \left(\frac{z_3}{z_1}(1 - \omega^b) + 1\right)\lambda_3 M_\varepsilon^{-1}, \\ z_2, & \text{if } \left(\frac{z_3}{z_1}(1 - \omega^b) + 1\right)\lambda_3 M_\varepsilon^{-1} \leq s < \left(\frac{z_3}{z_1}(1 - \omega^b) + 1\right)\lambda_3 M_\varepsilon^{-1} + \lambda_2 M_\varepsilon^{-1}, \\ z_1, & \text{if } \left(\frac{z_3}{z_1}(1 - \omega^b) + 1\right)\lambda_3 M_\varepsilon^{-1} + \lambda_2 M_\varepsilon^{-1} \leq s. \end{cases}$$

Now, let us consider suitable continuous approximations v_a^ε and v_b^ε of $\hat{v}_a^\varepsilon$ and $\hat{v}_b^\varepsilon$, which can be obtained in the same way as the one in Proposition 5.2.1. Let $w_a^\varepsilon, w_b^\varepsilon$ be the M_ε^{-1} -periodic extensions of v_a^ε and v_b^ε respectively, and define u_ε as in (5.15). It is easy to compute that

$$\begin{aligned} \int_0^{N^{-1}} \left(\int_0^x \hat{v}_a^\varepsilon(s) ds \right)^2 dx &\leq 3^{-1} (z_2^2 z_{21} \lambda_2^3 + z_3^2 z_{31} \lambda_3^3 - 3\lambda_2^3 f(\hat{y})), \\ \int_0^{N^{-1}} \left(\int_0^x \hat{v}_b^\varepsilon(s) ds \right)^2 dx &\leq 3^{-1} (z_2^2 z_{21} \lambda_2^3 + z_3^2 z_{31} \lambda_3^3 - 3\lambda_2^3 f(\hat{y})), \end{aligned}$$

respectively if $\hat{y}$ is smaller or bigger than $\sqrt{\frac{z_2 z_{21}}{z_3 z_{31}}}$. An argument as the one in Proposition 5.2.1 allows hence to prove

$$\begin{aligned} I^\varepsilon(u_\varepsilon) &\leq 3^{\frac{2}{3}} (E_0 + E_1)^{\frac{2}{3}} \left(z_2^2 z_{21} \lambda_2^3 + z_3^2 z_{31} \lambda_3^3 - 3\lambda_2^3 f(\hat{y}) \right)^{\frac{1}{3}} + c\varepsilon^\xi \\ &\leq 3^{\frac{2}{3}} \lambda_2 (E_0 + E_1)^{\frac{2}{3}} \left(z_2^2 z_{21} + z_3^2 z_{31} \hat{y}^3 - 3f(\hat{y}) \right)^{\frac{1}{3}} + c\varepsilon^\xi, \end{aligned}$$

for some $\xi > 0$, for some $\varepsilon_0 > 0$, and for every $\varepsilon \in (0, \varepsilon_0)$. Again, here we used the fact that, by construction, $\frac{\lambda_3}{\lambda_2} = \hat{y}$. The fact that $\hat{y}$ violates the inequality in (H8) yields

$$I^\varepsilon(u_\varepsilon) < A_0 \lambda_2 + B_0 \lambda_3 + c\varepsilon^\xi. \quad (5.101)$$

Furthermore, by arguing as in the proof of Proposition 5.2.1 we can prove estimates as the ones in (5.11). Therefore, by taking the lim sup in (5.101) and exploiting Lemma 5.4.1 we conclude the proof of the proposition. $\square$

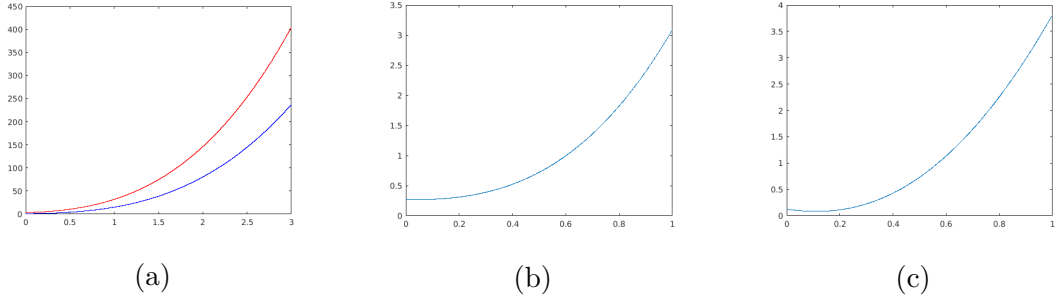


Figure 5.8: Verification of the hypotheses (H6)–(H8) for the examples of Section 5.6.1. Figure 5.8a is the graphical verification of (H6) for the two examples: in blue the case $z_2 = \frac{1}{3}$, in red the one with $z_2 = \frac{1}{2}$. Figure 5.8b and 5.8c verify (H7) and (H8) in the example with $z_2 = \frac{1}{3}$.

