



Mathematical justification of macroscopic models for diffusion MRI through the periodic unfolding method

Julien Coatléven

► To cite this version:

Julien Coatléven. Mathematical justification of macroscopic models for diffusion MRI through the periodic unfolding method. *Asymptotic Analysis*, 2015, 93 (3), pp.219-258. 10.3233/ASY-151294 . hal-01180148

HAL Id: hal-01180148

<https://hal.science/hal-01180148v1>

Submitted on 24 Jul 2015

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

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

Mathematical justification of macroscopic models for diffusion MRI through the periodic unfolding method

Julien Coatléven *

IFPEN 1 et 4 Avenue de Bois-Préau, 92852 Rueil-Malmaison, France

Abstract

Diffusion Magnetic Resonance Imaging (dMRI) is a promising tool to obtain useful information on cellular structure when applied to biological tissues. A coupled macroscopic model has been introduced recently through formal homogenization to model dMRI's signal attenuation. This model was based on a particular scaling of the permeability condition modeling cellular membranes. In this article, we explore all the possible scalings and mathematically justify the corresponding limit models, using the periodic unfolding method. We also illustrate through numerical simulations the respective behavior of the limit models when compared to dMRI measurements.

Keywords : Diffusion MRI, Homogenization, Bloch-Torrey equation, Periodic Unfolding, Imperfect transmission, trace jumps

1 Introduction

Diffusion Magnetic Resonance Imaging (dMRI) relies on the measurement of the diffusion of water molecules in the imaged sample. The potential clinical applications of diffusion MRI have multiplied during the last 25 years, making it a promising diagnosis tool. Indeed, when the imaged sample is biological tissue (for a survey, see [12]), it can be used to detect, for example, cerebral ischemia [19], demyelinating disorders [10], and tumors [16, 17, 18].

However, the link between the physiological modifications and the measurements has mostly remained qualitatively explained. In a recent paper [7], a macroscopic model has been introduced through formal homogenization, under the assumption that the media is periodic. It efficiently reproduces experimental dMRI's signal measurements and generalizes some phenomenological models that have appeared in the medical imaging literature (see [11]). It was based on a particular scaling of the permeability condition involved in the microscopic description of dMRI. Other scalings are possible and produce different limit models, which we want to identify as well. Our aim here is to provide a rigorous mathematical justification for all these limit models. The homogenization literature being particularly rich, many different techniques can be used to do so. We choose to use the unfolding operator, which was introduced first in [2] under the name of "dilation operator", and has been widely studied and used since (for a review, see [5]). There are of course other ways to mathematically justify our homogenization procedure, in particular one could use the notion of two-scale convergence (which is strongly related to the unfolding operator, see [1] and [5]), but the unfolding operator will prove to be a very natural way to treat the jump condition involved by the permeability condition of our model problem. We show that there are only five possible limit models, including of course the model of [7]. These limit models are similar to those that have already appeared in the literature when homogenizing elliptic problems with jump condition (see [14]-[15]) using two-scale convergence. Such problems have also been addressed through the periodic unfolding method in [9].

*julien.coatleven@ifpen.fr

The paper will be organized as follows. In a first section, we recall how dMRI is modeled through a two-compartment Bloch-Torrey equation with jump, and precise our periodicity hypothesis. Then, we define our model problem, which requires a careful treatment due to the presence of unbounded coefficients. In a second section, we state our main results : we present the homogenized macroscopic models we want to establish, and briefly display a few numerical results to compare them to the full model problem, emphasizing that the model introduced in [7] is indeed the most accurate one. The remaining sections are devoted to the homogenization proof. In the third section, we recall the main results of the periodic unfolding method. Then, in the fourth section we apply this theory to our model problem to rigorously establish our limit models.

2 Model problem

Water magnetization in presence of an applied magnetic field is modeled by a two-compartment Bloch-Torrey equation in $\mathbb{R}^d$ (d being the dimension, in practice $d = 2$ or 3), with the diffusion coefficient for water magnetization varying at the scale of biological cells. The biological cells' membrane is modeled by a permeability condition. We use $\mathbb{R}^d$ primarily to remain consistent with [7], where this choice is discussed from the physical and biological point of view. From the mathematical point of view, the results we obtain remain true for any Ω Lipschitz open subset of $\mathbb{R}^d$, with appropriate boundary conditions, such as Dirichlet or Neumann's (those boundary conditions, if not scaled, should simply be added to the homogenized problems in that case, the local cell problems remaining unchanged). However, both notations and proofs are more involved when dealing with bounded domains (see [4], [5] and [6]). Thus, using $\mathbb{R}^d$ also allow us to simplify the presentation.

We assume that the probed media is periodic (as explained in [7], it is enough to obtain the analytical form of the macroscopic equations) : $Y =]0, 1[^d$ denotes the reference periodic cell, and ε the size of the average square of tissue containing a single biological cell. Each periodicity cell is assumed to contain a biological cell, i.e. it can be divided into two connected parts : the cellular (or intra-cellular) domain Y_c (c stands for cell), the extra-cellular domain Y_e (e stands for extra-cellular), separated by the midline $\Gamma_m = \partial Y_c$ of the cellular membrane, assumed to be Lipschitz. The connec-

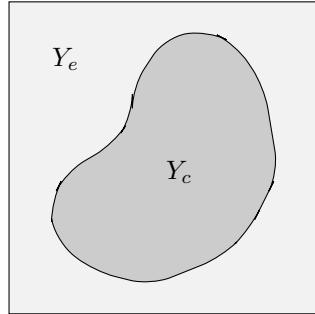


Figure 1: The periodic cell Y , containing a simplified biological cell without membrane

tivity assumption is made here in order to simplify some proofs and notations. As explained in [9], this assumption can easily be lifted. However, once this connectivity assumption made, we need to impose that Γ_m does not crosses ∂Y , if we want our model to remain realistic (the entire boundary of the biological cells must be recovered by a membrane). Of course, when dealing with several connected components, this hypothesis can be lifted, provided that when Γ_m crosses a face of Y , there exists another connected component of Y_c on the opposite face of Y to "close" the biological

cell by periodicity with a proper membrane. We denote :

$$\left\{ \begin{array}{l} \overline{\Omega}_e^\varepsilon = \bigcup_{\xi \in \mathbb{Z}^d} \varepsilon(\xi + \overline{Y}_e) \\ \Omega_c^\varepsilon = \bigcup_{\xi \in \mathbb{Z}^d} \varepsilon(\xi + Y_c) \\ \Omega_{ext}^\varepsilon = \Omega_e^\varepsilon \cup \Omega_c^\varepsilon \\ \Gamma_m^{\varepsilon, \xi} = \varepsilon(\xi + \Gamma_m) \end{array} \right. \quad (1)$$

The diffusion coefficient σ is assumed to be periodic, with period Y , on which it is (for simplicity) assumed to be defined piecewise as follows :

$$\sigma = \left\{ \begin{array}{ll} \sigma_c & \text{in } Y_c \\ \sigma_e & \text{in } Y_e \end{array} \right. \quad (2)$$

Then, the water magnetization M_ε is governed by the following Bloch-Torrey equation with jump :

$$\boxed{\begin{array}{ll} \frac{\partial M_\varepsilon(x, t)}{\partial t} + \imath \mathbf{q} \cdot \mathbf{x} f(t) M_\varepsilon(x, t) - \text{div}(\sigma_\varepsilon(x) \nabla M_\varepsilon(x, t)) = 0 & \text{in } \Omega_{ext}^\varepsilon \times]0, T[\\ \sigma_\varepsilon \nabla M_\varepsilon \cdot \nu|_{\Gamma_m^{\varepsilon, \xi}} = \kappa^\varepsilon [M_\varepsilon]|_{\Gamma_m^{\varepsilon, \xi}} & \forall \xi \in \mathbb{Z}^d \\ [\sigma_\varepsilon \nabla M_\varepsilon \cdot \nu]|_{\Gamma_m^{\varepsilon, \xi}} = 0 & \forall \xi \in \mathbb{Z}^d \\ M_\varepsilon(\cdot, 0) = M_{init} & \text{in } \Omega_{ext}^\varepsilon \end{array}} \quad (3)$$

where ν is the normal to $\Gamma_m^{\varepsilon, \xi}$, outgoing for Ω_{ext}^ε , κ^ε is the permeability coefficient, $\sigma_\varepsilon = \sigma\left(\frac{x}{\varepsilon}\right)$, $\imath = \sqrt{-1}$, M_{init} is the initial magnetization, $\mathbf{q}$, constant vector of $\mathbb{R}^d$, is the amplitude of the applied magnetic field and f its time profile. If we denote $M_{\varepsilon, \alpha}$ the restriction of M_ε to $\Omega_\alpha^\varepsilon$, then the trace jump is defined by :

$$[M_\varepsilon]|_{\Gamma_m^{\varepsilon, \xi}} = M_{\varepsilon, e}|_{\Gamma_m^{\varepsilon, \xi}} - M_{\varepsilon, c}|_{\Gamma_m^{\varepsilon, \xi}}$$

The three coefficients σ_e , σ_c and κ^ε are assumed to be strictly positive and satisfy almost everywhere, for some $0 < \sigma^- < \sigma^+ < +\infty$:

$$\sigma^- \leq \sigma_c \leq \sigma^+ \quad \text{and} \quad \sigma^- \leq \sigma_e \leq \sigma^+$$

Notice that Ω_{ext}^ε is not a connected domain, but can be decomposed into two subdomains, corresponding to the intra-cellular region and the extra-cellular region, and that the common boundary of Ω_c^ε and Ω_e^ε is the union of the mid-lines of the biological cell membranes :

$$\partial\Omega_c^\varepsilon = \partial\Omega_c^\varepsilon \cap \partial\Omega_e^\varepsilon = \partial\Omega_e^\varepsilon = \bigcup_{\xi \in \mathbb{Z}^d} \Gamma_m^{\varepsilon, \xi}$$

Problem (3) involves an unbounded coefficient, namely $\mathbf{q} \cdot \mathbf{x}$. Thus, its analysis is not entirely classical. In particular, it is not obvious in which sense its solution is to be understood. To make it clearer, we define $\widetilde{M}_\varepsilon$ almost everywhere on $\mathbb{R}^d \times]0, T[$ by :

$$\widetilde{M}_\varepsilon(x, t) = M_\varepsilon(x, t) e^{\imath \mathbf{q} \mathbf{x} \cdot \mathbf{n} \int_0^t f(s) ds} \quad (4)$$

denoting $\mathbf{q} = q\mathbf{n}$, where $\mathbf{n}$ is a unitary vector. Then, at least formally, we have :

$$\frac{\partial \widetilde{M}_\varepsilon}{\partial t} = \left(\frac{\partial M_\varepsilon}{\partial t} + \imath q \mathbf{n} \cdot \mathbf{x} f(t) M_\varepsilon \right) e^{\imath \mathbf{q} \mathbf{n} \cdot \mathbf{x} \int_0^t f(s) ds}$$

and :

$$\begin{aligned}\nabla \widetilde{M}_\varepsilon &= \left(\nabla M_\varepsilon + \imath q \mathbf{n} M_\varepsilon \left(\int_0^t f(s) ds \right) \right) e^{\imath q \mathbf{n} \cdot \mathbf{x} \int_0^t f(s) ds} \\ \operatorname{div}(\sigma_\varepsilon \nabla \widetilde{M}_\varepsilon) &= \left(\operatorname{div}(\sigma_\varepsilon \nabla M_\varepsilon) + \imath q \sigma_\varepsilon \nabla M_\varepsilon \cdot \mathbf{n} \left(\int_0^t f(s) ds \right) \right) e^{\imath q \mathbf{n} \cdot \mathbf{x} \int_0^t f(s) ds} \\ &\quad + \operatorname{div} \left(\imath q \sigma_\varepsilon \widetilde{M}_\varepsilon \cdot \mathbf{n} \left(\int_0^t f(s) ds \right) \right)\end{aligned}$$

Then, as :

$$\begin{aligned}&\imath q \sigma_\varepsilon \nabla M_\varepsilon \cdot \mathbf{n} \left(\int_0^t f(s) ds \right) e^{\imath q \mathbf{n} \cdot \mathbf{x} \int_0^t f(s) ds} \\ &= \imath q \sigma_\varepsilon \nabla \widetilde{M}_\varepsilon \cdot \mathbf{n} \left(\int_0^t f(s) ds \right) + q^2 \sigma_\varepsilon \widetilde{M}_\varepsilon \left(\int_0^t f(s) ds \right)^2\end{aligned}$$

we get :

$$\begin{aligned}\operatorname{div}(\sigma_\varepsilon \nabla \widetilde{M}_\varepsilon) &= \operatorname{div}(\sigma_\varepsilon \nabla M_\varepsilon) e^{\imath q \mathbf{n} \cdot \mathbf{x} \int_0^t f(s) ds} + \imath q \sigma_\varepsilon \nabla \widetilde{M}_\varepsilon \cdot \mathbf{n} \left(\int_0^t f(s) ds \right) \\ &\quad + q^2 \sigma_\varepsilon \left(\int_0^t f(s) ds \right)^2 \widetilde{M}_\varepsilon + \operatorname{div} \left(\imath q \sigma_\varepsilon \widetilde{M}_\varepsilon \cdot \mathbf{n} \left(\int_0^t f(s) ds \right) \right)\end{aligned}$$

and then, denoting

$$F(t) = \int_0^t f(s) ds$$

we finally obtain :

$$\boxed{\begin{aligned}&\frac{\partial \widetilde{M}_\varepsilon}{\partial t} - \operatorname{div} \left(\sigma_\varepsilon \left(\nabla \widetilde{M}_\varepsilon - \imath q \mathbf{n} F(t) \widetilde{M}_\varepsilon \right) \right) + \imath q \sigma_\varepsilon \nabla \widetilde{M}_\varepsilon \cdot \mathbf{n} F(t) \\ &+ q^2 \sigma_\varepsilon F(t)^2 \widetilde{M}_\varepsilon = 0 && \text{in } \Omega_{ext}^\varepsilon \times]0, T[\\ &\sigma_\varepsilon \left(\nabla \widetilde{M}_\varepsilon - \imath q \mathbf{n} F(t) \widetilde{M}_\varepsilon \right) \cdot \nu|_{\Gamma_m^\varepsilon, \xi} = \kappa^\varepsilon \left[\widetilde{M}_\varepsilon \right]|_{\Gamma_m^\varepsilon, \xi} && \forall \xi \in \mathbb{Z}^d \\ &\left[\sigma_\varepsilon \left(\nabla \widetilde{M}_\varepsilon - \imath q \mathbf{n} F(t) \widetilde{M}_\varepsilon \right) \cdot \nu \right]|_{\Gamma_m^\varepsilon, \xi} = 0 && \forall \xi \in \mathbb{Z}^d \\ &\widetilde{M}_\varepsilon(\cdot, 0) = M_{init} && \text{in } \Omega_{ext}^\varepsilon\end{aligned}} \quad (5)$$

While it may seem a more complicated problem, notice that it does not involve unbounded coefficients anymore. Though, we can hope to apply classical results to it. We denote $X'_\varepsilon = \widetilde{H}^{-1}(\Omega_c^\varepsilon) \times \widetilde{H}^{-1}(\Omega_e^\varepsilon)$ the dual space of $X_\varepsilon = H^1(\Omega_c^\varepsilon) \times H^1(\Omega_e^\varepsilon)$, $H_\varepsilon = L^2(\Omega_{ext}^\varepsilon)$ and :

$$W(0, T, X_\varepsilon) = \left\{ U \in L^2(0, T, X_\varepsilon) \mid \partial_t U \in L^2(0, T, X'_\varepsilon) \right\}$$

It is classical that $W(0, T, X_\varepsilon) \hookrightarrow C^0(0, T, H_\varepsilon)$. We have the following well-posedness result for $\widetilde{M}_\varepsilon$ (we refer the reader to [7] for the proof):

Theorem 2.1. *Let $f \in L^\infty([0, T])$ and $M_{init} \in X_\varepsilon$. Then, there exists a unique solution $\widetilde{M}_\varepsilon \in W(0, T, X_\varepsilon)$ of problem (5). Moreover, we have the estimates:*

$$\|\widetilde{M}_\varepsilon\|_{L^\infty(0, T, H_\varepsilon)} \leq e^{CT} \|M_{init}\|_{H_\varepsilon} \quad (6)$$

and

$$\|\widetilde{M}_\varepsilon\|_{L^2(0, T, X_\varepsilon)} \leq C' (1 + e^{CT}) \|M_{init}\|_{H_\varepsilon} \quad (7)$$

for some $C, C' > 0$ independent on ε .

The regularity of $\widetilde{M}_\varepsilon$ provides us precious information on the regularity we can expect for M_ε . Obviously, we have that :

$$M_\varepsilon \in L^2(0, T, X_\varepsilon) \cap C^0(0, T, H_\varepsilon)$$

as :

$$M_\varepsilon = \widetilde{M}_\varepsilon e^{-\imath q \mathbf{n} \cdot \mathbf{x} F(t)} \quad \text{and} \quad \nabla M_\varepsilon = \nabla \widetilde{M}_\varepsilon e^{-\imath q \mathbf{n} \cdot \mathbf{x} F(t)} - \imath q \mathbf{n} F(t) M_\varepsilon$$

The time regularity of M_ε is less obvious. We have :

$$\frac{\partial \widetilde{M}_\varepsilon}{\partial t} e^{-\imath q \mathbf{n} \cdot \mathbf{x} F(t)} = \frac{\partial M_\varepsilon}{\partial t} + \imath q \mathbf{n} \cdot \mathbf{x} f(t) M_\varepsilon$$

implying that :

$$\frac{\partial M_\varepsilon}{\partial t} + \imath q \mathbf{n} \cdot \mathbf{x} f(t) M_\varepsilon \in L^2(0, T, X'_\varepsilon)$$

However, we do not a priori have the regularity of $\partial_t M_\varepsilon$ or $\imath q \mathbf{n} \cdot \mathbf{x} f(t) M_\varepsilon$ separately. In the following, we denote :

$$\partial_t^{q, f} U = \partial_t U + \imath q \mathbf{n} \cdot \mathbf{x} f(t) U$$

and

$$W^{q, f}(0, T, X_\varepsilon) = \left\{ U \in L^2(0, T, X_\varepsilon) \mid \partial_t^{q, f} U \in L^2(0, T, X'_\varepsilon) \right\}$$

To any $U \in W^{q, f}(0, T, X_\varepsilon)$, we can associate $\widetilde{U}$ defined by :

$$\widetilde{U} = U e^{\imath q \mathbf{n} \cdot \mathbf{x} F(t)}$$

and we have $\widetilde{U} \in W(0, T, X_\varepsilon)$, with :

$$\|\widetilde{U}\|_{L^2(0, T, X_\varepsilon)}^2 + \|\partial_t \widetilde{U}\|_{L^2(0, T, X'_\varepsilon)}^2 \leq (1 + q \|F\|_{L^\infty(0, T)}^2) \|U\|_{L^2(0, T, X_\varepsilon)}^2 + \|\partial_t^{q, f} U\|_{L^2(0, T, X'_\varepsilon)}^2 \quad (8)$$

Conversely, to any $\widetilde{U} \in W(0, T, X_\varepsilon)$, we can associate $U \in W^{q, f}(0, T, X_\varepsilon)$ defined by:

$$U = \widetilde{U} e^{-\imath q \mathbf{n} \cdot \mathbf{x} F(t)}$$

and we have :

$$\|U\|_{L^2(0, T, X_\varepsilon)}^2 + \|\partial_t^{q, f} U\|_{L^2(0, T, X'_\varepsilon)}^2 \leq (1 + q \|F\|_{L^\infty(0, T)}^2) \|\widetilde{U}\|_{L^2(0, T, X_\varepsilon)}^2 + \|\partial_t \widetilde{U}\|_{L^2(0, T, X'_\varepsilon)}^2 \quad (9)$$

Thus, there exists a bijection between the two spaces, and $W^{q, f}(0, T, X_\varepsilon)$ is a Hilbert space for the norm :

$$\|U\|_{W^{q, f}}^2 = \|U\|_{L^2(0, T, X_\varepsilon)}^2 + \|\partial_t^{q, f} U\|_{L^2(0, T, X'_\varepsilon)}^2$$

Moreover, this implies that $W^{q, f}(0, T, X_\varepsilon) \hookrightarrow C^0(0, T, H_\varepsilon)$. It is consequently natural to define variational solutions to (3) as functions $M_\varepsilon \in W^{q, f}(0, T, X_\varepsilon)$ satisfying :

$$\left\{ \begin{array}{l} \langle \partial_t^{q, f} M_\varepsilon(t), V \rangle_{X'_\varepsilon, X_\varepsilon} + (\sigma_\varepsilon \nabla M_\varepsilon, \nabla V)_{H_\varepsilon} \\ + \sum_{\xi \in \mathbb{Z}^d} \kappa^\varepsilon \left([M_\varepsilon(t)]|_{\Gamma_m^{\varepsilon, \xi}}, [V]|_{\Gamma_m^{\varepsilon, \xi}} \right)_{L^2(\Gamma_m^{\varepsilon, \xi})} = 0 \quad \text{in } \mathcal{D}'([0, T]), \quad \forall V \in X_\varepsilon \\ M_\varepsilon(0) = M_{init} \end{array} \right. \quad (10)$$

We then immediately have :

Theorem 2.2. *Let $f \in L^\infty([0, T])$ and $M_{init} \in X_\varepsilon$. Then, there exists a unique solution $M_\varepsilon \in W^{q,f}(0, T, X_\varepsilon)$ of problem (10). Moreover, we have the estimates:*

$$\|M_\varepsilon\|_{L^\infty(0, T, H_\varepsilon)} \leq e^{CT} \|M_{init}\|_{H_\varepsilon} \quad (11)$$

and

$$\|M_\varepsilon\|_{L^2(0, T, X_\varepsilon)} \leq \sqrt{(1 + q\|F\|_{L^\infty(0, T)}^2)} C' (1 + e^{CT}) \|M_{init}\|_{H_\varepsilon} \quad (12)$$

for some $C, C' > 0$ independent on ε .

