



The Multiscale Hybrid Mixed Method in General Polygonal Meshes

Gabriel R Barrenechea, Fabrice Jaillet, Diego Paredes, Frédéric Valentin

► To cite this version:

Gabriel R Barrenechea, Fabrice Jaillet, Diego Paredes, Frédéric Valentin. The Multiscale Hybrid Mixed Method in General Polygonal Meshes. [Research Report] LNCC - Laboratorio Nacional de Computacao Cientifica. 2019. hal-02054681

HAL Id: hal-02054681

<https://inria.hal.science/hal-02054681>

Submitted on 2 Mar 2019

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.

THE MULTISCALE HYBRID MIXED METHOD IN GENERAL POLYGONAL MESHES

GABRIEL R. BARRENECHEA, FABRICE JAILLET, DIEGO PAREDES, AND FRÉDÉRIC VALENTIN

ABSTRACT. This work extends the general form of the Multiscale Hybrid-Mixed (MHM) method for the second-order Laplace (Darcy) equation to general non-conforming polygonal meshes. The main properties of the MHM method, i.e., stability, optimal convergence, and local conservation, are proven independently of the geometry of the elements used for the first level mesh. More precisely, it is proven that piecewise polynomials of degree k and $k + 1$, $k \geq 0$, for the Lagrange multipliers (flux), along with continuous piecewise polynomial interpolations of degree $k + 1$ posed on second-level sub-meshes are stable if the latter is fine enough with respect to the mesh for the Lagrange multiplier. We provide an explicit sufficient condition for this restriction. Also, we prove that the error converges with order $k + 1$ and $k + 2$ in the broken H^1 and L^2 norms, respectively, under usual regularity assumptions, and that such estimates also hold for non-convex; or even non-simply connected elements. Numerical results confirm the theoretical findings and illustrate the gain that the use of multiscale functions provides.

1. INTRODUCTION

Multiscale finite element methods have undergone an intense development in the last decades, both in theoretical and practical aspects, for their capacity to be accurate on coarse meshes without assuming neither scale separation nor periodicity of the exact solution, and to be prompt to leverage the new generation of massively parallel computers. Such properties have made multiscale numerical methods attractive to deal with real industrial applications compared to, for instance, more classical homogenization strategies, although its basics remain an important and challenging subject in the applied mathematics field. Starting with the original work [5] for the one-dimensional Laplace problem, the MsFEM was extended and analyzed for multi-dimensional problems in [27, 16]. Since, a large variety of approaches has been proposed leading to the VMS [28] and RFB [33, 12] methods, the HMM [15], and the LOD method [30], just to cite a few.

Recently, hybridisation has been used as a platform to develop multiscale methods. The Multiscale Hybrid-Mixed (MHM for short) method appears as a result of a hybrid formulation that starts at the continuous level posed on a *coarse* partition. Then, a decomposition of the exact solution is obtained in terms of a variable defined in the skeleton (the flux) and a constant per element of the coarse partition. The global problem needs for its construction the solution of local problems that fulfill the role of upscaling the under-mesh structures. Introduced and analysed in [24, 2, 31] for the Laplace (Darcy) equation, the MHM method has been further extended to other elliptic problems in [25, 23] as well as to mixed and hyperbolic models in [3] and [29], respectively. See also [26] for an abstract setting for the MHM method. This method shares similarities, but also fundamental differences, with related approaches such as the Multiscale Mortar method [4], the DEM [19], and the HDG method [11]. Also, in its fully discrete version, and according to the choices made for the interpolation spaces, the underlying

algebraic system associated to MHM method may be identified as being part of the family of FETI domain decomposition methods [20].

All the methods mentioned above are defined on conforming meshes, either simplicial, quadrilateral, or prismatic. The need for conformity of the meshes just mentioned makes some situations like non-matching interfaces resulting from, for example, moving domains, multi-physics mesh gluing, or front tracking, far from trivial. Due to this, among other reasons, in the last few years there has been an explosive development of numerical methods defined on general polyhedral meshes. Just to name a few, the Virtual Finite Element (VEM) [7, 6] uses virtual basis functions while keeping the conformity of the approach, and the discontinuous Galerkin methods uses unmapped polynomials on polygonal meshes [9, 8]. Also, the HHO method [18] appears after a hybridisation process, and as such, it has been recently linked to the HDG method in [10]. Motivated by the need to glue non-conforming meshes in a domain decomposition setting, the works [1, 21] propose a discretisation posed in polygonal elements, where the coupling is done by means of a Robin interface condition.

The main purpose of this work is to use the flexibility given by the MHM approach to extend its use to polygonal meshes. In fact, the first step in the derivation of the method is a new weak formulation having as unknowns the fluxes on the edges of the coarse partition and the mean value of the solution in each element. This formulation can be derived independently of the geometry of the elements. Thus, besides stating this fact rigorously, the bulk of the work is devoted to the proposal of a stable finite element method for this new weak problem. For this, we need two main ingredients, namely, a finite element space to approximate the fluxes over the edges (i.e., the Lagrange multipliers) and a finite element method to approximate the solution of the local problems. For the former we propose to use discontinuous polynomials of degree k or $k + 1$ ($k \geq 0$) in a sub-partition of the skeleton, while we use continuous Lagrangian polynomials of degree $k + 1$ on a conforming sub-mesh for the latter. It is important to mention that the meshes used in each element don't need to match at the interfaces. We focus in a diffusion equation in two space dimensions, but the results can be extended, without major complications, to three space dimensions and more complicated operators.

Other than the extension of the MHM framework to polygonal meshes, we now highlight the main contributions of this work:

- the method, and its analysis, has been extended to allowing discontinuous spaces for approximating the fluxes. This is new even for the MHM method proposed on simplicial meshes;
- the possibility of using equal order approximations for both the flux and the primal variable is also new, even in the case of conforming simplicial meshes. It is important to mention that this possibility does not require any extra stabilisation term to be added to the formulation. This comes at the price of supposing that the subelement-meshes are finer than the sub-partitions of the skeleton, but an explicit sufficient condition is given for this;
- error estimates, qualitatively similar (super-convergence) to the ones given in the last six rows of the summary given in [10, Table 1] are proven without post-processing procedures;
- convergence has been proven without the need for the coarse partition to be refined. It is enough that the sub-partitions to approximate the Lagrange multiplier and the sub-mesh used to solve the local problems get refined.

This work is outlined as follows. We end this section by presenting the model, notations and some preliminary results. Section 2 is dedicated to the presentation of the MHM method and important considerations on the second-level sub-meshes. Well-posedness and error analysis are

the subjects of Section 3 and Section 4, respectively. Section 5 assesses theoretical results through two numerical tests, followed by conclusions in Section 6. Two technical lemmas instrumental to the proof of the existence of a Fortin operator are left to the appendix.

1.1. Notations and the model problem. Let $\Omega \subset \mathbb{R}^2$ be an open, bounded, polygonal domain with a Lipschitz boundary $\partial\Omega$. Given $f \in L^2(\Omega)$ and $g \in H^{\frac{1}{2}}(\partial\Omega)$, this work aims at approximating the following boundary value problem: Find $u \in H^1(\Omega)$ such that $u|_{\partial\Omega} = g$ and

$$(1.1) \quad \int_{\Omega} A \nabla u \cdot \nabla v dx = \int_{\Omega} f v dx \quad \text{for all } v \in H_0^1(\Omega).$$

Here, $A \in L^\infty(\Omega)^{2 \times 2}$ is a symmetric matrix and may involve multiscale features. It is supposed to be uniformly elliptic in Ω . More precisely, we will assume that there exist positive constants $A_{\min}$ and $A_{\max}$ such that

$$(1.2) \quad A_{\min} |\boldsymbol{\xi}|^2 \leq \boldsymbol{\xi}^T A(\mathbf{x}) \boldsymbol{\xi} \leq A_{\max} |\boldsymbol{\xi}|^2 \quad \text{for all } \boldsymbol{\xi} \in \mathbb{R}^2 \text{ and } \mathbf{x} \in \Omega,$$

where $|\cdot|$ stands for the Euclidean norm in $\mathbb{R}^2$. We shall also make use of the following value

$$(1.3) \quad \omega := \frac{A_{\max}}{A_{\min}},$$

and note that if the entries of A are constant functions, then ω is simply the condition number of A . Above, and hereafter, we adopt standard notation for Sobolev and Lebesgue spaces aligned with, e.g., [17].

1.2. Partitions, triangulations, and spaces. We start by introducing $\mathcal{P}$, a collection of closed, bounded, disjoint polygons, denoted K , such that $\bar{\Omega} = \cup_{K \in \mathcal{P}} K$. The shape of the polygons K is, a priori, arbitrary, but we will suppose they satisfy a minimal angle condition (see Assumption 1 below for a more precise statement). The diameter of K is $\mathcal{H}_K$ and we denote $\mathcal{H} = \max_{K \in \mathcal{P}} \mathcal{H}_K$. We will not suppose necessarily that $\mathcal{H}$ tends to zero, but it may do so in many cases of practical interest. For each $K \in \mathcal{P}$, $\mathbf{n}^K$ denotes the unit outward normal to ∂K . We also introduce $\partial\mathcal{P}$, the set of boundaries ∂K , with $K \in \mathcal{P}$, and $\mathcal{E}$ the set of the edges of $\mathcal{P}$, that is

$$(1.4) \quad \mathcal{E} := \{E = K \cap K' \text{ or } K \cap \partial\Omega : K, K' \in \mathcal{P}, \text{ and it is not reduced to a single point}\}.$$

Associated to the partition $\mathcal{P}$, for $m \geq 0$ (where $H^0 = L^2$, as usual) we define the function spaces

$$(1.5) \quad H^m(\mathcal{P}) := \{v : v|_K \in H^m(K) \quad \text{for all } K \in \mathcal{P}\},$$

$$(1.6) \quad L^2(\mathcal{E}) := \{q : q|_E \in L^2(E) \quad \text{for all } E \in \mathcal{E}\},$$

with norms

$$(1.7) \quad \|v\|_{m,\mathcal{P}} := \left\{ \sum_{K \in \mathcal{P}} \|v\|_{m,K}^2 \right\}^{\frac{1}{2}} \quad \text{and} \quad \|q\|_{0,\mathcal{E}} := \left\{ \sum_{E \in \mathcal{E}} \|q\|_{0,E}^2 \right\}^{\frac{1}{2}}.$$

In addition, the following spaces will be useful in what follows

$$(1.8) \quad V := H^1(\mathcal{P}) \quad \text{with norm } \|v\|_V := \left\{ \sum_{K \in \mathcal{P}} \|v\|_{1,K}^2 \right\}^{\frac{1}{2}},$$

$$(1.9) \quad V_0 := \{v \in L^2(\Omega) : v|_K \in \mathbb{P}_0(K) \quad \text{for all } K \in \mathcal{P}\},$$

$$(1.10) \quad \Lambda := \{\boldsymbol{\tau} \cdot \mathbf{n}^K|_{\partial K} : \boldsymbol{\tau} \in H(\text{div}, \Omega) \quad \text{for all } K \in \mathcal{P}\},$$

and the space Λ is endowed with the norm

$$(1.11) \quad \|\mu\|_\Lambda := \inf_{\substack{\tau \in H(\text{div}, \Omega) \\ \tau \cdot \mathbf{n}^K = \mu \text{ on } \partial K, K \in \mathcal{P}}} \|\tau\|_{\text{div}, \Omega}.$$

We also denote by $(\cdot, \cdot)_D$ the $L^2(D)$ -inner product (we don't make a distinction between vector-valued and scalar-valued functions). We define the product on $\mathcal{P}$ as follows

$$(1.12) \quad (v, w)_\mathcal{P} := \sum_{K \in \mathcal{P}} (v, w)_K,$$

and the broken product on $\partial\mathcal{P}$ as

$$(1.13) \quad \langle \mu, v \rangle_{\partial\mathcal{P}} := \sum_{K \in \mathcal{P}} \langle \mu, v \rangle_{\partial K},$$

where the product on ∂K is the duality pairing between $H^{-\frac{1}{2}}(\partial K)$ and $H^{\frac{1}{2}}(\partial K)$. Following closely the arguments given in [2] we can prove that

$$(1.14) \quad \frac{\sqrt{2}}{2} \|\mu\|_\Lambda \leq \sup_{v \in V} \frac{\langle \mu, v \rangle_{\partial\mathcal{P}}}{\|v\|_V} \leq \|\mu\|_\Lambda,$$

and above and hereafter we lighten the notation and understand the supremum to be taken over sets excluding the zero function, even though this is not specifically indicated.

In Section 2.1 a characterisation of the weak solution of (1.1) will be derived using the partition $\mathcal{P}$ described above. That characterisation will be the starting point of the finite element method analysed in this work, which will need some further notations and partitions. More precisely, we introduce two partitions which do not coincide, but are not independent. We start describing the discretisation of the set of edges $\mathcal{E}$. For this we introduce $\mathcal{E}_H$, a partition of $\mathcal{E}$ for which each $E \in \mathcal{E}$ is split into segments F of length $H_F \leq H := \max_{F' \in \mathcal{E}_H} H_{F'}$. We will not assume the segments are of equal length, but we will require that neighbouring edges are not too dissimilar. More precisely, we impose the following assumption on $\mathcal{E}_H$:

Assumption (A1): The mesh $\mathcal{E}_H$ is such that in every $K \in \mathcal{P}$ a shape regular simplicial triangulation $\Xi_H(K)$ of K can be built in such a way that its trace on ∂K coincides with $\mathcal{E}_H$.

The triangulation $\Xi_H(K)$ will be useful in the definition of the method, but not explicitly used. More precisely, we make the following definitions:

- for each $F \in \mathcal{E}_H$, we denote by κ_F^K (or simply κ_F , when it is clear from the context to which element it does belong) the only element in $\Xi_H(K)$ such that $F = \kappa_F^K \cap \partial K$. Just to simplify the presentation, we will suppose that, for two different $F, F' \in \mathcal{E}_H$, $\kappa_F^K \neq \kappa_{F'}^K$;
- the triangulation $\Xi_H := \cup_{K \in \mathcal{P}} \Xi_H(K)$ will be referred to as *virtual triangulation*. This (conforming) partition will be very useful in the proofs below, but not used explicitly in the implementation of the method;
- for each $K \in \mathcal{P}$, we introduce a shape regular family of simplicial triangulations $\{\mathcal{T}_h^K\}_{h>0}$ built in the following way:
 - (1) first, on each $K \in \mathcal{P}$, the triangulation $\Xi_H(K)$ is refined once using a red refinement. The resulting triangulation is called *minimal triangulation*;
 - (2) then, for each K , the family $\{\mathcal{T}_h^K\}_{h>0}$ is formed by regular refinements of the minimal triangulation.

The diameter of $\mathfrak{T} \in \mathcal{T}_h^K$ is denoted by $h_{\mathfrak{T}}$, and $h := \max_{K \in \mathcal{P}} \max_{\mathfrak{T} \in \mathcal{T}_h^K} h_{\mathfrak{T}}$, denote $\mathcal{T}_h := \cup_{K \in \mathcal{P}} \mathcal{T}_h^K$, and, define the broken space

$$H(\text{div}, \mathcal{T}_h) := \{\boldsymbol{\tau} : \boldsymbol{\tau}|_{\mathfrak{T}} \in H(\text{div}, \mathfrak{T}), \forall \boldsymbol{\tau} \in \mathcal{T}_h\}.$$

It is important to remark that, if $E = K \cap K' \in \mathcal{E}$, then the traces of the two neighbouring triangulations $\mathcal{T}_h^K$ and $\mathcal{T}_h^{K'}$ do not need to coincide.

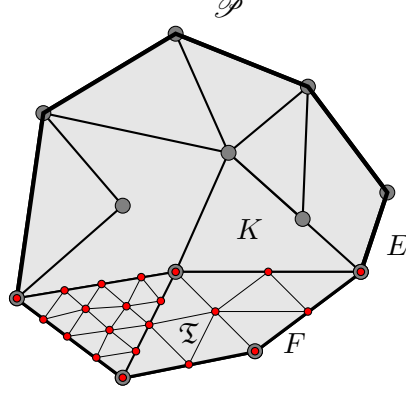


FIGURE 1. A domain partitioned with non-conforming polygonal elements K . Observe the sub-meshes discretising two different elements K with different granularity. The red dots represent the degrees of freedom associated with the sub-meshes and the gray dots with the mesh skeleton.

