

New primal-dual weak Galerkin finite element methods for convection-diffusion problems

Waixiang Cao^{a,1}, Chunmei Wang^{b,*,2}

^a School of Mathematical Sciences, Beijing Normal University, Beijing 100875, China

^b Department of Mathematics & Statistics, Texas Tech University, Lubbock, TX 79409, USA

ARTICLE INFO

Article history:

Received 7 June 2020

Received in revised form 10 October 2020

Accepted 9 December 2020

Available online xxxx

Keywords:

Primal-dual

Weak Galerkin

Finite element methods

Convection-diffusion

Weak gradient

Polygonal or polyhedral meshes

ABSTRACT

This article devises a new primal-dual weak Galerkin finite element method for the convection-diffusion equation. Optimal order error estimates are established for the primal-dual weak Galerkin approximations in various discrete norms and the standard L^2 norms. A series of numerical experiments are conducted and reported to verify the theoretical findings.

Crown Copyright © 2020 Published by Elsevier B.V. on behalf of IMACS. All rights reserved.

1. Introduction

This paper is concerned with new development of numerical methods for the convection-diffusion equations. For simplicity, we consider the model problem that seeks an unknown function u satisfying

$$\begin{aligned} -\nabla \cdot (a \nabla u + \mathbf{b}u) &= f, & \text{in } \Omega, \\ u &= g_1, & \text{on } \Gamma_D, \\ (a \nabla u + \mathbf{b}u) \cdot \mathbf{n} &= g_2, & \text{on } \Gamma_N, \end{aligned} \quad (1.1)$$

where $\Omega \subset \mathbb{R}^d$ ($d = 2, 3$) is an open bounded polygonal ($d = 2$) or polyhedral ($d = 3$) domain with Lipschitz continuous boundary $\partial\Omega$, Γ_D is the Dirichlet boundary, $\Gamma_N = \partial\Omega \setminus \Gamma_D$ is the Neumann boundary, and $\mathbf{n}$ is the unit outward normal direction to the Neumann boundary Γ_N . We assume that the convection vector $\mathbf{b} \in [L^\infty(\Omega)]^d$ is bounded, and the diffusion tensor $a = \{a_{ij}\}_{d \times d}$ is symmetric and positive definite in the sense that there exists a constant $\alpha > 0$, such that

$$\xi^T a \xi \geq \alpha \xi^T \xi, \quad \forall \xi \in \mathbb{R}^d.$$

* Corresponding author.

E-mail addresses: caowx@bnu.edu.cn (W. Cao), chunmei.wang@ttu.edu (C. Wang).

¹ The research of Waixiang Cao was partially supported by NSFC grant No. 11871106.

² The research of Chunmei Wang was partially supported by National Science Foundation Award DMS-1849483.

Furthermore, we assume that the diffusion tensor a and the convection tensor $\mathbf{b}$ are uniformly piecewise continuous functions.

The convection-diffusion equations arise in many areas of science and engineering. Readers are referred to the “Introduction” Section in [16] and the references cited therein for a detailed description of the convection-diffusion equations.

The weak Galerkin (WG) finite element method was first introduced by Wang and Ye in [14] for second order elliptic equations, and later was widely used for solving various partial differential equations, e.g., [1,13,15,7,17,8,11,9,12,10]. Recently, the authors in [3] have developed a new numerical scheme, called “primal-dual weak Galerkin (PDWG) finite element method” for the second order elliptic problem in non-divergence form. PDWG uses the weak Galerkin strategy to construct the *discrete weak Hessian* operator in the weak formulation of the model PDEs, and further seeks a discontinuous function which minimizes a stabilizer defined on the boundary of each element with the constraint given by the weak formulation of the model PDEs weakly defined on each element. The Euler-Lagrange method was employed to solve the constrained minimization problem leading to the primal-dual weak Galerkin finite element method, which has been further studied in [4,5,16,6,2]. The primal-dual weak Galerkin finite element method has shown the promising features as a discretization approach due to: (1) it works well for a wide class of PDE problems for which no traditional variational formulations are available; and (2) it is applicable to virtually any PDE problems where the inf-sup condition is satisfied.

Using the usual integration by parts one may derive a weak formulation for the model problem (1.1) as follows: Find $u \in H^1(\Omega)$ satisfying $u|_{\Gamma_D} = g_1$ and $(a\nabla u + \mathbf{b}u) \cdot \mathbf{n}|_{\Gamma_N} = g_2$, such that

$$\begin{aligned} & \int_T (a\nabla u + \mathbf{b}u) \cdot \nabla w dT - \int_{\partial T} (a\nabla u + \mathbf{b}u) \cdot \mathbf{n} w ds \\ &= \int_T f w dT, \quad \forall T \subset \Omega, w \in H^1(T). \end{aligned} \quad (1.2)$$

The PDWG numerical scheme developed in this paper is based on the weak formulation (1.2) for the convection-diffusion model problem (1.1). The gradient operator is the principal player in (1.2) so that a reconstructed gradient (i.e., weak gradient) is crucial in the PDWG finite element scheme. In contrast, the PDWG finite element method developed in [16] was based on a weak form principled by the operator $\mathcal{L} = \nabla \cdot (a\nabla)$ so that a reconstructed weak $\mathcal{L}$ played a key role in the construction of the numerical scheme. The two numerical methods are thus sharply different from each other, and each has its own advantage in theory and practical computation.

The rest of the paper is organized as follows. In Section 2, we present our primal-dual weak Galerkin scheme for the model problem (1.1) based on the weak formulation (1.2). In Section 3, we shall establish a result on the solution existence and uniqueness for the numerical method. Section 4 is devoted to the establishment of the property of mass conservation. The error equations for the primal-dual weak Galerkin algorithm are derived in Section 5. Sections 6–7 are devoted to the establishment of some optimal order error estimates for the PDWG solution in discrete norms as well as the usual L^2 -norm. Finally, various numerical examples are presented in the last section to support our theoretical findings.

Throughout this paper, we adopt standard notations for Sobolev spaces such as $W^{m,p}(D)$ on sub-domain $D \subset \Omega$ equipped with the norm $\|\cdot\|_{m,p,D}$ and the semi-norm $|\cdot|_{m,p,D}$. When $D = \Omega$, we omit the index D ; and if $p = 2$, we set $W^{m,p}(D) = H^m(D)$, $\|\cdot\|_{m,p,D} = \|\cdot\|_{m,D}$, and $|\cdot|_{m,p,D} = |\cdot|_{m,D}$, and if $m = 0$, $p = 2$, we set $\|\cdot\|_{m,p,D} = \|\cdot\|_D$.

2. Numerical algorithm

Let $\mathcal{T}_h$ be a partition of the domain Ω into polygons in 2D or polyhedra in 3D which is shape regular in the sense of [13]. Denote by $\mathcal{E}_h$ the set of all edges or flat faces in $\mathcal{T}_h$ and $\mathcal{E}_h^0 = \mathcal{E}_h \setminus \partial\Omega$ the set of all interior edges or flat faces. Denote by h_T the meshsize of $T \in \mathcal{T}_h$ and $h = \max_{T \in \mathcal{T}_h} h_T$ the meshsize for the partition $\mathcal{T}_h$.