3 Main results

In practice, the biological cell size is extremely small compared to the representative size of the probed tissue. A direct simulation involving the fine scale on the entire probed domain is consequently extremely costly to perform. Our goal is to get a limit problem which does not see the small scale, in order to diminish the computational effort needed, as well as to give an explanation for the macroscopic behavior observed from experimental measurements of water magnetization during diffusion MRI. To do so, we need to precise our hypothesis on κ^ε . We assume that κ^ε is of the form $\kappa^\varepsilon = \varepsilon^p \sigma_m$, where $\sigma_m > 0$ is a constant independent on ε and $p \in \mathbb{Z}$. Depending on the values of p , several macroscopic models will arise from homogenization. We claim that there are exactly five limit models, which we describe now.

3.1 The five macroscopic models

We will see in the proof that there are two sub-families of macroscopic models, those who involve several limit functions, which corresponds to $p \geq 1$, and those who involve a single one.

We introduce the spaces $X = H^1(\mathbb{R})^2$, its dual X' , $H = L^2(\mathbb{R}^d)^2$, and :

$$W(0, T, X) = \left\{ U \in L^2(0, T, X) \mid \partial_t U \in L^2(0, T, X') \right\}$$

$$W^{q,f}(0, T, X) = \left\{ U \in L^2(0, T, X) \mid \partial_t^{q,f} U \in L^2(0, T, X') \right\}$$

Proceeding as in the previous section, it is obvious that these two spaces are in bijection, thus $W^{q,f}(0, T, X)$ is a Hilbert space for the norm :

$$\|U\|_{W^{q,f}}^2 = \|U\|_{L^2(0, T, X)}^2 + \|\partial_t^{q,f} \tilde{U}\|_{L^2(0, T, X')}^2$$

and $W^{q,f}(0, T, X) \hookrightarrow C^0(0, T, H)$.

Now for $p \geq 1$, we claim that we have the following result :

Theorem 3.1. *We denote $w_{i,\alpha}$ the solutions for $i = 1, 2$ of the cell problems :*

$$\left\{ \begin{array}{ll} -\operatorname{div}_y(\sigma_\alpha(\nabla_y w_{i,\alpha} + e_i)) = 0 & \text{in } Y_\alpha \\ \sigma_e \nabla_y w_{i,e} \cdot \nu + \sigma_e e_i \cdot \nu = 0 & \text{on } \Gamma_m \\ \sigma_c \nabla_y w_{i,c} \cdot \nu + \sigma_c e_i \cdot \nu = 0 & \text{on } \Gamma_m \\ w_{i,e} - Y_e - \text{periodic} \end{array} \right. \quad (13)$$

We define the associated homogenized tensors by :

$$D_{\alpha,ij} = \frac{1}{|Y_\alpha|} \int_{Y_\alpha} \sigma_\alpha(\nabla w_{j,\alpha} + e_j) \cdot (\nabla w_{i,\alpha} + e_i) \quad (14)$$

and the coefficients η_α by :

$$\eta_c = \frac{\sigma_m |\Gamma_m|}{|Y_c|} \quad \text{and} \quad \eta_e = \frac{\sigma_m |\Gamma_m|}{|Y_e|} \quad (15)$$

Then, the solution M_ε of (10) satisfies, up to a subsequence :

$$M_\varepsilon \rightharpoonup \frac{|Y_e|}{|Y|} M_{0,e} + \frac{|Y_c|}{|Y|} M_{0,c} \quad \text{weakly in } L^2(0, T, L^2(\mathbb{R}^d)) \quad (16)$$

where if $p = 1$, $M_0 = (M_{0,e}, M_{0,c}) \in W^{q,f}(0, T, X)^2$ satisfies for all $V = (V_e, V_c) \in X^2$

$$\begin{aligned} & \left\langle \partial_t^{q,f} M_{0,e}, V_e \right\rangle_{H^{-1}, H^1} + (D_e \nabla_x M_{0,e}, \nabla V_e)_{L^2} + (\eta_e (M_{0,e} - M_{0,c}), V_e)_{L^2} = 0 \\ & \text{in } \mathcal{D}'([0, T]), \quad M_{0,e}(\cdot, 0) = M_{init} \text{ in } \mathbb{R}^d \\ & \left\langle \partial_t^{q,f} M_{0,c}, V_c \right\rangle_{H^{-1}, H^1} + (D_c \nabla_x M_{0,c}, \nabla V_c)_{L^2} + (\eta_c (M_{0,c} - M_{0,e}), V_c)_{L^2} = 0 \\ & \text{in } \mathcal{D}'([0, T]), \quad M_{0,c}(\cdot, 0) = M_{init} \text{ in } \mathbb{R}^d \end{aligned} \quad (17)$$

and if $p \geq 2$, $M_0 = (M_{0,e}, M_{0,c}) \in W^{q,f}(0, T, X)^2$ satisfies for all $V = (V_e, V_c) \in X^2$

$$\begin{aligned} & \left\langle \partial_t^{q,f} M_{0,e}, V_e \right\rangle_{H^{-1}, H^1} + (D_e \nabla_x M_{0,e}, \nabla V_e)_{L^2} = 0 \quad \text{in } \mathcal{D}'([0, T]) \\ & M_{0,e}(\cdot, 0) = M_{init} \quad \text{in } \mathbb{R}^d \\ & \left\langle \partial_t^{q,f} M_{0,c}, V_c \right\rangle_{H^{-1}, H^1} + (D_c \nabla_x M_{0,c}, \nabla V_c)_{L^2} = 0 \quad \text{in } \mathcal{D}'([0, T]) \\ & M_{0,c}(\cdot, 0) = M_{init} \quad \text{in } \mathbb{R}^d \end{aligned} \quad (18)$$

For $p \leq 0$, we claim that we have the following result :

Theorem 3.2. For $p \leq 0$, the solution M_ε of (10) satisfies, up to a subsequence:

$$M_\varepsilon \rightharpoonup M_0 \quad \text{weakly in } L^\infty(0, T, L^2(\mathbb{R}^d)) \quad (19)$$

where $M_0 \in W^{q,f}(0, T, H^1(\mathbb{R}^d))$ satisfies for all $V \in H^1(\mathbb{R}^d)$:

$$\begin{aligned} & \left\langle \partial_t^{q,f} M_0, V \right\rangle_{H^{-1}, H^1} + (D \nabla_x M_0, \nabla V)_{L^2} = 0 \quad \text{in } \mathcal{D}'([0, T]) \\ & M_0(\cdot, 0) = M_{init} \quad \text{in } \mathbb{R}^d \end{aligned}$$

where only the homogenized tensor changes with the values of $p \leq 0$. For $p = 0$, it is given by :

$$D_{\alpha,ij} = \sum_{\alpha \in \{c,e\}} \int_{Y_\alpha} \sigma_\alpha (\nabla w_{j,\alpha} + e_j) \cdot (\nabla w_{i,\alpha} + e_i) \quad (20)$$

where the $w_{i,\alpha}$ are solutions for $i = 1, 2$ of the cell problems :

$$\begin{cases} -\operatorname{div}_y (\sigma_\alpha (\nabla_y w_{i,\alpha} + e_i)) = 0 & \text{in } Y_\alpha \\ \sigma_e \nabla_y w_{i,e} \cdot \nu + \sigma_e e_i \cdot \nu = 0 & \text{on } \Gamma_m \\ \sigma_c \nabla_y w_{i,c} \cdot \nu + \sigma_c e_i \cdot \nu = 0 & \text{on } \Gamma_m \\ w_{i,e} \text{ } Y_e - \text{periodic} \end{cases} \quad (21)$$

For $p = -1$, it is given by :

$$D_{\alpha,ij} = \sum_{\alpha \in \{c,e\}} \int_{Y_\alpha} \sigma_\alpha (\nabla w_{j,\alpha} + e_j) \cdot (\nabla w_{i,\alpha} + e_i) + \sigma_m \int_{\Gamma_m} (w_{j,e} - w_{j,c})(w_{i,e} - w_{i,c}) \quad (22)$$

where the $w_{i,\alpha}$ are solutions for $i = 1, 2$ of the cell problems :

$$\begin{cases} -\operatorname{div}_y(\sigma_\alpha(\nabla_y w_{i,\alpha} + e_i)) = 0 & \text{in } Y_\alpha \\ \sigma_\alpha \nabla_y w_{i,\alpha} \cdot \nu + \sigma_\alpha e_i \cdot \nu = \sigma_m (w_{i,e} - w_{i,c}) & \text{on } \Gamma_m \\ \sigma_e \nabla_y w_{i,e} \cdot \nu + \sigma_e e_i \cdot \nu = \sigma_c \nabla_y w_{i,c} \cdot \nu + \sigma_c e_i \cdot \nu & \text{on } \Gamma_m \\ w_{i,e} \text{ } Y_e \text{ - periodic} \end{cases} \quad (23)$$

For $p \leq -2$, it is given by :

$$D_{ij} = \sum_{\alpha \in \{c,e\}} \int_{Y_\alpha} \sigma_\alpha (\nabla w_{j,\alpha} + e_j) \cdot (\nabla w_{i,\alpha} + e_i) \quad (24)$$

where the $w_{i,\alpha}$ are solutions for $i = 1, 2$ of the cell problems :

$$\begin{cases} -\operatorname{div}_y(\sigma_\alpha(\nabla_y w_{i,\alpha} + e_i)) = 0 & \text{in } Y_\alpha \\ w_{i,e} = w_{i,c} & \text{on } \Gamma_m \\ \sigma_e \nabla_y w_{i,e} \cdot \nu + \sigma_e e_i \cdot \nu = \sigma_c \nabla_y w_{i,c} \cdot \nu + \sigma_c e_i \cdot \nu & \text{on } \Gamma_m \\ w_{i,e} \text{ } Y_e \text{ - periodic} \end{cases} \quad (25)$$

Before turning to the proof of these two results, we emphasize their practical interest by studying the properties and numerical behavior of the limit models.

3.2 Some remarks on the macroscopic models

As we have assumed that Y_c does not touch the boundary of Y the homogenized systems (17) and (18) can be simplified since in that case $D_c = 0$. This can be easily seen by checking that $w_{i,c} = -y_i$ and therefore $\nabla_y w_{i,c} + e_i = 0$. Consequently :

$$\begin{aligned} & \left\langle \partial_t^{q,f} M_{0,e}, V_e \right\rangle_{H^{-1}, H^1} + (D_e \nabla_x M_{0,e}, \nabla V_e)_{L^2} + (\eta_e (M_{0,e} - M_{0,c}), V_e)_{L^2} = 0 \\ & \text{in } \mathcal{D}'([0, T]), \quad M_{0,e}(\cdot, 0) = M_{init} \text{ in } \mathbb{R}^d \\ & \left\langle \partial_t^{q,f} M_{0,c}, V_c \right\rangle_{H^{-1}, H^1} + (\eta_c (M_{0,c} - M_{0,e}), V_c)_{L^2} = 0 \\ & \text{in } \mathcal{D}'([0, T]), \quad M_{0,c}(\cdot, 0) = M_{init} \text{ in } \mathbb{R}^d \end{aligned}$$

(26)

for $p = 1$, and

$$\begin{aligned} & \left\langle \partial_t^{q,f} M_{0,e}, V_e \right\rangle_{H^{-1}, H^1} + (D_e \nabla_x M_{0,e}, \nabla V_e)_{L^2} = 0 & \text{in } \mathcal{D}'([0, T]) \\ & M_{0,e}(\cdot, 0) = M_{init} & \text{in } \mathbb{R}^d \\ & \left\langle \partial_t^{q,f} M_{0,c}, V_c \right\rangle_{H^{-1}, H^1} = 0 & \text{in } \mathcal{D}'([0, T]) \\ & M_{0,c}(\cdot, 0) = M_{init} & \text{in } \mathbb{R}^d \end{aligned}$$

(27)

for $p = 2$.

Let us also remark that since the boundary of Y is contained into the boundary of Y_e , the homogenized tensor D_e is positive definite as soon as σ_e is also positive definite [3]. Now, using the same change of unknowns than in the previous section, and reproducing the analysis of theorems (2.2) and (2.1), using the now established positivity of the tensors, one can obtain the well-posedness of the macroscopic models. For $p \geq 1$, the well-posedness is obtained in the space $W^{q,f}(0, T, \tilde{X}) \cap C^0(0, T, H)$, where $\tilde{X} = H^1 \times L^2$, using the density of X in $\tilde{X}$. In the case where D_c is also a positive definite matrix (which can be the case if we lift the connection hypothesis on Y_c), then the well-posedness is obtained in $W^{q,f}(0, T, X)$ (see [7] for details). For $p \leq 0$, the well-posedness is obtained in $W^{q,f}(0, T, H^1)$.

Assume now that our solutions are regular and decreasing enough, so that every term has a classical meaning (including the one involving an unbounded coefficient). Then, the equation for $M_{0,c}$ can be explicitly solved. For $p = 1$, we get :

$$M_{0,c}(x, t) = M_{init}G_c(x, t, 0) + \int_0^t G_c(x, t, s)M_{0,e}(x, s)ds \quad (28)$$

while for $p = 2$ we have :

$$M_{0,c}(x, t) = M_{init}G_c(x, t, 0) \quad (29)$$

where we have denoted :

$$G_c(x, t, s) = \exp\left(-\int_s^t (\imath \mathbf{q} \cdot \mathbf{x} f(r) + \eta_c)dr\right) \quad (30)$$

Thus, we can decouple the system (26) (setting $Q = \mathbf{q} \cdot \mathbf{x}$):

$$\left\{ \begin{array}{l} \frac{\partial M_{0,e}}{\partial t} + (\imath Q f(t) + \eta_e)M_{0,e} - \text{div}_x(D_e \nabla_x M_{0,e}) - \eta_e \int_0^t G_c(t, s)M_{0,e}(s)ds \\ = \eta_e M_{init}G_c(t) \quad \text{in } \mathbb{R}^d \times]0, T[\\ M_{0,e}(\cdot, 0) = M_{init} \quad \text{in } \mathbb{R}^d \\ M_{0,c} = M_{init}G_c(t, 0) + \int_0^t G_c(t, s)M_{0,e}(s)ds \quad \text{in } \mathbb{R}^d \times]0, T[\end{array} \right. \quad (31)$$

and we can also rewrite (27) :

$$\left\{ \begin{array}{l} \frac{\partial M_{0,e}}{\partial t} + \imath Q f(t)M_{0,e} - \text{div}_x(D_e \nabla_x M_{0,e}) = 0 \quad \text{in } \mathbb{R}^d \times]0, T[\\ M_{0,e}(\cdot, 0) = M_{init} \quad \text{in } \mathbb{R}^d \\ M_{0,c} = M_{init}G_c \quad \text{in } \mathbb{R}^d \times]0, T[\end{array} \right. \quad (32)$$

The first equation of (31) emphasizes the fact that this is the only macroscopic model, among those we have obtained, which will behave quite differently from a Bloch-Torrey equation. In particular, the presence of the integral with respect to time will give birth to memory effects for $M_{0,e}$.

3.3 Numerical exploration of the macroscopic models for diffusion MRI

Now we compare our macroscopic models to the full two-compartment model, with realistic values of the physical parameters, corresponding to practical situations of diffusion MRI.

The initial magnetization M_{init} is set to a known value, and the most commonly used normalized time profile of the magnetic field f is the Pulsed Gradient Spin Echo (PGSE) :

$$f(t) = \begin{cases} 1 & \text{for } 0 \leq t \leq \delta \\ 0 & \text{for } \delta \leq t \leq \Delta \\ -1 & \text{for } \Delta \leq t \leq \Delta + \delta \end{cases}$$

where $\delta, \Delta > 0$. Notice that

$$\int_0^{\Delta+\delta} f(s)ds = 0$$

We denote $T_E = \Delta + \delta$ the so-called time-echo. Several values $q = \gamma g$ are used for a fixed direction of the magnetic field gradient, where g is the gradient amplitude and γ is a fixed positive number. The signal or measurement is obtained at the time-echo T_E , for several directions g , and is given by :

$$\tilde{S}_\varepsilon(g) = \int_{\mathbb{R}^d} M_\varepsilon(x, T_E) dx$$

This signal is then normalized by dividing by $\int_{\mathbb{R}^d} M_{init}$. The normalized signal will be denoted by $S_\varepsilon(g)$:

$$S_\varepsilon(g) = \frac{\int_{\mathbb{R}^d} M_\varepsilon(x, T_E) dx}{\int_{\mathbb{R}^d} M_{init}}$$

From theorems 3.1 and 3.2, we know that in the case $p \geq 1$, we have :

$$S_\varepsilon(g) \rightarrow \left(\frac{|Y_e|}{|Y|} \int_{\mathbb{R}^d} M_{0,e}(x, T_E) dx + \frac{|Y_c|}{|Y|} \int_{\mathbb{R}^d} M_{0,c}(x, T_E) dx \right) \left(\int_{\mathbb{R}^d} M_{init} \right)^{-1}$$

while in the case $p \leq 0$, we have :

$$S_\varepsilon(g) \rightarrow \left(\int_{\mathbb{R}^d} M_0(x, T_E) dx \right) \left(\int_{\mathbb{R}^d} M_{init} \right)^{-1}$$

Naturally, we will compare the results to these quantities, which will be the signals $S^q(g)$ for the different values of the parameter q .