Associated to $\mathcal{E}_H$ and $\mathcal{T}_h^K$, for $k \geq 0$, we introduce the following finite element spaces:

$$\begin{aligned} V_h &:= \prod_{K \in \mathcal{P}} V_h(K) \quad \text{where} \quad V_h(K) := \{v_h \in C^0(K) : v_h|_{\mathfrak{T}} \in \mathbb{P}_{k+1}(\mathfrak{T}), \forall \mathfrak{T} \in \mathcal{T}_h^K\}, \\ (1.15) \quad \tilde{V}_h &:= \prod_{K \in \mathcal{P}} \tilde{V}_h(K) \quad \text{where} \quad \tilde{V}_h(K) := V_h(K) \cap L_0^2(K), \\ \Lambda_H &:= \{\mu_H \in \Lambda : \mu_H|_F \in \mathbb{P}_\ell(F), \forall F \in \mathcal{E}_H\} \quad \text{with } \ell = k \text{ or } \ell = k+1. \end{aligned}$$

We also introduce a projection onto Λ_H . We start by defining, for every $F \in \mathcal{E}_H$ the projection $\Pi_F^\ell : L^2(F) \rightarrow \mathbb{P}_\ell(F)$ as

$$(1.16) \quad (\Pi_F^\ell(\mu), q)_F = (\mu, q)_F \quad \text{for all } q \in \mathbb{P}_\ell(F),$$

and then we define the global projection $\Pi^\ell : L^2(\mathcal{E}) \rightarrow \Lambda_H$ by $\Pi^\ell(\mu)|_F = \Pi_F^\ell(\mu)$ for every $\mu \in L^2(\mathcal{E})$. In addition, we introduce $\mathcal{C}_h : H^1(\Omega) \rightarrow V_h$, a variant of the Cl  ment interpolation operator defined locally. That is, for every $v \in V$ we define $\mathcal{C}_h(v)|_K = \mathcal{C}_h^K(v)$, where $\mathcal{C}_h^K : H^1(K) \rightarrow V_h(K)$ is the usual Cl  ment interpolation operator. This mapping satisfies the following (see [17]): there exists $C > 0$, depending only on the shape of the elements $\mathfrak{T} \in \mathcal{T}_h^K$ such that, for all $v \in H^1(K)$ and all $\mathfrak{T} \in \mathcal{T}_h^K$,

$$(1.17) \quad \|\mathcal{C}_h^K(v)\|_{1,\mathfrak{T}} \leq C \|v\|_{1,\omega_{\mathfrak{T}}},$$

$$(1.18) \quad \|v - \mathcal{C}_h^K(v)\|_{0,\mathfrak{T}} \leq Ch |v|_{1,\omega_{\mathfrak{T}}},$$

where $\omega_{\mathfrak{T}} := \{\mathfrak{T}' \in \mathcal{T}_h^K : \mathfrak{T} \cap \mathfrak{T}' \neq \emptyset\}$.

Finally, in what follows C will denote a positive constant whose value does not depend on any mesh size, or the shape of $K \in \mathcal{P}$. The constant C is only allowed to depend on the shape of the elements of $\mathcal{T}_h$, and its value may change whenever it is written in two different locations.

Remark 1.1. *Assumption (A1) implies a restriction on the shape of K , and/or how refined the partition $\mathcal{E}_H$ is. More precisely, if the diameter of the polygons $K \in \mathcal{P}$ tends to zero, and their smallest angle is not bounded below, then the triangulation $\Xi_H(K)$ can not be built, as the smallest angle in $\Xi_H(K)$ would degenerate as well. On the other hand, even if the smallest angle of $\mathcal{P}$ is uniformly bounded below, if the polygons $K \in \mathcal{P}$ are very anisotropic, or irregular, then $\mathcal{E}_H$ needs to be fine enough so $\Xi_H(K)$ can be built.*

Remark 1.2. *The restriction of Ξ_H and $\mathcal{T}_h$ to be simplicial is made only for simplicity. The results presented in this work can be extended without major complications to the case in which those meshes are built using quadrilaterals. In addition, the way the triangulations $\mathcal{T}_h^K$ are built has been done mostly to make the presentation of the method (and the proofs) clearer. All the results presented below follow, with minor modifications, if we suppose that $\mathcal{E}_H$ and $\mathcal{T}_h^K$ are independent, but linked by the following assumption: For every $E = K \cap K' \in \mathcal{E}$ and every $F \in \mathcal{E}_H, F \subseteq E$, there exist two pairs of triangles $\mathfrak{T}_1^K, \mathfrak{T}_2^K \in \mathcal{T}_h^K$ and $\mathfrak{T}_1^{K'}, \mathfrak{T}_2^{K'} \in \mathcal{T}_h^{K'}$ such that they share one node and*

$$\left(\partial\mathfrak{T}_1^K \cup \partial\mathfrak{T}_2^K\right) \cap E \subseteq F \quad \text{and} \quad \left(\partial\mathfrak{T}_1^{K'} \cup \partial\mathfrak{T}_2^{K'}\right) \cap E \subseteq F.$$

This restriction requires the triangulations $\mathcal{T}_h^K$ to be fine enough such that for every $E \in \Lambda_H$, there are at least two triangles in $\mathcal{T}_h^K$ whose edges lying on ∂K are totally included in E . With this assumption the proofs of stability and convergence done in the subsequent sections follow almost exactly in the same way.

2. THE MHM METHOD

2.1. A characterisation of the exact solution. A fundamental point of the MHM method is a characterisation of u , weak solution of (1.1), as a function of a pair (λ, u_0) , solution of a mixed hybrid problem. To derive this characterisation we define the mappings $T \in \mathcal{L}(\Lambda, V)$ and $\hat{T} \in \mathcal{L}(L^2(\Omega), V)$ as follows:

- for all $\mu \in \Lambda$, $T\mu \in H^1(K) \cap L_0^2(K)$ is the unique solution of

$$(2.1) \quad \int_K A \nabla T\mu \cdot \nabla v \, dx = -\langle \mu, v \rangle_{\partial K} \quad \text{for all } v \in H^1(K) \cap L_0^2(K), \quad \forall K \in \mathcal{P};$$

- for all $q \in L^2(\Omega)$, $\hat{T}q \in H^1(K) \cap L_0^2(K)$ is the unique solution of

$$(2.2) \quad \int_K A \nabla \hat{T}q \cdot \nabla v \, dx = \int_K qv \, dx \quad \text{for all } v \in H^1(K) \cap L_0^2(K), \quad \forall K \in \mathcal{P}.$$

Following very closely the derivation from [2] the solution of (1.1) can be written as follows

$$(2.3) \quad u = u_0 + T\lambda + \hat{T}f,$$

where $(\lambda, u_0) \in \Lambda \times V_0$ solves the following mixed problem: Find $(\lambda, u_0) \in \Lambda \times V_0$ such that

$$(2.4) \quad \begin{aligned} a(\lambda, \mu) + b(\mu, u_0) &= -\langle \mu, \hat{T}f \rangle_{\partial\mathcal{P}} + \langle \mu, g \rangle_{\partial\Omega} \quad \text{for all } \mu \in \Lambda, \\ b(\lambda, v_0) &= (f, v_0)_{\mathcal{P}} \quad \text{for all } v_0 \in V_0, \end{aligned}$$

and the continuous bilinear forms $a(\cdot, \cdot)$ and $b(\cdot, \cdot)$ are given by

$$(2.5) \quad a : \Lambda \times \Lambda \rightarrow \mathbb{R} \quad a(\lambda, \mu) := \langle \mu, T\lambda \rangle_{\partial \mathcal{P}},$$

$$(2.6) \quad b : \Lambda \times V_0 \rightarrow \mathbb{R} \quad b(\mu, v_0) := \langle \mu, v_0 \rangle_{\partial \mathcal{P}}.$$

The well-posedness of (2.4) is stated next.

Theorem 2.1. *Let $\mathcal{N}$ be the space defined by*

$$(2.7) \quad \mathcal{N} := \{\mu \in \Lambda : b(\mu, v_0) = 0 \text{ for all } v_0 \in V_0\},$$

then there exist positive constants α (depending on $A_{\min}$) and β such that

$$(2.8) \quad a(\mu, \mu) \geq \alpha \|\mu\|_{\Lambda}^2 \text{ for all } \mu \in \mathcal{N},$$

$$(2.9) \quad \sup_{\mu \in \Lambda} \frac{b(\mu, v_0)}{\|\mu\|_{\Lambda}} \geq \beta \|v_0\|_V \text{ for all } v_0 \in V_0.$$

Consequently, (2.4) has a unique solution $(\lambda, u_0) \in \Lambda \times V_0$. Moreover, u given in (2.3) satisfies (1.1) and $A\nabla u \cdot \mathbf{n}^K|_{\partial K} = -\lambda$ for all $K \in \mathcal{P}$.

Proof. Regarding (2.8), the proof follows exactly as in [2]. As for (2.9), we omit the details since it is very similar to the one of Theorem 3.1 below. $\square$

We finish this section by two quick remarks on the solution (λ, u_0) of (2.4). First, the dual variable $\sigma := A\nabla u$ belongs to $H(\text{div}, \Omega)$ since the flux $\lambda = -\sigma \cdot \mathbf{n}^K|_{\partial K} \in \Lambda$ for all $K \in \mathcal{P}$. So, the relaxation of the continuity of u does not affect the continuity of the fluxes. In addition, since $T\lambda$ and $\hat{T}f$ have zero mean value in every $K \in \mathcal{P}$ the following holds

$$u_0|_K = \frac{1}{|K|} \int_K u \, dx.$$

2.2. The method. We start by defining the discrete equivalents of the operators defined in (2.1)-(2.2). Using the finite element spaces defined in (1.15) we introduce the following approximate mappings:

- for all $\mu \in \Lambda$, $T_h\mu \in \tilde{V}_h$ is the unique solution of

$$(2.10) \quad \int_K A\nabla T_h\mu \cdot \nabla v_h \, dx = -\langle \mu, v_h \rangle_{\partial K} \text{ for all } v_h \in \tilde{V}_h(K), \quad \forall K \in \mathcal{P};$$

- for all $q \in L^2(\Omega)$, $\hat{T}_h q \in \tilde{V}_h$ is the unique solution of

$$(2.11) \quad \int_K A\nabla \hat{T}_h q \cdot \nabla v_h \, dx = \int_K q v_h \, dx \text{ for all } v_h \in \tilde{V}_h(K), \quad \forall K \in \mathcal{P}.$$

Using the mappings (2.10)-(2.11) and the following approximate bilinear form

$$(2.12) \quad a_h : \Lambda \times \Lambda \rightarrow \mathbb{R} \quad \text{where} \quad a_h(\lambda, \mu) = \langle \mu, T_h\lambda \rangle_{\partial \mathcal{P}},$$

the MHM method associated to (2.4) reads: Find $(\lambda_H, u_0^h) \in \Lambda_H \times V_0$ such that

$$(2.13) \quad \begin{aligned} a_h(\lambda_H, \mu_H) + b(\mu_H, u_0^h) &= -\langle \mu_H, \hat{T}_h f \rangle_{\partial \mathcal{P}} + \langle \mu_H, g \rangle_{\partial \Omega} \text{ for all } \mu_H \in \Lambda_H, \\ b(\lambda_H, v_0) &= (f, v_0)_{\mathcal{P}} \text{ for all } v_0 \in V_0. \end{aligned}$$

The approximate solution u_h is given by

$$(2.14) \quad u_h := u_0^h + T_h\lambda_H + \hat{T}_h f.$$

Remark 2.1. *The fact that the exact flux λ is locally conservative is inherited by its discrete counterpart λ_H in each $K \in \mathcal{P}$. In fact, from the second equation in (2.13) we get*

$$\int_{\partial K} \lambda_H ds = \int_K f dx \quad \text{for all } K \in \mathcal{P}.$$

On the other hand, unless a mixed finite element method is used as a second order solver (see [14] for an example) in the second level mesh, this is not guaranteed. More precisely, for $\mathfrak{T} \in \mathcal{T}_h^K$ we, in general, have

$$-\int_{\mathfrak{T}} \nabla \cdot \sigma_h dx \neq \int_{\mathfrak{T}} f dx,$$

where $\sigma_h := A \nabla u_h$. Finally, it is worth mentioning that the choice of numerical method to define T_h and $\hat{T}_h$ is virtually unlimited. In this paper we have restricted the presentation to a Galerkin method, but other choices, such as stabilised, enriched, or DG-related methods, just to name a few, are also possible, leading to similar theoretical results.

Remark 2.2. *It is interesting to remark that $w_h := T_h \mu + \hat{T}_h f \in \tilde{V}_h(K)$ is the unique solution of the problem*

$$\int_K A \nabla (T_h \mu + \hat{T}_h f) \cdot \nabla v_h dx = -\langle \mu, v_h \rangle_{\partial K} + \int_K f v_h dx = \int_K A \nabla (T \mu + \hat{T} f) \cdot \nabla v_h dx,$$

for all $v_h \in \tilde{V}_h(K)$. Thus, using Cea's Lemma, the following estimate follows

$$(2.15) \quad \|T \mu + \hat{T} f - w_h\|_{1,K} \leq C_K \inf_{v_h \in \tilde{V}_h(K)} \|T \mu + \hat{T} f - v_h\|_{1,K},$$

where C_K depends on ratio ω given in (1.3) for each K , but is independent of h, H , or $\mathcal{H}$. This fact will be of paramount importance in the proof of optimal convergence of the method, even in the case the polygons $K \in \mathcal{P}$ are not convex.

Lemma 2.3. *There exist constants C , independent of $h, H, \mathcal{H}$, and A , such that*

$$(2.16) \quad \|T_h \mu\|_V \leq C A_{min}^{-1} \|\mu\|_{\Lambda} \quad \text{and} \quad \|T \mu\|_V \leq C A_{min}^{-1} \|\mu\|_{\Lambda},$$

$$(2.17) \quad \|\hat{T}_h f\|_V \leq C A_{min}^{-1} \|f\|_{0,\Omega} \quad \text{and} \quad \|\hat{T} f\|_V \leq C A_{min}^{-1} \|f\|_{0,\Omega}.$$

Proof. Let $\mu \in \Lambda$, from (2.10) and (1.14) we get

$$\begin{aligned} C A_{min} \|T_h \mu\|_V^2 &\leq \sum_{K \in \mathcal{P}} \int_K A \nabla T_h \mu \cdot \nabla T_h \mu dx = - \sum_{K \in \mathcal{P}} \langle \mu, T_h \mu \rangle_{\partial K} \\ &\leq \sup_{v_h \in V_h} \frac{\langle \mu, v_h \rangle_{\partial \mathcal{P}}}{\|v_h\|_V} \|T_h \mu\|_V \\ &\leq \|\mu\|_{\Lambda} \|T_h \mu\|_V, \end{aligned}$$

and the result involving T follows using an analogous argument which proves (2.16). Next, let $f \in L^2(\Omega)$. From (2.11) and using the Cauchy-Schwarz inequality, we get

$$\begin{aligned}
C A_{\min} \|\hat{T}_h f\|_V^2 &\leq \sum_{K \in \mathcal{P}} \int_K A \nabla \hat{T}_h f \cdot \nabla \hat{T}_h f \, dx = \sum_{K \in \mathcal{P}} (f, \hat{T}_h f)_K \\
&\leq \sum_{K \in \mathcal{P}} \|f\|_{0,K} \|\hat{T}_h f_K\|_{0,K} \\
&\leq \left(\sum_{K \in \mathcal{P}} \|f\|_{0,K}^2 \right)^{\frac{1}{2}} \left(\sum_{K \in \mathcal{P}} \|\hat{T}_h f\|_{0,K}^2 \right)^{\frac{1}{2}} \\
&\leq \|f\|_{0,\Omega} \|\hat{T}_h f\|_V,
\end{aligned}$$

and the result involving $\hat{T}$ follows using an analogous argument which gives (2.17). $\square$

Remark 2.4. Observe that the right-hand side of (2.13) may be rewritten using the following equivalence

$$(2.18) \quad -\langle \mu, \hat{T}_h f \rangle_{\partial K} = \int_K A \nabla T_h \mu \cdot \nabla \hat{T}_h f \, dx = \int_K T_h \mu \, f \, dx,$$

for all $\mu \in \Lambda$ and $K \in \mathcal{P}$. Also, if $f \in \mathbb{P}_0(K)$ then $\hat{T}_h f|_K = 0$ and the right-hand side of (2.13) simplifies.