By a weak function on $T \in \mathcal{T}_h$ we mean a triplet $v = \{v_0, v_b, v_n\}$ such that $v_0 \in L^2(T)$, $v_b \in L^2(\partial T)$ and $v_n \in L^2(\partial T)$, where ∂T is the boundary of T . The first and the second components, namely v_0 and v_b , should be understood as the value of v in the interior and on the boundary of T respectively. The third component v_n refers to the value of $(a\nabla v + \mathbf{b}v) \cdot \mathbf{n}$ on ∂T . Note that v_b and v_n may not necessarily be the trace of v_0 and $(a\nabla v_0 + \mathbf{b}v_0) \cdot \mathbf{n}$ on ∂T . Denote by $\mathcal{W}(T)$ the space of all weak functions on T ; i.e.,

$$\mathcal{W}(T) = \{v = \{v_0, v_b, v_n\} : v_0 \in L^2(T), v_b \in L^2(\partial T), v_n \in L^2(\partial T)\}. \quad (2.1)$$

The weak gradient of $v \in \mathcal{W}(T)$, denoted by $\nabla_w v$, is defined as a linear functional on $[H^1(T)]^d$ such that

$$(\nabla_w v, \boldsymbol{\psi})_T = -(v_0, \nabla \cdot \boldsymbol{\psi})_T + (v_b, \boldsymbol{\psi} \cdot \mathbf{n})_{\partial T},$$

for all $\boldsymbol{\psi} \in [H^1(T)]^d$. Denote by $P_r(T)$ the space of polynomials on T with degree $r \geq 0$. A discrete version of $\nabla_w v$, denoted by $\nabla_{w,r,T} v$, is defined as the unique vector-valued polynomial in $[P_r(T)]^d$ satisfying

$$(\nabla_{w,r,T} v, \boldsymbol{\psi})_T = -(v_0, \nabla \cdot \boldsymbol{\psi})_T + (v_b, \boldsymbol{\psi} \cdot \mathbf{n})_{\partial T}, \quad \forall \boldsymbol{\psi} \in [P_r(T)]^d. \quad (2.2)$$

For smooth v_0 , we have from the usual integration by parts that

$$(\nabla_{w,T} v, \psi)_T = (\nabla v_0, \psi)_T - \langle v_0 - v_b, \psi \cdot \mathbf{n} \rangle_{\partial T}, \quad \forall \psi \in [P_r(T)]^d. \quad (2.3)$$

For any given integer $k \geq 1$, denote by $W_k(T)$ the local discrete weak function space; i.e.,

$$W_k(T) = \{ \{v_0, v_b, v_n\} : v_0 \in P_k(T), v_b \in P_k(e), v_n \in P_l(e), e \subset \partial T \},$$

where $l = k - 1$ or $l = k$. Patching $W_k(T)$ over all the elements $T \in \mathcal{T}_h$ through a common value v_b and $\pm v_n$ on the interior interface $\mathcal{E}_h^0$, we arrive at a global weak finite element space W_h ; i.e.,

$$W_h = \{ \{v_0, v_b, v_n\} : \{v_0, v_b, v_n\}|_T \in W_k(T), \forall T \in \mathcal{T}_h \}.$$

Denote by W_h^0 the subspace of W_h with homogeneous Dirichlet and Neumann boundary conditions; i.e.,

$$W_h^0 = \{ \{v_0, v_b, v_n\} \in W_h : v_b = 0 \text{ on } \Gamma_D, v_n = 0 \text{ on } \Gamma_N \}. \quad (2.4)$$

Next, let M_h be the finite element space consisting of piecewise polynomials of degree k ; i.e.,

$$M_h = \{ \sigma : \sigma|_T \in P_k(T), \forall T \in \mathcal{T}_h \}. \quad (2.5)$$

Remark 2.1. The finite element space M_h in (2.5) can also be constructed by using piecewise polynomials of degree $k - 1$ in the forthcoming numerical scheme. All the mathematical results to be presented in this paper can be extended to the case of $k - 1$ without any difficulty.

For simplicity, for any $v = \{v_0, v_b, v_n\} \in W_h$, denote by $\nabla_w v$ the discrete weak gradient $\nabla_{w,k-1,T} v$ computed by using (2.2) on each element T ; i.e.,

$$(\nabla_w v)|_T = \nabla_{w,k-1,T}(v|_T), \quad v \in W_h.$$

Let us introduce the following bilinear forms:

$$\begin{aligned} s(u, v) &= \sum_{T \in \mathcal{T}_h} h_T^{-3} \langle u_0 - u_b, v_0 - v_b \rangle_{\partial T} \\ &\quad + h_T^{-1} \langle (a \nabla u_0 + \mathbf{b} u_0) \cdot \mathbf{n} - u_n, (a \nabla v_0 + \mathbf{b} v_0) \cdot \mathbf{n} - v_n \rangle_{\partial T}, \\ b(u, \lambda) &= \sum_{T \in \mathcal{T}_h} (a \nabla_w u + \mathbf{b} u_0, \nabla \lambda)_T - \langle u_n, \lambda \rangle_{\partial T}, \\ c(\lambda, \sigma) &= \tau_1 \sum_{T \in \mathcal{T}_h} h_T^2 (\nabla \lambda, \nabla \sigma)_T + \tau_2 \sum_{T \in \mathcal{T}_h} h_T^4 \sum_{i,j=1}^d (\partial_{ij}^2 \lambda, \partial_{ij}^2 \sigma)_T, \end{aligned}$$

where $u, v \in W_h$ and $\lambda, \sigma \in M_h$, $\tau_1 \geq 0$ and $\tau_2 \geq 0$ are two mesh-independent parameters.

Let $k \geq 1$ and $T \in \mathcal{T}_h$. Denote by $Q_0^{(k)}$ the L^2 projection operator onto $P_k(T)$. For each edge or face $e \subset \partial T$, denote by $Q_b^{(k)}$ and $Q_n^{(l)}$ the L^2 projection operators onto $P_k(e)$ and $P_l(e)$, respectively. For any $w \in H^1(\Omega)$, denote by $Q_h w$ the L^2 projection onto the weak finite element space W_h such that on each element T ,

$$Q_h w = \{ Q_0^{(k)} w, Q_b^{(k)} w, Q_n^{(l)} ((a \nabla w + \mathbf{b} w) \cdot \mathbf{n}) \}.$$

Denote by $Q_h^{(k-1)}$ the L^2 projection operator onto the space $[P_{k-1}(T)]^d$.

The numerical scheme for the convection-diffusion problem (1.1) based on the variational formulation (1.2) can be stated as follows:

Primal-Dual Weak Galerkin Algorithm 2.1. Find $(u_h; \lambda_h) \in W_h \times M_h$ satisfying $u_b = Q_b^{(k)} g_1$ on Γ_D and $u_n = Q_n^{(l)} g_2$ on Γ_N , such that

$$s(u_h, v) + b(v, \lambda_h) = 0, \quad \forall v \in W_h^0, \quad (2.6)$$

$$-c(\lambda_h, \sigma) + b(u_h, \sigma) = (f, \sigma), \quad \forall \sigma \in M_h. \quad (2.7)$$

Remark 2.2. For the case of $l = k$, one may take $\tau_1 = \tau_2 = 0$ and thus $c(\lambda_h, \sigma) = 0$; for the case of $l = k - 1$ and $k = 1$, one may take $\tau_1 > 0$ and $\tau_2 = 0$; for the case of $l = k - 1$ and $k \geq 2$, one would take $\tau_1 = 0$ and $\tau_2 > 0$, as suggested by the mathematical theory.

3. Solution existence and uniqueness

For the sake of analysis, in what follows of this paper, we assume that the diffusion tensor a and the convection tensor $\mathbf{b}$ in the convection-diffusion equation (1.1) are piecewise constants in Ω with respect to the finite element partition $\mathcal{T}_h$. However, the analysis can be extended to the case that a and $\mathbf{b}$ are piecewise smooth functions without any difficulty.

The L^2 projection operators Q_h and $Q_h^{(k-1)}$ satisfy the following commutative property [13]:

$$\nabla_w(Q_h w) = Q_h^{(k-1)}(\nabla w), \quad \forall w \in H^1(T). \quad (3.1)$$

In the finite element spaces W_h and M_h , we introduce the following seminorms:

$$\|v\|_{W_h} = s(v, v)^{\frac{1}{2}}, \quad v \in W_h; \quad (3.2)$$

$$\|\sigma\|_{M_h} = c(\sigma, \sigma)^{\frac{1}{2}}, \quad \sigma \in M_h. \quad (3.3)$$

Lemma 3.1 (Generalized inf-sup condition). For any $\lambda \in M_h$, there exists a $v \in W_h^0$ satisfying

$$b(v, \lambda) \geq \begin{cases} \frac{1}{2} \|\lambda\|^2, & l = k, \\ \frac{1}{2} \|\lambda\|^2 - \beta h^2 \|\nabla \lambda\|^2, & k = 1, l = k - 1, \\ \frac{1}{2} \|\lambda\|^2 - \beta h^4 |\lambda|_2^2, & k \geq 2, l = k - 1, \end{cases} \quad (3.4)$$

for some constant $\beta > 0$.