We perform the numerical simulations with initial data $M_{init} = 1$. Remark that this case does not strictly enters our previous setting : indeed, such an initial condition is not $L^2(\mathbb{R}^d)$. However, it can be easily shown that the solution of (3) for this initial condition is quasi-periodic in space, with period ε and quasi-periodic coefficient $\exp(-\imath \varepsilon q_i \int_0^t f(s)ds)$ in each direction e_i . Thus, we can solve (3) on a single periodic cell in that case, allowing us to perform numerical simulations (the domain now being bounded). The normalized signal on the whole space is then defined as :

$$S_\varepsilon(g) = \lim_{N \rightarrow +\infty} \frac{\int_{\Omega_N} M_\varepsilon(x, T_E) dx}{\int_{\Omega_N} M_{init}}$$

where Ω_N is the union of N cells, and it is easy to see that we get, for a constant M_{init} :

$$S_\varepsilon(g) = \frac{\int_{Y_\varepsilon} M_\varepsilon(x, T_E) dx}{|Y_\varepsilon| M_{init}}$$

where Y_ε denotes a single periodicity cell. For the homogenized problem, the same holds with for instance Y as periodic cell, the problem being quasi-periodic for any period, coefficients being constant in space.

For meshing simplicity, we perform numerical simulations in dimension 2 and we use circular biological cells, of radius R_m . The values we use for the other coefficients are in the range of values experimentally measured for the biological tissues involved in diffusion MRI. It is common in the diffusion MRI community not to display S_ε as a function of g , but as a function of the so called b-value :

$$b(q) = \|q\|^2 \int_0^{T_E} \left(\int_0^t f(s)ds \right)^2 dt$$

(see [7] for details about the reason why this quantity is used). We display a comparison of $\log S(b(q))$ for the different models, comparing them to the two-compartment model, on figures 3 and 4. We

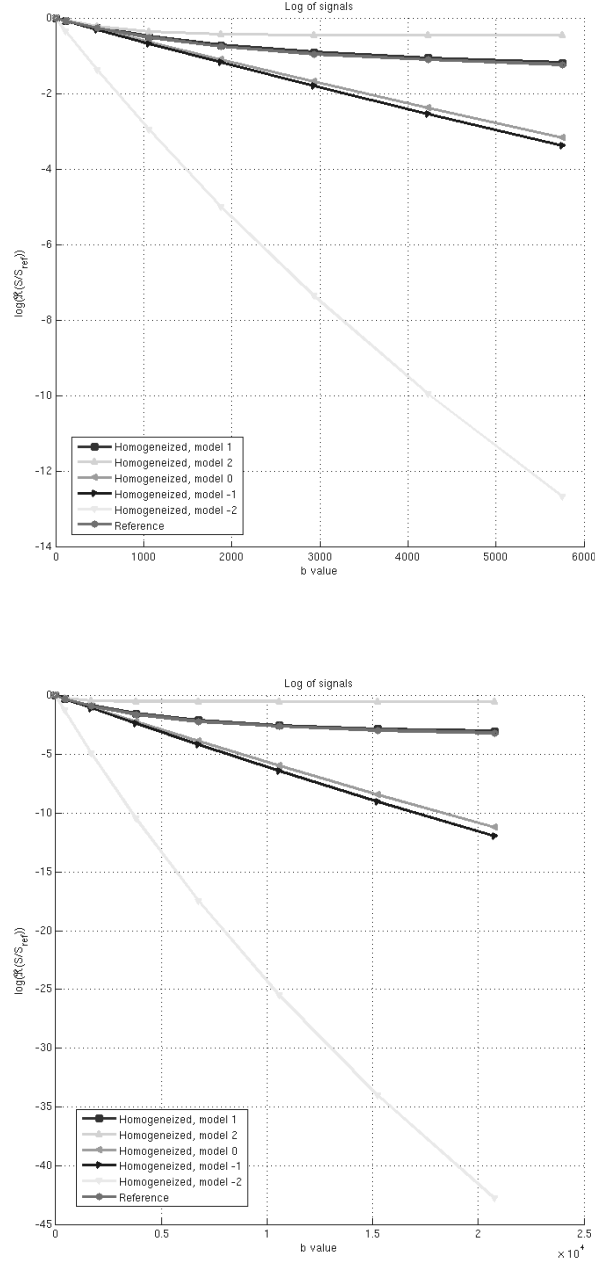


Figure 2: Log of the normalized signals, $\sigma_e = \sigma_c = 3e - 3mm^2s^{-1}$, $\kappa^\varepsilon = 5e - 5ms^{-1}$, $Rm = 0.45$, $\delta = 3.5ms$, $\Delta = 5, 15, 25, 35ms$

clearly see on these figures that one of the five models seems to reproduce the behavior of the signal much better than the four other, in all the tested situations. This model is the limit model when $p = 1$, i.e. the coupled macroscopic model (17). In particular, it is the only model that accurately reproduces the 'curvature' of the signal : the signal obtained from the full model is not a straight

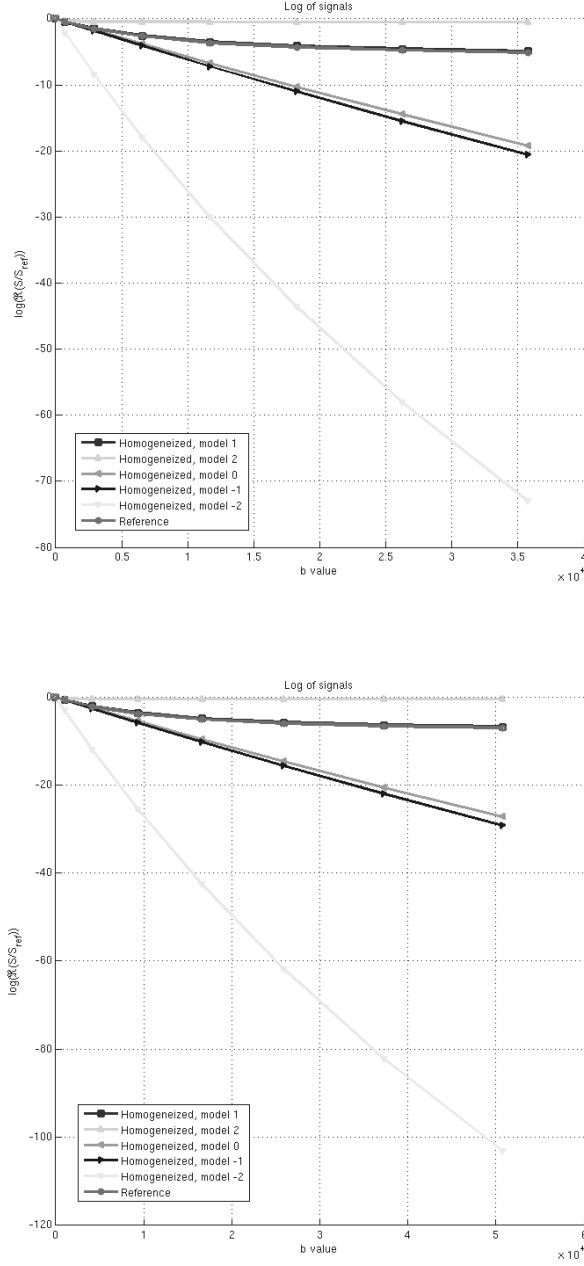


Figure 3: Log of the normalized signals, $\sigma_e = \sigma_c = 3e - 3mm^2s^{-1}$, $\kappa^e = 5e - 5ms^{-1}$, $Rm = 0.45$, $\delta = 3.5ms$, $\Delta = 5, 15, 25, 35ms$

line, as a single Bloch-Torrey equation with homogeneous coefficients would produce (see [7]). The models with $p \leq 0$, can consequently not reproduce such behavior, while for $p = 2$ the 'curvature' is clearly not the right one. In fact, this 'curvature' corresponds to the macroscopic counterpart of the cell membranes. As the model for $p = 1$ is the only one which transports on the equation the

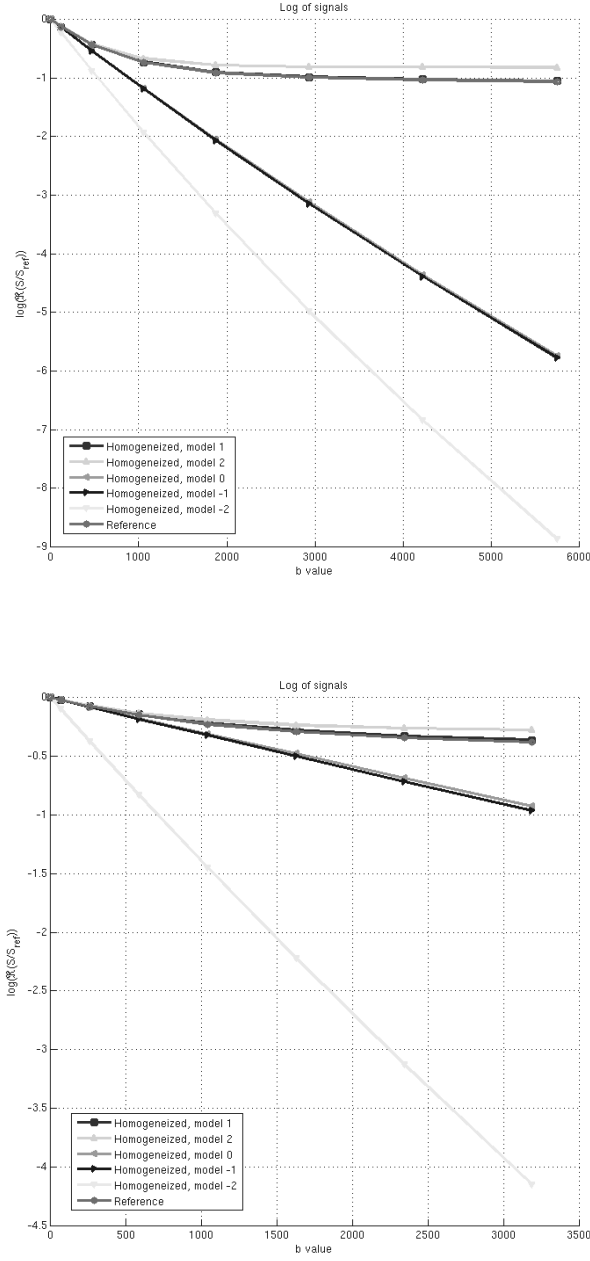


Figure 4: Log of the normalized signals, $\sigma_e = 3e - 3mm^2s^{-1}$, $\sigma_c = 1e - 3mm^2s^{-1}$, $\kappa^e = 1e - 5ms^{-1}$, $Rm = 0.375$, $\delta = 3.5ms$, $\Delta = 5ms$ and $\sigma_e = 4e - 3mm^2s^{-1}$, $\sigma_c = 1e - 3mm^2s^{-1}$, $\kappa^e = 1e - 5ms^{-1}$, $Rm = 0.49$, $\delta = 2.5ms$, $\Delta = 5ms$

effect of the membranes through the coupling coefficients, it seems logical that it is indeed the most accurate. For a more thorough numerical study of the limit model when $p = 1$, we refer the reader to [13].

4 Some results from the unfolding method in homogenization

In this section, we present the tools that will be needed to rigorously establish the macroscopic models by homogenization of the two compartment model. The results of this section are extracted from [4], [5], [6] and [9], with the slight difference that we consider here the time dependent case. As the time variable plays the role of a parameter, the extension is straightforward, which is why we do not recall any proof here. However, we felt that it would be more convenient for the reader to sum up these results here.

4.1 The periodic unfolding operator

We first need to introduce some notations to describe our periodic domain. We denote $[z]_Y$ the unique integer combination of the vectors e_i of the canonical basis of $\mathbb{R}^d$ such that $z - [z]_Y$ belongs to Y , and set $\{z\}_Y = z - [z]_Y$. Then, for each $x \in \mathbb{R}^d$, we have :

$$x = \varepsilon \left(\left[\frac{x}{\varepsilon} \right]_Y + \left\{ \frac{x}{\varepsilon} \right\}_Y \right)$$

The periodic unfolding operators $\mathcal{T}_\varepsilon^\alpha$ for $\alpha \in \{c, e\}$ are defined for any ϕ Lebesgue measurable on $\Omega_\alpha^\varepsilon \times]0, T[$, by :

$$\mathcal{T}_\varepsilon^\alpha(\phi)(x, y, t) = \phi \left(\varepsilon \left[\frac{x}{\varepsilon} \right]_Y + \varepsilon y, t \right) \quad \text{for a.e. } (x, y, t) \in \mathbb{R}^d \times Y_\alpha \times]0, T[\quad (33)$$

While ϕ is defined only on $\Omega_\alpha^\varepsilon \times]0, T[$, $\mathcal{T}_\varepsilon^\alpha(\phi)$ is defined on $\mathbb{R}^d \times Y_\alpha \times]0, T[$, where it is Lebesgue-measurable. It is obvious from the definition that, for v and w Lebesgue measurable on $\Omega_\alpha^\varepsilon \times]0, T[$:

$$\mathcal{T}_\varepsilon^\alpha(vw) = \mathcal{T}_\varepsilon^\alpha(v)\mathcal{T}_\varepsilon^\alpha(w) \quad (34)$$

Convergence properties of the unfolding operator with respect to ε , are expressed through the mean value operator $\mathcal{M}_{Y_\alpha} : L^p(\mathbb{R}^d \times Y_\alpha) \rightarrow L^p(\mathbb{R}^d)$, for $p \in [1, +\infty]$, defined by :

$$\mathcal{M}_{Y_\alpha}(\Phi)(x) = \frac{1}{|Y_\alpha|} \int_{Y_\alpha} \Phi(x, y) dy \quad (35)$$

for which using Hölder's inequality, one can immediately see that for $\Phi \in L^\beta(0, T, L^2(\mathbb{R}^d \times Y_\alpha))$ and $\beta \in \{2, \infty\}$:

$$\|\mathcal{M}_{Y_\alpha}(\Phi)\|_{L^\beta(0, T, L^2(\mathbb{R}^d))} \leq |Y_\alpha|^{-1/2} \|\Phi\|_{L^\beta(0, T, L^2(\mathbb{R}^d \times Y_\alpha))} \quad (36)$$

Theorem 4.1. *For f measurable on $Y_\alpha \times]0, T[$, extended by Y periodicity on all Ω_α , we define the sequence :*

$$f_\varepsilon(x, t) = f \left(\frac{x}{\varepsilon}, t \right) \quad \text{a.e. for } x \in \Omega_\alpha^\varepsilon \quad (37)$$

Then,

$$\mathcal{T}_\varepsilon^\alpha(f_\varepsilon)(x, y, t) = f(y, t) \quad \text{for a.e. } (x, y, t) \in \mathbb{R}^d \times Y_\alpha \times]0, T[$$

The operator $\mathcal{T}_\varepsilon^\alpha$ is linear and continuous from $L^2(0, T, L^2(\Omega_\alpha^\varepsilon))$ to $L^2(0, T, L^2(\mathbb{R}^d \times Y_\alpha))$ or from $L^\infty(0, T, L^2(\Omega_\alpha^\varepsilon))$ to $L^\infty(0, T, L^2(\mathbb{R}^d \times Y_\alpha))$. For any $\phi \in L^1(0, T, \Omega_\alpha^\varepsilon)$ and any $w \in L^2(0, T, L^2(\Omega_\alpha^\varepsilon))$, we have :

$$(i) \quad \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_\alpha} \mathcal{T}_\varepsilon^\alpha(\phi)(x, y, t) dx dy dt = \int_0^T \int_{\Omega_\alpha^\varepsilon} \phi(x, t) dx dt$$

$$(ii) \quad \|\mathcal{T}_\varepsilon^\alpha(\phi)\|_{L^2(0, T, L^2(\mathbb{R}^d \times Y_\alpha))} = |Y|^{1/2} \|\phi\|_{L^2(0, T, L^2(\Omega_\alpha^\varepsilon))}$$

Similarly, for any $\phi \in L^\infty(0, T, \Omega_\alpha^\varepsilon)$ and any $w \in L^\infty(0, T, L^2(\Omega_\alpha^\varepsilon))$, we have :

$$(i) \quad \frac{1}{|Y|} \int_{\mathbb{R}^d \times Y_\alpha} \mathcal{T}_\varepsilon^\alpha(\phi)(x, y, t) dx dy = \int_{\Omega_\alpha^\varepsilon} \phi(x, t) dx$$

$$(ii) \quad \|\mathcal{T}_\varepsilon^\alpha(\phi)\|_{L^\infty(0, T, L^2(\mathbb{R}^d \times Y_\alpha))} = |Y|^{1/2} \|\phi\|_{L^\infty(0, T, L^2(\Omega_\alpha^\varepsilon))}$$

In the remaining of this section, we will use $\beta \in \{2, \infty\}$ to state simultaneously the results for L^2 in time functions and L^∞ in time functions, in order to avoid repetitions. To do so, we introduce the notation " β -weak convergence", to be understood as weak convergence in the case $\beta = 2$, and weak* convergence in the case $\beta = \infty$.

Theorem 4.2. *The unfolding operator $\mathcal{T}_\varepsilon^\alpha$ has the following convergence properties :*

(i) For $w \in L^\beta(0, T, L^2(\mathbb{R}^d))$,

$$\mathcal{T}_\varepsilon^\alpha(w) \rightarrow w \quad \text{strongly in } L^\beta(0, T, L^2(\mathbb{R}^d \times Y_\alpha))$$

and

$$\chi_{\varepsilon, \alpha} w \rightharpoonup \frac{|Y_\alpha|}{|Y|} w \quad \beta\text{-weakly in } L^\beta(0, T, L^2(\mathbb{R}^d))$$

where $\chi_{\varepsilon, \alpha}$ is the characteristic function of $\Omega_\alpha^\varepsilon$.