3. WELL-POSEDNESS

We address the well-posedness of the MHM method given in (2.13). The main ingredient is the construction of a Fortin operator. This will be detailed for the choice $\ell = k$, and sketched for the choice $\ell = k + 1$.

3.1. The $\mathbb{P}_k(F) \times \mathbb{P}_{k+1}(\mathfrak{T})$ element. We start considering the case $\ell = k$. The following result ensures the existence of a Fortin operator acting on V with image in V_h , and using the space Λ_H . This result will be key to the well-posedness of (2.13).

Lemma 3.1. *Let us assume Assumption (A1). Then, there exists a mapping $\Pi_h : V \rightarrow V_h$ such that, for all $v \in V$:*

$$(3.1) \quad \int_F \Pi_h(v) \mu_H \, ds = \int_F v \mu_H \, ds \quad \text{for all } \mu_H \in \Lambda_H \quad \text{and } F \in \mathcal{E}_H,$$

$$(3.2) \quad \|\Pi_h(v)\|_V \leq C \|v\|_V,$$

where $C > 0$ does not depend on h, H , or the shape of $K \in \mathcal{P}$.

Proof. We will prove the result for the case in which the mesh $\mathcal{T}_h^K$ is the minimal mesh allowed. That is, $\mathcal{T}_h^K$ is obtained by performing only one red refinement of the mesh $\Xi_H(K)$. Since the discrete spaces $V_h(K)$ associated to further refinements of this mesh include the one associated to the minimal mesh, the stability proved below will also apply to those choices.

Let $K \in \mathcal{P}$, and let us define the mapping $\rho_h^K : H^1(K) \rightarrow V_h(K)$ as follows. For every $F \in \mathcal{E}_H \cap \partial K$, there are exactly two neighboring triangles $\mathfrak{T}_1, \mathfrak{T}_2 \in \mathcal{T}_h^K$ with at least one edge contained in F . Let e_1 and e_2 be these edges, so that $F = e_1 \cup e_2$. Let us denote by $\mathbf{x}_1, \dots, \mathbf{x}_{k+1}$ the position of $k+1$ degrees of freedom of $V_h(K)$ in F° (the interior of F), and let $\varphi_1, \dots, \varphi_{k+1} \in V_h(K)$ be the basis functions of $V_h(K)$ associated to these degrees of

freedom. In addition, we fix a basis $\{\hat{\mu}_1, \dots, \hat{\mu}_{k+1}\}$ of $\mathbb{P}_k(0, 1)$ satisfying $|\hat{\mu}_i(\hat{x})| \leq 1$ in $(0, 1)$ for $i = 1, \dots, k+1$. Then, we map $(0, 1)$ onto F to build a basis $\{\mu_1, \dots, \mu_{k+1}\}$ of $\mathbb{P}_k(F)$, and then we have that $|\mu_i(\mathbf{x})| \leq 1$ for all $\mathbf{x} \in F$.

Next, let $(\alpha_1^F, \dots, \alpha_{k+1}^F)^T$ be the solution of the linear system

$$(3.3) \quad \sum_{j=1}^{k+1} \int_F \mu_i(s) \varphi_j(s) ds \alpha_j^F = \int_F v(s) \mu_i(s) ds \quad i = 1, \dots, k+1.$$

This system can be written in matrix form as

$$(3.4) \quad \mathbb{A} \boldsymbol{\alpha}^F = \left[\int_F v(s) \mu_i(s) ds \right]_{i=1}^{k+1},$$

where $\boldsymbol{\alpha}^F = (\alpha_1^F, \dots, \alpha_{k+1}^F)^T$, and

$$(3.5) \quad \mathbb{A} = (a_{ij})_{i,j=1,\dots,k+1} \quad , \quad a_{ij} = \int_F \mu_i(s) \varphi_j(s) ds.$$

We notice that, changing variables $\mathbb{A} = H_F \hat{\mathbb{A}}$, where $\hat{\mathbb{A}} = (\hat{a}_{ij})_{i,j=1,\dots,k+1}$ and

$$(3.6) \quad \hat{a}_{ij} = \int_0^1 \hat{\mu}_i(\hat{s}) \hat{\varphi}_j(\hat{s}) d\hat{s},$$

where $\{\hat{\varphi}_1, \dots, \hat{\varphi}_{k+1}\}$ is the corresponding basis in $(0, 1)$. We notice that we can always choose the position of $\mathbf{x}_1, \dots, \mathbf{x}_{k+1}$ in such a way that this basis in $(0, 1)$ remains unchanged, so $\hat{\mathbb{A}}$ is independent of F . Due to Lemma 6.1 (see the Appendix) $\hat{\mathbb{A}}$ is invertible, which also shows the invertibility of $\mathbb{A}$.

With these ingredients we define

$$(3.7) \quad \rho_h^K(v) := \sum_{F \in \mathcal{E}_H \cap \partial K} \rho_F^K(v) \quad \text{where} \quad \rho_F^K(v) = \sum_{i=1}^{k+1} \alpha_i^F \varphi_i,$$

and, for each F , $\alpha_i^F, i = 1, \dots, k+1$, solve (3.3).

We now analyse the stability of ρ_h^K . First,

$$(3.8) \quad \mathbb{A} \boldsymbol{\alpha} = \left[\int_F v(s) \mu_i(s) ds \right]_{i=1}^{k+1} \implies \boldsymbol{\alpha}^F = H_F^{-1} \hat{\mathbb{A}}^{-1} \left[\int_F v(s) \mu_i(s) ds \right]_{i=1}^{k+1}.$$

Thus, using (3.8) and since $\hat{\mathbb{A}}$ does not depend on F , and using Cauchy-Schwarz's inequality and $|\mu_i| \leq 1$ on $F \subset \mathcal{E}_H$ we arrive at

$$(3.9) \quad \begin{aligned} \sum_{i=1}^{k+1} (\alpha_i^F)^2 &\leq C H_F^{-2} \sum_{i=1}^{k+1} \left(\int_F v(s) \mu_i(s) ds \right)^2 \\ &\leq C H_F^{-2} \|v\|_{0,F}^2 \max_{i=1,\dots,k+1} \|\mu_i\|_{0,F}^2 \\ &\leq C H_F^{-2} H_F \|v\|_{0,F}^2 \\ &\leq C H_F^{-1} \|v\|_{0,F}^2, \end{aligned}$$

where $C > 0$ depends only on $\hat{\mathbb{A}}$, but not on h, H , or the shape or size of K , but may depend on the degree k .

Hence, recalling that κ_F is the unique triangle in $\Xi_H(K)$ that has F as one of its edges, using (3.9) and a local trace inequality in κ_F we arrive at

$$\begin{aligned}
\|\rho_h^K(v)\|_{0,K}^2 &= \sum_{F \in \mathcal{E}_H \cap \partial K} \|\rho_F^K(v)\|_{0,\mathfrak{T}_1 \cup \mathfrak{T}_2}^2 \\
&= \sum_{F \in \mathcal{E}_H \cap \partial K} \int_{\mathfrak{T}_1 \cup \mathfrak{T}_2} \left(\sum_{i=1}^{k+1} \alpha_i^F \varphi_i \right)^2 \\
&\leq C \sum_{F \in \mathcal{E}_H \cap \partial K} |\mathfrak{T}_1 \cup \mathfrak{T}_2| \sum_{i=1}^{k+1} (\alpha_i^F)^2 \\
&\leq C \sum_{F \in \mathcal{E}_H \cap \partial K} |\mathfrak{T}_1 \cup \mathfrak{T}_2| H_F^{-1} \|v\|_{0,F}^2 \\
(3.10) \quad &\leq C \sum_{F \in \mathcal{E}_H \cap \partial K} \left\{ \|v\|_{0,\kappa_F}^2 + h_{\mathfrak{T}_1}^2 |v|_{1,\kappa_F}^2 \right\},
\end{aligned}$$

where we have used the relation between H and h and the regularity of $\mathcal{T}_h^K$ in the last step. Analogously, an inverse inequality and (3.10) gives

$$(3.11) \quad |\rho_h^K(v)|_{1,K}^2 = \sum_{F \in \mathcal{E}_H \cap \partial K} |\rho_F^K(v)|_{1,\mathfrak{T}_1 \cup \mathfrak{T}_2}^2 \leq C \sum_{F \in \mathcal{E}_H \cap \partial K} \left(h_{\mathfrak{T}_1}^{-2} \|v\|_{0,\kappa_F}^2 + |v|_{1,\kappa_F}^2 \right).$$

We now define the Fortin operator as follows:

$$(3.12) \quad \Pi_h(v)|_K := \mathcal{C}_h^K(v) + \rho_h^K(v - \mathcal{C}_h^K(v)).$$

The proof of (3.1) follows by noting that, thanks to (3.3), for each $F \in \mathcal{E}_H \cap \partial K$ and any $\mu_H \in \Lambda_H$, the following holds

$$\begin{aligned}
\int_F \Pi_h(v) \mu_H ds &= \int_F \mathcal{C}_h^K(v) \mu_H ds + \int_F \rho_h^K(v - \mathcal{C}_h^K(v)) \mu_H ds \\
&= \int_F \mathcal{C}_h^K(v) \mu_H ds + \int_F (v - \mathcal{C}_h^K(v)) \mu_H ds \\
&= \int_F v \mu_H ds.
\end{aligned}$$

To prove (3.2) we use (3.10), (3.11), (1.17) and (1.18) and obtain

$$\begin{aligned}
\|\Pi_h(v)\|_{1,K}^2 &\leq C \left(\|\mathcal{C}_h^K(v)\|_{1,K}^2 + \|\rho_h^K(v - \mathcal{C}_h^K(v))\|_{1,K}^2 \right) \\
&\leq C \left(\|v\|_{1,K}^2 + \sum_{F \in \mathcal{E}_H \cap \partial K} h_{\mathfrak{T}_1}^{-2} \|v - \mathcal{C}_h^K(v)\|_{0,\kappa_F}^2 + |v - \mathcal{C}_h^K(v)|_{1,\kappa_F}^2 \right) \\
&\leq C \|v\|_{1,K}^2,
\end{aligned}$$

and the proof is finished adding over $K \in \mathcal{P}$. $\square$

We are ready to prove the well-posedness of the MHM method. This is stated in the next result.

Theorem 3.1. *Let us suppose that Assumption (A1) is satisfied. Then,*

a) *There exists $\beta_0 > 0$ such that, for all $v_0 \in V_0$:*

$$(3.13) \quad \sup_{\mu_H \in \Lambda_H} \frac{b(\mu_H, v_0)}{\|\mu_H\|_\Lambda} = \sup_{\mu_H \in \Lambda_H} \frac{\langle \mu_H, v_0 \rangle_{\partial \mathcal{D}}}{\|\mu_H\|_\Lambda} \geq \beta_0 \|v_0\|_V.$$

b) *There exists $\alpha_0 > 0$ such that*

$$(3.14) \quad a_h(\mu_H, \mu_H) \geq \alpha_0 \|\mu_H\|_\Lambda^2 \quad \text{for all } \mu_H \in \mathcal{N}_H,$$

where $\mathcal{N}_H$ is the discrete kernel of $b(\cdot, \cdot)$, that is,

$$(3.15) \quad \mathcal{N}_H := \{\mu_H \in \Lambda_H : b(\mu_H, v_0) = 0 \quad \text{for all } v_0 \in V_0\}.$$

Thus, (2.13) is well-posed.

Proof. We start proving (3.13). Over the partition Ξ_H we consider the space

$$(3.16) \quad X_H := \{\tau_H \in H(\text{div}, \Omega) : \tau_H|_\kappa \in RT_0(\kappa) \quad \text{for all } \kappa \in \Xi_H\}.$$

That is, the global Raviart-Thomas space of the lowest order defined in Ξ_H . Let now $v_0 \in V_0$. Then, there exists $\tilde{\tau}_H \in X_H$ such that $\nabla \cdot \tilde{\tau}_H = v_0$ in Ω and $\beta_0 \|\tilde{\tau}_H\|_{\text{div}, \Omega} \leq \|v_0\|_{0, \Omega} = \|v_0\|_V$, where β_0 does not depend on $\mathcal{H}$, H , or h . Thus

$$(3.17) \quad \beta_0 \|\tilde{\tau}_H\|_{\text{div}, \Omega} \|v_0\|_V \leq \int_\Omega \nabla \cdot \tilde{\tau}_H v_0 \, dx = \sum_{K \in \mathcal{D}} \langle \tilde{\tau}_H \cdot \mathbf{n}^K, v_0 \rangle_{\partial K},$$

where we have also used the fact that v_0 is continuous in every $K \in \mathcal{D}$ and the normal component of $\tilde{\tau}_H$ is continuous across every internal edge of Ξ . Thus, defining $\tilde{\mu}_H|_F = \tilde{\tau}_H \cdot \mathbf{n}^K|_F$ on every $F \in \mathcal{E}_H$, and using the definition of the norm in Λ , we arrive at

$$(3.18) \quad \beta_0 \|\tilde{\mu}_H\|_\Lambda \|v_0\|_V \leq \sum_{K \in \mathcal{D}} \langle \tilde{\tau}_H \cdot \mathbf{n}^K, v_0 \rangle_{\partial K} = \sum_{K \in \mathcal{D}} \langle \tilde{\mu}_H, v_0 \rangle_{\partial K} = b(\tilde{\mu}_H, v_0),$$

which proves (3.13).

To prove the ellipticity (3.14), let $\mu_H \in \mathcal{N}_H$. Then, for all $v_h \in V_h$ (not necessarily having zero mean value in K) the following equality holds

$$(3.19) \quad \int_K A \nabla T_h \mu_H \cdot \nabla v_h \, dx = -\langle \mu_H, v_h \rangle_{\partial K}.$$

Next, (1.14) and Lemma 3.1 imply the existence of a constant $C > 0$, independent of $h, H, \mathcal{H}$, and the shape of the elements in $\mathcal{D}$, such that

$$(3.20) \quad \begin{aligned} \frac{\sqrt{2}}{2} \|\mu_H\|_\Lambda &\leq \sup_{v \in V} \frac{\langle \mu_H, v \rangle_{\partial \mathcal{D}}}{\|v\|_V} \\ &\leq C \sup_{v \in V} \frac{\langle \mu_H, \Pi_h v \rangle_{\partial \mathcal{D}}}{\|\Pi_h v\|_V} \\ &\leq C \sup_{v_h \in V_h} \frac{\langle \mu_H, v_h \rangle_{\partial \mathcal{D}}}{\|v_h\|_V} \\ &= C \sup_{v_h \in \tilde{V}} \frac{-\sum_{K \in \mathcal{D}} \int_K A \nabla T_h \mu_H \cdot \nabla v_h \, dx}{\|v_h\|_V} \\ &\leq C \left\{ \sum_{K \in \mathcal{D}} A_{\max}^K \|\nabla T_h \mu_H\|_{0, K}^2 \right\}^{\frac{1}{2}}, \end{aligned}$$

where we also used (3.19). Hence, using the definition of $a_h(\cdot, \cdot)$ we arrive at

$$(3.21) \quad a_h(\mu_H, \mu_H) = \langle \mu_H, T_h \mu_H \rangle_{\partial \mathcal{P}} = \sum_{K \in \mathcal{P}} \int_K A \nabla T_h \mu_H \cdot \nabla T_h \mu_H \, dx \geq A_{\min} \|\nabla T_h \mu_H\|_{0, \mathcal{P}}^2 \geq C \omega^{-1} \|\mu_H\|_{\Lambda}^2,$$

which proves (3.14) with $\alpha_0 = C \omega$, and ω given in (1.3). This last result implies the stability and well-posedness of the discrete problem (2.13). $\square$

3.2. The $\mathbb{P}_{k+1}(F) \times \mathbb{P}_{k+1}(\mathfrak{T})$ element. Method (2.13) can also be implemented when the space of Lagrange multipliers Λ_H is built using polynomials of degree $\ell = k + 1, k \geq 0$. The main difference for this case resides on the minimal mesh. More precisely, we need to consider the following cases:

- For $k = 0, 1$: The mesh $\mathcal{T}_h^K$ needs to be the result of at least two red refinements of the mesh $\Xi_H(K)$.
- For $k \geq 2$: The mesh $\mathcal{T}_h^K$ needs to be the result of at least one red refinement of the mesh $\Xi_H(K)$.