5.6.1 Two examples

An easy example where hypotheses (H1)–(H8) hold is when

$$W(s) = (s - 1)^2(s + 1)^2(s - 3^{-1})^2.$$

Indeed, in this case $E_0 \approx 1.054$, $E_1 \approx 0.165$, $A_0 \approx 0.718$, $B_0 \approx 1.883$, $z_1 = -1$, $z_2 = \frac{1}{3}$, $z_3 = 1$. It is trivial to check that in this context (H1)–(H5) hold. Hypotheses (H6)–(H8) are here verified graphically (cf. Figure 5.8). It can be proved that (H6)–(H8) hold for W of the form

$$W(s) = (s - 1)^2(s + 1)^2(s - z_2)^2,$$

whenever $z_2 \in (-0.49, 0.49)$. The bound on z_2 is not sharp.

On the other hand, let us consider

$$W(s) = (s - 1)^2(s + 1)^2(s - 2^{-1})^2,$$

where, in our notation, $E_0 \approx 1.406$, $E_1 \approx 0.073$, $A_0 \approx 1.186$, $B_0 \approx 2.143$, $z_1 = -1$, $z_2 = \frac{1}{2}$, $z_3 = 1$. In this case (H1)–(H6) hold (cf. Figure 5.8a). However, hypotheses (H7) and (H8) fail respectively in a neighbourhood of $\hat{y}_7 = 0.585$ and $\hat{y}_8 = 0.204$. Here, it is energetically very cheap to pass from z_2 to z_3 so other microstructures are energetically favourable for $\frac{\lambda_3}{\lambda_2}$ close to $\hat{y}_7$ or $\hat{y}_8$. Nonetheless, as $z_3 \leq 3|z_1|$, thanks to Theorem 5.0.2 we can still select minimizing sequences by adding vanishing interfacial energy.

Bibliography

- [1] R. Abeyaratne and J.K. Knowles. On the driving traction acting on a surface of strain discontinuity in a continuum. *Journal of the Mechanics and Physics of Solids*, 38(3):345–360, 1990.
- [2] R. Abeyaratne and J.K. Knowles. A continuum model of a thermoelastic solid capable of undergoing phase transitions. *Journal of the Mechanics and Physics of Solids*, 41(3):541–571, 1993.
- [3] G. Alberti, S. Bianchini, and G. Crippa. Structure of level sets and Sard-type properties of Lipschitz maps. *Ann. Sc. Norm. Super. Pisa Cl. Sci. (5)*, 12(4):863–902, 2013.
- [4] L. Ambrosio, N. Fusco, and D. Pallara. *Functions of bounded variation and free discontinuity problems*. Oxford Mathematical Monographs. The Clarendon Press, Oxford University Press, New York, 2000.
- [5] G Arlt. Twinning in ferroelectric and ferroelastic ceramics: stress relief. *Journal of Materials Science*, 25(6):2655–2666, 1990.
- [6] J. M. Ball and R. D. James. A characterization of plane strain. *Proc. Roy. Soc. London Ser. A*, 432(1884):93–99, 1991.
- [7] J.M. Ball. Convexity conditions and existence theorems in nonlinear elasticity. *Arch. Rational Mech. Anal.*, 63(4):337–403, 1976/77.
- [8] J.M. Ball. A version of the fundamental theorem for Young measures. In *PDEs and continuum models of phase transitions (Nice, 1988)*, volume 344 of *Lecture Notes in Phys.*, pages 207–215. Springer, Berlin, 1989.
- [9] J.M. Ball. Some open problems in elasticity. In *Geometry, mechanics, and dynamics*, pages 3–59. Springer, New York, 2002.