(ii) Let $(w_\varepsilon)_\varepsilon$ be a sequence in $L^\beta(0, T, L^2(\mathbb{R}^d))$ such that :

$$w_\varepsilon \rightarrow w \quad \text{strongly in } L^\beta(0, T, L^2(\mathbb{R}^d))$$

then

$$\mathcal{T}_\varepsilon^\alpha(w_\varepsilon) \rightarrow w \quad \text{strongly in } L^\beta(0, T, L^2(\mathbb{R}^d \times Y_\alpha))$$

(iii) Let $(w_\varepsilon)_\varepsilon$ be a bounded sequence in $L^\beta(0, T, L^2(\Omega_\alpha^\varepsilon))$, i.e. there exists $C > 0$ such that :

$$\|w_\varepsilon\|_{L^\beta(0, T, L^2(\Omega_\alpha^\varepsilon))} \leq C$$

Then the corresponding sequence $(\mathcal{T}_\varepsilon^\alpha(w_\varepsilon))_\varepsilon$ is bounded in $L^\beta(0, T, L^2(\mathbb{R}^d \times Y_\alpha))$. Moreover, if

$$\mathcal{T}_\varepsilon^\alpha(w_\varepsilon) \rightharpoonup \hat{w} \quad \beta\text{-weakly in } L^\beta(0, T, L^2(\mathbb{R}^d \times Y_\alpha))$$

then

$$\chi_{\varepsilon, \alpha} w_\varepsilon \rightharpoonup \frac{|Y_\alpha|}{|Y|} \mathcal{M}_{Y_\alpha}(\hat{w}) \quad \beta\text{-weakly in } L^\beta(0, T, L^2(\mathbb{R}^d))$$

(iv) If $\mathcal{T}_\varepsilon^\alpha(w_\varepsilon) \rightharpoonup \hat{w}$ in $L^\beta(0, T, L^2(\mathbb{R}^d \times Y_\alpha))$, then

$$\|\hat{w}\|_{L^\beta(0, T, L^2(\mathbb{R}^d \times Y_\alpha))} \leq \liminf_{\varepsilon \rightarrow 0} |Y|^{1/2} \|w_\varepsilon\|_{L^\beta(0, T, L^2(\Omega_\alpha^\varepsilon))} \quad (38)$$

(v) Let $f \in L^\beta(0, T, L^2(Y_\alpha))$, and f_ε be the sequence defined in (37). Then:

$$\chi_{\varepsilon, \alpha} f_\varepsilon \rightharpoonup \frac{|Y_\alpha|}{|Y|} \mathcal{M}_{Y_\alpha}(f) \quad \beta\text{-weakly in } L^\beta(0, T, L^2(\mathbb{R}^d)) \quad (39)$$

Now, we introduce the local average operator $\mathcal{M}_\varepsilon^\alpha : L^2(0, T, L^2(\Omega_\alpha^\varepsilon)) \rightarrow L^2(0, T, L^2(\mathbb{R}^d))$, which is nothing but the average of a function on each periodicity cell. It is defined as follows :

$$\mathcal{M}_\varepsilon^\alpha(\phi)(x, t) = \frac{1}{\varepsilon^d |Y_\alpha|} \int_{\varepsilon([\frac{x}{\varepsilon}]_Y + Y_\alpha)} \phi(\zeta, t) d\zeta \quad \text{for a.e. } (x, t) \in \mathbb{R}^d \times]0, T[$$

Using the obvious change of variable in each cell, we have for a.e. $(x, t) \in \mathbb{R}^d \times]0, T[$:

$$\begin{aligned} \mathcal{M}_\varepsilon^\alpha(\phi)(x, t) &= \frac{1}{\varepsilon^d |Y_\alpha|} \int_{\varepsilon([\frac{x}{\varepsilon}]_Y + Y_\alpha)} \phi(\zeta) d\zeta = \frac{1}{\varepsilon^d |Y_\alpha|} \int_{Y_\alpha} \phi(\varepsilon [\frac{x}{\varepsilon}]_Y + \varepsilon y) \varepsilon^d dy \\ &= \mathcal{M}_{Y_\alpha}(\mathcal{T}_\varepsilon^\alpha(\phi))(x, t) \end{aligned}$$

We also have on $\mathbb{R}^d \times Y_\alpha$:

$$\mathcal{T}_\varepsilon^\alpha(\mathcal{M}_\varepsilon^\alpha(\phi)) = \mathcal{M}_\varepsilon^\alpha(\phi)$$

as :

$$\left[\frac{\varepsilon [\frac{x}{\varepsilon}]_Y + \varepsilon y}{\varepsilon} \right]_Y = \left[\left[\frac{x}{\varepsilon} \right]_Y + y \right]_Y = \left[\frac{x}{\varepsilon} \right]_Y$$

4.2 The unfolding operator and gradients

First, remark that for $w \in L^2(0, T, H^1(\Omega_\alpha^\varepsilon))$, we have :

$$\nabla_y \mathcal{T}_\varepsilon^\alpha(w) = \varepsilon \mathcal{T}_\varepsilon^\alpha(\nabla w) \quad \text{for a.e. } (x, y, t) \in \mathbb{R}^d \times Y_\alpha \times]0, T[\quad (40)$$

Consequently, from point (iii) of theorem 4.1 we immediately deduce that $\mathcal{T}_\varepsilon^\alpha$ maps $L^2(0, T, H^1(\Omega_\alpha^\varepsilon))$ into $L^2(0, T, L^2(\Omega, H^1(Y_\alpha)))$. The next theorem, extracted from [6] and [9], describes the convergence properties of the unfolding operators acting on gradients :

Theorem 4.3. *Let $(w_\varepsilon)_\varepsilon$ be a sequence of $L^2(0, T, H^1(\Omega_\alpha^\varepsilon))$ such that :*

$$\|w_\varepsilon\|_{L^2(0, T, H^1(\Omega_\alpha^\varepsilon))} \leq C \quad (41)$$

Then there exists w in $L^2(0, T, H^1(\mathbb{R}^d))$ such that, up to a subsequence :

$$\mathcal{T}_\varepsilon^\alpha(w_\varepsilon) \rightharpoonup w \quad \text{weakly in } L^2(0, T, L^2(\Omega, H^1(Y_\alpha))) \quad (42)$$

Moreover, there exists $\hat{w} \in L^2(0, T, L^2(\Omega, H^1(Y_\alpha)))$ such that $\mathcal{M}_{Y_\alpha}(\hat{w}) = 0$ and, up to a subsequence:

$$\begin{aligned} (i) \quad & \mathcal{T}_\varepsilon^\alpha(\nabla w_\varepsilon) \rightharpoonup \nabla w + \nabla_y \hat{w} \quad \text{weakly in } L^2(0, T, L^2(\mathbb{R}^d \times Y_\alpha)) \\ (ii) \quad & \frac{1}{\varepsilon} (\mathcal{T}_\varepsilon^\alpha(w_\varepsilon) - \mathcal{M}_\varepsilon^\alpha(w_\varepsilon)) \rightharpoonup \hat{w} + y_\alpha \cdot \nabla w \quad \text{weakly in } L^2(0, T, L^2(\mathbb{R}^d, H^1(Y_\alpha))) \end{aligned} \quad (43)$$

where :

$$y_\alpha = y - \mathcal{M}_{Y_\alpha}(y)$$

If $\alpha = c$, then $\hat{w} \in L^2(0, T, L^2(\Omega, H_\#^1(Y_e)))$.

4.3 The unfolding operator on the boundary

Denote

$$\Gamma_\varepsilon = \bigcup_{\xi \in \mathbb{Z}^d} \varepsilon(\xi + \Gamma_m)$$

For any φ Lebesgue measurable on $\Gamma_\varepsilon \times]0, T[$, the boundary unfolding operator $\mathcal{T}_\varepsilon^\Gamma$ is defined by:

$$\mathcal{T}_\varepsilon^\Gamma(\varphi)(x, y, t) = \varphi\left(\varepsilon \left[\frac{x}{\varepsilon}\right]_Y + \varepsilon y\right) \quad \text{for a.e. } (x, y, t) \in \mathbb{R}^d \times \Gamma_m \times]0, T[\quad (44)$$

As before, if (φ, ψ) are two functions Lebesgue-measurable on $\Gamma_\varepsilon \times]0, T[$, we have:

$$\mathcal{T}_\varepsilon^\Gamma(\varphi\psi) = \mathcal{T}_\varepsilon^\Gamma(\varphi)\mathcal{T}_\varepsilon^\Gamma(\psi) \quad (45)$$

One can immediately remark that for $\phi \in L^2(0, T, H^1(\Omega_\alpha^\varepsilon))$, $\alpha \in \{c, e\}$, one has simply:

$$\mathcal{T}_\varepsilon^\Gamma(\phi) = \mathcal{T}_\varepsilon^\alpha(\phi)|_{\Gamma_\varepsilon} \quad (46)$$

i.e. $\mathcal{T}_\varepsilon^\Gamma(\phi)$ is the trace of $\mathcal{T}_\varepsilon^\alpha(\phi)$ on Γ_ε . It follows from the trace theorem that there exists a constant $C > 0$ independent on ε such that :

$$\|\mathcal{T}_\varepsilon^\Gamma(\phi)\|_{L^2(0, T, L^2(\mathbb{R}^d, H^{1/2}(\Gamma_m)))} \leq C \|\mathcal{T}_\varepsilon^\alpha(\phi)\|_{L^2(0, T, L^2(\mathbb{R}^d, H^1(Y_\alpha)))} \quad (47)$$

Theorem 4.4. *The boundary unfolding operator $\mathcal{T}_\varepsilon^\Gamma$ as the following properties :*

(i) *For any $\varphi \in L^1(0, T, L^1(\Gamma_\varepsilon))$, we have :*

$$\int_0^T \int_{\Gamma_\varepsilon} \varphi d\mu(x) dt = \frac{1}{\varepsilon|Y|} \int_0^T \int_{\mathbb{R}^d \times \Gamma_m} \mathcal{T}_\varepsilon^\Gamma(\varphi)(x, y, t) dx d\mu(y) dt \quad (48)$$

(ii) For any $\varphi \in L^2(0, T, L^2(\Gamma_\varepsilon))$

$$\|\mathcal{T}_\varepsilon^\Gamma(\varphi)\|_{L^2(0, T, L^2(\mathbb{R}^d \times \Gamma_m))} \leq \varepsilon^{1/2} |Y|^{1/2} \|\varphi\|_{L^2(0, T, L^2(\Gamma_\varepsilon))}$$

(iii) Let $(w_\varepsilon)_\varepsilon$ be a $L^2(0, T, H^1(\Omega_\alpha^\varepsilon))$ sequence and φ_ε a $L^2(0, T, L^2(\Gamma_\varepsilon))$ sequence such that :

$$\begin{aligned} \mathcal{T}_\varepsilon^\alpha(w_\varepsilon) &\rightarrow w \quad \text{weakly in } L^2(0, T, L^2(\mathbb{R}^d, H^1(Y_\alpha))) \\ \mathcal{T}_\varepsilon^\Gamma(\varphi_\varepsilon) &\rightarrow \varphi \quad \text{strongly in } L^2(0, T, L^2(\mathbb{R}^d \times \Gamma_m)) \end{aligned}$$

Then :

$$\varepsilon \int_{\Gamma_\varepsilon} \varphi_\varepsilon w_\varepsilon d\mu(x) dt \rightarrow \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times \Gamma_m} \varphi(x, y, t) w(x, y, t) dx d\mu(y) dt$$

(iv) Let $w_\varepsilon \in L^2(0, T, H^1(\Omega_{ext}^\varepsilon))$. Then :

$$\frac{1}{\varepsilon |Y|} \int_0^T \int_{\mathbb{R}^d \times \Gamma_m} (\mathcal{T}_\varepsilon^e(w_\varepsilon) - \mathcal{T}_\varepsilon^c(w_\varepsilon))^2 \leq \int_0^T \int_{\Gamma_\varepsilon} (w_{\varepsilon, e} - w_{\varepsilon, c})^2$$

We conclude this section by a useful result, especially designed for handling trace jumps between Ω_e^ε and Ω_c^ε . It is essentially a rewriting of a result of [9], in a form that is well suited for handling our particular problem :

Theorem 4.5. Let $w_\varepsilon = (w_{\varepsilon, e}, w_{\varepsilon, c})$ be a sequence of $L^2(0, T, H^1(\Omega_{ext}^\varepsilon)) = L^2(0, T, H^1(\Omega_e^\varepsilon)) \times L^2(0, T, H^1(\Omega_c^\varepsilon))$ such that there exists $\gamma \geq 1/2$ and $C > 0$ independent on ε such that:

$$\|w_\varepsilon\|_{L^2(0, T, H^1(\Omega_{ext}^\varepsilon))} \leq C \quad \text{and} \quad \|w_{\varepsilon, e} - w_{\varepsilon, c}\|_{L^2(0, T, L^2(\Gamma_\varepsilon))} \leq C\varepsilon^\gamma$$

Let $(\varphi_\varepsilon)_\varepsilon$ be a sequence in $L^2(0, T, H^1(\mathbb{R}^d))$. Assume that there exists

$$\varphi \in L^2(0, T, L^2(\mathbb{R}^d, H^1(\Gamma_m)))$$

such that:

$$\mathcal{T}_\varepsilon^\Gamma(\varphi_\varepsilon) \rightarrow \varphi \quad \text{strongly in } L^2(0, T, L^2(\mathbb{R}^d, H^1(Y_\alpha)))$$

Then, there exists $\hat{w}_c \in L^2(0, T, H^1(Y_c))$, $\hat{w}_e \in L^2(0, T, H_\sharp^1(Y_e))$ and $\mu \in L^2(0, T, L^2(\mathbb{R}^d))$ such that:

$$\int_0^T \int_{\Gamma_\varepsilon} \varphi_\varepsilon (w_{\varepsilon, e} - w_{\varepsilon, c}) \rightarrow \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times \Gamma_m} \varphi (\hat{w}_e - \hat{w}_c + \mu)$$

Proof. Using (45), we apply point (i) of theorem 4.4 to the product $\varphi_\varepsilon w_{\varepsilon, \alpha}$, which gives, using (46):

$$\int_0^T \int_{\Gamma_\varepsilon} \varphi_\varepsilon w_{\varepsilon, \alpha} d\mu(x) dt = \frac{1}{\varepsilon |Y|} \int_0^T \int_{\mathbb{R}^d \times \Gamma_m} \mathcal{T}_\varepsilon^\Gamma(\varphi_\varepsilon)(x, y, t) \mathcal{T}_\varepsilon^\alpha(w_\varepsilon)(x, y, t) dx d\mu(y) dt$$

This can be rewritten :

$$\begin{aligned} \int_0^T \int_{\Gamma_\varepsilon} \varphi_\varepsilon w_{\varepsilon, \alpha} d\mu(x) dt &= \frac{|\Gamma_m|}{\varepsilon |Y|} \int_0^T \int_{\mathbb{R}^d} \mathcal{M}_{\Gamma_m}(\mathcal{T}_\varepsilon^\Gamma(\varphi_\varepsilon))(x, t) \mathcal{M}_\varepsilon^\alpha(w_\varepsilon)(x, t) dx dt \\ &+ \frac{1}{\varepsilon |Y|} \int_0^T \int_{\mathbb{R}^d \times \Gamma_m} \mathcal{T}_\varepsilon^\Gamma(\varphi_\varepsilon)(x, y, t) (\mathcal{T}_\varepsilon^\alpha(w_\varepsilon) - \mathcal{M}_\varepsilon^\alpha(w_\varepsilon))(x, y, t) dx d\mu(y) dt \end{aligned}$$

Thus, we deduce that :

$$\begin{aligned} \int_0^T \int_{\Gamma_\varepsilon} \varphi_\varepsilon(w_{\varepsilon,e} - w_{\varepsilon,c}) &= \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d} \frac{1}{\varepsilon} (\mathcal{M}_\varepsilon^c(w_\varepsilon) - \mathcal{M}_\varepsilon^e(w_\varepsilon)) \mathcal{M}_{\Gamma_m}(\mathcal{T}_\varepsilon^\Gamma(\varphi_\varepsilon)) dx dt \\ &+ \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times \Gamma_m} \left[\frac{1}{\varepsilon} (\mathcal{T}_\varepsilon^e(w_\varepsilon) - \mathcal{M}_\varepsilon^e(w_\varepsilon)) - \frac{1}{\varepsilon} (\mathcal{T}_\varepsilon^c(w_\varepsilon) - \mathcal{M}_\varepsilon^c(w_\varepsilon)) \right] \mathcal{T}_\varepsilon^\Gamma(\varphi_\varepsilon) dx d\mu(y) dt \end{aligned}$$

We know from theorem 4.3 that :

$$\frac{1}{\varepsilon} (\mathcal{T}_\varepsilon^\alpha(w_\varepsilon) - \mathcal{M}_\varepsilon^\alpha(w_\varepsilon)) \rightharpoonup \hat{w}_\alpha + y_\alpha \cdot \nabla w \quad \text{weakly in } L^2(0, T, L^2(\mathbb{R}^d, H^1(Y_\alpha)))$$

and thus, applying the trace theorem (in the y variable) :

$$\frac{1}{\varepsilon} (\mathcal{T}_\varepsilon^\alpha(w_\varepsilon) - \mathcal{M}_\varepsilon^\alpha(w_\varepsilon)) \rightharpoonup \hat{w}_\alpha + y_\alpha \cdot \nabla w \quad \text{weakly in } L^2(0, T, L^2(\mathbb{R}^d \times \Gamma_m))$$

and consequently :

$$\frac{1}{\varepsilon} (\mathcal{T}_\varepsilon^e(w_\varepsilon) - \mathcal{M}_\varepsilon^e(w_\varepsilon)) - \frac{1}{\varepsilon} (\mathcal{T}_\varepsilon^c(w_\varepsilon) - \mathcal{M}_\varepsilon^c(w_\varepsilon)) \rightharpoonup \hat{w}_e - \hat{w}_c + \left(\frac{1}{|Y_c|} \int_{Y_c} y - \frac{1}{|Y_e|} \int_{Y_e} y \right) \cdot \nabla w$$

weakly in $L^2(0, T, L^2(\mathbb{R}^d \times \Gamma_m))$. Moreover, from point (iv) of theorem 4.4 and the bound on the trace jump, we know that there exists $\mu_1 \in L^2(0, T, L^2(\mathbb{R}^d \times \Gamma_m))$ such that :

$$\frac{1}{\varepsilon} (\mathcal{T}_\varepsilon^e(w_\varepsilon) - \mathcal{T}_\varepsilon^c(w_\varepsilon)) \rightharpoonup \mu_1 \quad \text{weakly in } L^2(0, T, L^2(\mathbb{R}^d \times \Gamma_m))$$

Thus, combining these two weak convergence results, we deduce that $\frac{1}{\varepsilon} (\mathcal{M}_\varepsilon^e(w_\varepsilon) - \mathcal{M}_\varepsilon^c(w_\varepsilon))$, which is independent on y , is weakly convergent in $L^2(0, T, L^2(\mathbb{R}^d \times \Gamma_m))$, to some $\mu_2 \in L^2(0, T, L^2(\mathbb{R}^d))$ also independent on y . Thus, defining :

$$\mu = \left(\frac{1}{|Y_c|} \int_{Y_c} y - \frac{1}{|Y_e|} \int_{Y_e} y \right) \cdot \nabla w + \mu_2$$

we obtain the desired result. $\square$

5 Derivation of the macroscopic models by homogenization of the two compartment model

To ensure that we work in a pleasant mathematical framework, we will apply the periodic unfolding method to $\widetilde{M}_\varepsilon$ rather than to M_ε . Then, using the equivalence between these two families of functions, the homogenized model for the original unknown can be easily deduced. Using the results of the periodic unfolding method presented in section 4, establishing rigorously the five limit models for the unknown $\widetilde{M}_\varepsilon$ is quite straightforward.