We notice that for $k \geq 2$ the same situation as for the one treated in the last section is assumed, while for the lowest order case some extra mesh refinement is required to allow for the proof of stability. Then, with these considerations on the meshes used, for $k \geq 0$, we consider here the following finite element spaces

$$(3.22) \quad V_h := \prod_{K \in \mathcal{P}} V_h(K) \quad \text{where} \quad V_h(K) := \{v_h \in C^0(K) : v_h|_{\mathfrak{T}} \in \mathbb{P}_{k+1}(\mathfrak{T}), \forall \mathfrak{T} \in \mathcal{T}_h^K\},$$

$$(3.23) \quad \Lambda_H := \{\mu_H \in \Lambda : \mu_H|_F \in \mathbb{P}_{k+1}(F), \forall F \in \mathcal{E}_H\}.$$

A close inspection of the results proven in the last section shows that the only difference lies in the proof of existence of a Fortin operator, that is, Lemma 3.1. Its proof shows that the only difference that appears when considering $\ell = k + 1$ lies on the invertibility of the matrix $\mathbb{A}$ built in (3.5), which follows in this case from Lemma 6.2 (see the appendix for the proof), taking into consideration the minimal mesh requirement needed when $k = 0, 1$. Then, the stability of the discrete scheme can be proved using Lemma 3.1.

4. CONVERGENCE

The results in this section treat, in a unified manner, both cases $\ell = k$ and $\ell = k + 1$. We start by presenting an interpolation estimate for the space Λ_H . This estimate, in addition to extending [32, Lemma 9] to polygonal meshes, gives an error estimate depending on the parameter H , i.e. the discretisation parameter associated to the mesh $\mathcal{E}_H$, and not depending on $\mathcal{H}$, associated to the mesh $\mathcal{P}$.

Lemma 4.1. *Suppose $w \in H^{\ell+2}(\mathcal{P}) \cap H_0^1(\Omega)$, $A \nabla w \in H^{\ell+1}(\mathcal{P})$, with $\ell \geq 0$, and $A \nabla w \in H(\text{div}, \Omega)$. Let $\mu \in \Lambda$ be defined by $\mu|_E := (A \nabla w|_K \cdot \mathbf{n}^K)|_E$ for each $E \in \mathcal{E}$. Then, there exists a positive constant C , independent of h , H , $\mathcal{H}$ and A , such that*

$$(4.1) \quad \inf_{\mu_H \in \Lambda_H} \|\mu - \mu_H\|_{\Lambda} \leq C H^{\ell+1} |A \nabla w|_{\ell+1, \mathcal{P}},$$

where Λ_H is given in (1.15).

Proof. Let $w \in H^{\ell+2}(\mathcal{P})$ and $E \in \mathcal{E}$. Let $K \in \mathcal{P}$ be such that $F \subset \partial K$. We define $\chi_E := A \nabla w \cdot \mathbf{n}_E \in H^{\ell+1}(\mathcal{P})$ where $\mathbf{n}_E$ must be understood as the trivial extension of the normal vector $\mathbf{n}^K|_E$ to E to a constant function in the whole of K , with $|\mathbf{n}_E| = 1$. Observe that $\mu := \chi_E|_F \in L^2(F)$ for each $F \subset E \in \mathcal{E}$.

Now, let us recall that κ_F denotes the unique element in $\Xi_H(K)$ such that $\kappa_F \cap \partial K = F$. Let $\hat{\mathfrak{T}}$ be the standard reference element with vertices $(0,0)$, $(1,0)$ and $(0,1)$, and let $\mathcal{F}_F : \hat{\mathfrak{T}} \rightarrow \kappa_F$ be the invertible affine transformation such that $\mathcal{F}_F(\hat{F}) = F$, where $\hat{F} = [0,1]$. First, we observe that

$$(4.2) \quad \widehat{\Pi_F^\ell \mu} = \Pi_{\hat{F}}^\ell \hat{\mu}.$$

So

$$(4.3) \quad \begin{aligned} \int_F (\mu - \Pi_F^\ell \mu) v \, ds &= H_F \int_{\hat{F}} (\hat{\mu} - \widehat{\Pi_F^\ell \mu}) \hat{v} \, d\hat{s} \\ &= H_F \int_{\hat{F}} (\hat{\mu} - \Pi_{\hat{F}}^\ell \hat{\mu}) \hat{v} \, d\hat{s}. \end{aligned}$$

In addition, a scaling argument (see, e.g., [17]) and the mesh regularity of $\Xi_H(K)$ give

$$(4.4) \quad |\hat{z}|_{\ell+1, \hat{\mathfrak{T}}} \leq C H_F^\ell |z|_{\ell+1, \kappa_F} \quad \text{and} \quad |\hat{v}|_{1, \hat{\mathfrak{T}}} \leq C |v|_{1, \kappa_F},$$

for all $z \in H^{\ell+1}(\kappa_F)$ and $v \in H^1(\kappa_F)$. Now, using [13, Lemma 3] and (4.4), we get, for all $v \in V$,

$$(4.5) \quad \begin{aligned} \int_{\hat{F}} (\hat{\mu} - \Pi_{\hat{F}}^\ell \hat{\mu}) \hat{v} \, d\hat{s} &\leq C |\hat{\chi}_E|_{\ell+1, \hat{\mathfrak{T}}} |\hat{v}|_{1, \hat{\mathfrak{T}}} \\ &\leq C H_F^\ell |\chi_E|_{\ell+1, \kappa_F} |v|_{1, \kappa_F}, \end{aligned}$$

where C is a positive constant that only depends on the shape of κ_F .

Next, we define $\tilde{\mu}_H := \Pi_F^\ell \mu$, and in each $F \subset \partial K$ and $K \in \mathcal{P}$. A change of variables, (4.5), (4.4), the regularity of $\Xi_H(K)$, and the fact that the elements κ_F do not overlap give

$$\begin{aligned} \langle \mu - \tilde{\mu}_H, v \rangle_{\partial K} &= \sum_{F \subset \partial K} \int_F (\mu - \tilde{\mu}_H) v \, ds = \sum_{F \subset \partial K} H_F \int_{\hat{F}} (\hat{\mu} - \widehat{\tilde{\mu}_H}) \hat{v} \, d\hat{s} \\ &\leq C \sum_{F \subset \partial K} H_F^{\ell+1} |\chi_E|_{\ell+1, \kappa_F} |v|_{1, \kappa_F} \\ &\leq C H^{\ell+1} \left\{ \sum_{F \subset \partial K} |\chi_E|_{\ell+1, \kappa_F}^2 \right\}^{\frac{1}{2}} \left\{ \sum_{F \subset \partial K} |v|_{1, \kappa_F}^2 \right\}^{\frac{1}{2}} \\ &\leq C H^{\ell+1} |A \nabla w|_{\ell+1, K} |v|_{1, K}, \end{aligned}$$

for all $v \in V$, where we used $|\chi_E|_{\ell+1, \kappa_F} = |A \nabla w \cdot \mathbf{n}_E|_{\ell+1, \kappa_F} \leq |A \nabla w|_{\ell+1, \kappa_F}$.

Finally, adding over $K \in \mathcal{P}$, and collecting the above results, we arrive at

$$b(\mu - \tilde{\mu}_H, v) = \langle \mu - \tilde{\mu}_H, v \rangle_{\partial \mathcal{P}} \leq C H^{\ell+1} |A \nabla w|_{\ell+1, \mathcal{P}} |v|_{1, \mathcal{P}} \leq C H^{\ell+1} |A \nabla w|_{\ell+1, \mathcal{P}} \|v\|_V,$$

which immediately leads to

$$\sup_{v \in V} \frac{b(\mu - \tilde{\mu}_H, v)}{\|v\|_V} \leq C H^{\ell+1} |A \nabla w|_{\ell+1, \mathcal{P}},$$

and (4.1) follows from (1.14). $\square$

We are ready to present the main convergence result for the present method. From now on, we will assume that the solution u of (1.1) is such that all the norms on the right-hand side of the estimates are finite.

Theorem 4.1. *There exists $C > 0$, independent of h, H , and $\mathcal{H}$, such that*

$$(4.6) \quad \|u_0 - u_0^h\|_V + \|\lambda - \lambda_H\|_\Lambda \leq C \left(h^{k+1} |u|_{k+2, \mathcal{P}} + H^{\ell+1} |A\nabla u|_{\ell+1, \mathcal{P}} \right).$$

In addition, if we denote $u_h := u_0^h + T_h \lambda_H + \hat{T}_h f$, then the following error estimate holds

$$(4.7) \quad \|u - u_h\|_V \leq C \left(h^{k+1} |u|_{k+2, \mathcal{P}} + H^{\ell+1} |A\nabla u|_{\ell+1, \mathcal{P}} \right).$$

Proof. Let $\Lambda_H^* \subset \Lambda_H$ be defined by

$$\Lambda_H^* := \left\{ \mu_H \in \Lambda_H : \int_{\partial K} \mu_H = \int_K f \text{ for all } K \in \mathcal{P} \right\},$$

and let $\mu_H \in \Lambda_H^*$ be arbitrary. Observing that $\lambda_H - \mu_H \in \mathcal{N}_H$, we get from Theorem 3.1-(b), (2.4) and (2.13) that

$$\begin{aligned} \alpha_0 \|\lambda_H - \mu_H\|_\Lambda^2 &\leq a_h(\lambda_H - \mu_H, \lambda_H - \mu_H) \\ &= a_h(\lambda_H, \lambda_H - \mu_H) - a(\lambda, \lambda_H - \mu_H) + a_h(\lambda - \mu_H, \lambda_H - \mu_H) \\ &\quad + a(\lambda, \lambda_H - \mu_H) - a_h(\lambda, \lambda_H - \mu_H) \\ &\leq C \left(\|\lambda - \mu_H\|_\Lambda + \|(T - T_h)\lambda + (\hat{T} - \hat{T}_h)f\|_V \right) \|\lambda_H - \mu_H\|_\Lambda, \end{aligned}$$

where we used the stability of T_h given in (2.16). Using that $u = u_0 + T\lambda + \hat{T}f \in H^{k+2}(\mathcal{P})$ and (2.15) we get

$$(4.8) \quad \|T\lambda + \hat{T}f - (T_h\lambda + \hat{T}_h f)\|_V \leq Ch^{k+1} |u|_{k+2, \mathcal{P}},$$

which gives

$$\begin{aligned} \|\lambda_H - \mu_H\|_\Lambda &\leq C \left(\|\lambda - \mu_H\|_\Lambda + \|(T - T_h)\lambda + (\hat{T} - \hat{T}_h)f\|_V \right) \\ &\leq C \left(\|\lambda - \mu_H\|_\Lambda + h^{k+1} |u|_{k+2, \mathcal{P}} \right). \end{aligned}$$

Next, we observe that the second equation of (2.4) gives

$$(4.9) \quad \int_{\partial K} \Pi^\ell \lambda v_0 dx = \sum_{F \subset \partial K} \int_F \Pi^\ell \lambda v_0 dx = \sum_{F \subset \partial K} \int_F \lambda v_0 ds = \int_K f v_0 dx \quad \text{for all } v_0 \in V_0,$$

and then $\Pi^\ell \lambda \in \Lambda_H^*$. Thus, setting $\mu_H = \Pi^\ell \lambda$ and using Lemma 4.1 we get

$$\|\lambda_H - \mu_H\|_\Lambda \leq C \left(h^{k+1} |u|_{k+2, \mathcal{P}} + H^{\ell+1} |A\nabla u|_{\ell+1, \mathcal{P}} \right).$$

Finally, the triangle inequality and Lemma 4.1 lead to

$$(4.10) \quad \|\lambda - \lambda_H\|_\Lambda \leq C \left(h^{k+1} |u|_{k+2, \mathcal{P}} + H^{\ell+1} |A\nabla u|_{\ell+1, \mathcal{P}} \right).$$

From Theorem 3.1 item (a), there exists $\xi_H \in \Lambda_H$, with $\|\xi_H\|_\Lambda = 1$, such that

$$\begin{aligned} \beta_0 \|u_0^h - u_0\|_V &\leq b(u_0^h - u_0, \xi_H) \\ &= -a_h(\lambda_H, \xi_H) + a(\lambda, \xi_H) + \langle \xi_H, (\hat{T}_h - \hat{T})f \rangle_{\partial \mathcal{P}} \\ &= -a_h(\lambda_H - \lambda, \xi_H) + a(\lambda, \xi_H) - a_h(\lambda, \xi_H) + \langle \xi_H, (\hat{T}_h - \hat{T})f \rangle_{\partial \mathcal{P}} \\ &\leq C \left(\|\lambda_H - \lambda\|_\Lambda + \|(T_h - T)\lambda + (\hat{T}_h - \hat{T})f\|_V \right), \end{aligned}$$

where we used (2.16), (2.4) and (2.13) once more. From (4.10) and (4.8), the following estimate holds

$$(4.11) \quad \|u_0^h - u_0\|_V \leq C \left(h^{k+1} |u|_{k+2, \mathcal{P}} + H^{\ell+1} |A \nabla u|_{\ell+1, \mathcal{P}} \right),$$

and adding (4.10) and (4.11) the result (4.6) follows. To prove (4.7), we see that

$$\begin{aligned} \|u - u_h\|_V &\leq \|u_0 - u_0^h\|_V + \|T\lambda + \hat{T}f - (T_h\lambda_H + \hat{T}_h f)\|_V \\ &\leq \|u_0 - u_0^h\|_V + \|T_h(\lambda - \lambda_H)\|_V + \|T\lambda + \hat{T}f - (T_h\lambda + \hat{T}_h f)\|_V, \end{aligned}$$

and the proof is finished applying (2.16), (4.6) and (4.8). $\square$

4.1. An error estimate for $\|\sigma - \sigma_h\|_{div, \Omega}$. In addition to the convergence result just proved, the following result states that the MHM method also produces an accurate discrete solution of the dual variable $\sigma := A \nabla u$ in the following $H(div, \mathcal{T}_h)$ norm

$$(4.12) \quad \|\tau\|_{div, \mathcal{T}_h}^2 := \sum_{K \in \mathcal{P}} \sum_{\mathfrak{T} \in \mathcal{T}_h^K} \|\tau\|_{div, \mathfrak{T}}^2 \quad \text{for all } \tau \text{ such that } \tau|_{\mathfrak{T}} \in H(div, \mathfrak{T}).$$

In the proof of the result below we will use, for every $K \in \mathcal{P}$, the following space

$$(4.13) \quad \mathbf{W}_h(K) := \{\tau_h \in H(div, \Omega) : \tau_h|_{\mathfrak{T}} \in RT_k(\mathfrak{T}), \forall \mathfrak{T} \in \mathcal{T}_h^K\},$$

where $RT_k(\mathfrak{T})$ stands for the local Raviart-Thomas space of order k in T .