Proof. Consider the auxiliary problem of seeking w such that

$$\begin{aligned} -\nabla \cdot (a \nabla w + \mathbf{b} w) &= \lambda, & \text{in } \Omega, \\ w &= 0, & \text{on } \Gamma_D, \\ (a \nabla w + \mathbf{b} w) \cdot \mathbf{n} &= 0, & \text{on } \Gamma_N. \end{aligned} \quad (3.5)$$

Assume that the auxiliary problem (3.5) has the H^2 -regularity property in the sense that there exists a constant C satisfying

$$\|w\|_2 \leq C \|\lambda\|. \quad (3.6)$$

By taking $v = Q_h w = \{Q_0^{(k)} w, Q_b^{(k)} w, Q_n^{(l)}((a \nabla w + \mathbf{b} w) \cdot \mathbf{n})\} \in W_h^0$ in $b(v, \lambda)$, we have from (2.2) and the usual integration by parts that

$$\begin{aligned} b(v, \lambda) &= b(Q_h w, \lambda) \\ &= \sum_{T \in \mathcal{T}_h} (a \nabla_w Q_h w + \mathbf{b} Q_0^{(k)} w, \nabla \lambda)_T - \langle Q_n^{(l)}((a \nabla w + \mathbf{b} w) \cdot \mathbf{n}), \lambda \rangle_{\partial T} \\ &= \sum_{T \in \mathcal{T}_h} (a Q_h^{(k-1)}(\nabla w) + \mathbf{b} Q_0^{(k)} w, \nabla \lambda)_T - \langle Q_n^{(l)}((a \nabla w + \mathbf{b} w) \cdot \mathbf{n}), \lambda \rangle_{\partial T} \\ &= \sum_{T \in \mathcal{T}_h} (a \nabla w + \mathbf{b} w, \nabla \lambda)_T - \langle Q_n^{(l)}((a \nabla w + \mathbf{b} w) \cdot \mathbf{n}), \lambda \rangle_{\partial T} \\ &= \sum_{T \in \mathcal{T}_h} -(\nabla \cdot (a \nabla w + \mathbf{b} w), \lambda)_T - \langle (Q_n^{(l)} - I)((a \nabla w + \mathbf{b} w) \cdot \mathbf{n}), \lambda \rangle_{\partial T} \\ &= \|\lambda\|^2 - \sum_{T \in \mathcal{T}_h} \langle (Q_n^{(l)} - I)((a \nabla w + \mathbf{b} w) \cdot \mathbf{n}), (I - Q_n^{(l)}) \lambda \rangle_{\partial T}, \end{aligned} \quad (3.7)$$

where we have used the first equation of (3.5), (3.1), and the property of the L^2 projection $Q_n^{(l)}$.

We shall discuss the estimate of the term $\sum_{T \in \mathcal{T}_h} \langle (Q_n^{(l)} - I)((a \nabla w + \mathbf{b} w) \cdot \mathbf{n}), (I - Q_n^{(l)}) \lambda \rangle_{\partial T}$ in various situations. For the case of $l = k$, we have

$$\sum_{T \in \mathcal{T}_h} \langle (Q_n^{(l)} - I)((a \nabla w + \mathbf{b} w) \cdot \mathbf{n}), (I - Q_n^{(l)}) \lambda \rangle_{\partial T} = 0,$$

which, together with (3.7), gives (3.4) for the case of $l = k$. For the case of $l = k - 1$, using the Cauchy-Schwarz inequality and the trace inequality (6.1) gives

$$\begin{aligned}
 & \left| \sum_{T \in \mathcal{T}_h} \langle (Q_n^{(l)} - I)((a \nabla w + \mathbf{b}w) \cdot \mathbf{n}), (I - Q_n^{(l)})\lambda \rangle_{\partial T} \right| \\
 & \leq \left(\sum_{T \in \mathcal{T}_h} \|(Q_n^{(l)} - I)((a \nabla w + \mathbf{b}w) \cdot \mathbf{n})\|_{\partial T}^2 \right)^{\frac{1}{2}} \left(\sum_{T \in \mathcal{T}_h} \|(I - Q_n^{(l)})\lambda\|_{\partial T}^2 \right)^{\frac{1}{2}} \\
 & \leq C \left(\sum_{T \in \mathcal{T}_h} h_T^{-1} \|(Q_0^{(l)} - I)((a \nabla w + \mathbf{b}w))\|_T^2 \right. \\
 & \quad \left. + h_T \|(Q_0^{(l)} - I)(a \nabla w + \mathbf{b}w)\|_{1,T}^2 \right)^{\frac{1}{2}} \\
 & \quad \left(\sum_{T \in \mathcal{T}_h} h_T^{-1} \|(I - Q_0^{(l)})\lambda\|_T^2 + h_T \|(I - Q_0^{(l)})\lambda\|_{1,T}^2 \right)^{\frac{1}{2}} \\
 & \leq \begin{cases} Ch \|\nabla \lambda\| \|\mathbf{w}\|_2, & k = 1, l = k - 1, \\ Ch^2 |\lambda|_2 \|\mathbf{w}\|_2, & k \geq 2, l = k - 1. \end{cases}
 \end{aligned} \tag{3.8}$$

Substituting (3.8) into (3.7) and using the Young's inequality and the H^2 -regularity property (3.6) gives

$$\begin{aligned}
 |b(v, \lambda)| & \geq \|\lambda\|^2 - \epsilon \|\mathbf{w}\|_2^2 - C\epsilon^{-1} \begin{cases} h^2 \|\nabla \lambda\|^2, & k = 1, l = k - 1 \\ h^4 |\lambda|_2^2, & k \geq 2, l = k - 1 \end{cases} \\
 & \geq (1 - \epsilon C) \|\lambda\|^2 - C\epsilon^{-1} \begin{cases} h^2 \|\nabla \lambda\|^2, & k = 1, l = k - 1 \\ h^4 |\lambda|_2^2, & k \geq 2, l = k - 1 \end{cases} \\
 & \geq \frac{1}{2} \|\lambda\|^2 - \beta \begin{cases} h^2 \|\nabla \lambda\|^2, & k = 1, l = k - 1, \\ h^4 |\lambda|_2^2, & k \geq 2, l = k - 1, \end{cases}
 \end{aligned}$$

where $\epsilon > 0$ is a parameter satisfying $1 - \epsilon C \geq \frac{1}{2}$, and $\beta = C\epsilon^{-1} > 0$. This completes the proof of (3.4) for the case of $l = k - 1$ and further completes the proof of the lemma. $\square$

Theorem 3.2. The primal-dual weak Galerkin algorithm (2.6)-(2.7) has a unique solution.

Proof. It suffices to prove that the homogeneous problem of (2.6)-(2.7) has only trivial solution. To this end, we assume $f = 0$, $g_1 = 0$ and $g_2 = 0$. By letting $v = u_h$ and $\sigma = \lambda_h$ in (2.6)-(2.7), we have from the difference of (2.6)-(2.7) that

$$s(u_h, u_h) + c(\lambda_h, \lambda_h) = 0,$$

which implies $u_0 = u_b$ and $(a \nabla u_0 + \mathbf{b}u_0) \cdot \mathbf{n} = u_n$ on each ∂T ; and $c(\lambda_h, \lambda_h) = 0$. From $c(\lambda_h, \lambda_h) = 0$ we have $\nabla \lambda_h = 0$ on each element $T \in \mathcal{T}_h$ if $\tau_1 > 0$ and $\partial_{ij}^2 \lambda_h = 0$ for $i, j = 1, \dots, d$ on each element $T \in \mathcal{T}_h$ if $\tau_2 > 0$, which shows that $c(\lambda_h, \sigma) = 0$ for all $\sigma \in M_h$.