We start by recalling that $\widetilde{M}_\varepsilon$ is given as the element of $W(0, T, X_\varepsilon)$ that satisfies, for all $V \in X_\varepsilon$:

$$\begin{aligned} &\left| \begin{aligned} &\frac{d}{dt} \left\langle \widetilde{M}_\varepsilon(t), V \right\rangle_{X'_\varepsilon, X_\varepsilon} + \left(\sigma_\varepsilon(\nabla \widetilde{M}_\varepsilon(t) - \imath q \mathbf{n} F(t) \widetilde{M}_\varepsilon), \nabla V \right)_{H_\varepsilon} \\ &+ \imath q F(t) \left(\sigma_\varepsilon \nabla \widetilde{M}_\varepsilon(t) \cdot \mathbf{n}, V \right)_{H_\varepsilon} + q^2 F(t)^2 \left(\sigma_\varepsilon \widetilde{M}_\varepsilon(t), V \right)_{H_\varepsilon} \\ &+ \sum_{\xi \in \mathbb{Z}^d} \kappa^\varepsilon \left(\left[\widetilde{M}_\varepsilon(t) \right]_{\Gamma_m^{\varepsilon, \xi}}, [V]_{\Gamma_m^{\varepsilon, \xi}} \right)_{L^2(\Gamma_m^{\varepsilon, \xi})} = 0 \quad \text{in } \mathcal{D}'([0, T]) \\ &\widetilde{M}_\varepsilon(0) = M_{init} \end{aligned} \right. \end{aligned} \quad (49)$$

As $W(0, T, X_\varepsilon) \hookrightarrow C^0(0, T, H_\varepsilon)$, for any function $\Psi \in C^1([0, T], H^1(\Omega_{ext}^\varepsilon)) = (\Psi_e, \Psi_c) \in C^1([0, T], H^1(\Omega_e^\varepsilon)) \times C^1([0, T], H^1(\Omega_c^\varepsilon))$, we obtain after integrating by parts :

$$\begin{aligned}
& \int_0^T \int_{\Omega_{ext}^\varepsilon} \left[-\widetilde{M}_\varepsilon(x, t) \frac{\partial \Psi(x, t)}{\partial t} + \sigma_\varepsilon \nabla \widetilde{M}_\varepsilon(x, t) \nabla \Psi(x, t) \right] dx dt \\
& + \int_0^T \int_{\Gamma_\varepsilon} \kappa^\varepsilon(\widetilde{M}_{\varepsilon, e} - \widetilde{M}_{\varepsilon, c})(\Psi_e - \Psi_c) d\mu(x) dt \\
& + \int_0^T \int_{\Omega_{ext}^\varepsilon} \left[\imath q F(t) \nabla \widetilde{M}_\varepsilon(x, t) \cdot \mathbf{n} \Psi(x, t) - \imath q F(t) \widetilde{M}_\varepsilon(x, t) \mathbf{n} \cdot \nabla \Psi(x, t) \right] dx dt \\
& = \int_{\Omega_{ext}^\varepsilon} M_{init}(x) \Psi(x, 0) dx - \int_{\Omega_{ext}^\varepsilon} M_\varepsilon(x, T) \Psi(x, T) dx
\end{aligned} \tag{50}$$

and in particular, for any function

$$\begin{aligned}
\Psi \in C_c^1([0, T], H^1(\Omega_{ext})) = (\Psi_e, \Psi_c) \in C_c^1([0, T], H^1(\Omega_e)) \times C_c^1([0, T], H^1(\Omega_c)) \\
& \int_0^T \int_{\Omega_{ext}^\varepsilon} \left[-\widetilde{M}_\varepsilon(x, t) \frac{\partial \Psi(x, t)}{\partial t} + \sigma_\varepsilon \nabla \widetilde{M}_\varepsilon(x, t) \nabla \Psi(x, t) \right] dx dt \\
& + \int_0^T \int_{\Gamma_\varepsilon} \kappa^\varepsilon(\widetilde{M}_{\varepsilon, e} - \widetilde{M}_{\varepsilon, c})(\Psi_e - \Psi_c) d\mu(x) dt \\
& + \int_0^T \int_{\Omega_{ext}^\varepsilon} \left[\imath q F(t) \nabla \widetilde{M}_\varepsilon(x, t) \cdot \mathbf{n} \Psi(x, t) - \imath q F(t) \widetilde{M}_\varepsilon(x, t) \mathbf{n} \cdot \nabla \Psi(x, t) \right] dx dt \\
& = \int_{\Omega_{ext}^\varepsilon} M_{init}(x) \Psi(x, 0) dx
\end{aligned} \tag{51}$$

Recall that we have set $\kappa^\varepsilon = \varepsilon^p \sigma_m$. Using the above formulation, we now rigorously derive five macroscopic models, according to the values of p . We start by the simplest case, when $p \geq 1$, and then we treat the three remaining ones.

5.1 The macroscopic models when $p \geq 1$

The aim of this subsection is to prove the following theorem :

Theorem 5.1. *Let $p \geq 1$ and let $\widetilde{M}_\varepsilon = (\widetilde{M}_{\varepsilon, e}, \widetilde{M}_{\varepsilon, c})$ be the unique solution of the two-compartment model (3), with $f \in L^\infty(]0, T[)$ and $\widetilde{M}_{init} \in L^2(\mathbb{R}^d)$. Then, there exists :*

$$(\widetilde{M}_{0, e}, \widetilde{M}_{1, e}) \in L^2(0, T, H^1(\mathbb{R}^d)) \times L^2(0, T, L^2(\Omega, H_\#^1(Y_e)))$$

and

$$(\widetilde{M}_{0, c}, \widetilde{M}_{1, c}) \in L^2(0, T, L^2(\mathbb{R}^d)) \times L^2(0, T, L^2(\Omega, H^1(Y_c)))$$

such that, up to a subsequence :

$$\left| \begin{aligned}
& \widetilde{M}_\varepsilon \rightharpoonup \frac{|Y_e|}{|Y|} \widetilde{M}_{0, e} + \frac{|Y_c|}{|Y|} \widetilde{M}_{0, c} \quad \text{weakly in } L^2(0, T, L^2(\mathbb{R}^d)) \\
& \mathcal{T}_\varepsilon^\alpha(\widetilde{M}_{\varepsilon, \alpha}) \rightharpoonup \widetilde{M}_{0, \alpha} \quad \text{weakly in } L^2(0, T, L^2(\Omega, H^1(Y_\alpha))), \quad \alpha \in \{c, e\} \\
& \mathcal{T}_\varepsilon^\alpha(\nabla \widetilde{M}_{\varepsilon, \alpha}) \rightharpoonup \nabla \widetilde{M}_{0, \alpha} + \nabla_y \widetilde{M}_{1, \alpha} \quad \text{weakly in } L^2(0, T, L^2(\mathbb{R}^d \times Y_\alpha)), \quad \alpha \in \{c, e\}
\end{aligned} \right. \tag{52}$$

and the four functions $(\widetilde{M}_{0,e}, \widetilde{M}_{1,e})$ and $(\widetilde{M}_{0,c}, \widetilde{M}_{1,c})$ are solutions of the following variational formulations, $\forall \Psi = (\Psi_e, \Psi_c) \in C_c^1([0, T[, H^1(\Omega_e)) \times C_c^1([0, T[, H^1(\Omega_c))$ and $\forall \Phi = (\Phi_e, \Phi_c) \in L^2(0, T, L^2(\Omega, H_\#^1(Y_e))) \times L^2(0, T, L^2(\Omega, H^1(Y_c)))$:

$$\begin{aligned}
& \left| \sum_{\alpha \in \{e, c\}} \left(\frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_\alpha} \sigma_\alpha(y) \left[\nabla \widetilde{M}_{0,\alpha}(x, t) + \nabla_y \widetilde{M}_{1,\alpha}(x, y, t) \right] [\nabla \Psi_\alpha(x, t) + \nabla_y \Phi_\alpha(x, y, t)] dx dy dt \right. \right. \\
& - \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_\alpha} \imath q F(t) \sigma_\alpha(y) \widetilde{M}_{0,\alpha}(x, t) \mathbf{n} \cdot [\nabla \Psi_\alpha(x, t) + \nabla_y \Phi_\alpha(x, y, t)] dx dy dt \\
& + \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_\alpha} \imath q F(t) \sigma_\alpha(y) \left[\nabla \widetilde{M}_{0,\alpha}(x, t) + \nabla_y \widetilde{M}_{1,\alpha}(x, y, t) \right] \cdot \mathbf{n} \Psi_\alpha(x, t) dx dy dt \\
& + \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_\alpha} q^2 F(t)^2 \sigma_\alpha(y) \widetilde{M}_{0,\alpha}(x, t) \Psi_\alpha(x, t) dx dy dt \\
& \left. - \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_\alpha} \widetilde{M}_{0,\alpha} \frac{\partial \Psi_\alpha}{\partial t}(x, t) dx dy dt \right) + \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times \Gamma_m} \sigma_m \left[\widetilde{M}_{0,e} - \widetilde{M}_{0,c} \right] [\Psi_e - \Psi_c] dx d\mu(y) dt \\
& = \sum_{\alpha \in \{e, c\}} \frac{1}{|Y|} \int_{\mathbb{R}^d \times Y_\alpha} M_{init}(x) \Psi_\alpha(x, 0)
\end{aligned} \tag{53}$$

in the case $p = 1$ and :

$$\begin{aligned}
& \left| \sum_{\alpha \in \{e, c\}} \left(\frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_\alpha} \sigma_\alpha(y) \left[\nabla \widetilde{M}_{0,\alpha}(x, t) + \nabla_y \widetilde{M}_{1,\alpha}(x, y, t) \right] [\nabla \Psi_\alpha(x, t) + \nabla_y \Phi_\alpha(x, y, t)] dx dy dt \right. \right. \\
& - \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_\alpha} \imath q F(t) \sigma_\alpha(y) \widetilde{M}_{0,\alpha}(x, t) \mathbf{n} \cdot [\nabla \Psi_\alpha(x, t) + \nabla_y \Phi_\alpha(x, y, t)] dx dy dt \\
& + \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_\alpha} \imath q F(t) \sigma_\alpha(y) \left[\nabla \widetilde{M}_{0,\alpha}(x, t) + \nabla_y \widetilde{M}_{1,\alpha}(x, y, t) \right] \cdot \mathbf{n} \Psi_\alpha(x, t) dx dy dt \\
& + \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_\alpha} q^2 F(t)^2 \sigma_\alpha(y) \widetilde{M}_{0,\alpha}(x, t) \Psi_\alpha(x, t) dx dy dt \\
& \left. - \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_\alpha} \widetilde{M}_{0,\alpha} \frac{\partial \Psi_\alpha}{\partial t}(x, t) dx dy dt \right) = \sum_{\alpha \in \{e, c\}} \frac{1}{|Y|} \int_{\mathbb{R}^d \times Y_\alpha} M_{init}(x) \Psi_\alpha(x, 0)
\end{aligned} \tag{54}$$

in the case $p \geq 2$.

Proof. The main steps of the proof are identical for the two situations. In fact, the only difference will come from the treatment of the boundary term in (51). From the bound (12) of theorem 2.1, we know that there exists $C > 0$ independent on ε such that, for $\alpha \in \{c, e\}$:

$$\|\widetilde{M}_{\varepsilon, \alpha}\|_{L^2(0, T, H^1(\Omega_\alpha^\varepsilon))} \leq e^{CT} \|M_{init}\|_{L^2(\mathbb{R}^d)}$$

Thus, we deduce from theorem 4.3 the existence of :

$$(\widetilde{M}_{0,e}, \widetilde{M}_{1,e}) \in L^2(0, T, H^1(\mathbb{R}^d)) \times L^2(0, T, L^2(\mathbb{R}^d, H_\#^1(Y_e)))$$

and

$$(\widetilde{M}_{0,c}, \widetilde{M}_{1,c}) \in L^2(0, T, H^1(\mathbb{R}^d)) \times L^2(0, T, L^2(\mathbb{R}^d, H^1(Y_c)))$$

such that :

$$\begin{aligned}
\mathcal{T}_\varepsilon^e(\widetilde{M}_{\varepsilon, e}) &\rightharpoonup \widetilde{M}_{0, e} \quad \text{weakly in } L^2(0, T, L^2(\mathbb{R}^d, H^1(Y_e))) \\
\mathcal{T}_\varepsilon^e(\nabla \widetilde{M}_{\varepsilon, e}) &\rightharpoonup \nabla \widetilde{M}_{0, e} + \nabla_y \widetilde{M}_{1, e} \quad \text{weakly in } L^2(0, T, L^2(\mathbb{R}^d \times Y_e)) \\
\mathcal{T}_\varepsilon^c(\widetilde{M}_{\varepsilon, c}) &\rightharpoonup \widetilde{M}_{0, c} \quad \text{weakly in } L^2(0, T, L^2(\mathbb{R}^d, H^1(Y_c))) \\
\mathcal{T}_\varepsilon^c(\nabla \widetilde{M}_{\varepsilon, c}) &\rightharpoonup \nabla \widetilde{M}_{0, c} + \nabla_y \widetilde{M}_{1, c} \quad \text{weakly in } L^2(0, T, L^2(\mathbb{R}^d \times Y_c))
\end{aligned}$$

We write :

$$\widetilde{M}_\varepsilon = \chi_{\varepsilon,e} \widetilde{M}_{\varepsilon,e} + \chi_{\varepsilon,c} \widetilde{M}_{\varepsilon,c}$$

Then, point (iii) of theorem 4.2 immediately gives, as $\widetilde{M}_{0,\alpha}$ is independent on y :

$$\chi_{\varepsilon,\alpha} \widetilde{M}_{\varepsilon,\alpha} \rightharpoonup \frac{|Y_\alpha|}{|Y|} \widetilde{M}_{0,\alpha} \quad \text{weakly in } L^2(0, T, L^2(\mathbb{R}^d))$$

and we consequently have as announced :

$$\widetilde{M}_\varepsilon \rightharpoonup \frac{|Y_e|}{|Y|} \widetilde{M}_{0,e} + \frac{|Y_c|}{|Y|} \widetilde{M}_{0,c} \quad \text{weakly in } L^2(0, T, L^2(\mathbb{R}^d))$$

Choosing test functions $\Psi_\alpha \in C_c^1([0, T], H^1(\mathbb{R}^d))$, we pass to the limit in the preceding expressions. For the volumic terms, from theorem 4.1 and formula (34), point (ii) of theorem 4.2 and the above weak convergences, we deduce :

$$\begin{aligned} & \int_0^T \int_{\Omega_{\varepsilon,xt}^\varepsilon} \left[\sigma_\varepsilon \nabla \widetilde{M}_\varepsilon(x, t) \nabla \Psi(x, t) - \widetilde{M}_\varepsilon(x, t) \frac{\partial \Psi(x, t)}{\partial t} \right] dx dt \\ & \rightarrow \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_e} \sigma_e(y) \left[\nabla \widetilde{M}_{0,e}(x, t) + \nabla_y \widetilde{M}_{1,e}(x, y, t) \right] \nabla \Psi_e(x, t) dx dy dt \\ & \quad - \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_e} \widetilde{M}_{0,e}(x, t) \frac{\partial \Psi_e}{\partial t}(x, t) dx dy dt \\ & \quad + \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_c} \sigma_c(y) \left[\nabla \widetilde{M}_{0,c}(x, t) + \nabla_y \widetilde{M}_{1,c}(x, y, t) \right] \nabla \Psi_c(x, t) dx dy dt \\ & \quad - \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_c} \widetilde{M}_{0,c}(x, t) \frac{\partial \Psi_c}{\partial t}(x, t) dx dy dt \end{aligned}$$

then

$$\begin{aligned} & \int_0^T \int_{\Omega_{\varepsilon,xt}^\varepsilon} \imath q F(t) \sigma_\varepsilon \nabla \widetilde{M}_\varepsilon(x, t) \cdot \mathbf{n} \Psi(x, t) dx dt = \\ & \rightarrow \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_e} \imath q F(t) \sigma_e(y) \left[\nabla \widetilde{M}_{0,e}(x, t) + \nabla_y \widetilde{M}_{1,e}(x, y, t) \right] \cdot \mathbf{n} \Psi_e(x, t) dx dy dt \\ & \quad + \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_c} \imath q F(t) \sigma_c(y) \left[\nabla \widetilde{M}_{0,c}(x, t) + \nabla_y \widetilde{M}_{1,c}(x, y, t) \right] \cdot \mathbf{n} \Psi_c(x, t) dx dy dt \end{aligned}$$

and

$$\begin{aligned} & - \int_0^T \int_{\Omega_{\varepsilon,xt}^\varepsilon} \imath q F(t) \sigma_\varepsilon \widetilde{M}_\varepsilon(x, t) \mathbf{n} \cdot \nabla \Psi(x, t) dx dt = \\ & \rightarrow - \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_e} \imath q F(t) \sigma_e(y) \widetilde{M}_{0,e}(x, t) \mathbf{n} \cdot \nabla \Psi_e(x, t) dx dy dt \\ & \quad - \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_c} \imath q F(t) \sigma_c(y) \widetilde{M}_{0,c}(x, t) \mathbf{n} \cdot \nabla \Psi_c(x, t) dx dy dt \end{aligned}$$

For the boundary terms, we will use results obtained through the periodic unfolding operator on the boundary. From point (ii) of theorem 4.2, we deduce that $\mathcal{T}_\varepsilon^\Gamma(\Psi_\alpha) \rightarrow \Psi_\alpha$ strongly in $L^2(0, T, L^2(\mathbb{R}^d, H^1(Y_\alpha)))$. As $\mathcal{T}_\varepsilon^\Gamma(\Psi_\alpha) = \mathcal{T}_\varepsilon^\Gamma(\Psi_\alpha)|_{y \in \Gamma_m}$, the continuity of the trace operator implies

the strong convergence of $\mathcal{T}_\varepsilon^\Gamma(\Psi_\alpha)$ towards Ψ_α in $L^2(0, T, L^2(\mathbb{R}^d, H^{1/2}))$ and thus, using the above weak convergences and theorem 4.4, we get in the case $p = 1$:

$$\begin{aligned} & \int_0^T \int_{\Gamma_\varepsilon} \varepsilon \sigma_m(\widetilde{M}_{\varepsilon,e} - \widetilde{M}_{\varepsilon,c})(\Psi_e - \Psi_c) d\mu(x) dt \\ & \rightarrow \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times \Gamma_m} \sigma_m(\widetilde{M}_{0,e} - \widetilde{M}_{0,c})(\Psi_e - \Psi_c) dx d\mu(y) dt \end{aligned}$$

and $\int_0^T \int_{\Gamma_\varepsilon} \varepsilon^p \sigma_m(\widetilde{M}_{\varepsilon,e} - \widetilde{M}_{\varepsilon,c})(\Psi_e - \Psi_c) d\mu(x) dt \rightarrow 0$ in the case $p > 1$. Finally, for the second member, we have :

$$\begin{aligned} \int_{\Omega_\alpha^\varepsilon} M_{init}(x) \Psi_\alpha(x, 0) dx &= \frac{1}{|Y|} \int_{\mathbb{R}^d \times Y_\alpha} \mathcal{T}_\varepsilon^\alpha(M_{init})(x, y) \mathcal{T}_\varepsilon^\alpha(\Psi_\alpha(\cdot, 0))(x, y) dx dy \\ &\rightarrow \frac{1}{|Y|} \int_{\mathbb{R}^d \times Y_\alpha} M_{init}(x) \Psi_\alpha(x, 0) dx dy \end{aligned}$$

Consequently, we have obtained the result for test functions independent on y . Now, we use as test function $\Psi^\varepsilon = (\Psi_e^\varepsilon, \Psi_c^\varepsilon)$, where $\Psi_\alpha^\varepsilon = \varepsilon \psi_\alpha(x, t) \Phi_\alpha(\frac{x}{\varepsilon})$, with $\psi_\alpha \in \mathcal{D}([0, T] \times \mathbb{R}^d)$ and where $\Phi_e \in H_\#^1(Y_e)$ is Y -periodic and $\Phi_c \in H^1(Y_c)$ is periodically repeated on each Y -translated of Y_c . Then, due to the additional factor ε , all the terms (including the second member) except the ones involving $\nabla_x \Psi^\varepsilon$ will vanish when taking the limit $\varepsilon \rightarrow 0$. For this term, we have :

$$\nabla \Psi_\alpha(x, t) = \varepsilon \nabla \psi_\alpha(x, t) \Phi_\alpha\left(\frac{x}{\varepsilon}\right) + \psi_\alpha(x, t) \nabla_y \Phi_\alpha\left(\frac{x}{\varepsilon}\right)$$