- [10] J.M. Ball and Y. Sengül. Quasistatic nonlinear viscoelasticity and gradient flows. *J. Dynam. Differential Equations*, 27(3-4):405–442, 2015.
- [11] J.M. Ball and C. Carstensen. Hadamard’s compatibility condition for microstructures. *In progress*.
- [12] J.M. Ball and C. Carstensen. Nonclassical austenite-martensite interfaces. *Le Journal de Physique IV*, 7(C5):35–40, 1997.
- [13] J.M. Ball and C. Carstensen. Compatibility conditions for microstructures and the austenite–martensite transition. *Materials Science and Engineering: A*, 273:231–236, 1999.
- [14] J.M. Ball and C. Carstensen. Interaction of martensitic microstructures in adjacent grains. *arXiv preprint arXiv:1708.03286*, 2017.
- [15] J.M. Ball, C. Chu, and R.D. James. Hysteresis during stress-induced variant rearrangement. *Le Journal de Physique IV*, 5(C8):C8–245, 1995.
- [16] J.M. Ball and E.C.M. Crooks. Local minimizers and planar interfaces in a phase-transition model with interfacial energy. *Calc. Var. Partial Differential Equations*, 40(3-4):501–538, 2011.
- [17] J.M. Ball, P.J. Holmes, R.D. James, R.L. Pego, and P.J. Swart. On the dynamics of fine structure. *J. Nonlinear Sci.*, 1(1):17–70, 1991.
- [18] J.M. Ball and R.D. James. Fine phase mixtures as minimizers of energy. *Arch. Rational Mech. Anal.*, 100(1):13–52, 1987.
- [19] J.M. Ball and R.D. James. Proposed experimental tests of a theory of fine microstructure and the two-well problem. *Phil. Trans. R. Soc. Lond. A*, 338(1650):389–450, 1992.
- [20] J.M. Ball and K. Koumatos. An investigation of non-planar austenite-martensite interfaces. *Math. Models Methods Appl. Sci.*, 24(10):1937–1956, 2014.
- [21] J.M. Ball and K. Koumatos. Quasiconvexity at the boundary and the nucleation of austenite. *Arch. Ration. Mech. Anal.*, 219(1):89–157, 2016.
- [22] K. Bhattacharya. Self-accommodation in martensite. *Arch. Rational Mech. Anal.*, 120(3):201–244, 1992.

- [23] K. Bhattacharya. *Microstructure of martensite*. Oxford Series on Materials Modelling. Oxford University Press, Oxford, 2003. Why it forms and how it gives rise to the shape-memory effect.
- [24] F. Boyer and P. Fabrie. *Mathematical tools for the study of the incompressible Navier-Stokes equations and related models*, volume 183 of *Applied Mathematical Sciences*. Springer, New York, 2013.
- [25] A. Braides. Γ -convergence for beginners, volume 22 of *Oxford Lecture Series in Mathematics and its Applications*. Oxford University Press, Oxford, 2002.
- [26] X. Chen, V. Srivastava, V. Dabade, and R.D. James. Study of the *cofactor conditions*: conditions of supercompatibility between phases. *J. Mech. Phys. Solids*, 61(12):2566–2587, 2013.
- [27] X. Chen, N. Tamura, A. MacDowell, and R.D. James. In-situ characterization of highly reversible phase transformation by synchrotron x-ray laue microdiffraction. *Applied Physics Letters*, 108(21):211902, 2016.
- [28] C. Chluba, W. Ge, R.L. de Miranda, J. Strobel, L. Kienle, E. Quandt, and M. Wuttig. Ultralow-fatigue shape memory alloy films. *Science*, 348(6238):1004–1007, 2015.
- [29] M. Cicalese, E.N. Spadaro, and C.I. Zeppieri. Asymptotic analysis of a second-order singular perturbation model for phase transitions. *Calc. Var. Partial Differential Equations*, 41(1-2):127–150, 2011.
- [30] S. Conti and G. Dolzmann. On the theory of relaxation in nonlinear elasticity with constraints on the determinant. *Arch. Ration. Mech. Anal.*, 217(2):413–437, 2015.
- [31] S. Conti, G. Dolzmann, and S. Müller. The div-curl lemma for sequences whose divergence and curl are compact in $W^{-1,1}$. *C. R. Math. Acad. Sci. Paris*, 349(3-4):175–178, 2011.
- [32] S. Conti, I. Fonseca, and G. Leoni. A Γ -convergence result for the two-gradient theory of phase transitions. *Comm. Pure Appl. Math.*, 55(7):857–936, 2002.
- [33] S. Conti and B. Schweizer. Rigidity and gamma convergence for solid-solid phase transitions with $SO(2)$ invariance. *Comm. Pure Appl. Math.*, 59(6):830–868, 2006.