Using (2.7), (2.3) and the usual integration by parts, we have

$$\begin{aligned}
 0 & = b(u_h, \sigma) \\
 & = \sum_{T \in \mathcal{T}_h} (a \nabla u_h + \mathbf{b}u_0, \nabla \sigma)_T - \langle u_n, \sigma \rangle_{\partial T} \\
 & = \sum_{T \in \mathcal{T}_h} (\nabla u_0, a \nabla \sigma)_T - \langle u_0 - u_b, a \nabla \sigma \cdot \mathbf{n} \rangle_{\partial T} - (\nabla \cdot (\mathbf{b}u_0), \sigma)_T \\
 & \quad + \langle \mathbf{b}u_0 \cdot \mathbf{n}, \sigma \rangle_{\partial T} - \langle u_n, \sigma \rangle_{\partial T} \\
 & = \sum_{T \in \mathcal{T}_h} -(\nabla \cdot (a \nabla u_0), \sigma)_T + \langle a \nabla u_0 \cdot \mathbf{n}, \sigma \rangle_{\partial T} - \langle u_0 - u_b, a \nabla \sigma \cdot \mathbf{n} \rangle_{\partial T} \\
 & \quad - (\nabla \cdot (\mathbf{b}u_0), \sigma)_T + \langle \mathbf{b}u_0 \cdot \mathbf{n}, \sigma \rangle_{\partial T} - \langle u_n, \sigma \rangle_{\partial T}
 \end{aligned}$$

$$\begin{aligned}
&= \sum_{T \in \mathcal{T}_h} -(\nabla \cdot (a \nabla u_0 + \mathbf{b} u_0), \sigma)_T - \langle u_0 - u_b, a \nabla \sigma \cdot \mathbf{n} \rangle_{\partial T} \\
&\quad + \langle (a \nabla u_0 + \mathbf{b} u_0) \cdot \mathbf{n} - u_n, \sigma \rangle_{\partial T} \\
&= \sum_{T \in \mathcal{T}_h} -(\nabla \cdot (a \nabla u_0 + \mathbf{b} u_0), \sigma)_T,
\end{aligned}$$

where we used $u_0 = u_b$ and $(a \nabla u_0 + \mathbf{b} u_0) \cdot \mathbf{n} = u_n$ on each ∂T . This gives $\nabla \cdot (a \nabla u_0 + \mathbf{b} u_0) = 0$ on each element $T \in \mathcal{T}_h$ by taking $\sigma = \nabla \cdot (a \nabla u_0 + \mathbf{b} u_0)$. From $(a \nabla u_0 + \mathbf{b} u_0) \cdot \mathbf{n} = u_n$ on each ∂T , and $a \nabla u_0 + \mathbf{b} u_0 \in H(\text{div}; T)$, we obtain $a \nabla u_0 + \mathbf{b} u_0 \in H(\text{div}; \Omega)$ and further $\nabla \cdot (a \nabla u_0 + \mathbf{b} u_0) = 0$ in Ω . Using $g_1 = 0$ on Γ_D and $u_0 = u_b$ on each ∂T , gives $u_0 = 0$ on Γ_D . Using $g_2 = 0$ on Γ_N and $(a \nabla u_0 + \mathbf{b} u_0) \cdot \mathbf{n} = u_n$ on each ∂T , yields $(a \nabla u_0 + \mathbf{b} u_0) \cdot \mathbf{n} = 0$ on Γ_N . Therefore, from the solution uniqueness of the PDE problem, we have $u_0 \equiv 0$ in Ω . We further obtain $u_b \equiv 0$, $u_n \equiv 0$ and thus $u_h \equiv 0$ in Ω .

From $u_h \equiv 0$ in Ω , (2.6) can be simplified as follows

$$b(v, \lambda_h) = 0, \quad \forall v \in W_h^0.$$

From Lemma 3.1, there exists a $v \in W_h^0$, satisfying

$$0 = b(v, \lambda_h) \geq \begin{cases} \frac{1}{2} \|\lambda_h\|^2, & l = k, \\ \frac{1}{2} \|\lambda_h\|^2 - \beta h^2 \|\nabla \lambda_h\|^2, & k = 1, l = k - 1, \\ \frac{1}{2} \|\lambda_h\|^2 - \beta h^4 |\lambda_h|_2^2, & k \geq 2, l = k - 1, \end{cases} \quad (3.9)$$

for some constant $\beta > 0$. For the case of $l = k$, it follows from (3.9) that $\lambda_h \equiv 0$ in Ω . Note that when $l = k - 1$ and $k = 1$, we take $\tau_1 > 0$ and $\tau_2 = 0$; when $l = k - 1$ and $k \geq 2$, we take $\tau_1 = 0$ and $\tau_2 > 0$. Thus, for the case of $l = k - 1$, using $c(\lambda_h, \lambda_h) \equiv 0$ gives $\nabla \lambda_h = 0$ on each $T \in \mathcal{T}_h$ for $k = 1$; and $\partial_{ij}^2 \lambda_h = 0$ for any $i, j = 1, \dots, d$ on each $T \in \mathcal{T}_h$ for $k \geq 2$, which, combined with (3.9), yields $\lambda_h \equiv 0$ in Ω for the case of $l = k - 1$. This completes the proof of this theorem. $\square$

4. Mass conservation

The first equation in the convection-diffusion model problem (1.1) can be rewritten in a conservative form; i.e.,

$$-\nabla \cdot \mathbf{F} = f, \quad (4.1)$$

$$\mathbf{F} = a \nabla u + \mathbf{b} u. \quad (4.2)$$

On each element $T \in \mathcal{T}_h$, integrating (4.1) over T gives the integral formulation of the mass conservation; i.e.,

$$-\int_{\partial T} \mathbf{F} \cdot \mathbf{n} ds = \int_T f dT. \quad (4.3)$$

We claim that the numerical solution arising from the primal-dual weak Galerkin scheme (2.6)-(2.7) for the convection-diffusion model problem (1.1) retains the mass conservation property (4.3) locally on each element $T \in \mathcal{T}_h$ with a numerical flux $\mathbf{F}_h$. To this end, for any given element $T \in \mathcal{T}_h$, choosing the test function σ in (2.7) such that $\sigma = 1$ on T and $\sigma = 0$ elsewhere, yields

$$\begin{aligned}
& -\tau_1 h_T^2 (\nabla \lambda_h, \nabla 1)_T - \tau_2 h_T^4 \sum_{i,j=1}^d (\partial_{ij}^2 \lambda_h, \partial_{ij}^2 1)_T + (a \nabla_w u_h + \mathbf{b} u_0, \nabla 1)_T - \langle u_n, 1 \rangle_{\partial T} \\
&= (f, 1)_T,
\end{aligned}$$

which can be simplified as follows

$$-\langle u_n \mathbf{n} \cdot \mathbf{n}, 1 \rangle_{\partial T} = (f, 1)_T.$$

This implies that the primal-dual weak Galerkin algorithm (2.6)-(2.7) conserves mass with a numerical flux given by

$$\mathbf{F}_h|_{\partial T} = u_n \mathbf{n}.$$

It is easy to check that

$$\mathbf{F}_h|_{\partial T_1} \cdot \mathbf{n}_{T_1} + \mathbf{F}_h|_{\partial T_2} \cdot \mathbf{n}_{T_2} = 0, \quad \text{on } e = \partial T_1 \cap \partial T_2,$$

where $\mathbf{n}_{T_1}$ and $\mathbf{n}_{T_2}$ are the unit outward normal directions along the interior edge or flat face $e = \partial T_1 \cap \partial T_2$ pointing exterior to T_1 and T_2 , respectively. This indicates the continuity of the numerical flux $\mathbf{F}_h$ along the normal direction on each interior edge or flat face $e \in \mathcal{E}_h^0$.

The result can be summarized as follows.