Then, we get, using theorem 4.1 :

$$\begin{aligned} & \int_0^T \int_{\Omega_{\varepsilon x t}^\varepsilon} \sigma_\varepsilon \nabla \widetilde{M}_\varepsilon(x, t) \nabla \Psi^\varepsilon(x, t) dx dt \\ &= \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_e} \sigma_e \mathcal{T}_\varepsilon^e(\nabla \widetilde{M}_{\varepsilon,e})(x, y, t) [\varepsilon \mathcal{T}_\varepsilon^e(\nabla \psi_e(x, t)) \Phi_e(y) \\ & \quad + \mathcal{T}_\varepsilon^e(\psi_e(x, t)) \nabla_y \Phi_e(y)] dx dy dt \\ &+ \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_c} \sigma_c \mathcal{T}_\varepsilon^c(\nabla \widetilde{M}_{\varepsilon,c})(x, y, t) [\varepsilon \mathcal{T}_\varepsilon^c(\nabla \psi_c(x, t)) \Phi_c(y) \\ & \quad + \mathcal{T}_\varepsilon^c(\psi_c(x, t)) \nabla_y \Phi_c(y)] dx dy dt \end{aligned}$$

Passing to the limit yields for this term :

$$\begin{aligned} & \int_0^T \int_{\Omega_{\varepsilon x t}^\varepsilon} \sigma_\varepsilon \nabla \widetilde{M}_\varepsilon(x, t) \nabla \Psi^\varepsilon(x, t) dx dt \\ & \rightarrow \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_e} \sigma_e(y) \left[\nabla \widetilde{M}_{0,e}(x, t) + \nabla_y \widetilde{M}_{1,e}(x, y, t) \right] \Psi_e(x, t) \nabla_y \Phi_e(y) dx dy dt \\ & + \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_c} \sigma_c(y) \left[\nabla \widetilde{M}_{0,c}(x, t) + \nabla_y \widetilde{M}_{1,c}(x, y, t) \right] \Psi_c(x, t) \nabla_y \Phi_c(y) dx dy dt \end{aligned}$$

In the same way, we have :

$$- \int_0^T \int_{\Omega_{\varepsilon x t}^\varepsilon} \iota q F(t) \sigma_\varepsilon \widetilde{M}_\varepsilon(x, t) \mathbf{n} \cdot \nabla \Psi^\varepsilon(x, t) dx dt =$$

$$\begin{aligned}
& \rightarrow -\frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_e} \imath q F(t) \sigma_e(y) \widetilde{M}_{0,e}(x,t) \mathbf{n} \cdot \Psi_e(x,t) \nabla_y \Phi_e(y) dx dy dt \\
& - \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_c} \imath q F(t) \sigma_c(y) \widetilde{M}_{0,c}(x,t) \mathbf{n} \cdot \Psi_c(x,t) \nabla_y \Phi_c(y) dx dy dt
\end{aligned}$$

and the sum of these two terms is equal to zero as all the other terms vanish. Using the density of $\mathcal{D}([0, T] \times \mathbb{R}^d) \times H_{\#}^1(Y_e)$ in $L^2(0, T, L^2(\mathbb{R}^d, H_{\#}^1(Y_e)))$ and of $\mathcal{D}([0, T] \times \mathbb{R}^d) \times H^1(Y_c)$ in $L^2(0, T, L^2(\mathbb{R}^d, H^1(Y_c)))$, we deduce that for any $\Phi_e \in L^2(0, T, L^2(\mathbb{R}^d, H_{\#}^1(Y_e)))$ and any $\Phi_c \in L^2(0, T, L^2(\mathbb{R}^d, H^1(Y_c)))$, we have :

$$\begin{aligned}
& \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_e} \sigma_e(y) \left[\nabla \widetilde{M}_{0,e}(x,t) + \nabla_y \widetilde{M}_{1,e}(x,y,t) \right] \nabla_y \Phi_e(x,y,t) dx dy dt \\
& + \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_c} \sigma_c(y) \left[\nabla \widetilde{M}_{0,c}(x,t) + \nabla_y \widetilde{M}_{1,c}(x,y,t) \right] \nabla_y \Phi_c(x,y,t) dx dy dt \\
& - \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_e} \imath q F(t) \sigma_e(y) \widetilde{M}_{0,e}(x,t) \mathbf{n} \cdot \nabla_y \Phi_e(x,y,t) dx dy dt \\
& - \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_c} \imath q F(t) \sigma_c(y) \widetilde{M}_{0,c}(x,t) \mathbf{n} \cdot \nabla_y \Phi_c(x,y,t) dx dy dt = 0
\end{aligned}$$

Adding this to the previous result, we obtain (53) and (54). $\square$

Now, we identify the corresponding macroscopic models :

Theorem 5.2. *Let $w_{i,\alpha}$ be solutions for $i = 1, 2$ of the cell problems :*

$$\begin{cases}
-div_y(\sigma_{\alpha}(\nabla_y w_{i,\alpha} + e_i)) = 0 & \text{in } Y_{\alpha} \\
\sigma_e \nabla_y w_{i,e} \cdot \nu + \sigma_e e_i \cdot \nu = 0 & \text{on } \Gamma_m \\
\sigma_c \nabla_y w_{i,c} \cdot \nu + \sigma_c e_i \cdot \nu = 0 & \text{on } \Gamma_m \\
w_{i,e} \text{ } Y_e - \text{periodic}
\end{cases} \quad (21)$$

Then, defining the homogenized tensors by :

$$D_{\alpha,ij} = \frac{1}{|Y_{\alpha}|} \int_{Y_{\alpha}} \sigma_{\alpha}(\nabla w_{j,\alpha} + e_j) \cdot (\nabla w_{i,\alpha} + e_i) \quad (14)$$

and the coefficients η_{α} by :

$$\eta_c = \frac{\sigma_m |\Gamma_m|}{|Y_c|} \quad \text{and} \quad \eta_e = \frac{\sigma_m |\Gamma_m|}{|Y_e|} \quad (15)$$

the variational problem (53) implies that $\widetilde{M}_0 = (\widetilde{M}_{0,e}, \widetilde{M}_{0,c}) \in W(0, T, X)$ satisfies for any $V \in X$

the macroscopic model :

$$\begin{aligned}
& \left\langle \partial_t \widetilde{M}_{0,e}, V_e \right\rangle_{H^{-1}, H^1} + \left(D_e \nabla \widetilde{M}_{0,e}, \nabla V_e \right)_{L^2} + \left(\eta_e (\widetilde{M}_{0,e} - \widetilde{M}_{0,c}), V_e \right)_{L^2} \\
& - \left(\imath q F(t) D_e \widetilde{M}_{0,e} \mathbf{n}, \nabla V_e \right)_{L^2} + \left(\imath q F(t) D_e \nabla \widetilde{M}_{0,e} \cdot \mathbf{n}, V_e \right)_{L^2} \\
& + \left(q^2 F(t)^2 D_e \mathbf{n} \cdot \mathbf{n} \widetilde{M}_{0,e}, V_e \right)_{L^2} = 0 \quad \text{in } \mathcal{D}'([0, T]) \\
& \widetilde{M}_{0,e}(\cdot, 0) = M_{init} \quad \text{in } \mathbb{R}^d \\
& \left\langle \partial_t \widetilde{M}_{0,c}, V_c \right\rangle_{H^{-1}, H^1} + \left(D_c \nabla_x \widetilde{M}_{0,c}, \nabla V_c \right)_{L^2} + \left(\eta_c (\widetilde{M}_{0,c} - \widetilde{M}_{0,e}), V_c \right)_{L^2} \\
& - \left(\imath q F(t) D_c \widetilde{M}_{0,c} \mathbf{n}, \nabla V_c \right)_{L^2} + \left(\imath q F(t) D_c \nabla \widetilde{M}_{0,c} \cdot \mathbf{n}, V_c \right)_{L^2} \\
& + \left(q^2 F(t)^2 D_c \mathbf{n} \cdot \mathbf{n} \widetilde{M}_{0,c}, V_c \right)_{L^2} = 0 \quad \text{in } \mathcal{D}'([0, T]) \\
& \widetilde{M}_{0,c}(\cdot, 0) = M_{init} \quad \text{in } \mathbb{R}^d
\end{aligned} \tag{55}$$

The variational problem (54) implies that $\widetilde{M}_0 = (\widetilde{M}_{0,e}, \widetilde{M}_{0,c}) \in W(0, T, X)$ satisfies for any $V \in X$ the macroscopic model :

$$\begin{aligned}
& \left\langle \partial_t \widetilde{M}_{0,e}, V_e \right\rangle_{H^{-1}, H^1} + \left(D_e \nabla \widetilde{M}_{0,e}, \nabla V_e \right)_{L^2} - \left(\imath q F(t) D_e \widetilde{M}_{0,e} \mathbf{n}, \nabla V_e \right)_{L^2} \\
& + \left(\imath q F(t) D_e \nabla \widetilde{M}_{0,e} \cdot \mathbf{n}, V_e \right)_{L^2} + \left(q^2 F(t)^2 D_e \mathbf{n} \cdot \mathbf{n} \widetilde{M}_{0,e}, V_e \right)_{L^2} = 0 \quad \text{in } \mathcal{D}'([0, T]) \\
& \widetilde{M}_{0,e}(\cdot, 0) = M_{init} \quad \text{in } \mathbb{R}^d \\
& \left\langle \partial_t \widetilde{M}_{0,c}, V_c \right\rangle_{H^{-1}, H^1} + \left(D_c \nabla_x \widetilde{M}_{0,c}, \nabla V_c \right)_{L^2} - \left(\imath q F(t) D_c \widetilde{M}_{0,c} \mathbf{n}, \nabla V_c \right)_{L^2} \\
& + \left(\imath q F(t) D_c \nabla \widetilde{M}_{0,c} \cdot \mathbf{n}, V_c \right)_{L^2} + \left(q^2 F(t)^2 D_c \mathbf{n} \cdot \mathbf{n} \widetilde{M}_{0,c}, V_c \right)_{L^2} = 0 \quad \text{in } \mathcal{D}'([0, T]) \\
& \widetilde{M}_{0,c}(\cdot, 0) = M_{init} \quad \text{in } \mathbb{R}^d
\end{aligned} \tag{18}$$

In both case, we have :

$$\widetilde{M}_{1,\alpha} = \sum_{i=1}^d w_{i,\alpha} \left(\frac{\partial \widetilde{M}_{0,\alpha}}{\partial x_i} - \imath q F(t) \mathbf{n}_i \widetilde{M}_{0,\alpha} \right) \quad \text{in }]0, T[\times \mathbb{R}^d \times Y_\alpha, \quad \alpha \in \{e, c\}$$

Proof. We detail the proof only for $p = 1$, the proof being identical for $p \geq 2$. As we have seen in the proof of theorem 5.1, we have for any $\psi_\alpha \in \mathcal{D}([0, T] \times \Omega)$, $\Phi_e \in H_\#^1(Y_e)$ and $\Phi_c \in H^1(Y_c)$:

$$\begin{aligned}
& \frac{1}{|Y|} \int_0^T \int_{\Omega \times Y_e} \sigma_e \left[\nabla \widetilde{M}_{0,e}(x, t) + \nabla_y \widetilde{M}_{1,e}(x, y, t) \right] \Psi_e(x, t) \nabla_y \Phi_e(y) dx dy dt \\
& + \frac{1}{|Y|} \int_0^T \int_{\Omega \times Y_c} \sigma_c \left[\nabla \widetilde{M}_{0,c}(x, y) + \nabla_y \widetilde{M}_{1,c}(x, y, t) \right] \Psi_c(x, t) \nabla_y \Phi_c(y) dx dy dt \\
& - \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_e} \imath q F(t) \sigma_e(y) \widetilde{M}_{0,e}(x, t) \mathbf{n} \cdot \Psi_e(x, t) \nabla_y \Phi_e(y) dx dy dt
\end{aligned}$$

$$-\frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_c} \imath q F(t) \sigma_c(y) \widetilde{M}_{0,c}(x, t) \mathbf{n} \cdot \Psi_c(x, t) \nabla_y \Phi_c(y) dx dy dt = 0$$

This implies that, in the distributional sense, we have, for any $\Phi_e \in H_{\sharp}^1(Y_e)$:

$$\int_{Y_e} \sigma_e \left[\nabla \widetilde{M}_{0,e}(x, t) + \nabla_y \widetilde{M}_{1,e}(x, y, t) - \imath q F(t) \sigma_e(y) \widetilde{M}_{0,e}(x, t) \mathbf{n} \right] \nabla_y \Phi_e(y) dy = 0$$

and for any $\Phi_c \in H^1(Y_c)$:

$$\int_{Y_c} \sigma_c \left[\nabla \widetilde{M}_{0,c}(x, t) + \nabla_y \widetilde{M}_{1,c}(x, y, t) - \imath q F(t) \sigma_c(y) \widetilde{M}_{0,c}(x, t) \mathbf{n} \right] \nabla_y \Phi_c(y) dy = 0$$

which is the classical variational formulation of the problem, for a.e. $(x, t) \in \Omega \times]0, T[$:

$$\left\{ \begin{array}{ll} -\operatorname{div}_y(\sigma_\alpha(\nabla_y \widetilde{M}_{1,\alpha}(x, y, t) + \nabla_x \widetilde{M}_{0,\alpha}(x, t) - \imath q F(t) \widetilde{M}_{0,\alpha}(x, t) \mathbf{n})) = 0 & \text{in } Y_\alpha \\ \sigma_e \nabla_y \widetilde{M}_{1,e}(x, y, t) \cdot \nu + \sigma_e \nabla_x \widetilde{M}_{0,e}(x, t) \cdot \nu - \imath q F(t) \widetilde{M}_{0,e}(x, t) \mathbf{n} \cdot \nu = 0 & \text{on } \Gamma_m \\ \sigma_c \nabla_y \widetilde{M}_{1,c}(x, y, t) \cdot \nu + \sigma_c \nabla_x \widetilde{M}_{0,c}(x, t) \cdot \nu - \imath q F(t) \widetilde{M}_{0,c}(x, t) \mathbf{n} \cdot \nu = 0 & \text{on } \Gamma_m \\ \widetilde{M}_{1,e} \text{ } Y_e \text{ - periodic} \end{array} \right.$$

which provides the desired formula for $\widetilde{M}_{1,\alpha}$. Now, for any $\Psi_\alpha \in C_c^1([0, T[, H^1(\mathbb{R}^d))$, we have, for $\alpha \in \{e, c\}$, using the independence on y of $M_{0,\alpha}$:

$$\begin{aligned} & \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_\alpha} \sigma_\alpha \left[\nabla \widetilde{M}_{0,\alpha}(x, t) + \nabla_y \widetilde{M}_{1,\alpha}(x, y, t) \right] \nabla \Psi_\alpha(x, t) dx dy dt \\ &= \frac{|Y_\alpha|}{|Y|} \int_0^T \int_{\mathbb{R}^d} D_\alpha \nabla \widetilde{M}_{0,\alpha} \nabla \Psi_\alpha \\ & - \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_\alpha} \sigma_\alpha \imath q F(t) \left(\sum_j^d \nabla_y w_{j,\alpha} n_j \right) \cdot \widetilde{M}_{0,\alpha}(x, t) \nabla \Psi_\alpha(x, t) \end{aligned}$$

where we have used the fact that :

$$D_{\alpha,ij} = \frac{1}{|Y_\alpha|} \int_{Y_\alpha} \sigma_\alpha(\nabla w_{j,\alpha} + e_j) \cdot (\nabla w_{i,\alpha} + e_i) = \frac{1}{|Y_\alpha|} \int_{Y_\alpha} \sigma_\alpha(\nabla w_{j,\alpha} + e_j) \cdot e_i$$

Thus we get :

$$\begin{aligned} & \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_\alpha} \sigma_\alpha \left[\nabla \widetilde{M}_{0,\alpha}(x, t) + \nabla_y \widetilde{M}_{1,\alpha}(x, y, t) \right] \nabla \Psi_\alpha(x, t) dx dy dt \\ & - \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_\alpha} \imath q F(t) \sigma_\alpha(y) \widetilde{M}_{0,\alpha}(x, t) \mathbf{n} \cdot \nabla \Psi_\alpha(x, t) dx dy dt \\ &= \frac{|Y_\alpha|}{|Y|} \int_0^T \int_{\mathbb{R}^d} D_\alpha \nabla \widetilde{M}_{0,\alpha} \nabla \Psi_\alpha - \frac{|Y_\alpha|}{|Y|} \int_0^T \int_{\mathbb{R}^d} \imath q F(t) \widetilde{M}_{0,\alpha} D_\alpha \mathbf{n} \nabla \Psi_\alpha \end{aligned}$$

In the same way, we have:

$$\frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_\alpha} \imath q F(t) \sigma_\alpha(y) \left[\nabla \widetilde{M}_{0,\alpha}(x, t) + \nabla_y \widetilde{M}_{1,\alpha}(x, y, t) \right] \cdot \mathbf{n} \Psi_\alpha(x, t) dx dy dt$$

$$\begin{aligned}
&= \frac{|Y_\alpha|}{|Y|} \int_0^T \int_{\mathbb{R}^d} \imath q F(t) D_\alpha \nabla \widetilde{M}_{0,\alpha} \cdot \mathbf{n} \Psi_\alpha \\
&+ \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_\alpha} q^2 F(t)^2 \sigma_\alpha \left(\sum_{j=1}^d \nabla_y w_{j,\alpha} n_j \right) \cdot \mathbf{n} \widetilde{M}_{0,\alpha} \Psi_\alpha
\end{aligned}$$

Thus we have, as $\mathbf{n} \cdot \mathbf{n} = 1$:

$$\begin{aligned}
&\frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_\alpha} \imath q F(t) \sigma_\alpha(y) \left[\nabla \widetilde{M}_{0,\alpha}(x, t) + \nabla_y \widetilde{M}_{1,\alpha}(x, y, t) \right] \cdot \mathbf{n} \Psi_\alpha(x, t) dx dy dt \\
&+ \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_\alpha} q^2 F(t)^2 \sigma_\alpha(y) \widetilde{M}_{0,\alpha}(x, t) \Psi_\alpha(x, t) dx dy dt \\
&= \frac{|Y_\alpha|}{|Y|} \int_0^T \int_{\mathbb{R}^d} \imath q F(t) D_\alpha \nabla \widetilde{M}_{0,\alpha} \cdot \mathbf{n} \Psi_\alpha + \frac{|Y_\alpha|}{|Y|} \int_0^T \int_{\mathbb{R}^d} q^2 F(t)^2 D_\alpha \mathbf{n} \cdot \mathbf{n} \widetilde{M}_{0,\alpha} \Psi_\alpha
\end{aligned}$$