Theorem 4.2. *Let us assume that $f \in H^{k+1}(\mathcal{T}_h)$, that the families of partitions $\{\mathcal{T}_h^K\}_{h>0}$ are quasi-uniform, and that, for all $K \in \mathcal{P}$ and all $F \in \mathcal{E}_H \cap \partial K$,*

$$H \leq C h_{min} := \min\{h_{\mathfrak{T}} : \mathfrak{T} \in \mathcal{T}_h^K, K \in \mathcal{P}\},$$

for some $C > 0$. Denoting $\sigma_h := A \nabla (T_h \lambda_H + \hat{T}_h f)$, there exists $C > 0$, independent of any mesh size, and the shape of K , such that

$$(4.14) \quad \|\sigma - \sigma_h\|_{div, \mathcal{T}_h} \leq C \left((H^\ell + h^k) \|u\|_{k+2, \mathcal{P}} + h^{k+1} \|f\|_{k+1, \mathcal{T}_h} \right).$$

Proof. Since $\|\sigma - \sigma_h\|_{0, \Omega}$ has been bounded in Theorem 4.1, we only bound the difference $\nabla \cdot \sigma - \nabla \cdot \sigma_h$. Let $\bar{\sigma}_h|_K := \mathcal{R}_K(\sigma)$, where $\mathcal{R}_K$ stands for the Raviart-Thomas interpolation operator with values in $\mathbf{W}_h(K)$. Using well-known properties of $\mathcal{R}_K$ (see, e.g., [17]) we get, for every $\mathfrak{T} \in \mathcal{T}_h^K$ and every $K \in \mathcal{P}$, that

$$(4.15) \quad \begin{aligned} \|\sigma - \bar{\sigma}_h\|_{0, \mathfrak{T}} &\leq C h_{\mathfrak{T}}^{k+1} \|\sigma\|_{k+1, \mathfrak{T}}, \\ \|\nabla \cdot \sigma - \nabla \cdot \bar{\sigma}_h\|_{0, \mathfrak{T}} &\leq C h_{\mathfrak{T}}^{k+1} |\nabla \cdot \sigma|_{k+1, \mathfrak{T}} = C h_{\mathfrak{T}}^{k+1} |f|_{k+1, \mathfrak{T}}. \end{aligned}$$

Thus, using the triangle inequality, (4.15), an inverse inequality in each $\mathfrak{T}$, and Theorem 4.1, we arrive at

$$\begin{aligned}
\|\nabla \cdot \boldsymbol{\sigma} - \nabla \cdot \boldsymbol{\sigma}_h\|_{0,\mathcal{T}_h}^2 &\leq 2 \sum_{K \in \mathcal{P}} \sum_{\mathfrak{T} \in \mathcal{T}_h^K} \left(\|\nabla \cdot \boldsymbol{\sigma} - \nabla \cdot \bar{\boldsymbol{\sigma}}_h\|_{0,\mathfrak{T}}^2 + \|\nabla \cdot \bar{\boldsymbol{\sigma}}_h - \nabla \cdot \boldsymbol{\sigma}_h\|_{0,\mathfrak{T}}^2 \right) \\
&\leq C \sum_{K \in \mathcal{P}} \sum_{\mathfrak{T} \in \mathcal{T}_h^K} \left(h_{\mathfrak{T}}^{2k+2} |f|_{k+1,\mathfrak{T}}^2 + h_{\mathfrak{T}}^{-2} \|\bar{\boldsymbol{\sigma}}_h - \boldsymbol{\sigma}_h\|_{0,\mathfrak{T}}^2 \right) \\
&\leq Ch^{2k+2} |f|_{k+1,\mathcal{T}_h}^2 + Ch_{\min}^{-2} \left(\|\boldsymbol{\sigma}_h - \boldsymbol{\sigma}\|_{0,\Omega}^2 + \|\boldsymbol{\sigma} - \bar{\boldsymbol{\sigma}}_h\|_{0,\Omega}^2 \right) \\
&\leq Ch^{2k+2} |f|_{k+1,\mathcal{T}_h}^2 + Ch_{\min}^{-2} \left([h^{2k+2} |u|_{k+2,\mathcal{P}} + H^{2\ell+2} |A \nabla u|_{\ell+1,\mathcal{P}}^2] + h^{2k+2} |\boldsymbol{\sigma}|_{k+1,\mathcal{P}}^2 \right) \\
(4.16) \quad &\leq C \left(h^{2k+2} |f|_{k+1,\mathcal{T}_h}^2 + h^{2k} |u|_{k+2,\mathcal{P}} + H^{2\ell} |A \nabla u|_{\ell+1,\mathcal{P}}^2 + h^{2k} |\boldsymbol{\sigma}|_{k+1,\mathcal{P}}^2 \right),
\end{aligned}$$

which finishes the proof. $\square$

Remark 4.2. *If the second level problems are written in mixed form, then they can alternatively be solved using a mixed inf-sup stable method. As a result, $\nabla \cdot \boldsymbol{\sigma}_h - \nabla \cdot \bar{\boldsymbol{\sigma}}_h = 0$ in every $T \in \mathcal{T}_h^K$, and then the following improved error estimate can be obtained*

$$(4.17) \quad \|\boldsymbol{\sigma} - \boldsymbol{\sigma}_h\|_{div,\Omega} \leq C \left(h^{k+1} \|f\|_{k+1,\mathcal{P}} + (H^{\ell+1} + h^{k+1}) \|u\|_{k+2,\mathcal{P}} \right),$$

requiring only the mesh regularity. In addition, it is worth mentioning that in the above case, $\boldsymbol{\sigma}_h \in H(\text{div}, \Omega)$. This idea was explored for simplicial elements in [14].

4.2. Error estimates for $\|u - u_h\|_{0,\Omega}$. In order to prove a higher order of convergence in the $L^2(\Omega)$ norm for u , we first make the following remark: if every $K \in \mathcal{P}$ is convex and $A \in W^{1,\infty}(K)^{2 \times 2}$, then the solution of

$$(4.18) \quad -\nabla \cdot (A \nabla \phi) = g \quad \text{in } K, \quad A \nabla \phi \cdot \mathbf{n}^K = 0 \quad \text{on } \partial K,$$

with $\int_K g \, dx = 0$, belongs to $H^2(K)$ and satisfies

$$(4.19) \quad \|\phi\|_{2,K} \leq C \|g\|_{0,K},$$

where C depends only on $A_{\min}$ and $\|A\|_{1,\infty,K}$ (see [22]). In particular, this implies that there exists $C > 0$, independent of the shape and size of K , such that

$$(4.20) \quad \|\hat{T}q\|_{2,K} \leq C \|q\|_{0,K},$$

for all $q \in L^2(\mathcal{P})$. As a consequence, arguing as in Remark 2.2, the application of Aubin-Nitsche's Lemma gives

$$(4.21) \quad \|(T - T_h)\mu + (\hat{T} - \hat{T}_h)q\|_{0,K} \leq Ch \|\nabla((T - T_h)\mu + (\hat{T} - \hat{T}_h)q)\|_{0,K},$$

for all $(\mu, q) \in \Lambda \times L^2(\mathcal{P})$, where $C > 0$ is independent of $h, H, \mathcal{H}$, and the shape of K . With these ingredients we now present a first error estimate in the $L^2(\Omega)$ norm for which we impose the hypothesis of convexity of both Ω and $K \in \mathcal{P}$.

Theorem 4.3. *Let us assume that Ω , and every $K \in \mathcal{P}$, are convex. Then, there exists $C > 0$, independent of h, H , and $\mathcal{H}$, such that*

$$(4.22) \quad \|u - u_h\|_{0,\Omega} \leq C \left(h^{k+2} + H^{\ell+2} \right) \|u\|_{k+2,\mathcal{P}}.$$

Proof. First, we observe that (2.18) also holds if one replaces T_h and $\hat{T}_h$ by T and $\hat{T}$, respectively. Next, define $e = u - u_h$ and let w satisfy the following elliptic problem

$$(4.23) \quad -\nabla \cdot (A \nabla w) = e \quad \text{in } \Omega, \quad \text{and} \quad w = 0 \quad \text{on } \partial\Omega.$$

Problem (4.23) is well-posed, w belongs to $H_0^1(\Omega) \cap H^2(\Omega)$ and satisfies the following bound

$$(4.24) \quad \|w\|_{2,\mathcal{D}} \leq C \|e\|_{0,\Omega}.$$

As it was done in Section 2.1, $w = w_0 + T\gamma + \hat{T}e$, where $(\gamma, w_0) \in \Lambda \times V_0$ satisfies

$$(4.25) \quad \begin{aligned} a(\gamma, \mu) + b(\mu, w_0) &= \int_{\Omega} T\mu e \, dx \quad \text{for all } \mu \in \Lambda, \\ b(\gamma, v_0) &= \int_{\Omega} e v_0 \, dx \quad \text{for all } v_0 \in V_0. \end{aligned}$$

Now, let $(\gamma_H, w_0^h) \in \Lambda_H \times V_0$ be the following discrete approximation of (4.25), i.e.,

$$\begin{aligned} a_h(\gamma_H, \mu_H) + b(\mu_H, w_0^h) &= \int_{\Omega} T_h \mu_H e \, dx \quad \text{for all } \mu_H \in \Lambda_H, \\ b(\gamma_H, v_0) &= \int_{\Omega} e v_0 \, dx \quad \text{for all } v_0 \in V_0. \end{aligned}$$

Using Theorem 4.1 and (4.24) we obtain the following error estimate

$$(4.26) \quad \|\gamma - \gamma_H\|_{\Lambda} + \|w_0 - w_0^h\|_V \leq C(h + H) |w|_{2,\mathcal{D}} \leq C(h + H) \|e\|_{0,\Omega}.$$

From the definitions of the second-level local solvers T and T_h , the following Galerkin orthogonality property holds

$$(4.27) \quad \int_K A \nabla(T_h \lambda) \cdot \nabla(T_h \gamma_H) \, dx = -\langle \lambda, T_h \gamma_H \rangle_{\partial K} = \int_K A \nabla(T \lambda) \cdot \nabla(T_h \gamma_H) \, dx,$$

similarly, from the local solvers $\hat{T}$ and $\hat{T}_h$, we get

$$(4.28) \quad \int_K A \nabla(\hat{T}_h f) \cdot \nabla(T_h \gamma_H) \, dx = (f, T_h \gamma_H)_K = \int_K A \nabla(\hat{T} f) \cdot \nabla(T_h \gamma_H) \, dx.$$

Then, using (4.25) we arrive at

$$\begin{aligned} \|e\|_{0,\Omega}^2 &= (u - u_h, e)_{\mathcal{D}} \\ &= (u_0 - u_0^h + T\lambda - T_h \lambda_H, e)_{\mathcal{D}} + (\hat{T}f - \hat{T}_h f, e)_{\mathcal{D}} \\ &= b(\gamma_H, u_0 - u_0^h) + (T\lambda - T_h \lambda_H, e)_{\mathcal{D}} + (T\lambda_H - T_h \lambda_H + \hat{T}f - \hat{T}_h f, e)_{\mathcal{D}} \\ &= b(\gamma_H, u_0 - u_0^h) + a(\lambda - \lambda_H, \gamma) + b(\lambda - \lambda_H, w_0) + (T\lambda_H - T_h \lambda_H + \hat{T}f - \hat{T}_h f, e)_{\mathcal{D}} \\ &= \underbrace{b(\gamma_H, u_0 - u_0^h) + a(\lambda - \lambda_H, \gamma)}_{(I)} + \underbrace{(T\lambda_H - T_h \lambda_H + \hat{T}f - \hat{T}_h f, e)_{\mathcal{D}}}_{(II)} \end{aligned}$$

where we used $b(\lambda - \lambda_H, w_0) = 0$. Let us estimate (I). Since $\gamma_H \in \Lambda_H$, from (2.4) and (2.13), and Cauchy-Schwarz's inequality it holds

$$\begin{aligned} b(\gamma_H, u_0 - u_0^h) + a(\lambda - \lambda_H, \gamma) &= -a(\gamma_H, \lambda) + a_h(\gamma_H, \lambda_H) - \langle \gamma_H, \hat{T}f - \hat{T}_h f \rangle_{\partial\mathcal{D}} + a(\lambda - \lambda_H, \gamma) \\ &= a(\gamma - \gamma_H, \lambda - \lambda_H) + a_h(\gamma_H, \lambda_H) - a(\gamma_H, \lambda_H) - \langle \gamma_H, \hat{T}f - \hat{T}_h f \rangle_{\partial\mathcal{D}} \\ &\leq \underbrace{C \|\gamma - \gamma_H\|_{\Lambda} \|\lambda - \lambda_H\|_{\Lambda}}_{(I1)} + \underbrace{(\nabla(T - T_h)\lambda_H + \nabla(\hat{T} - \hat{T}_h)f, A \nabla T \gamma_H)_{\mathcal{D}}}_{(I2)}, \end{aligned}$$

where we used (2.1)-(2.2) and symmetry of $a(\cdot, \cdot)$ and $a_h(\cdot, \cdot)$ in the last inequality. The first term (I1) is bounded using (4.26) and Theorem 4.1 as follows

(4.29)

$$\|\gamma - \gamma_H\|_\Lambda \|\lambda - \lambda_H\|_\Lambda \leq C(h + H) \|e\|_{0,\Omega} \|\lambda - \lambda_H\|_\Lambda \leq C \left(h^{k+2} + H^{\ell+2} \right) \|u\|_{k+2,\mathcal{D}} \|e\|_{0,\Omega}.$$

Next, using (4.27)-(4.28), Cauchy-Schwarz inequality and the stability of T and T_h in (2.16), term (I2) reads

$$\begin{aligned} & (\nabla(T - T_h)\lambda_H + \nabla(\hat{T} - \hat{T}_h)f, A \nabla T \gamma_H)_{\mathcal{D}} \\ & \leq (A \nabla(T - T_h)\lambda + A \nabla(\hat{T} - \hat{T}_h)f, \nabla(T - T_h)\gamma_H)_{\mathcal{D}} + (A \nabla(T - T_h)(\lambda_H - \lambda), \nabla(T - T_h)\gamma_H)_{\mathcal{D}} \\ & \leq C \left(\|\nabla(T - T_h)\lambda + \nabla(\hat{T} - \hat{T}_h)f\|_{0,\mathcal{D}} + \|\nabla(T - T_h)(\lambda_H - \lambda)\|_{0,\mathcal{D}} \right) \|\nabla(T - T_h)\gamma_H\|_{0,\mathcal{D}} \\ & \leq C \left(\|\nabla(T - T_h)\lambda + \nabla(\hat{T} - \hat{T}_h)f\|_{0,\mathcal{D}} + \|\lambda_H - \lambda\|_\Lambda \right) \|\nabla(T - T_h)\gamma_H\|_{0,\mathcal{D}} \\ & \leq C \left(\|\nabla(T - T_h)\lambda + \nabla(\hat{T} - \hat{T}_h)f\|_{0,\mathcal{D}} + \|\lambda_H - \lambda\|_\Lambda \right) \left(\|\nabla(T - T_h)(\gamma_H - \gamma)\|_{0,\mathcal{D}} + \|\nabla(T - T_h)\gamma\|_{0,\mathcal{D}} \right) \\ & \leq C \left(\|\nabla(T - T_h)\lambda + \nabla(\hat{T} - \hat{T}_h)f\|_{0,\mathcal{D}} + \|\lambda_H - \lambda\|_\Lambda \right) \left(\|\gamma_H - \gamma\|_\Lambda + \|\nabla(T - T_h)\gamma\|_{0,\mathcal{D}} \right) \\ & \leq C \left(\|\nabla(T - T_h)\lambda + \nabla(\hat{T} - \hat{T}_h)f\|_{0,\mathcal{D}} + \|\lambda_H - \lambda\|_\Lambda \right) \left(\|\gamma_H - \gamma\|_\Lambda + \|\nabla(T - T_h)\gamma + \nabla(\hat{T} - \hat{T}_h)e\|_{0,\mathcal{D}} \right. \\ & \quad \left. + \|\nabla(\hat{T} - \hat{T}_h)e\|_{0,\mathcal{D}} \right). \end{aligned}$$

Now, (2.15) and (4.24) give

$$(4.30) \quad \|\nabla((T - T_h)\gamma + (\hat{T} - \hat{T}_h)e)\|_{0,\mathcal{D}} \leq C h |w|_{2,\mathcal{D}} \leq C h \|e\|_{0,\Omega}.$$

In addition, standard Galerkin error estimates over each $K \in \mathcal{D}$ and (4.20) lead to

$$(4.31) \quad \|\nabla(\hat{T} - \hat{T}_h)e\|_{0,\mathcal{D}} \leq C h |\hat{T}e|_{2,\mathcal{D}} \leq C h \|e\|_{0,\Omega}.$$

Thus, thanks to (2.15), Theorem 4.1, (4.26), (4.30), and (4.31) we get the following bound for (I2)

$$\begin{aligned} (\nabla(T - T_h)\lambda_H + \nabla(\hat{T} - \hat{T}_h)f, A \nabla T \gamma_H)_{\mathcal{D}} & \leq C \left(h^{k+1} + H^{\ell+1} \right) \|u\|_{k+2,\mathcal{D}} \left((h + H) \|e\|_{0,\Omega} + h \|e\|_{0,\Omega} \right) \\ & \leq C \left(h^{k+2} + H^{\ell+2} \right) \|u\|_{k+2,\mathcal{D}} \|e\|_{0,\Omega}. \end{aligned}$$