Theorem 4.1. Let $(u_h = \{u_0, u_b, u_n\}; \lambda_h)$ be the numerical solution of the convection-diffusion model problem (1.1) arising from the primal-dual weak Galerkin finite element method (2.6)–(2.7). Define a numerical flux function as follows:

$$\mathbf{F}_h|_{\partial T} := u_n \mathbf{n}, \quad \text{on } \partial T, \quad T \in \mathcal{T}_h.$$

Then, the numerical flux approximation $\mathbf{F}_h$ is continuous across each interior edge or flat face $e \in \mathcal{E}_h^0$ in the normal direction, and satisfies the following mass conservation property; i.e.,

$$-\int_{\partial T} \mathbf{F}_h \cdot \mathbf{n} ds = \int_T f dT.$$

5. Error equations

Let u and $(u_h; \lambda_h) \in W_h \times M_h$ be the exact solution of (1.1) and the PDWG solution arising from the numerical scheme (2.6)–(2.7), respectively. Denote by $\mathcal{Q}_h^{(k)}$ the L^2 projection onto the finite element space M_h . Note that the exact solution of the Lagrange multiplier λ is 0. Define two error functions by

$$e_h = u_h - Q_h u, \quad (5.1)$$

$$\epsilon_h = \lambda_h - \mathcal{Q}_h^{(k)} \lambda = \lambda_h. \quad (5.2)$$

Lemma 5.1. The error functions e_h and ϵ_h defined in (5.1)–(5.2) satisfy the following error equations for the primal-dual WG finite element scheme (2.6)–(2.7); i.e.,

$$s(e_h, v) + b(v, \epsilon_h) = -s(Q_h u, v), \quad \forall v \in W_h^0, \quad (5.3)$$

$$-c(\epsilon_h, \sigma) + b(e_h, \sigma) = \ell_u(\sigma), \quad \forall \sigma \in M_h, \quad (5.4)$$

where

$$\ell_u(\sigma) = \begin{cases} 0, & l = k, \\ \sum_{T \in \mathcal{T}_h} \langle (Q_n^{(l)} - I)((a \nabla u + \mathbf{b}u) \cdot \mathbf{n}), \sigma \rangle_{\partial T}, & l = k - 1. \end{cases} \quad (5.5)$$

Proof. Note that the exact solution of the Lagrange multiplier λ is 0. Subtracting $s(Q_h u, v)$ from both sides of (2.6) yields

$$s(u_h - Q_h u, v) + b(v, \lambda_h - \mathcal{Q}_h^{(k)} \lambda) = -s(Q_h u, v), \quad \forall v \in W_h^0.$$

This completes the proof of (5.3). Next, for any $\sigma \in M_h$, we have

$$\begin{aligned} b(Q_h u, \sigma) &= \sum_{T \in \mathcal{T}_h} (a \nabla_w Q_h u + \mathbf{b} Q_0^{(k)} u, \nabla \sigma)_T - \langle Q_n^{(l)} ((a \nabla u + \mathbf{b}u) \cdot \mathbf{n}), \sigma \rangle_{\partial T} \\ &= \sum_{T \in \mathcal{T}_h} (a \mathcal{Q}_h^{(k-1)} \nabla u + \mathbf{b} Q_0^{(k)} u, \nabla \sigma)_T - \langle Q_n^{(l)} ((a \nabla u + \mathbf{b}u) \cdot \mathbf{n}), \sigma \rangle_{\partial T} \\ &= \sum_{T \in \mathcal{T}_h} (a \nabla u + \mathbf{b}u, \nabla \sigma)_T - \langle Q_n^{(l)} ((a \nabla u + \mathbf{b}u) \cdot \mathbf{n}), \sigma \rangle_{\partial T} \\ &= \sum_{T \in \mathcal{T}_h} -(\nabla \cdot (a \nabla u + \mathbf{b}u), \sigma)_T + \langle (a \nabla u + \mathbf{b}u) \cdot \mathbf{n}, \sigma \rangle_{\partial T} \\ &\quad - \langle Q_n^{(l)} ((a \nabla u + \mathbf{b}u) \cdot \mathbf{n}), \sigma \rangle_{\partial T} \\ &= \sum_{T \in \mathcal{T}_h} (f, \sigma)_T - \sum_{T \in \mathcal{T}_h} \langle (Q_n^{(l)} - I)((a \nabla u + \mathbf{b}u) \cdot \mathbf{n}), \sigma \rangle_{\partial T}, \end{aligned}$$

where we have used the operator identity (3.1), the usual integration by parts, and the first equation of (1.1). Note that for the case of $l = k$, we have $\sum_{T \in \mathcal{T}_h} \langle (Q_n^{(l)} - I)((a \nabla u + \mathbf{b}u) \cdot \mathbf{n}), \sigma \rangle_{\partial T} = 0$. Combining the above with (2.7) yields (5.4). This completes the proof of the lemma. $\square$

6. Residual error estimates

Recall that $\mathcal{T}_h$ is a shape-regular finite element partition of the domain Ω . For any $T \in \mathcal{T}_h$ and $\varphi \in H^1(T)$, the following trace inequality holds true [13]:

$$\|\varphi\|_{\partial T}^2 \leq C(h_T^{-1}\|\varphi\|_T^2 + h_T\|\nabla\varphi\|_T^2). \quad (6.1)$$

If φ is a polynomial on the element $T \in \mathcal{T}_h$, then from the inverse inequality (see also [13]) we have

$$\|\varphi\|_{\partial T}^2 \leq Ch_T^{-1}\|\varphi\|_T^2. \quad (6.2)$$

Lemma 6.1. [13] Let $\mathcal{T}_h$ be a finite element partition of the domain Ω satisfying the shape regularity assumptions given in [13]. Then, for any $0 \leq p \leq 2$, $1 \leq m \leq k$, one has

$$\sum_{T \in \mathcal{T}_h} h_T^{2p} \|u - Q_0^{(m)} u\|_{p,T}^2 \leq Ch^{2(m+1)} \|u\|_{m+1}^2, \quad (6.3)$$

$$\sum_{T \in \mathcal{T}_h} h_T^{2p} \|\nabla u - Q_h^{(m-1)} \nabla u\|_{p,T}^2 \leq Ch^{2m} \|u\|_{m+1}^2, \quad (6.4)$$

$$\sum_{T \in \mathcal{T}_h} h_T^{2p} \|u - Q_h^{(m)} u\|_{p,T}^2 \leq Ch^{2(m+1)} \|u\|_{m+1}^2. \quad (6.5)$$

Theorem 6.2. Let u and $(u_h; \lambda_h) \in W_h \times M_h$ be the exact solution of (1.1) and PDWG solution of (2.6)-(2.7), respectively. Assume that the exact solution u of (1.1) is sufficiently regular such that $u \in H^{k+1}(\Omega)$. Then, there exists a constant C such that the following error estimate holds true:

$$\|u_h - Q_h u\|_{W_h} + \|\lambda_h - Q_h^{(k)} \lambda\|_{M_h} \leq \begin{cases} Ch^{k-1} \|u\|_{k+1}, & l = k, \\ C(1 + \tau_1^{-\frac{1}{2}})h^{k-1} \|u\|_{k+1}, & k = 1, l = k-1, \\ C(1 + \tau_2^{-\frac{1}{2}})h^{k-1} \|u\|_{k+1}, & k \geq 2, l = k-1. \end{cases} \quad (6.6)$$

Proof. By choosing $v = e_h$ and $\sigma = \epsilon_h$ in (5.3)-(5.4), we have from the difference of (5.3) and (5.4) that

$$s(e_h, e_h) + c(\epsilon_h, \epsilon_h) = -s(Q_h u, e_h) - \ell_u(\epsilon_h). \quad (6.7)$$

Recall that

$$\begin{aligned} & s(Q_h u, e_h) \\ &= \sum_{T \in \mathcal{T}_h} h_T^{-3} \langle Q_0^{(k)} u - Q_b^{(k)} u, e_0 - e_b \rangle_{\partial T} + \sum_{T \in \mathcal{T}_h} h_T^{-1} \langle (a \nabla Q_0^{(k)} u + \mathbf{b} Q_0^{(k)} u) \cdot \mathbf{n} \\ & \quad - Q_n^{(l)} ((a \nabla u + \mathbf{b} u) \cdot \mathbf{n}), (a \nabla e_0 + \mathbf{b} e_0) \cdot \mathbf{n} - e_n \rangle_{\partial T}. \end{aligned} \quad (6.8)$$