Finally, we have :

$$-\frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_\alpha} \widetilde{M}_{0,\alpha}(x, t) \frac{\partial \Psi_\alpha}{\partial t}(x, t) dx dy dt = -\frac{|Y_\alpha|}{|Y|} \int_0^T \int_{\mathbb{R}^d} \widetilde{M}_{0,\alpha}(x, t) \frac{\partial \Psi_\alpha}{\partial t}(x, t) dx dt$$

and

$$\begin{aligned}
&\frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times \Gamma_m} \sigma_m \left[\widetilde{M}_{0,e} - \widetilde{M}_{0,c} \right] [\Psi_e - \Psi_c] dx d\mu(y) dt \\
&= \frac{\sigma_m |\Gamma_m|}{|Y|} \int_0^T \int_{\mathbb{R}^d} \left[\widetilde{M}_{0,e} - \widetilde{M}_{0,c} \right] [\Psi_e - \Psi_c] dx dt
\end{aligned}$$

Taking alternatively $\Psi_e = 0$ or $\Psi_c = 0$, we obtain, for any $\Psi_e \in C^1([0, T[, H^1(\mathbb{R}^d))$ and any $\Psi_c \in C^1([0, T[, H^1(\mathbb{R}^d))$, multiplying respectively by $\frac{|Y|}{|Y_e|}$ and $\frac{|Y|}{|Y_c|}$:

$$\begin{aligned}
&\int_0^T \int_{\mathbb{R}^d} D_\alpha \nabla \widetilde{M}_{0,\alpha} \nabla \Psi_\alpha - \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d} \imath q F(t) \widetilde{M}_{0,\alpha} D_\alpha \mathbf{n} \nabla \Psi_\alpha \\
&+ \int_0^T \int_{\mathbb{R}^d} \imath q F(t) D_\alpha \nabla \widetilde{M}_{0,\alpha} \cdot \mathbf{n} \Psi_\alpha + \int_0^T \int_{\mathbb{R}^d} q^2 F(t)^2 D_\alpha \mathbf{n} \cdot \mathbf{n} \widetilde{M}_{0,\alpha} \Psi_\alpha \\
&- \int_0^T \int_{\mathbb{R}^d} \widetilde{M}_{0,\alpha} \frac{\partial \Psi_\alpha}{\partial t} dx dt + \frac{\sigma_m |\Gamma_m|}{|Y_\alpha|} \int_0^T \int_{\mathbb{R}^d} \left[\widetilde{M}_{0,\alpha} - \widetilde{M}_{0,\beta} \right] \Psi_\alpha dx dt = \int_{\mathbb{R}^d} M_{init} \Psi_\alpha dx
\end{aligned}$$

for $\alpha \in \{e, c\}$, with $\beta = c$ if $\alpha = e$ and $\beta = e$ if $\alpha = c$. Integrating by parts in time, we obtain the desired result. $\square$

It is clear that we can immediately deduce theorem 3.1 from theorem 5.2, using the change of unknowns :

$$M_{i,\alpha} = \widetilde{M}_{i,\alpha} e^{-\imath q \mathbf{x} \cdot \mathbf{n} F(t)} \quad M_\varepsilon = \widetilde{M}_\varepsilon e^{-\imath q \mathbf{x} \cdot \mathbf{n} F(t)}$$

5.2 The macroscopic models when $p \leq 0$

The aim of this subsection is to prove the following theorem :

Theorem 5.3. *Let $p = 0$ or $p = -1$ and let $\widetilde{M}_\varepsilon = (\widetilde{M}_{\varepsilon,e}, \widetilde{M}_{\varepsilon,c})$ be the unique solution of the two-compartment model (3), with $f \in L^\infty(]0, T[)$ and $\widetilde{M}_{init} \in L^2(\mathbb{R}^d)$. Then, there exists :*

$$(\widetilde{M}_0, \widetilde{M}_{1,e}) \in L^2(0, T, H^1(\mathbb{R}^d)) \times L^2(0, T, L^2(\mathbb{R}^d, H_\sharp^1(Y_e)))$$

and

$$\widetilde{M}_{1,c} \in L^2(0, T, L^2(\mathbb{R}^d, H^1(Y_c)))$$

such that :

$$\left\{ \begin{array}{l} \widetilde{M}_\varepsilon \rightharpoonup \widetilde{M}_0 \quad \text{weakly in } L^2(0, T, L^2(\mathbb{R}^d)) \\ \mathcal{T}_\varepsilon^\alpha(\widetilde{M}_{\varepsilon,\alpha}) \rightharpoonup \widetilde{M}_0 \quad \text{weakly in } L^2(0, T, L^2(\Omega, H^1(Y_\alpha))), \quad \alpha \in \{c, e\} \\ \mathcal{T}_\varepsilon^\alpha(\nabla \widetilde{M}_{\varepsilon,\alpha}) \rightharpoonup \nabla \widetilde{M}_0 + \nabla_y \widetilde{M}_{1,\alpha} \quad \text{weakly in } L^2(0, T, L^2(\mathbb{R}^d \times Y_\alpha)), \quad \alpha \in \{c, e\} \end{array} \right. \quad (56)$$

and the three functions $\widetilde{M}_0, \widetilde{M}_{1,e}$ and $\widetilde{M}_{1,c}$ are solutions of the following variational formulations, $\forall \Psi \in C_c^1([0, T], H^1(\mathbb{R}^d))$ and $\forall \Phi = (\Phi_e, \Phi_c) \in L^2(0, T, L^2(\mathbb{R}^d, H_\sharp^1(Y_e))) \times L^2(0, T, L^2(\mathbb{R}^d, H^1(Y_c)))$:

$$\left\{ \begin{array}{l} \sum_{\alpha \in \{e, c\}} \left(\frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_\alpha} \sigma_\alpha(y) [\nabla \widetilde{M}_0(x, t) + \nabla_y \widetilde{M}_{1,\alpha}(x, y, t)] [\nabla \Psi(x, t) + \nabla_y \Phi_\alpha(x, y, t)] dx dy dt \right. \\ + \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_\alpha} \imath q F(t) \sigma_\alpha(y) [\nabla \widetilde{M}_0(x, t) + \nabla_y \widetilde{M}_{1,\alpha}(x, y, t)] \cdot \mathbf{n} \Psi(x, t) dx dy dt \\ - \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_\alpha} \imath q F(t) \sigma_\alpha(y) \widetilde{M}_0(x, t) \mathbf{n} \cdot [\nabla \Psi(x, t) + \nabla_y \Phi_\alpha(x, y, t)] dx dy dt \\ \left. + \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_\alpha} q^2(t) F(t) \sigma_\alpha(y) \widetilde{M}_0(x, t) \Psi(x, t) dx dy dt \right) \\ - \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y} \widetilde{M}_0 \frac{\partial \Psi}{\partial t}(x, t) dx dy dt = \frac{1}{|Y|} \int_{\mathbb{R}^d \times Y} \widetilde{M}_{init}(x) \Psi(x, 0) \end{array} \right. \quad (57)$$

in the case $p = 0$ and :

$$\left\{ \begin{array}{l} \sum_{\alpha \in \{e, c\}} \left(\frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_\alpha} \sigma_\alpha(y) [\nabla \widetilde{M}_0(x, t) + \nabla_y \widetilde{M}_{1,\alpha}(x, y, t)] [\nabla \Psi(x, t) + \nabla_y \Phi_\alpha(x, y, t)] dx dy dt \right. \\ + \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_\alpha} \imath q F(t) \sigma_\alpha(y) [\nabla \widetilde{M}_0(x, t) + \nabla_y \widetilde{M}_{1,\alpha}(x, y, t)] \cdot \mathbf{n} \Psi(x, t) dx dy dt \\ - \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_\alpha} \imath q F(t) \sigma_\alpha(y) \widetilde{M}_0(x, t) \mathbf{n} \cdot [\nabla \Psi(x, t) + \nabla_y \Phi_\alpha(x, y, t)] dx dy dt \\ \left. + \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_\alpha} q^2(t) F(t) \sigma_\alpha(y) \widetilde{M}_0(x, t) \Psi(x, t) dx dy dt \right) \\ - \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y} \widetilde{M}_0 \frac{\partial \Psi}{\partial t}(x, t) dx dy dt + \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times \Gamma_m} \sigma_m [\widetilde{M}_{1,e} - \widetilde{M}_{1,c} + \mu] [\Phi_e - \Phi_c] dx d\mu(y) dt \\ = \frac{1}{|Y|} \int_{\mathbb{R}^d \times Y} \widetilde{M}_{init}(x) \Psi(x, 0) \end{array} \right. \quad (58)$$

in the case $p = -1$, for some $\mu \in L^2(0, T, L^2(\mathbb{R}^d))$. Let $p \leq -2$ and let $\widetilde{M}_\varepsilon = (\widetilde{M}_{\varepsilon,e}, \widetilde{M}_{\varepsilon,c})$ be the unique solution of the two-compartment model (3), with $f \in L^\infty(]0, T[)$ and $\widetilde{M}_{init} \in L^2(\mathbb{R}^d)$. Then, there exists :

$$(\widetilde{M}_0, \widetilde{M}_1) \in L^2(0, T, H^1(\mathbb{R}^d)) \times L^2(0, T, L^2(\mathbb{R}^d, H_\sharp^1(Y)))$$

such that :

$$\left\{ \begin{array}{l} \widetilde{M}_\varepsilon \rightharpoonup \widetilde{M}_0 \quad \text{weakly in } L^2(0, T, L^2(\mathbb{R}^d)) \\ \mathcal{T}_\varepsilon^\alpha(\widetilde{M}_{\varepsilon, \alpha}) \rightharpoonup \widetilde{M}_0 \quad \text{weakly in } L^2(0, T, L^2(\Omega, H^1(Y_\alpha))), \quad \alpha \in \{c, e\} \\ \mathcal{T}_\varepsilon^\alpha(\nabla \widetilde{M}_{\varepsilon, \alpha}) \rightharpoonup \nabla \widetilde{M}_0 + \nabla_y \widetilde{M}_1 \quad \text{weakly in } L^2(0, T, L^2(\mathbb{R}^d \times Y_\alpha)), \quad \alpha \in \{c, e\} \end{array} \right. \quad (59)$$

and the two functions $(\widetilde{M}_0, \widetilde{M}_1)$ are solutions of the following variational formulations, $\forall \Psi \in C_c^1([0, T], H^1(\mathbb{R}^d))$ and $\forall \Phi \in L^2(\mathbb{R}^d, H_\#^1(Y))$:

$$\left\{ \begin{array}{l} \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y} \sigma(y) [\nabla \widetilde{M}_0(x, t) + \nabla_y \widetilde{M}_1(x, y, t)] [\nabla \Psi(x, t) + \nabla_y \Phi(x, y)] dx dy dt \\ + \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y} \imath q F(t) \sigma(y) [\nabla \widetilde{M}_0(x, t) + \nabla_y \widetilde{M}_1(x, y, t)] \cdot \mathbf{n} \Psi(x, t) dx dy dt \\ - \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y} \imath q F(t) \sigma_e(y) \widetilde{M}_0(x, t) \mathbf{n} \cdot [\nabla \Psi(x, t) + \nabla_y \Phi_\alpha(x, y, t)] dx dy dt \\ + \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y} q^2(t) F(t) \sigma(y) \widetilde{M}_0(x, t) \Psi(x, t) dx dy dt \\ - \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y} \widetilde{M}_0 \frac{\partial \Psi}{\partial t}(x, t) dx dy dt = \frac{1}{|Y|} \int_{\mathbb{R}^d \times Y} \widetilde{M}_{init}(x) \Psi(x, 0) \end{array} \right. \quad (60)$$

in the case $p \leq -2$.

Proof. The main difference with the proof of theorem 5.1 comes from the treatment of the boundary term, thus we will only detail this part. In the three cases, we have again, from the bound (12) of theorem 2.1 that for $\alpha \in \{c, e\}$, that for some $C > 0$ independent on ε :

$$\|\widetilde{M}_{\varepsilon, \alpha}\|_{L^2(0, T, H^1(\Omega_\alpha^\varepsilon))} \leq e^{CT} \|\widetilde{M}_{init}\|_{L^2(\mathbb{R}^d)}$$

Moreover, taking $V = M_\varepsilon$ in (49) and integrating over time, we easily obtain in this case that :

$$\int_0^T \int_{\Gamma_\varepsilon} (\widetilde{M}_{\varepsilon, e} - \widetilde{M}_{\varepsilon, c})^2 \leq \frac{\varepsilon^{-p} C(\sigma^+, q, \|F\|_{L^\infty([0, T][])})}{\sigma_m} \|\widetilde{M}_\varepsilon\|_{L^2(0, T, X_\varepsilon)} \quad (61)$$

We deduce again from theorem 4.3 the existence of :

$$(\widetilde{M}_{0, e}, \widetilde{M}_{1, e}) \in L^2(0, T, H^1(\mathbb{R}^d)) \times L^2(0, T, L^2(\mathbb{R}^d, H_\#^1(Y_e)))$$

and

$$(\widetilde{M}_{0, c}, \widetilde{M}_{1, c}) \in L^2(0, T, H^1(\mathbb{R}^d)) \times L^2(0, T, L^2(\mathbb{R}^d, H^1(Y_c)))$$

such that :

$$\begin{aligned} \mathcal{T}_\varepsilon^e(\widetilde{M}_{\varepsilon, e}) &\rightharpoonup \widetilde{M}_{0, e} \quad \text{weakly in } L^2(0, T, L^2(\mathbb{R}^d, H^1(Y_e))) \\ \mathcal{T}_\varepsilon^e(\nabla \widetilde{M}_{\varepsilon, e}) &\rightharpoonup \nabla \widetilde{M}_{0, e} + \nabla_y \widetilde{M}_{1, e} \quad \text{weakly in } L^2(0, T, L^2(\mathbb{R}^d \times Y_e)) \\ \mathcal{T}_\varepsilon^c(\widetilde{M}_{\varepsilon, c}) &\rightharpoonup \widetilde{M}_{0, c} \quad \text{weakly in } L^2(0, T, L^2(\Omega, H^1(Y_c))) \\ \mathcal{T}_\varepsilon^c(\nabla \widetilde{M}_{\varepsilon, c}) &\rightharpoonup \nabla \widetilde{M}_{0, c} + \nabla_y \widetilde{M}_{1, c} \quad \text{weakly in } L^2(0, T, L^2(\mathbb{R}^d \times Y_c)) \end{aligned}$$

Now, for any function $\Psi = (\Psi_e, \Psi_c) \in C_c^1([0, T], H^1(\Omega_e)) \times C_c^1([0, T], H^1(\Omega_c))$, using the periodic unfolding operator on the boundary and theorem 4.4, we obtain:

$$\int_0^T \int_{\Gamma_\varepsilon} \varepsilon \sigma_m (\widetilde{M}_{\varepsilon, e} - \widetilde{M}_{\varepsilon, c}) (\Psi_e - \Psi_c) d\mu(x) dt$$

$$= \left(\frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times \Gamma_m} \sigma_m \left(\mathcal{T}_\varepsilon^\Gamma(\widetilde{M}_{\varepsilon,e}) - \mathcal{T}_\varepsilon^\Gamma(\widetilde{M}_{\varepsilon,c}) \right) (\mathcal{T}_\varepsilon^e(\Psi_e) - \mathcal{T}_\varepsilon^c(\Psi_c)) dx d\mu(y) dt \right)$$

Proceeding as in the proof of theorem 5.1 we get once again:

$$\begin{aligned} & \int_0^T \int_{\Gamma_\varepsilon} \varepsilon \sigma_m(\widetilde{M}_{\varepsilon,e} - \widetilde{M}_{\varepsilon,c})(\Psi_e - \Psi_c) d\mu(x) dt \\ & \rightarrow \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times \Gamma_m} \sigma_m(\widetilde{M}_{0,e} - \widetilde{M}_{0,c})(\Psi_e - \Psi_c) dx d\mu(y) dt \end{aligned}$$

Consequently, as :

$$\begin{aligned} & \int_0^T \int_{\Gamma_\varepsilon} \varepsilon^p \sigma_m(\widetilde{M}_{\varepsilon,e} - \widetilde{M}_{\varepsilon,c})(\Psi_e - \Psi_c) d\mu(x) dt \\ & = \varepsilon^{p-1} \left(\int_0^T \int_{\Gamma_\varepsilon} \varepsilon \sigma_m(\widetilde{M}_{\varepsilon,e} - \widetilde{M}_{\varepsilon,c})(\Psi_e - \Psi_c) d\mu(x) dt \right) \end{aligned}$$

if

$$\frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times \Gamma_m} \sigma_m(\widetilde{M}_{0,e} - \widetilde{M}_{0,c})(\Psi_e - \Psi_c) dx d\mu(y) dt \neq 0$$

then, as $p \leq 0$:

$$\left| \int_0^T \int_{\Gamma_\varepsilon} \varepsilon^p \sigma_m(\widetilde{M}_{\varepsilon,e} - \widetilde{M}_{\varepsilon,c})(\Psi_e - \Psi_c) d\mu(x) dt \right| \longrightarrow +\infty$$

However, as we have from (51):

$$\begin{aligned} & \int_0^T \int_{\Omega_{\varepsilon,xt}^\varepsilon} \left[\sigma_\varepsilon \nabla \widetilde{M}_\varepsilon(x, t) \nabla \Psi(x, t) + \imath q F(t) \nabla \widetilde{M}_\varepsilon(x, t) \cdot \mathbf{n} \Psi(x, t) \right. \\ & \quad \left. - \imath q F(t) \widetilde{M}_\varepsilon(x, t) \nabla \Psi(x, t) \cdot \mathbf{n} \right] dx dt \\ & + \int_0^T \int_{\Omega_{\varepsilon,xt}^\varepsilon} \left[q^2 F(t)^2 \sigma_\varepsilon \widetilde{M}_\varepsilon(x, t) \Psi(x, t) - \widetilde{M}_\varepsilon(x, t) \frac{\partial \Psi(x, t)}{\partial t} \right] dx dt \\ & + \int_0^T \int_{\Gamma_\varepsilon} \sigma_m \varepsilon^p (\widetilde{M}_{\varepsilon,e} - \widetilde{M}_{\varepsilon,c})(\Psi_e - \Psi_c) d\mu(x) dt = \int_{\Omega_{\varepsilon,xt}^\varepsilon} \widetilde{M}_{init}(x) \Psi(x, 0) dx \end{aligned}$$

using Cauchy-Schwarz inequality we deduce that there exists some $C(\sigma^+, q, \|F\|_{L^\infty([0, T])}) > 0$ independent on ε such that :

$$\begin{aligned} & \left| \int_0^T \int_{\Gamma_\varepsilon} \varepsilon^p \sigma_m(\widetilde{M}_{\varepsilon,e} - \widetilde{M}_{\varepsilon,c})(\Psi_e - \Psi_c) d\mu(x) dt \right| \\ & \leq C(\sigma^+, q, \|F\|_{L^\infty([0, T])}) \|\widetilde{M}_\varepsilon\|_{L^2(0, T, X_\varepsilon)} \|\Psi\|_{L^2(0, T, X)} \end{aligned}$$

and thus, we deduce from the bound (12) of theorem 2.1 that :

$$\lim_{\varepsilon \rightarrow 0} \left| \int_0^T \int_{\Gamma_\varepsilon} \varepsilon^p \sigma_m(\widetilde{M}_{\varepsilon,e} - \widetilde{M}_{\varepsilon,c})(\Psi_e - \Psi_c) d\mu(x) dt \right| < +\infty$$

Consequently :

$$\frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times \Gamma_m} \sigma_m(\widetilde{M}_{0,e} - \widetilde{M}_{0,c})(\Psi_e - \Psi_c) dx d\mu(y) dt$$

$$= \frac{|\Gamma_m| \sigma_m}{|Y|} \int_0^T \int_{\mathbb{R}^d} \sigma_m(\widetilde{M}_{0,e} - \widetilde{M}_{0,c})(\Psi_e - \Psi_c) dx d\mu(y) dt = 0$$

for all $\Psi = (\Psi_e, \Psi_c) \in C_c^1([0, T[, H^1(\Omega_e)) \times C_c^1([0, T[, H^1(\Omega_c))$, from which we deduce that $\widetilde{M}_{0,e} = \widetilde{M}_{0,c}$. Consequently, we can restrict the test functions to all $\Psi \in C_c^1([0, T[, H^1(\mathbb{R}^d))$. Then, the boundary term vanishes, because of the continuity in trace of Ψ , and we obtain the first part of the result, by going to the limit in the variational formulation through the periodic unfolding operators, as in the proof of theorem 5.1. Now, we use the test function $\Psi_\alpha^\varepsilon = \varepsilon \psi_\alpha(x, t) \Phi_\alpha(\frac{x}{\varepsilon})$, with $\psi_\alpha \in \mathcal{D}([0, T[\times \Omega)$ and where $\Phi_e \in H_\#^1(Y_e)$ is Y -periodic and $\Phi_c \in H^1(Y_c)$ is periodically repeated on each Y -translated of Y_c . Again, due to the additional factor ε , all the volumic terms (including the second member) except the one involving $\nabla_x \Psi^\varepsilon$ will vanish when taking the limit $\varepsilon \rightarrow 0$. For the boundary term, we must distinguish according to the value of p . For $p = 0$, as $\widetilde{M}_{0,e} = \widetilde{M}_{0,c}$, we have in fact shown that:

$$\int_0^T \int_{\widehat{\Gamma}_\varepsilon} \varepsilon \sigma_m(\widetilde{M}_{\varepsilon,e} - \widetilde{M}_{\varepsilon,c})(\Psi_e \Phi_{f,\varepsilon} - \Psi_c \Phi_{c,\varepsilon}) d\mu(x) dt \rightarrow 0$$

Proceeding as in the proof of theorem 5.1, using the unfolding operator on the boundary through points (i) and (iii) of theorem 4.4, the density of $\mathcal{D}([0, T[\times \mathbb{R}^d) \times H_\#^1(Y_e)$ in $L^2(0, T, L^2(\Omega, H_\#^1(Y_e)))$ and of $\mathcal{D}([0, T[\times \mathbb{R}^d) \times H^1(Y_c)$ in $L^2(0, T, L^2(\Omega, H^1(Y_c)))$, and the previous result, we easily obtain (57).