We next bound (II). Using (4.21) and the stability of T and T_h given in (2.16), we get

$$(4.32) \quad \|(T - T_h)(\lambda_H - \lambda)\|_{0,\Omega} \leq C h \|\nabla((T - T_h)(\lambda_H - \lambda))\|_{0,\mathcal{D}} \leq C h \|\lambda_H - \lambda\|_\Lambda,$$

and similarly using (4.21) we get the bound

$$\begin{aligned} (4.33) \quad \|(T - T_h)\lambda + (\hat{T} - \hat{T}_h)f\|_{0,\Omega} & \leq C h \|\nabla((T - T_h)\lambda + (\hat{T} - \hat{T}_h)f)\|_{0,\mathcal{D}} \\ & \leq C h^{k+2} \|T\lambda + \hat{T}f\|_{k+2,\mathcal{D}} \\ & \leq C h^{k+2} \|u\|_{k+2,\mathcal{D}}. \end{aligned}$$

So, (4.32), (4.33), and Theorem 4.1 give

$$\begin{aligned}
(T\lambda_H - T_h\lambda_H + \hat{T}f - \hat{T}_hf, e)_{\mathcal{P}} &\leq \|T\lambda_H - T_h\lambda_H + \hat{T}f - \hat{T}_hf\|_{0,\Omega} \|e\|_{0,\Omega} \\
&\leq \left(\|(T - T_h)(\lambda_H - \lambda)\|_{0,\Omega} + \|(T - T_h)\lambda + (\hat{T} - \hat{T}_h)f\|_{0,\Omega} \right) \|e\|_{0,\Omega} \\
&\leq C h \left(\|\lambda_H - \lambda\|_{\Lambda} + h^{k+1} \|u\|_{k+2,\mathcal{P}} \right) \|e\|_{0,\Omega} \\
&\leq C h \left(\left(H^{\ell+1} + h^{k+1} \right) \|u\|_{k+2,\mathcal{P}} + h^{k+1} \|u\|_{k+2,\mathcal{P}} \right) \|e\|_{0,\Omega} \\
&\leq C \left(H^{\ell+2} + h^{k+2} \right) \|u\|_{k+2,\mathcal{P}} \|e\|_{0,\Omega}.
\end{aligned}$$

Finally, the desired result (4.22) follows gathering all previous estimates and adding up (I) and (II). $\square$

In all the results proven so far, the diameter of the elements $K \in \mathcal{P}$ does not need to tend to zero to obtain optimal order error estimates. On the other hand, in the proof of the last result the convexity of the elements K was necessary. If the hypothesis of convexity is relaxed a similar result can be proven, now under the assumption that the elements K do shrink in size, and the boundness of the Poincaré constant holds in the sense of [34]. This is stated in the next result.

Theorem 4.4. *Assume that Ω is convex. Then, there exists $C > 0$, independent of h , H and $\mathcal{H}$, such that*

$$(4.34) \quad \|u - u_h\|_{0,\Omega} \leq C \mathcal{H} \left(h^{k+1} + H^{\ell+1} \right) \|u\|_{k+2,\mathcal{P}}.$$

Proof. First, observe that the convexity of $K \in \mathcal{P}$ is employed in the proof of Theorem 4.3 to obtain the estimates (4.31), (4.32) and (4.33). Thereby, we can avoid such an assumption and still derive a similar estimate by using the following generalised Poincaré inequality (see [34])

$$(4.35) \quad \|v\|_{0,K} \leq C \mathcal{H}_K \|\nabla v\|_{0,K} \quad \text{for all } v \in H^1(K) \cap L_0^2(K),$$

where C is independent of $\mathcal{H}_K$.

To obtain an analogue of (4.31), we use that $\hat{T}e|_K \in L_0^2(K)$ and (2.2), and we get the following bound

$$\begin{aligned}
\|\nabla \hat{T}e\|_{0,\mathcal{P}}^2 &\leq C \|A^{1/2} \nabla \hat{T}e\|_{0,\mathcal{P}}^2 \\
&= C (A \nabla \hat{T}e, \nabla \hat{T}e)_{\mathcal{P}} \\
&= C (e, \hat{T}e)_{\mathcal{P}} \\
&\leq C \sum_{K \in \mathcal{P}} \|e\|_{0,K} \|\hat{T}e\|_{0,K} \\
&\leq C \sum_{K \in \mathcal{P}} \mathcal{H}_K \|e\|_{0,K} \|\nabla \hat{T}e\|_{0,K} \\
&\leq C \mathcal{H} \|e\|_{0,\Omega} \|\nabla \hat{T}e\|_{0,\mathcal{P}},
\end{aligned}$$

where we used the generalised Poincaré inequality (4.35), and then

$$\|\nabla \hat{T}e\|_{0,\mathcal{P}} \leq C \mathcal{H} \|e\|_{0,\mathcal{P}}.$$

Following analogous steps, we get

$$\|\nabla \hat{T}_h e\|_{0,\mathcal{P}} \leq C \mathcal{H} \|e\|_{0,\mathcal{P}},$$

and, thus, from the triangle inequality it holds

$$(4.36) \quad \|\nabla(\hat{T} - \hat{T}_h)e\|_{0,\mathcal{P}} \leq C \mathcal{H} \|e\|_{0,\mathcal{P}}.$$

Proceeding in a very similar manner, we derive the following estimates mimicking (4.32) and (4.33)

$$(4.37) \quad \|(T - T_h)(\lambda_H - \lambda)\|_{0,\Omega} \leq C \mathcal{H} \|\lambda_H - \lambda\|_{\Lambda},$$

and

$$(4.38) \quad \|(T - T_h)\lambda + (\hat{T} - \hat{T}_h)f\|_{0,\Omega} \leq C \mathcal{H} h^{k+1} \|u\|_{k+2,\mathcal{P}}.$$

Thus, following otherwise the same steps as in the proof of the last theorem, we arrive that

$$\begin{aligned} (\nabla(T - T_h)\lambda_H + \nabla(\hat{T} - \hat{T}_h)f, A \nabla T \gamma_H)_{\mathcal{P}} &\leq C \left(h^{k+1} + H^{\ell+1}\right) \|u\|_{k+2,\mathcal{P}} \left(h + H + \mathcal{H}\right) \|e\|_{0,\Omega}, \\ &\leq C \mathcal{H} \left(h^{k+1} + H^{\ell+1}\right) \|u\|_{k+2,\mathcal{P}} \|e\|_{0,\Omega}, \end{aligned}$$

and

$$(T\lambda_H - T_h\lambda_H + \hat{T}f - \hat{T}_hf, e)_{\mathcal{P}} \leq C \mathcal{H} \left(h^{k+1} + H^{\ell+1}\right) \|u\|_{k+2,\mathcal{P}} \|e\|_{0,\Omega}.$$

Thus, reproducing the exact same steps of the proof of Theorem 4.3, we get

$$\|e\|_{0,\Omega}^2 \leq C \mathcal{H} \left(h^{k+1} + H^{\ell+1}\right) \|u\|_{k+2,\mathcal{P}} \|e\|_{0,\Omega},$$

and the proof is finished dividing by $\|e\|_{0,\Omega}$. $\square$

Remark 4.3. We finish this section by summarizing in Table 1 the approximation error for the primal (u) and dual (σ) variables in the $L^2(\Omega)$, V , and $H(\text{div}, \mathcal{T}_h)$ norms, when convex and non-convex elements are used.

	Convex K			Non-Convex K		
$k \geq 0$	$\ \cdot\ _{0,\Omega}$	$\ \cdot\ _V$	$\ \cdot\ _{\text{div}, \mathcal{T}_h}$	$\ \cdot\ _{0,\Omega}$	$\ \cdot\ _V$	$\ \cdot\ _{\text{div}, \mathcal{T}_h}$
$\mathbb{P}_k(F) \times \mathbb{P}_{k+1}(\mathfrak{T})$	H^{k+2}	H^{k+1}	H^k	$\mathcal{H}H^{k+1}$	H^{k+1}	H^k
$\mathbb{P}_{k+1}(F) \times \mathbb{P}_{k+1}(\mathfrak{T})$	$h^{k+2} + H^{k+3}$	$h^{k+1} + H^{k+2}$	$h^k + H^{k+1}$	$\mathcal{H}(h^{k+1} + H^{k+2})$	$h^{k+1} + H^{k+2}$	$h^k + H^{k+1}$

TABLE 1. Error estimates in the $L^2(\Omega)$, V and $H(\text{div}, \mathcal{T}_h)$ norms for a partition $\mathcal{P}$ built up on convex and non-convex elements $K \in \mathcal{P}$.

In view of Table 1, special attention on sub-mesh refinement must be made in the $\mathbb{P}_{k+1}(F) \times \mathbb{P}_{k+1}(\mathfrak{T})$ case so as not to “pollute” first-level convergences. Also, observe that rates are similar to the ones reported in [10, Table 1]. Interestingly, in the present approach there is no need to post-process the variables, nor to add additional stabilising terms to the formulation in order to prove these orders of convergence.

5. NUMERICAL VALIDATION

In this section we present two sets of numerical experiments illustrating the performance of the MHM method proposed in this work. The goal of the first one is to check the results given by the error analysis using a smooth analytical solution, while the second test case aims at showing the robustness of the method when applied to a problem with a highly contrasting coefficient.

Since for most of the results the diameter of the elements $K \in \mathcal{P}$ does not need to decrease in order to have a converging method, we shall distinguish two kinds of convergences, namely,

- $\mathcal{H} \rightarrow 0$: the *mesh-based* convergence;
- $H \rightarrow 0$ with $\mathcal{H}$ fixed: the *space-based* convergence.

Second-level meshes are made of triangles that respect the requirement for the well-posedness of the MHM method. Their diameter h tends to zero in both mesh-based and space-based convergence validations.

5.1. An analytical smooth solution. We take Ω to be the unit square and set the coefficient A as the identity matrix. We prescribe homogeneous Dirichlet boundary conditions and the right-hand side in such a way that the exact solution of (1.1) is given by

$$u(x, y) = \sin(2\pi x) \sin(2\pi y).$$

We validate the error estimates using two distinct meshes, namely, conforming quadrilateral meshes and meshes where the elements are L -shaped. We first study the performance of the pair $\mathbb{P}_k(F) \times \mathbb{P}_{k+1}(\mathfrak{T})$ with $k \in \{0, 1, 2, 3\}$. For the mesh-based convergence we use $H = \mathcal{H}$, this is, we consider $F = E$ and do not divide the edges of the partitions any further (the sub-element meshes are the minimal triangulations allowed by the stability results presented in Section 3). The results for quadrilateral elements and $k = 0$ and $k = 3$ are depicted in Figure 2, where we can observe that all the errors tend to zero as predicted by the results in Section 4.

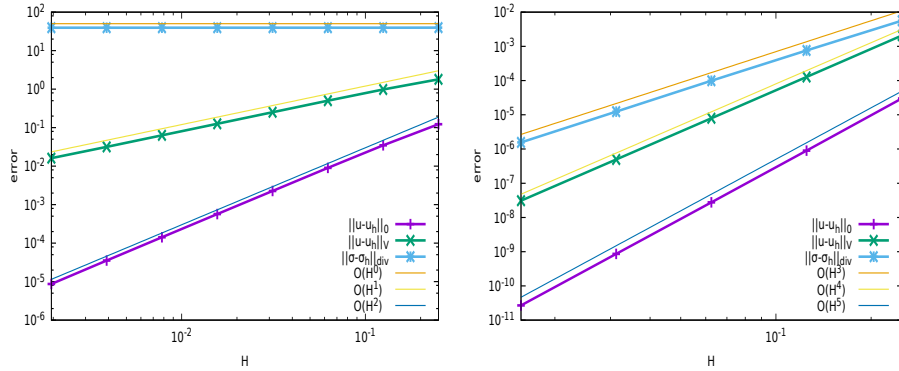


FIGURE 2. The mesh-based convergence history on quadrilateral elements for $k = 0$ (left) and $k = 3$ (right).

The same test, still using quadrilateral meshes, is repeated for $k = 1$ and $k = 2$ in Figure 3. In that figure we also report the results obtained for the space-based strategy. For this we fix the coarse mesh to have 16 squares and then the edges get refined in a structured way, with the implied refinement in the subelement-meshes $\mathcal{T}_h^K$. Interestingly, Figure 3 shows an (unexpected) extra $O(H^{1/2})$ in the convergence rate when the space-based approach is adopted. In fact, we observe a gain of one order of magnitude in accuracy between the space-based and mesh-based strategies. Other numerical experiments show the same behavior, but the proof of this fact

is lacking. We next perform the same study with non-convex L -shaped element meshes using $k = 2$.

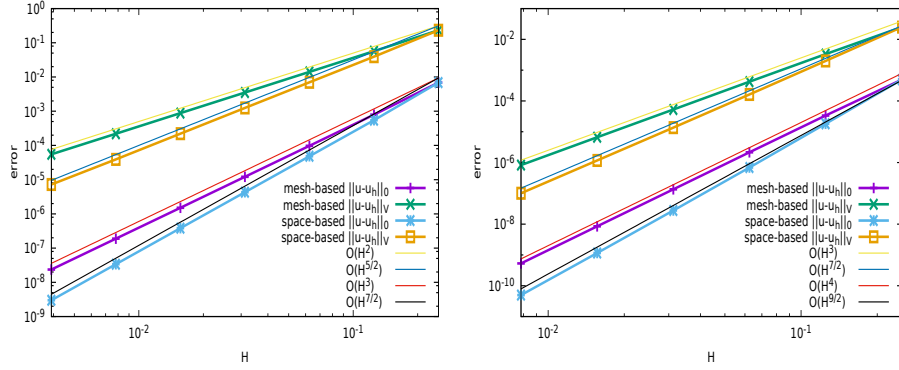


FIGURE 3. Comparison between mesh-based and space-based convergences on quadrilateral elements for $k = 1$ (left) and $k = 2$ (right) in the L^2 and V norms.

Figure 4 shows the isolines of the primal and dual variables for the L -shaped case. As predicted by the theory, all the errors tend to zero with optimal rates for the mesh-based strategy. The errors for the space-based strategy tend to zero with a $H^{\frac{1}{2}}$ extra rate (see Figure 5, left), which is especially noticeable when the results are depicted with respect to the degrees of freedom, as done in the right-hand side of Figure 5.

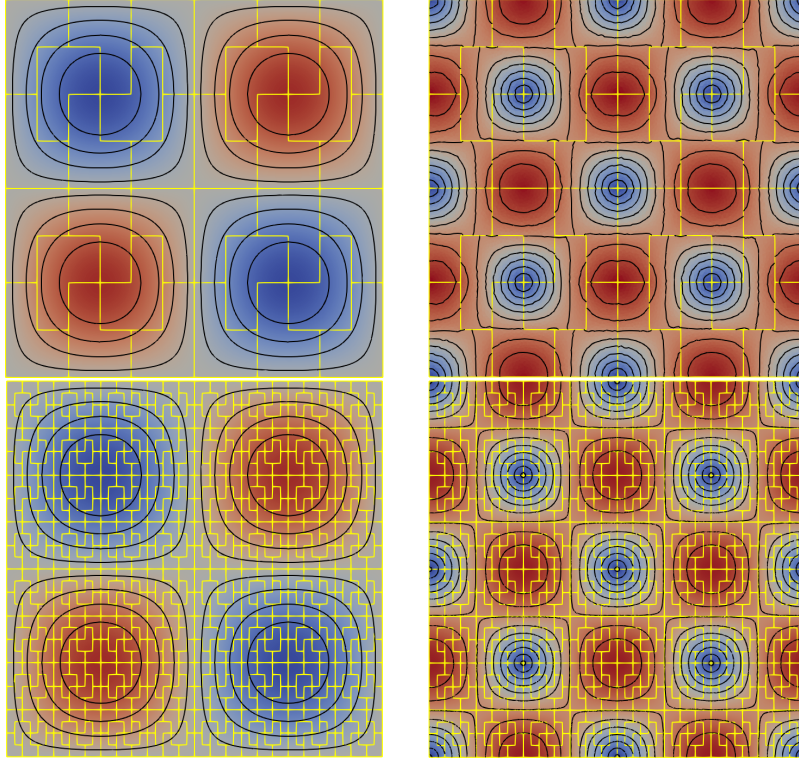


FIGURE 4. Sequence of two refined L -shaped meshes and isolines of u_h (left) and $|\sigma_h|$ (right).