The first term on the right-hand side of (6.8) can be estimated by using the Cauchy-Schwarz inequality, the trace inequality (6.1), and the estimate (6.3) with $m = k$ as follows

$$\begin{aligned} & \left| \sum_{T \in \mathcal{T}_h} h_T^{-3} \langle Q_0^{(k)} u - Q_b^{(k)} u, e_0 - e_b \rangle_{\partial T} \right| \\ &= \left| \sum_{T \in \mathcal{T}_h} h_T^{-3} \langle Q_0^{(k)} u - u, e_0 - e_b \rangle_{\partial T} \right| \\ &\leq \left(\sum_{T \in \mathcal{T}_h} h_T^{-3} \|u - Q_0^{(k)} u\|_{\partial T}^2 \right)^{\frac{1}{2}} \left(\sum_{T \in \mathcal{T}_h} h_T^{-3} \|e_0 - e_b\|_{\partial T}^2 \right)^{\frac{1}{2}} \\ &\leq C \left(\sum_{T \in \mathcal{T}_h} h_T^{-4} \|u - Q_0^{(k)} u\|_T^2 + h_T^{-2} \|u - Q_0^{(k)} u\|_{1,T}^2 \right)^{\frac{1}{2}} \|e_h\|_{W_h} \\ &\leq Ch^{k-1} \|u\|_{k+1} \|e_h\|_{W_h}. \end{aligned} \quad (6.9)$$

Similarly, the second term on the right-hand side of (6.8) has the following estimate

$$\left| \sum_{T \in \mathcal{T}_h} h_T^{-1} \langle (a \nabla Q_0^{(k)} u + \mathbf{b} Q_0^{(k)} u) \cdot \mathbf{n} - Q_n^{(l)} ((a \nabla u + \mathbf{b} u) \cdot \mathbf{n}), (a \nabla e_0 + \mathbf{b} e_0) \cdot \mathbf{n} - e_n \rangle_{\partial T} \right| \leq C h^{k-1} \|u\|_{k+1} \|e_h\|_{W_h}. \quad (6.10)$$

Substituting (6.9) and (6.10) into (6.8) gives

$$|s(Q_h u, e_h)| \leq C h^{k-1} \|u\|_{k+1} \|e_h\|_{W_h}. \quad (6.11)$$

We shall further discuss the second term on the right-hand side of (6.7). For the case of $l = k$, from (5.5), we have

$$\ell_u(\epsilon_h) = 0. \quad (6.12)$$

We now consider the case of $l = k - 1$. By denoting

$$F_u = a \nabla u + \mathbf{b} u,$$

and then using (5.5), the Cauchy-Schwarz inequality, the trace inequality (6.1), and the estimate (6.3) with $m = l = k - 1$, we have

$$\begin{aligned} |\ell_u(\epsilon_h)| &= |\ell_u(\epsilon_h - I_h^l \epsilon_h)| = \left| \sum_{T \in \mathcal{T}_h} \langle (Q_n^{(l)} - I)(F_u \cdot \mathbf{n}), \epsilon_h - I_h^l \epsilon_h \rangle_{\partial T} \right| \\ &\leq \left(\sum_{T \in \mathcal{T}_h} \|(Q_n^{(l)} - I)(F_u \cdot \mathbf{n})\|_{\partial T}^2 \right)^{\frac{1}{2}} \left(\sum_{T \in \mathcal{T}_h} \|\epsilon_h - I_h^l \epsilon_h\|_{\partial T}^2 \right)^{\frac{1}{2}} \\ &\leq C \left(\sum_{T \in \mathcal{T}_h} h_T^{-1} \|(Q_0^{(l)} - I)F_u\|_T^2 + h_T \|(Q_0^{(l)} - I)F_u\|_{1,T}^2 \right)^{\frac{1}{2}} \\ &\quad \left(\sum_{T \in \mathcal{T}_h} h_T^{-1} \|\epsilon_h - I_h^l \epsilon_h\|_T^2 + h_T \|\nabla(\epsilon_h - I_h^l \epsilon_h)\|_T^2 \right)^{\frac{1}{2}} \\ &\leq C h^l \|F_u\|_{l+1} \left(\sum_{T \in \mathcal{T}_h} \|\epsilon_h - I_h^l \epsilon_h\|_T^2 + h_T^2 \|\nabla(\epsilon_h - I_h^l \epsilon_h)\|_T^2 \right)^{\frac{1}{2}}, \end{aligned} \quad (6.13)$$

where $I_h^l \epsilon_h$ denotes the cell average and linear interpolation of ϵ_h on each element $T \in \mathcal{T}_h$ for $l = 0$ and $l \geq 1$, respectively. Choosing $l = k - 1$ in the above inequality and using the approximation property of the interpolation function yields

$$|\ell_u(\epsilon_h)| \leq \begin{cases} C \tau_1^{-\frac{1}{2}} h^{k-1} \|u\|_{k+1} \|e_h\|_{M_h}, & k = 1, l = k - 1, \\ C \tau_2^{-\frac{1}{2}} h^{k-1} \|u\|_{k+1} \|e_h\|_{M_h}, & k \geq 2, l = k - 1. \end{cases} \quad (6.14)$$

Substituting (6.11), (6.12), and (6.14) into (6.7) gives the error estimate (6.6). This completes the proof of the theorem. $\square$

Theorem 6.3. Under the assumption of Theorem 6.2, there exists a constant C such that the following error estimate holds true:

$$\left(\sum_{T \in \mathcal{T}_h} \|\nabla \cdot (a \nabla e_0 + \mathbf{b} e_0)\|_T^2 \right)^{\frac{1}{2}} \leq \begin{cases} C h^{k-1} \|u\|_{k+1}, & l = k, \\ C(1 + \tau_1^{\frac{1}{2}})(1 + \tau_1^{-\frac{1}{2}}) h^{k-1} \|u\|_{k+1}, & k = 1, l = k - 1, \\ C(1 + \tau_2^{\frac{1}{2}})(1 + \tau_2^{-\frac{1}{2}}) h^{k-1} \|u\|_{k+1}, & k \geq 2, l = k - 1. \end{cases} \quad (6.15)$$

Proof. From the error equation (5.4) we have

$$b(e_h, \sigma) = c(\epsilon_h, \sigma) + \ell_u(\sigma), \quad \forall \sigma \in M_h. \quad (6.16)$$

Recall that

$$\begin{aligned}
 b(e_h, \sigma) &= \sum_{T \in \mathcal{T}_h} (a \nabla_w e_h + \mathbf{b} e_0, \nabla \sigma)_T - \langle e_n, \sigma \rangle_{\partial T} \\
 &= \sum_{T \in \mathcal{T}_h} (a \nabla e_0 + \mathbf{b} e_0, \nabla \sigma)_T + \langle e_b - e_0, a \nabla \sigma \cdot \mathbf{n} \rangle_{\partial T} - \langle e_n, \sigma \rangle_{\partial T} \\
 &= - \sum_{T \in \mathcal{T}_h} (\nabla \cdot (a \nabla e_0 + \mathbf{b} e_0), \sigma)_T - \langle e_b - e_0, a \nabla \sigma \cdot \mathbf{n} \rangle_{\partial T} \\
 &\quad + \langle e_n - (a \nabla e_0 + \mathbf{b} e_0) \cdot \mathbf{n}, \sigma \rangle_{\partial T},
 \end{aligned} \tag{6.17}$$

where we have used (2.3) with $\psi = a \nabla \sigma$ and the usual integration by parts. Substituting (6.17) into (6.16) gives

$$\begin{aligned}
 &- \sum_{T \in \mathcal{T}_h} (\nabla \cdot (a \nabla e_0 + \mathbf{b} e_0), \sigma)_T \\
 &= c(\epsilon_h, \sigma) + \ell_u(\sigma) + \sum_{T \in \mathcal{T}_h} \langle e_0 - e_b, a \nabla \sigma \cdot \mathbf{n} \rangle_{\partial T} + \langle e_n - (a \nabla e_0 + \mathbf{b} e_0) \cdot \mathbf{n}, \sigma \rangle_{\partial T} \\
 &= J_1 + J_2 + J_3 + J_4,
 \end{aligned} \tag{6.18}$$

where J_i is defined accordingly for $i = 1, \dots, 4$.