Now, for $p = -1$, we have from theorem 4.1 and formula (45):

$$\mathcal{T}_\varepsilon^\Gamma(\Psi_\alpha \Phi_{\alpha,\varepsilon})(x, y, t) = \mathcal{T}_\varepsilon^\Gamma(\Psi_\alpha)(x, y, t) \Phi_\alpha(y) \longrightarrow \Psi_\alpha(x, t) \Psi_\alpha(y)$$

strongly in $L^2(0, T, L^2(\mathbb{R}^d \times \Gamma_m))$. Estimate (61) now allow us to use theorem 4.5, and thus there exists some $\mu \in L^2(0, T, L^2(\mathbb{R}^d))$:

$$\begin{aligned} & \int_0^T \int_{\Gamma_\varepsilon} \varepsilon^{-1} \sigma_m(\widetilde{M}_{\varepsilon,e} - \widetilde{M}_{\varepsilon,c})(\Psi_e^\varepsilon - \Psi_c^\varepsilon) d\mu(x) dt \\ &= \int_0^T \int_{\Gamma_\varepsilon} \sigma_m(\widetilde{M}_{\varepsilon,e} - \widetilde{M}_{\varepsilon,c})(\Psi_e \Phi_{f,\varepsilon} - \Psi_c \Phi_{c,\varepsilon}) d\mu(x) dt \\ &\longrightarrow \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times \Gamma_m} [\widetilde{M}_{1,e} - \widetilde{M}_{1,c} + \mu] [\Psi_e(x, t) \Phi_e(y) - \Psi_c(x, t) \Phi_c(y)] dx d\mu(y) dt \end{aligned}$$

that (58) directly follows by density.

Finally, for $p \leq -2$, we have :

$$\begin{aligned} & \int_0^T \int_{\Omega_{\varepsilon_{xt}}} \left[\sigma_\varepsilon \nabla \widetilde{M}_\varepsilon(x, t) \nabla \Psi(x, t) + \imath q F(t) \nabla \widetilde{M}_\varepsilon(x, t) \cdot \mathbf{n} \Psi(x, t) \right. \\ & \quad \left. - \imath q F(t) \widetilde{M}_\varepsilon(x, t) \nabla \Psi(x, t) \cdot \mathbf{n} \right] dx dt \\ &+ \int_0^T \int_{\Omega_{\varepsilon_{xt}}} \left[q^2 F(t)^2 \sigma_\varepsilon \widetilde{M}_\varepsilon(x, t) \Psi(x, t) - \widetilde{M}_\varepsilon(x, t) \frac{\partial \Psi(x, t)}{\partial t} \right] dx dt \\ &+ \int_0^T \int_{\Gamma_\varepsilon} \sigma_m^p(\widetilde{M}_{\varepsilon,e} - \widetilde{M}_{\varepsilon,c})(\Psi_e - \Psi_c) d\mu(x) dt = \int_{\Omega_{\varepsilon_{xt}}} \widetilde{M}_{init}(x) \Psi(x, 0) dx \end{aligned}$$

and consequently using Cauchy-Schwarz inequality and the bound on $\widetilde{M}_\varepsilon$ we deduce that :

$$\lim_{\varepsilon \rightarrow 0} \left| \int_0^T \int_{\Gamma_\varepsilon} \varepsilon^p \sigma_m(\widetilde{M}_{\varepsilon,e} - \widetilde{M}_{\varepsilon,c})(\Psi_e \Phi_{f,\varepsilon} - \Psi_c \Phi_{c,\varepsilon}) d\mu(x) dt \right| < +\infty$$

As $p+1 \leq -1$ and :

$$\begin{aligned} & \int_0^T \int_{\Gamma_\varepsilon} \varepsilon^p \sigma_m(\widetilde{M}_{\varepsilon,e} - \widetilde{M}_{\varepsilon,c})(\Psi_e \Phi_{f,\varepsilon} - \Psi_c \Phi_{c,\varepsilon}) d\mu(x) dt \\ &= \varepsilon^{p+1} \left(\int_0^T \int_{\Gamma_\varepsilon} \sigma_m \varepsilon^{-1} (\widetilde{M}_{\varepsilon,e} - \widetilde{M}_{\varepsilon,c})(\Psi_e \Phi_{f,\varepsilon} - \Psi_c \Phi_{c,\varepsilon}) d\mu(x) dt \right) \end{aligned}$$

this is possible if and only if $\widetilde{M}_{1,e} = \widetilde{M}_{1,c} + \mu$ in that case, with $\mu \in L^2(0, T, L^2(\mathbb{R}^d))$ independent on y . Then, taking $\Psi^\varepsilon = \Psi(x, t) \Phi_\varepsilon(x)$, with $\Psi \in \mathcal{D}([0, T] \times \mathbb{R}^d)$ and $\Phi \in H_\#^1(Y)$, we obtain, the boundary term being zero because of the trace continuity of Ψ^ε :

$$\begin{aligned} & \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y} \sigma(y) \left[\nabla \widetilde{M}_0(x, t) + \nabla_y \widetilde{M}_1(x, y, t) \right] \Psi(x, t) \nabla_y \Phi(y) dx dy dt \\ & - \frac{1}{|Y|} \int_0^T \int_{\mathbb{R}^d \times Y_e} \imath q F(t) \sigma_e(y) \widetilde{M}_0(x, t) \mathbf{n} \cdot \Psi(x, t) \nabla_y \Phi(y) dx dy dt = 0 \end{aligned}$$

and thus, by density (60) is proved. $\square$

Now, reinterpreting the variational formulations we have obtained, we can prove in exactly the same way than for $p \geq 1$, that :

Theorem 5.4. *The variational problems (57), (58) and (60) imply that $M_0 \in W^{q,f}(0, T, H^1(\mathbb{R}^d))$ satisfies for all $V \in H^1(\mathbb{R}^d)$:*

$$\begin{aligned} & \left\langle \partial_t \widetilde{M}_0, V_e \right\rangle_{H^{-1}, H^1} + \left(D \nabla_x \widetilde{M}_0, \nabla V \right)_{L^2} - \left(\imath q F(t) D \widetilde{M}_0 \mathbf{n}, \nabla V \right)_{L^2} \\ & + \left(\imath q F(t) D \nabla \widetilde{M}_0 \cdot \mathbf{n}, V_e \right)_{L^2} + \left(q^2 F(t)^2 D \mathbf{n} \cdot \mathbf{n} \widetilde{M}_0, V \right)_{L^2} = 0 \quad \text{in } \mathcal{D}'([0, T]) \\ & \widetilde{M}_0(\cdot, 0) = M_{init} \quad \text{in } \mathbb{R}^d \end{aligned} \quad (62)$$

where only the homogenized tensor changes from a variational problem to another. For (57), it is given by :

$$D_{\alpha, ij} = \sum_{\alpha \in \{c, e\}} \int_{Y_\alpha} \sigma_\alpha (\nabla w_{j, \alpha} + e_j) \cdot (\nabla w_{i, \alpha} + e_i) \quad (20)$$

where the $w_{i, \alpha}$ are solutions for $i = 1, 2$ of the cell problems :

$$\begin{cases} -\operatorname{div}_y (\sigma_\alpha (\nabla_y w_{i, \alpha} + e_i)) = 0 & \text{in } Y_\alpha \\ \sigma_e \nabla_y w_{i, e} \cdot \nu + \sigma_e e_i \cdot \nu = 0 & \text{on } \Gamma_m \\ \sigma_c \nabla_y w_{i, c} \cdot \nu + \sigma_c e_i \cdot \nu = 0 & \text{on } \Gamma_m \\ w_{i, e} \text{ } Y_e - \text{periodic} \end{cases} \quad (21)$$

For (58), it is given by :

$$D_{\alpha, ij} = \sum_{\alpha \in \{c, e\}} \int_{Y_\alpha} \sigma_\alpha (\nabla w_{j, \alpha} + e_j) \cdot (\nabla w_{i, \alpha} + e_i) + \sigma_m \int_{\Gamma_m} (w_{j, e} - w_{j, c})(w_{i, e} - w_{i, c}) \quad (22)$$

where the $w_{i, \alpha}$ are solutions for $i = 1, 2$ of the cell problems :

$$\begin{cases} -\operatorname{div}_y (\sigma_\alpha (\nabla_y w_{i, \alpha} + e_i)) = 0 & \text{in } Y_\alpha \\ \sigma_\alpha \nabla_y w_{i, \alpha} \cdot \nu + \sigma_\alpha e_i \cdot \nu = \sigma_m (w_{i, e} - w_{i, c}) & \text{on } \Gamma_m \\ \sigma_e \nabla_y w_{i, e} \cdot \nu + \sigma_e e_i \cdot \nu = \sigma_c \nabla_y w_{i, c} \cdot \nu + \sigma_c e_i \cdot \nu & \text{on } \Gamma_m \\ w_{i, e} \text{ } Y_e - \text{periodic} \end{cases} \quad (23)$$

For (60), it is given by :

$$D_{ij} = \sum_{\alpha \in \{c,e\}} \int_{Y_\alpha} \sigma_\alpha (\nabla w_{j,\alpha} + e_j) \cdot (\nabla w_{i,\alpha} + e_i) \quad (24)$$

where the $w_{i,\alpha}$ are solutions for $i = 1, 2$ of the cell problems :

$$\left\{ \begin{array}{ll} -\operatorname{div}_y(\sigma_\alpha(\nabla_y w_{i,\alpha} + e_i)) = 0 & \text{in } Y_\alpha \\ w_{i,e} = w_{i,c} & \text{on } \Gamma_m \\ \sigma_e \nabla_y w_{i,e} \cdot \nu + \sigma_e e_i \cdot \nu = \sigma_c \nabla_y w_{i,c} \cdot \nu + \sigma_c e_i \cdot \nu & \text{on } \Gamma_m \\ w_{i,e} \text{ } Y_e \text{ - periodic} & \end{array} \right. \quad (25)$$

Finally, for (57), we have :

$$\widetilde{M}_{1,\alpha} = \sum_{i=1}^d w_{i,\alpha} \left(\frac{\partial \widetilde{M}_0}{\partial x_i} - iqF(t) \mathbf{n}_i \widetilde{M}_0 \right) \quad \text{in } \mathbb{R}^d \times Y_\alpha, \quad \alpha \in \{e, c\}$$

while for (58) we have, for some $\mu \in L^2(0, T, L^2(\mathbb{R}^d))$:

$$\widetilde{M}_{1,e} = \sum_{i=1}^d w_{i,e} \left(\frac{\partial \widetilde{M}_0}{\partial x_i} - iqF(t) \mathbf{n}_i \widetilde{M}_0 \right) \quad \text{in }]0, T[\times \mathbb{R}^d \times Y_e$$

and

$$\widetilde{M}_{1,c} = \sum_{i=1}^d w_{i,c} \left(\frac{\partial \widetilde{M}_0}{\partial x_i} - iqF(t) \mathbf{n}_i \widetilde{M}_0 \right) - \mu \quad \text{in }]0, T[\times \mathbb{R}^d \times Y_c$$

and for (60) we have :

$$\widetilde{M}_1 = \sum_{i=1}^d w_i \left(\frac{\partial \widetilde{M}_0}{\partial x_i} - iqF(t) \mathbf{n}_i \widetilde{M}_0 \right) \quad \text{in }]0, T[\times \mathbb{R}^d \times Y$$

It is again clear that theorem 3.2 is a direct consequence from theorem 5.4, using the change of unknowns :

$$M_i = \widetilde{M}_i e^{-iq\mathbf{x} \cdot \mathbf{n} F(t)} \quad M_\varepsilon = \widetilde{M}_\varepsilon e^{-iq\mathbf{x} \cdot \mathbf{n} F(t)}$$

6 Conclusion

We have proposed a mathematical justification to the macroscopic model for dMRI introduced in [7], as well as for all the other possible limit problem for two-compartment Bloch-Torrey equations with a scaled permeability (as it is easy to notice that integer exponents are sufficient to cover all the possible cases). This leads to five macroscopic models, among which the coupled model of [7] seems to be the more efficient for modeling realistic dMRI measurements. These results can be easily generalized to situations which involves several permeability conditions, with potentially different scaling for each permeability coefficient, and thus cover many useful situations (for instance the case where a myelin layer covers the cell, with a permeability condition between the cell and the myelin as well as between the myelin and the extra-cellular fluid).

7 Acknowledgements

The author would like to thank Professor H. Haddard for our numerous discussions on the subject of this article and his precious advices. The author would also like to thank an anonymous reviewer for his/her corrections.

References

- [1] G. ALLAIRE, *Homogenization and two-scale convergence*, SIAM J. Math. Anal., 1992
- [2] T. ARBOGAST, J. DOUGLAS AND U. HORNING, *Derivation of the double porosity model of single phase flow via homogenization theory*, SIAM J. Math. Anal. 21, pp. 823-836, 1990
- [3] A. BENSOUSSAN, J.L. LIONS AND G. PAPANICOLAOU, *Asymptotic analysis for periodic structures, vol. 5 de Studies in Mathematics and its Applications*, North-Holland Publishing Co., Amsterdam, 1978
- [4] D. CIORANESCU, P. DONATO AND R. ZAKI, *The Periodic Unfolding Method in Perforated Domains*, Portugaliae Mathematica, 63, 4, pp. 467-496, 2012
- [5] D. CIORANESCU, A. DAMLAMIAN AND G. GRISO, *The Periodic Unfolding Method in Homogenization*, SIAM J. Math. Anal. 40, pp. 1585-1620, 2012
- [6] D. CIORANESCU, A. DAMLAMIAN AND P. DONATO AND G. GRISO AND R. ZAKI, *The Periodic Unfolding Method in Domains with Holes*, SIAM J. Math. Anal. 44, pp. 718-760, 2012
- [7] J. COATLÉVEN, H. HADDARD AND J-R. LI, *A new macroscopic model including membrane exchange for diffusion MRI*, SIAM J. Appl. Math., Vol. 74, No. 2, pp. 516-546, 2014
- [8] R. DAUTRAY AND J.-L. LIONS, *Mathematical analysis and numerical methods for science and technology, volume 5 : evolution problems*, Springer, 1993
- [9] P. DONATO, K. H. LE NGUYEN AND R. TARDIEU, *The periodic unfolding method for a class of imperfect transmission problems*, SIAM J. Math. Anal. Vol. 176, N. 6, pp. 891-927, 2011
- [10] M. A. HORSFIELD AND D. K. JONES, *Applications of diffusion-weighted and diffusion tensor mri to white matter diseases : a review*, NMR Biomed., vol. 15, no. 7-8, pp. 570-577, 2002
- [11] J. KARGER, H. PFEIFER AND W. HEINIK, *Principles and application of self-diffusion measurements by nuclear magnetic resonance*, Advances in magnetic resonance, vol. 12, pp. 1-89, 1988
- [12] D. LEBIHAN AND H. JOHANSEN-BERG, *Diffusion mri at 25: Exploring brain tissue structure and function*, NeuroImage, 2012
- [13] J-R. LI , H. NGUYEN , D. VAN NGUYEN , H. HADDAR , J. COATLÉVEN AND D. LEBIHAN, *Numerical study of a macroscopic finite pulse model of the diffusion MRI signal*, J. Magn. Reson. vol. 248, pp. 54-65, 2014
- [14] J-N. PERNIN, *Diffusion in composite solid: threshold phenomenon and homogenization*, International Journal of Engineering Science, 37 1597-1610, 1999
- [15] M. A. PETER AND MICHAEL BÖHM, *Different choices of scaling in homogenization of diffusion and interfacial exchange in a porous medium*, Math. Meth. Appl. Sci, 31:1257-1282, 2008

- [16] D. SCHNAPPAUFF AND M. ZEILE AND M. B. NIEDERHAGEN AND B. FLEIGE AND P.-U. TUNN AND B. HAMM AND O. DUDECK, *Diffusion-weighted echo-planar magnetic resonance imaging for the assessment of tumor cellularity in patients with soft-tissue sarcomas*, J. Magn. Reson. Imaging, vol. 29, no. 6, pp. 1355-1359, 2009
- [17] T. SUGAHARA, Y. KOROGE, M. KOCHI, I. IKUSHIMA, Y. SHIGEMATU, T. HIRAI, T. OKUDA, L. LIANG, Y. GE, Y. KOMOHARA, Y. USHIO AND M. TAKAHASHI, *Usefulness of diffusion-weighted mri with echo-planar technique in the evaluation of cellularity in gliomas*, J. Magn. Reson. Imaging, vol. 9, no. 1, pp. 53-60, 1999
- [18] Y. TSUSHIMA AND A. TAKAHASHI-TAKETOMI AND K. ENDO, *Magnetic resonance (mr) differential diagnosis of breast tumors using apparent diffusion coefficient (adc) on 1.5-t*, J. Magn. Reson. Imaging, vol. 30, no. 2, pp. 249-255, 2009
- [19] S. WARACH AND D. CHIEN AND W. LI, M. RONTAL AND R. R. EDELMAN, *Fast magnetic resonance diffusion-weighted imaging of acute human stroke*, Neurology, vol. 42, no. 9, 1992