- [34] B. Dacorogna. *Direct methods in the calculus of variations*, volume 78 of *Applied Mathematical Sciences*. Springer, New York, second edition, 2008.
- [35] B. Dacorogna and P. Marcellini. *Implicit partial differential equations*, volume 37 of *Progress in Nonlinear Differential Equations and their Applications*. Birkhäuser Boston, Inc., Boston, MA, 1999.
- [36] G. Dolzmann. *Variational methods for crystalline microstructure—analysis and computation*, volume 1803 of *Lecture Notes in Mathematics*. Springer-Verlag, Berlin, 2003.
- [37] G. Dolzmann and S. Müller. Microstructures with finite surface energy: the two-well problem. *Arch. Rational Mech. Anal.*, 132(2):101–141, 1995.
- [38] J.L. Ericksen. On the symmetry and stability of thermoelastic solids. *Journal of Applied Mechanics*, 45(4):740–744, 1978.
- [39] L.C. Evans. *Partial differential equations*, volume 19 of *Graduate Studies in Mathematics*. American Mathematical Society, Providence, RI, second edition, 2010.
- [40] L.C. Evans and R.F. Gariepy. *Measure theory and fine properties of functions*. Studies in Advanced Mathematics. CRC Press, Boca Raton, FL, 1992.
- [41] H. Federer. *Geometric measure theory*. Die Grundlehren der mathematischen Wissenschaften, Band 153. Springer-Verlag New York Inc., New York, 1969.
- [42] H. Flanders. Differentiation under the integral sign. *Amer. Math. Monthly*, 80:615–627; correction, *ibid.* 81 (1974), 145, 1973.
- [43] V. Girault and P.A. Raviart. *Finite element approximation of the Navier-Stokes equations*, volume 749 of *Lecture Notes in Mathematics*. Springer-Verlag, Berlin-New York, 1979.
- [44] M. Gromov. *Partial differential relations*, volume 9 of *Ergebnisse der Mathematik und ihrer Grenzgebiete (3) [Results in Mathematics and Related Areas (3)]*. Springer-Verlag, Berlin, 1986.
- [45] H. Gu, L. Bumke, C. Chluba, E. Quandt, and R.D. James. Phase engineering and supercompatibility of shape memory alloys. *Materials Today*, 2017.

- [46] H. Halkin. Implicit functions and optimization problems without continuous differentiability of the data. *SIAM J. Control*, 12:229–236, 1974. Collection of articles dedicated to the memory of Lucien W. Neustadt.
- [47] T. Inamura, H. Hosoda, and S. Miyazaki. Incompatibility and preferred morphology in the self-accommodation microstructure of β -titanium shape memory alloy. *Philosophical Magazine*, 93(6):618–634, 2013.
- [48] T. Iwaniec, G.C. Verchota, and A.L. Vogel. The failure of rank-one connections. *Arch. Ration. Mech. Anal.*, 163(2):125–169, 2002.
- [49] R.D. James. Taming the temperamental metal transformation. *Science*, 348(6238):968–969, 2015.
- [50] R.D. James. Materials from mathematics. <http://www.ams.org/CEB-2018-Master.pdf>, 2018.
- [51] R.D. James and Z. Zhang. A way to search for multiferroic materials with unlikely combinations of physical properties. In *Magnetism and structure in functional materials*, pages 159–175. Springer, 2005.
- [52] S.J. Kim and R. Abeyaratne. On the effect of the heat generated during a stress-induced thermoelastic phase transformation. *Contin. Mech. Thermodyn.*, 7(3):311–332, 1995.
- [53] D. Kinderlehrer and P. Pedregal. Characterizations of Young measures generated by gradients. *Arch. Rational Mech. Anal.*, 115(4):329–365, 1991.
- [54] D. Kinderlehrer and P. Pedregal. Gradient Young measures generated by sequences in Sobolev spaces. *J. Geom. Anal.*, 4(1):59–90, 1994.
- [55] J. Kristensen. A necessary and sufficient condition for lower semicontinuity. *Nonlinear Anal.*, 120:43–56, 2015.
- [56] N.H. Kuiper. On C^1 -isometric imbeddings. I, II. volume 17, pages 545–556, 683–689, 1955.
- [57] L. Modica and S. Mortola. Un esempio di Γ^- -convergenza. *Boll. Un. Mat. Ital. B (5)*, 14(1):285–299, 1977.
- [58] F. Morgan. *Geometric measure theory*. Academic Press, Inc., San Diego, CA, second edition, 1995. A beginner’s guide.