We shall estimate each term J_i in (6.18) respectively. With J_1 , we have for the case of $l = k$, $J_1 = 0$. For the case of $l = k - 1$ and $k = 1$, we have

$$\begin{aligned}
 J_1 &= \tau_1 \sum_{T \in \mathcal{T}_h} h_T^2 (\nabla \epsilon_h, \nabla \sigma)_T \\
 &\leq \left(\sum_{T \in \mathcal{T}_h} \tau_1 h_T^2 \|\nabla \epsilon_h\|_T^2 \right)^{\frac{1}{2}} \left(\sum_{T \in \mathcal{T}_h} \tau_1 h_T^2 \|\nabla \sigma\|_T^2 \right)^{\frac{1}{2}} \\
 &\leq C \tau_1^{\frac{1}{2}} \|\epsilon_h\|_{M_h} \left(\sum_{T \in \mathcal{T}_h} \|\sigma\|_T^2 \right)^{\frac{1}{2}},
 \end{aligned}$$

where we have used the Cauchy-Schwarz inequality and the inverse inequality. Similarly, for the case of $l = k - 1$ and $k \geq 2$, we have

$$J_1 \leq C \tau_2^{\frac{1}{2}} \|\epsilon_h\|_{M_h} \left(\sum_{T \in \mathcal{T}_h} \|\sigma\|_T^2 \right)^{\frac{1}{2}}.$$

As to the term J_2 , we have from (5.5) that for $l = k$, $\ell_u(\sigma) = 0$; for $l = k - 1$, we have, by following the same argument as that in (6.13)

$$\begin{aligned}
 |J_2| &= |\ell_u(\sigma)| \leq \left(\sum_{T \in \mathcal{T}_h} \|(Q_n^{(l)} - I)((a \nabla u + \mathbf{b} u) \cdot \mathbf{n})\|_{\partial T}^2 \right)^{\frac{1}{2}} \left(\sum_{T \in \mathcal{T}_h} \|\sigma\|_{\partial T}^2 \right)^{\frac{1}{2}} \\
 &\leq C h^{k-1} \|u\|_{k+1} \left(\sum_{T \in \mathcal{T}_h} \|\sigma\|_T^2 \right)^{\frac{1}{2}},
 \end{aligned}$$

where we used the Cauchy-Schwarz inequality, the estimate (6.3) with $m = l = k - 1$, and the trace inequalities (6.1) and (6.2). As to the term J_3 , we have

$$\begin{aligned}
 J_3 &\leq \left(\sum_{T \in \mathcal{T}_h} h_T^{-3} \|e_0 - e_b\|_{\partial T}^2 \right)^{\frac{1}{2}} \left(\sum_{T \in \mathcal{T}_h} h_T^3 \|a \nabla \sigma \cdot \mathbf{n}\|_{\partial T}^2 \right)^{\frac{1}{2}} \\
 &\leq C \|e_h\|_{W_h} \left(\sum_{T \in \mathcal{T}_h} h_T^2 \|a \nabla \sigma \cdot \mathbf{n}\|_T^2 \right)^{\frac{1}{2}} \\
 &\leq C \|e_h\|_{W_h} \left(\sum_{T \in \mathcal{T}_h} \|\sigma\|_T^2 \right)^{\frac{1}{2}},
 \end{aligned}$$

where we used the Cauchy-Schwarz inequality, the trace inequality (6.2) and the inverse inequality.

For the last term J_4 , we have

$$\begin{aligned} & \sum_{T \in \mathcal{T}_h} \langle e_n - (a \nabla e_0 + \mathbf{b} e_0) \cdot \mathbf{n}, \sigma \rangle_{\partial T} \\ & \leq \left(\sum_{T \in \mathcal{T}_h} h_T^{-1} \|e_n - (a \nabla e_0 + \mathbf{b} e_0) \cdot \mathbf{n}\|_{\partial T}^2 \right)^{\frac{1}{2}} \left(\sum_{T \in \mathcal{T}_h} h_T \|\sigma\|_{\partial T}^2 \right)^{\frac{1}{2}} \\ & \leq C \|e_h\|_{W_h} \left(\sum_{T \in \mathcal{T}_h} \|\sigma\|_T^2 \right)^{\frac{1}{2}}, \end{aligned}$$

where we used the Cauchy-Schwarz inequality and the trace inequality (6.2).

Substituting the above estimates for $J_i (i = 1, \dots, 4)$ into (6.18) and combining with (6.6) completes the proof of (6.15). $\square$

7. Error estimates in H^1 and L^2

Consider the dual problem of seeking an unknown function w such that

$$-\nabla \cdot (a \nabla w) + \mathbf{b} \cdot \nabla w = e_0, \quad \text{in } \Omega, \quad (7.1)$$

$$w = 0, \quad \text{on } \Gamma_D, \quad (7.2)$$

$$a \nabla w \cdot \mathbf{n} = 0, \quad \text{on } \Gamma_N, \quad (7.3)$$

for any given $e_0 \in L^2(\Omega)$. The problem (7.1)-(7.3) is said to be $H^{1+s} (\frac{1}{2} < s \leq 1)$ -regular in the sense that

$$\|w\|_{1+s} \leq C \|e_0\|. \quad (7.4)$$

Lemma 7.1. Let $e_h = \{e_0, e_b, e_n\}$ be the error function defined in (5.1). There holds

$$\|\nabla_w e_h - \nabla e_0\|_T \leq C h_T^{-\frac{1}{2}} \|e_0 - e_b\|_{\partial T}. \quad (7.5)$$

Proof. From (2.3), we have

$$(\nabla_w e_h - \nabla e_0, \boldsymbol{\psi})_T = -\langle e_0 - e_b, \boldsymbol{\psi} \cdot \mathbf{n} \rangle_{\partial T}, \quad \forall \boldsymbol{\psi} \in [P_{k-1}(T)]^d.$$

From the Cauchy-Schwarz inequality and the trace inequality (6.2), we thus have

$$\begin{aligned} \|\nabla_w e_h - \nabla e_0\|_T & \leq \sup_{\forall \boldsymbol{\psi} \in [P_{k-1}(T)]^d} \frac{\|e_0 - e_b\|_{\partial T} \|\boldsymbol{\psi} \cdot \mathbf{n}\|_{\partial T}}{\|\boldsymbol{\psi}\|_T} \\ & \leq \sup_{\forall \boldsymbol{\psi} \in [P_{k-1}(T)]^d} \frac{C h_T^{-\frac{1}{2}} \|e_0 - e_b\|_{\partial T} \|\boldsymbol{\psi}\|_T}{\|\boldsymbol{\psi}\|_T} \\ & \leq C h_T^{-\frac{1}{2}} \|e_0 - e_b\|_{\partial T}. \end{aligned}$$

This completes the proof of the lemma. $\square$

The following theorem presents the error estimate in the usual L^2 norm for the first component u_0 in the primal variable $u_h = \{u_0, u_b, u_n\}$ of the PDWG solution arising from the numerical scheme (2.6)-(2.7).

Theorem 7.2. Assume that the dual problem (7.1)-(7.3) has the H^{1+s} -regularity with a priori estimate (7.4) for $s \in (\frac{1}{2}, 1]$. There exists a constant C such that

$$\|e_0\| \leq \begin{cases} C h^{k+s} \|u\|_{k+1}, & l = k, \\ C (1 + \tau_1^{\frac{1}{2}} (1 + \tau_1^{-\frac{1}{2}})) h^k \|u\|_{k+1}, & k = 1, l = k - 1, \\ C (1 + \tau_2^{\frac{1}{2}} (1 + \tau_2^{-\frac{1}{2}})) h^{k+s} \|u\|_{k+1}, & k \geq 2, l = k - 1, \end{cases} \quad (7.6)$$

provided that the meshsize h is sufficiently small.