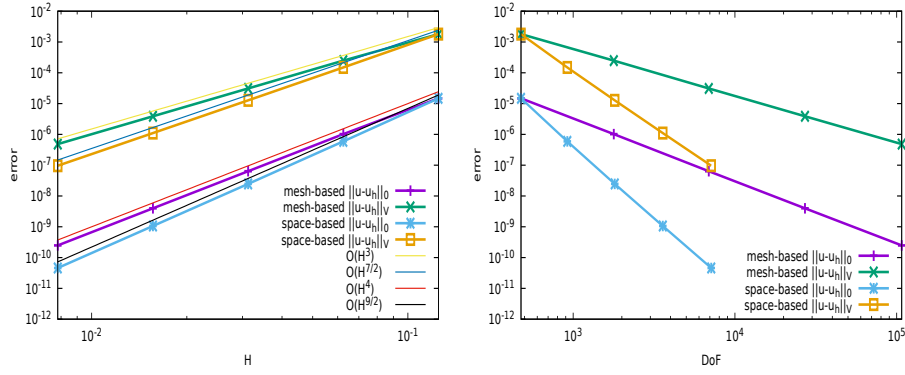


FIGURE 5. Comparison between mesh-based and space-based convergences on L-shaped elements for $k = 2$ in the L^2 and V norms with respect to H (left) and the DoF (right).

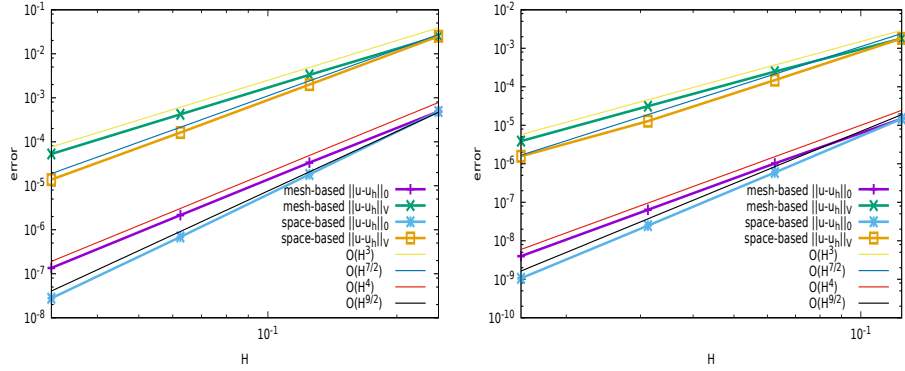


FIGURE 6. Comparison between mesh-based and space-based convergences on quadrangular (left) and L-shaped elements (right) using the $\mathbb{P}_{k+1}(F) \times \mathbb{P}_{k+1}(\mathfrak{T})$ element with $k = 1$ in the $L^2(\Omega)$ and $H^1(\mathcal{P})$ norms.

Now, we address the case $\mathbb{P}_{k+1}(F) \times \mathbb{P}_{k+1}(\mathfrak{T})$ with $k = 1$, again using both quadrilateral and L-shaped meshes. The results are depicted in Figure 6 where, again, we can observe that the errors tend to zero with rates according to the results in Section 4. In particular, the error for u_h in the $H^1(\mathcal{P})$ -norm tends to zero as H^3 while the error in the $L^2(\Omega)$ -norm shows an H^4 convergence rate. Also, once again, the space-based approach shows a faster convergence rate than the mesh-based one. More precisely, the errors $\|u - u_h\|_{1,\mathcal{P}}$ and $\|u - u_h\|_{0,\Omega}$ tend to zero with rates given by $H^{\frac{7}{2}}$ and $H^{\frac{9}{2}}$, respectively. According to the theory presented in the last section, the convergence order is related to $h^{k+1} + H^{\ell+1}$. So, in order to obtain the convergence orders reported in the present results, h needs to be small enough with respect to H .

5.2. Heterogeneous media cases. The domain corresponds to a unit square with 3×3 or 9×9 square inclusions in which the diffusion has different values. In Figure 7 we depict the case with 3 inclusions in each direction. The diffusion matrix is given by $A = \kappa I$ with $\kappa = 1$ in the blue region and $\kappa = 10^5$ in the green one, where I stands for the identity matrix. Homogeneous Dirichlet conditions are prescribed in the whole boundary, and $f = 1$ in Ω .

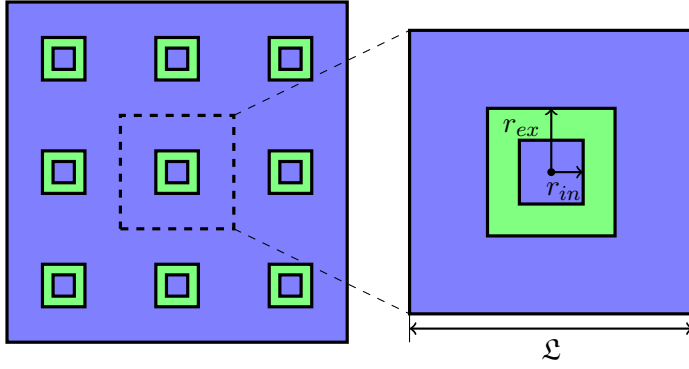


FIGURE 7. Sketch of the heterogeneous diffusive media with 3×3 inclusions. Here $r_{in} = \frac{1}{32}$, $r_{ex} = \frac{1}{16}$, $\mathfrak{L} = \frac{1}{3}$, and the diffusive coefficient is $\kappa = 1$ (blue) and $\kappa = 10^5$ (green).

Since the analytical solution is not known, we have computed a reference solution on a highly refined grid containing 524,288 triangular elements, and the approximate solution is computed using quadratic elements (the global linear system has approximately one million degrees of freedom). For the MHM method we have partitioned the domain using a mesh of 32 L -shaped elements and have used the $\mathbb{P}_1(F) \times \mathbb{P}_2(\mathfrak{T})$ element. The linear system solved by the MHM method has 624 unknowns, while each basis function has been computed solving an off-line local problem on a sub-element mesh of approximately 3,000 triangular elements. In Figure 8 we depict a basis function resulting from this process, where we can observe the influence of the physical coefficient on the shape of the function. It is interesting to notice that the elements do not follow the jump in the diffusion coefficient.

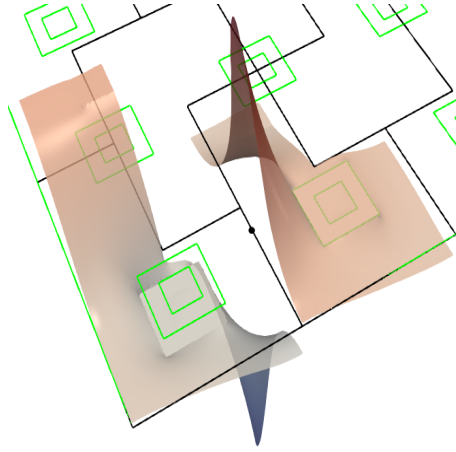


FIGURE 8. A multiscale basis function associated with a degree of freedom (black dot) on a L -shaped mesh with the $\mathbb{P}_1(F) \times \mathbb{P}_2(\mathfrak{T})$ element. We observe the influence of the high-contrast coefficients on the basis function.

For further comparison, we also compare the MHM solution to the solution of the Galerkin method in a mesh containing 512 quadratic elements (3,283 degrees of freedom). In the case in which there are 3 inclusions in each direction, we see in Figure 9 that the MHM's solution

reproduces the reference solution very accurately while the Galerkin method fails to do so. This is especially noticeable in the cross-sections shown in Figure 9.

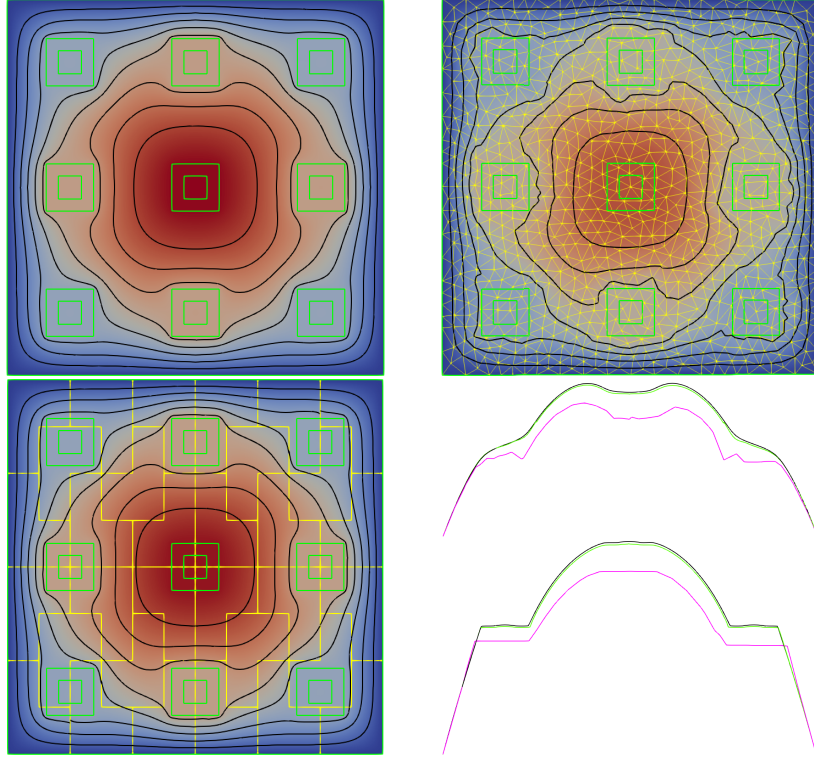


FIGURE 9. The 3×3 inclusion test case: Comparison between the reference solution (top left) and the Galerkin's solution on a conforming triangular mesh with $DoF = 3,283$ (top right) and the MHM's solution (bottom left) on a non-conforming L -shaped mesh with $DoF = 624$. Cross sections at $x = 0.25$ and $x = 0.5$ (bottom right) are shown illustrating the lack of accuracy of the Galerkin method compared to the MHM solution.

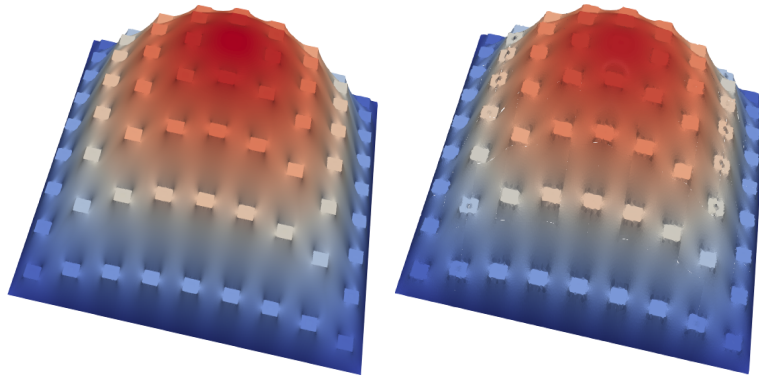


FIGURE 10. The 9×9 inclusions test case. Elevation of the reference solution (left) and the MHM's solution obtained on a 32 L -shaped element mesh (right) with $DoF = 624$.

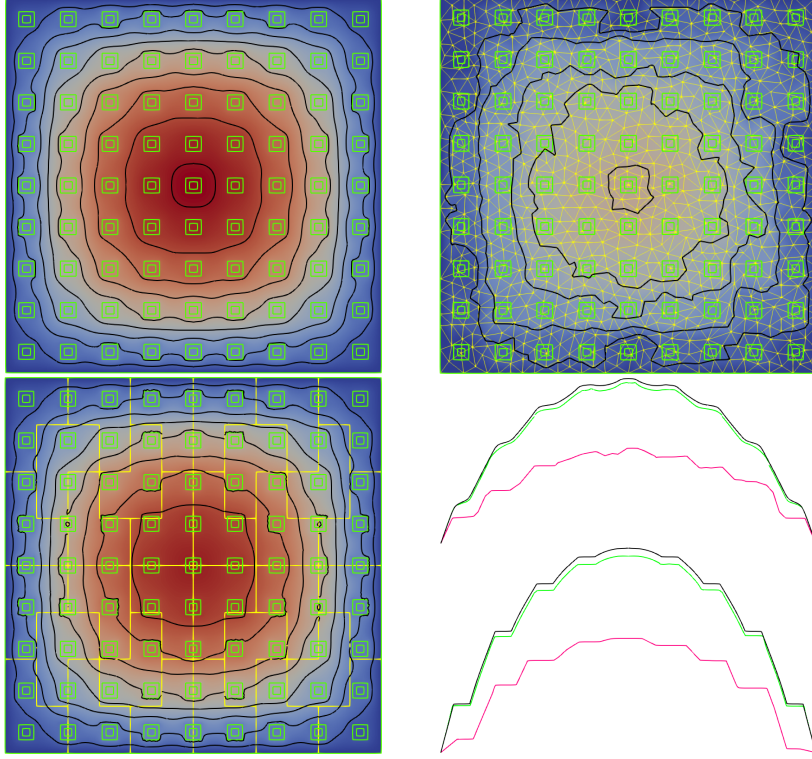


FIGURE 11. The 9×9 inclusion test case: Comparison between the reference solution (top left) and the Galerkin solution on a conform triangular mesh with $DoF = 3,283$ (top right), and the MHM solution (bottom left) on a non-conforming L -shaped mesh with $DoF = 624$. Cross sections at $x = 0.25$ and $x = 0.5$ (bottom right) show that the MHM method improves the accuracy when compared to the Galerkin method with polynomial basis functions.

We next address the case in which there are 9 inclusions in each direction. In Figure 10 we depict elevations of the reference and MHM solutions. We can observe that the use of the multiscale functions makes it possible to give an accurate approximation of the reference solution using a small number of degrees of freedom on the online phase. A finer comparison is shown in Figure 11 where isolines of the solution (along with the computation mesh for the Galerkin and MHM methods), and cross sections are depicted. We can observe, once again, the gain obtained when using multiscale functions, when compared to the standard Galerkin method.

6. CONCLUSION

In this work we have extended the MHM method to general polygonal meshes. Stability and optimal convergence has been proven even when the polynomial degree used to approximate the fluxes is the same as the one used to approximate the primal variable, and when the discrete space for the Lagrange multiplier is formed by discontinuous polynomials. Two converging scenarios appear. In the mesh-based approach, the polygonal mesh gets refined while keeping the same sub-partitions of the skeleton, and of the elements. Meanwhile, in the space-based approach the coarse mesh is fixed at the beginning of the calculation while the sub-partitions of the skeleton and the sub-element meshes get refined. Our numerical experiments show that the latter approach leads to a faster convergence than the former. A formal proof of this fact

is lacking, but this helps to understand the reasons why the method, supplied with an edge refinement strategy, has given satisfactory results in the past (see, e.g., [25]).

The numerical results show a very robust behaviour of the method, even in a case with high contrast. On the other hand, a more refined error analysis of this situation is needed, along with the full error analysis for the three-dimensional case. Also, more intensive, three-dimensional numerical validation is needed, including the possibility of using elements that are not simply connected (possibility allowed by the theory, but that we have not tested so far). This last case would be very attractive to treat cases such as PDEs posed on perforated, non-periodic, domains. These issues will be the subject of future research.

APPENDIX

Lemma 6.1. *Let $k \geq 0$, and let us define the following space*

$$\mathcal{X} = \left\{ q \in H_0^1(0,1) : q|_{(0,0.5)} \in \mathbb{P}_{k+1}(0,0.5) \text{ and } q|_{(0.5,1)} \in \mathbb{P}_{k+1}(0.5,1) \right\}.$$

Then, if μ satisfies

$$(6.1) \quad \mu \in \mathbb{P}_k(0,1) \text{ and } \int_0^1 \mu(x)q(x)dx = 0 \quad \text{for all } q \in \mathcal{X},$$

then $\mu = 0$ in $[0,1]$.