- [59] C.B. Morrey. Quasi-convexity and the lower semicontinuity of multiple integrals. *Pacific J. Math.*, 2:25–53, 1952.
- [60] S. Müller. Minimizing sequences for nonconvex functionals, phase transitions and singular perturbations. In *Problems involving change of type (Stuttgart, 1988)*, volume 359 of *Lecture Notes in Phys.*, pages 31–44. Springer, Berlin, 1990.
- [61] S. Müller. Singular perturbations as a selection criterion for periodic minimizing sequences. *Calc. Var. Partial Differential Equations*, 1(2):169–204, 1993.
- [62] S. Müller. Variational models for microstructure and phase transitions. In *Calculus of variations and geometric evolution problems (Cetraro, 1996)*, volume 1713 of *Lecture Notes in Math.*, pages 85–210. Springer, Berlin, 1999.
- [63] S. Müller and M.A. Sychev. Optimal existence theorems for nonhomogeneous differential inclusions. *J. Funct. Anal.*, 181(2):447–475, 2001.
- [64] S. Müller and V. Šverák. Attainment results for the two-well problem by convex integration. pages 239–251, 1996.
- [65] S. Müller and V. Šverák. Convex integration with constraints and applications to phase transitions and partial differential equations. *J. Eur. Math. Soc. (JEMS)*, 1(4):393–422, 1999.
- [66] J. Nash. C^1 isometric imbeddings. *Ann. of Math. (2)*, 60:383–396, 1954.
- [67] X. Ni, J.R. Greer, K. Bhattacharya, R.D. James, and X. Chen. Exceptional resilience of small-scale Au₃₀Cu₂₅Zn₄₅ under cyclic stress-induced phase transformation. *Nano letters*, 16(12):7621–7625, 2016.
- [68] R.A. Nicolaides and N.J. Walkington. Strong convergence of numerical solutions to degenerate variational problems. *Math. Comp.*, 64(209):117–127, 1995.
- [69] P. Pedregal. *Parametrized measures and variational principles*, volume 30 of *Progress in Nonlinear Differential Equations and their Applications*. Birkhäuser Verlag, Basel, 1997.
- [70] R.L. Pego. Phase transitions in one-dimensional nonlinear viscoelasticity: admissibility and stability. *Arch. Rational Mech. Anal.*, 97(4):353–394, 1987.
- [71] M. Pitteri. Reconciliation of local and global symmetries of crystals. *J. Elasticity*, 14(2):175–190, 1984.

- [72] M. Pitteri and G. Zanzotto. *Continuum models for phase transitions and twinning in crystals*, volume 19 of *Applied Mathematics (Boca Raton)*. Chapman & Hall/CRC, Boca Raton, FL, 2003.
- [73] J.G. Rešetnjak. Liouville’s conformal mapping theorem under minimal regularity hypotheses. *Sibirsk. Mat. Ž.*, 8:835–840, 1967.
- [74] Y. Sengul. *Well-posedness of dynamics of microstructure in solids*. PhD thesis, University of Oxford, 2010.
- [75] J. Simon. Sobolev, Besov and Nikoskii fractional spaces: imbeddings and comparisons for vector valued spaces on an interval. *Ann. Mat. Pura Appl. (4)*, 157:117–148, 1990.
- [76] Y. Song, X. Chen, V. Dabade, T.W. Shield, and R.D. James. Enhanced reversibility and unusual microstructure of a phase-transforming material. *Nature*, 502(7469):85, 2013.
- [77] M.A. Sychev. A new approach to Young measure theory, relaxation and convergence in energy. *Ann. Inst. H. Poincaré Anal. Non Linéaire*, 16(6):773–812, 1999.
- [78] L. Székelyhidi Jr. From isometric embeddings to turbulence. In *HCDTE lecture notes. Part II. Nonlinear hyperbolic PDEs, dispersive and transport equations*, volume 7 of *AIMS Ser. Appl. Math.*, page 63. Am. Inst. Math. Sci. (AIMS), Springfield, MO, 2013.
- [79] C. Truesdell. *Rational thermodynamics*. Springer-Verlag, New York, second edition, 1984. With an appendix by C. C. Wang, With additional appendices by 23 contributors.
- [80] A. Visintin. Stefan problem with a kinetic condition at the free boundary. *Ann. Mat. Pura Appl. (4)*, 146:97–122, 1987.
- [81] A. Visintin. *Models of phase transitions*, volume 28 of *Progress in Nonlinear Differential Equations and their Applications*. Birkhäuser Boston, Inc., Boston, MA, 1996.
- [82] M.S. Wechsler, D.S. Lieberman, and T.A. Read. On the theory of the formation of martensite. *Trans AIME*, 197:1503–1515, 1953.

- [83] W.Q. Xie. The Stefan problem with a kinetic condition at the free boundary. *SIAM J. Math. Anal.*, 21(2):362–373, 1990.
- [84] Z. Zhang, R.D. James, and S. Müller. Energy barriers and hysteresis in martensitic phase transformations. *Acta Materialia*, 57(15):4332–4352, 2009.