Proof. The case $k = 0$ will be treated first. For $k = 0$ let

$$q(x) = \begin{cases} x & \text{if } x \leq 0.5 \\ 1-x & \text{if } x > 0.5 \end{cases}.$$

Then, $q \in \mathcal{X}$. Thus, if $\mu \in \mathbb{R} \setminus \{0\}$, then

$$\int_0^1 \mu q(x)dx = \mu \int_0^1 q(x)dx = \frac{\mu}{4} \neq 0,$$

which contradicts the hypothesis.

Now, let $k \geq 1$, and let $\mu \in \mathbb{P}_k(0,1)$ satisfying (6.1). Let $q_1, q_2 \in \mathcal{X}$ be defined by

$$(6.2) \quad q_1(x) = \begin{cases} x(0.5-x) & \text{if } x \leq 0.5 \\ 0 & \text{if } x > 0.5 \end{cases}, \quad q_2(x) = \begin{cases} 0 & \text{if } x \leq 0.5 \\ (x-0.5)(1-x) & \text{if } x > 0.5 \end{cases}.$$

Since

$$(6.3) \quad \int_0^1 \mu(x)q_1(x)dx = \int_0^{0.5} \mu(x)x(0.5-x)dx = 0,$$

and $q_1(x) \geq 0$ in $(0,0.5)$, then μ changes sign in $(0,0.5)$, which implies that μ has at least one root $x_1 \in (0,0.5)$. Analogously, since

$$(6.4) \quad \int_0^1 \mu(x)q_2(x)dx = \int_{0.5}^1 \mu(x)(x-0.5)(1-x)dx = 0,$$

then μ has also a root $x_2 \in (0.5,1)$. So, if $k = 1$, then $\mu \in \mathbb{P}_1(0,1)$ has two roots in $(0,1)$, and thus it needs to be identically zero. If $k \geq 2$, we define the following functions in $\mathcal{X}$:

$$(6.5) \quad q_3(x) = q_1(x)(x-x_1) \quad \text{and} \quad q_4(x) = q_2(x)(x-x_2).$$

If μ does not have any other root in $(0,0.5)$, then, since μ and q_3 change signs at the exact same points in $(0,0.5)$, then

$$(6.6) \quad \int_0^1 \mu(x)q_3(x)dx = \int_0^{0.5} \mu(x)q_3(x)dx \neq 0,$$

which contradicts (6.1). In an analogous way, using the function q_4 , we deduce the existence of a fourth root $x_4 \in (0.5, 1)$. If $k = 3$, then μ is a quadratic polynomial with four distinct roots in $(0, 1)$, and then it needs to be identically zero.

For $k \geq 3$, following completely analogous steps to the ones just described, we deduce that μ needs to vanish at, at least, $2(k+1) - 2$ different points in $(0, 1)$, which, for a polynomial $\mu \in \mathbb{P}_k(0, 1)$, is impossible unless $\mu = 0$. This finishes the proof. $\square$

The following result extends the previous result to the $\mathbb{P}_{k+1}(F) \times \mathbb{P}_{k+1}(\mathfrak{T})$ case.

Lemma 6.2. *Let $k \geq 0$. Then, if μ satisfies*

$$(6.7) \quad \mu \in \mathbb{P}_{k+1}(0, 1) \text{ and } \int_0^1 \mu(x)q(x) dx = 0 \quad \text{for all } q \in \mathcal{X},$$

then $\mu = 0$ in $[0, 1]$, where the space $\mathcal{X}$ is defined as:

- For $k = 0$:

$\mathcal{X} = \{q \in H_0^1(0, 1) : \text{the restriction of } q \text{ to } (0, 1/4), (1/4, 1/2), (1/2, 3/4), \text{ and } (3/4, 1) \text{ belongs to } \mathbb{P}_1\}.$

- For $k = 1$:

$\mathcal{X} = \{q \in H_0^1(0, 1) : \text{the restriction of } q \text{ to } (0, 1/3), (1/3, 2/3), \text{ and } (2/3, 1) \text{ belongs to } \mathbb{P}_2\}.$

- For $k \geq 2$:

$$\mathcal{X} = \left\{ q \in H_0^1(0, 1) : q|_{(0, 0.5)} \in \mathbb{P}_{k+1}(0, 0.5) \text{ and } q|_{(0.5, 1)} \in \mathbb{P}_{k+1}(0.5, 1) \right\}.$$

Proof. We split the proof in the three cases described above:

$k = 0$: Let $\mu \in \mathbb{P}_1(0, 1)$ satisfying (6.1), and let $\ell_1, \ell_2 \in \mathcal{X}$ be defined as follows:

$$(6.8) \quad \ell_1(x) = \begin{cases} x & \text{if } x \leq 0.25 \\ 0.25 - x & \text{if } 0.25 \leq x \leq 0.5 \\ 0 & \text{if } 0.5 \leq x \end{cases}, \quad \ell_2(x) = \begin{cases} 0 & \text{if } x \leq 0.5 \\ x - 0.5 & \text{if } 0.5 \leq x \leq 0.75 \\ 1 - x & \text{if } 0.75 \leq x \end{cases}.$$

Since μ satisfies (6.1) then

$$(6.9) \quad \int_0^1 \mu(x)\ell_1(x) dx = \int_0^{0.5} \mu(x)\ell_1(x) dx = 0,$$

which implies that μ has a root in $(0, 0.5)$ since ℓ_1 is positive in that interval. Analogously, since

$$(6.10) \quad \int_0^1 \mu(x)\ell_2(x) dx = \int_{0.5}^1 \mu(x)\ell_2(x) dx = 0,$$

and ℓ_2 is positive in $(0.5, 1)$, we deduce that μ has a root also in $(0.5, 1)$. Since $\mu \in \mathbb{P}_1(0, 1)$ and has two different roots in $(0, 1)$, then it is identically zero.

$k = 1$: Let $\mu \in \mathbb{P}_2(0, 1)$ satisfying (6.1), and let $\ell_3, \ell_4, \ell_5 \in \mathcal{X}$ be defined as follows:

$$\ell_3(x) = \begin{cases} x(1/3 - x) & \text{if } x \leq 1/3 \\ 0 & \text{if } 1/3 \leq x \end{cases}, \quad \ell_4(x) = \begin{cases} (x - 1/3)(2/3 - x) & \text{if } 1/3 \leq x \leq 2/3 \\ 0 & \text{else} \end{cases}$$

$$\ell_5(x) = \begin{cases} 0 & \text{if } x \leq 2/3 \\ (x - 2/3)(1 - x) & \text{if } 2/3 \leq x \end{cases}.$$

Since μ satisfies (6.1) then

$$(6.11) \quad \int_0^1 \mu(x)\ell_3(x) dx = \int_0^{1/3} \mu(x)x(1/3 - x) dx = 0,$$

which, due to the positivity of ℓ_3 implies the existence of $x_1 \in (0, 1/3)$ such that $\mu(x_1) = 0$. In a completely analogous way, we deduce there exist $x_2 \in (1/3, 2/3)$ and $x_3 \in (2/3, 1)$ such that $\mu(x_2) = \mu(x_3) = 0$. Thus, μ must necessarily vanish.

$k \geq 2$: From the proof of Lemma 6.1, any function μ that satisfies (6.7) has at least $2k$ different roots in $(0, 1)$. Now, for $k \geq 2$, $2k \geq k + 2$, and thus, if $\mu \in \mathbb{P}_{k+1}(0, 1)$ satisfies (6.7) it will vanish at at least $k + 2$ different points in $(0, 1)$, which implies that μ is identically zero. This finishes the proof. $\square$

FUNDING

The third author was supported by FONDECYT Project 1181572. The fourth author was partially supported by CNPq/Brazil No. 301576/2013-0, EU H2020 Program and from MCTI/RNP-Brazil under the HPC4E Project, Grant Agreement No. 689772, and by CAPES/Brazil No. 88881.143295/2017-01 under the PHOTOM Project.

REFERENCES

- [1] Y. Achdou, C. Japhet, Y. Maday, and F. Nataf. A new cement to glue non-conforming grids with Robin interface conditions: the finite volume case. *Numer. Math.*, 92(4):593–620, 2002.
- [2] R. Araya, C. Harder, D. Paredes, and F. Valentin. Multiscale hybrid-mixed method. *SIAM J. Numer. Anal.*, 51(6):3505–3531, 2013.
- [3] R. Araya, C. Harder, A. Poza, and F. Valentin. Multiscale hybrid-mixed method for the Stokes and Brinkman equations – the method. *Comp. Meth. Appl. Mech. Eng.*, 324:29–53, 2017.
- [4] T. Arbogast, G. Pencheva, M. F. Wheeler, and I. Yotov. A multiscale mortar mixed finite element method. *SIAM Multiscale Modeling and Simulation*, 6:319–346, 2007.
- [5] I. Babuska and E. Osborn. Generalized finite element methods: Their performance and their relation to mixed methods. *SIAM J. Num. Anal.*, 20(3):510–536, 1983.
- [6] L. Beirão da Veiga, F. Brezzi, L. D. Marini, and A. Russo. The hitchhiker’s guide to the virtual element method. *Math. Models Methods Appl. Sci.*, 24(8):1541–1573, 2014.
- [7] L. Beirão da Veiga, F. Brezzi, A. Cangiani, G. Manzini, L.D. Marini, and A. Russo. Basic principles of virtual element methods. *Math. Models Methods Appl. Sci.*, 23(1):199–214, 2013.
- [8] A. Cangiani, E. H. Dong, A. Georgoulis, and P. Houston. *hp-Version Discontinuous Galerkin Methods on Polygonal and Polyhedral Meshes*. Springer, 2017.
- [9] A. Cangiani, E. Georgoulis, and P. Houston. *hp-version discontinuous Galerkin methods on polygonal and polyhedral meshes*. *Math. Models Methods Appl. Sci.*, 24(10):2009–2041, 2014.
- [10] B. Cockburn, D. Di Pietro, and A. Ern. Bridging the hybrid high-order and hybridizable discontinuous Galerkin methods. *ESAIM Math. Model. Numer. Anal.*, 50(3):635–650, 2016.
- [11] B. Cockburn, J. Gopalakrishnan, and R. Lazarov. Unified hybridization of discontinuous Galerkin, mixed, and continuous Galerkin methods for second order elliptic problems. *SIAM J. Numer. Anal.*, 47(2):1319–1365, 2009.
- [12] A.L.G.A. Coutinho, L. P. Franca, and F. Valentin. Numerical multiscale methods. *International Journal for Numerical Methods in Fluids*, 70(4):403–419, 2012.
- [13] M. Crouzeix and P.A. Raviart. Conforming and nonconforming finite element methods for solving the stationary stokes equations. *I. Revue française d’automatique, informatique, recherche opérationnelle. Mathématique*, 7(3):33–75, 1973.
- [14] O. Duran, P. R. B Devloo, A.T.A. Gomes, and F. Valentin. A multiscale hybrid method for Darcy’s problems using mixed finite element local solvers. Technical Report HAL-01868853, INRIA, 2018.
- [15] W. E and B. Engquist. The heterogeneous multi-scale methods. *Comm. Math. Sci.*, 1:87–133, 2003.
- [16] Y.R. Efendiev, T.Y. Hou, and X.H. Wu. Convergence of a nonconforming multiscale finite element method. *SIAM J. Numer. Anal.*, 37(3):888–910 (electronic), 2000.
- [17] A. Ern and J.-L. Guermond. *Theory and practice of finite elements*, volume 159 of *Applied Mathematical Sciences*. Springer-Verlag, New York, 2004.
- [18] A. Ern and D. D. A. Pietro. Arbitrary-order mixed methods for heterogeneous anisotropic diffusion on general meshes. *to appear in IMA (DOI: 10.1093/imanum/drw003)*, 2016.

- [19] C. Farhat, I. Harari, and L. P. Franca. The discontinuous enrichment method. *Comput. Methods Appl. Mech. Engrg.*, 190(48):6455–6479, 2001.
- [20] C. Farhat, J. Mandel, and F.-X. Roux. Optimal convergence properties of the feti domain decomposition method. *Comm. Numer. Methods Engrg.*, 115:365–385, 1994.
- [21] M. J. Gander, C. Japhet, Y. Maday, and F. Nataf. A new cement to glue nonconforming grids with Robin interface conditions: the finite element case. In *Domain decomposition methods in science and engineering*, volume 40 of *Lect. Notes Comput. Sci. Eng.*, pages 259–266. Springer, Berlin, 2005.
- [22] P. Grisvard. *Elliptic Problems in Non-Smooth Domains*. Pitman Publishing, 1985.
- [23] C. Harder, A.L. Madureira, and F. Valentin. A hybrid-mixed method for elasticity. *ESAIM: Math. Model. Num. Anal.*, 50(2):311–336, 2016.
- [24] C. Harder, D. Paredes, and F. Valentin. A family of multiscale hybrid-mixed finite element methods for the Darcy equation with rough coefficients. *J. Comput. Phys.*, 245:107–130, 2013.
- [25] C. Harder, D. Paredes, and F. Valentin. On a multiscale hybrid-mixed method for advective-reactive dominated problems with heterogenous coefficients. *SIAM Multiscale Model. and Simul.*, 13(2):491–518, 2015.
- [26] C. Harder and F. Valentin. Foundations of the MHM method. In G. R. Barrenechea, F. Brezzi, A. Cangiani, and E. H. Georgoulis, editors, *Building Bridges: Connections and Challenges in Modern Approaches to Numerical Partial Differential Equations*, Lecture Notes in Computational Science and Engineering. Springer, 2016.
- [27] T. Y. Hou, X. Wu, and Z. Cai. Convergence of a multiscale finite element method for elliptic problems with rapidly oscillating coefficients. *Math. Comp.*, 68(227):913–943, 1999.
- [28] T. J. R. Hughes. Multiscale phenomena: Green’s functions, the Dirichlet-to-Neumann formulation, subgrid scale models, bubbles and the origin of stabilized methods. 127:387–401, 1995.
- [29] S. Lanteri, D. Paredes, C. Scheid, and F. Valentin. The multiscale hybrid-mixed method for the Maxwell equations in heterogeneous media. *SIAM Multiscale Model. Simul.*, 16(4):1648–1683, 2018.
- [30] A. Malqvist and D. Peterseim. Localization of elliptic multiscale problems. *Math. Comp.*, 83(290):2583–2603, 2014.
- [31] D. Paredes, F. Valentin, and H. M. Versieux. On the robustness of multiscale hybrid-mixed methods. *Math. Comp.*, 86(304):525–548, 2017.
- [32] P.A. Raviart and J.M. Thomas. Primal hybrid finite element methods for 2nd order elliptic equations. *Math. Comp.*, 31(138):391–413, 1977.
- [33] G. Sangalli. Capturing small scales in elliptic problems using a Residual-Free Bubbles finite element method. *SIAM Multiscale Model. Simul.*, 1(3):485–503, 2003.
- [34] A. Veiser and R. Verfurth. Poincaré constants for finite element stars. *IMA Journal of Numerical Analysis*, 32(1):30–47.

(G.R.B.) DEPARTMENT OF MATHEMATICS AND STATISTICS, UNIVERSITY OF STRATHCLYDE, 26 RICHMOND STREET, GLASGOW G1 1XH, SCOTLAND. gabriel.barrenechea@strath.ac.uk

(F.J.) UNIVERSITÉ DE LYON, IUT LYON 1, CNRS, LIRIS, UMR5205, F-69622, VILLEURBANNE, FRANCE, fabrice.jaillet@liris.cnrs.fr

(D.P.) DEPARTAMENTO DE INGENIERÍA MATEMÁTICA, UNIVERSIDAD DE CONCEPCIÓN, CONCEPCIÓN, CHILE , dparedes@udec.cl

(F.V.) DEPARTAMENTO DE MÉTODOS MATEMÁTICOS E COMPUTACIONAIS, LNCC - LABORATÓRIO NACIONAL DE COMPUTAÇÃO CIENTÍFICA, AV. GETÚLIO VARGAS, 333, 25651-070 PETRÓPOLIS - RJ, BRAZIL, valentin@lncc.br