Proof. Testing (7.1) with e_0 on each element $T \in \mathcal{T}_h$, we obtain from the usual integration by parts that

$$\begin{aligned} \|e_0\|^2 &= \sum_{T \in \mathcal{T}_h} (-\nabla \cdot (a \nabla w) + \mathbf{b} \cdot \nabla w, e_0)_T \\ &= \sum_{T \in \mathcal{T}_h} (a \nabla w, \nabla e_0)_T - \langle a \nabla w \cdot \mathbf{n}, e_0 \rangle_{\partial T} + (\mathbf{b} \cdot \nabla w, e_0)_T \\ &= \sum_{T \in \mathcal{T}_h} (a \nabla w, \nabla e_0)_T - \langle a \nabla w \cdot \mathbf{n}, e_0 - e_b \rangle_{\partial T} + (\mathbf{b} \cdot \nabla w, e_0)_T, \end{aligned} \quad (7.7)$$

where we used $\sum_{T \in \mathcal{T}_h} \langle a \nabla w \cdot \mathbf{n}, e_b \rangle_{\partial T} = \langle a \nabla w \cdot \mathbf{n}, e_b \rangle_{\partial \Omega} = 0$ due to the facts that $a \nabla w \cdot \mathbf{n} = 0$ on Γ_N and $e_b = 0$ on Γ_D . It follows from (2.3) and (3.1) that

$$\begin{aligned} (a \nabla_w e_h, \nabla_w (Q_h w))_T &= (a \nabla_w e_h, Q_h^{(k-1)} \nabla w)_T \\ &= (a \nabla e_0, Q_h^{(k-1)} \nabla w)_T - \langle e_0 - e_b, a Q_h^{(k-1)} \nabla w \cdot \mathbf{n} \rangle_{\partial T} \\ &= (a \nabla e_0, \nabla w)_T - \langle e_0 - e_b, a Q_h^{(k-1)} \nabla w \cdot \mathbf{n} \rangle_{\partial T}, \end{aligned}$$

which gives

$$\begin{aligned} (a \nabla e_0, \nabla w)_T &= (a \nabla_w e_h, Q_h^{(k-1)} \nabla w)_T + \langle e_0 - e_b, a Q_h^{(k-1)} \nabla w \cdot \mathbf{n} \rangle_{\partial T} \\ &= (a \nabla_w e_h, \nabla w)_T + \langle e_0 - e_b, a Q_h^{(k-1)} \nabla w \cdot \mathbf{n} \rangle_{\partial T}. \end{aligned} \quad (7.8)$$

Substituting (7.8) into (7.7) leads to

$$\begin{aligned} \|e_0\|^2 &= \sum_{T \in \mathcal{T}_h} (a \nabla_w e_h, \nabla w)_T + (\mathbf{b} e_0, \nabla w)_T + \langle e_0 - e_b, a(Q_h^{(k-1)} - I) \nabla w \cdot \mathbf{n} \rangle_{\partial T} \\ &= b(e_h, w) + \sum_{T \in \mathcal{T}_h} \langle e_n, w \rangle_{\partial T} + \langle e_0 - e_b, a(Q_h^{(k-1)} - I) \nabla w \cdot \mathbf{n} \rangle_{\partial T} \\ &= c(\epsilon_h, Q_h^{(k)} w) + \ell_u(Q_h^{(k)} w) + b(e_h, (I - Q_h^{(k)}) w) \\ &\quad + \sum_{T \in \mathcal{T}_h} \langle e_0 - e_b, a(Q_h^{(k-1)} - I) \nabla w \cdot \mathbf{n} \rangle_{\partial T} \\ &= I_1 + I_2 + I_3 + I_4, \end{aligned} \quad (7.9)$$

where in the second step, we have used the fact that $\sum_{T \in \mathcal{T}_h} \langle e_n, w \rangle_{\partial T} = \langle e_n, w \rangle_{\partial \Omega} = 0$ due to the facts that $w = 0$ on Γ_D and $e_n = 0$ on Γ_N , and $I_i (i = 1, \dots, 4)$ is defined accordingly.

We shall estimate each of the four terms I_i for $i = 1, \dots, 4$ in (7.9). As to the term I_1 , for the case of $l = k$ where $\tau_1 = 0$ and $\tau_2 = 0$, we have

$$I_1 = c(\epsilon_h, Q_h^{(k)} w) = 0. \quad (7.10)$$

For the case of $l = k - 1$, we have, from the Cauchy-Schwarz inequality and (6.6),

$$\begin{aligned} I_1 &= \tau_1 \sum_{T \in \mathcal{T}_h} h_T^2 (\nabla \epsilon_h, \nabla Q_h^{(k)} w)_T + \tau_2 \sum_{i,j=1}^d \sum_{T \in \mathcal{T}_h} h_T^4 (\partial_{ij}^2 \epsilon_h, \partial_{ij}^2 (Q_h^{(k)} w))_T \\ &\leq \| \epsilon_h \|_{M_h} \left(\sum_{T \in \mathcal{T}_h} \left(\tau_1 h_T^2 \| \nabla Q_h^{(k)} w \|_T^2 + \tau_2 \sum_{i,j=1}^d h_T^4 \| \partial_{ij}^2 Q_h^{(k)} w \|_T^2 \right) \right)^{\frac{1}{2}} \\ &\leq \begin{cases} Ch \tau_1^{\frac{1}{2}} \| \epsilon_h \|_{M_h} \| w \|_1, & k = 1, l = k - 1, \\ Ch^{1+s} \tau_2^{\frac{1}{2}} \| \epsilon_h \|_{M_h} \| w \|_{1+s}, & k \geq 2, l = k - 1. \end{cases} \end{aligned} \quad (7.11)$$

Here in the last step, we have used the fact that $\tau_2 = 0$ for $k = 1$ and $\tau_1 = 0$ for $k \geq 2$, and the inverse inequality

$$|Q_h^{(k)} w|_2 \leq Ch^{s-1} |Q_h^{(k)} w|_{1+s} \leq Ch^{s-1} \| w \|_{1+s}, \quad \frac{1}{2} < s \leq 1.$$

As to the term I_2 , for the case of $l = k$, we have from (5.5) that

$$I_2 = \ell_u(Q_h^{(k)} w) = 0.$$

For the case of $l = k - 1$, by the same argument as what we did in (6.13), we have

$$\begin{aligned}
 |I_2| &= |\ell_u(\mathcal{Q}_h^{(k)} w)| \\
 &\leq Ch^{k-1} \|u\|_{k+1} \left(\sum_{T \in \mathcal{T}_h} \|(I - I_h^l) \mathcal{Q}_h^{(k)} w\|_T^2 + h_T^2 \|\nabla((I - I_h^l) \mathcal{Q}_h^{(k)} w)\|_T^2 \right)^{\frac{1}{2}} \\
 &\leq \begin{cases} Ch^k \|u\|_{k+1} \|w\|_1, & k = 1, l = k - 1, \\ Ch^{k+s} \|u\|_{k+1} \|w\|_{1+s}, & k \geq 2, l = k - 1. \end{cases}
 \end{aligned} \tag{7.12}$$

Here for any function v , $I_h^l v$ denotes the cell average and linear interpolation of v on each element $T \in \mathcal{T}_h$ for $l = 0$ and $l \geq 1$, respectively.

To estimate I_3 , we note that

$$\begin{aligned}
 I_3 &= \sum_{T \in \mathcal{T}_h} \langle (a \nabla_w e_h + \mathbf{b} e_0) \cdot \mathbf{n} - e_n, (I - \mathcal{Q}_h^{(k)}) w \rangle_{\partial T} \\
 &\quad - \sum_{T \in \mathcal{T}_h} \langle \nabla \cdot (a \nabla_w e_h + \mathbf{b} e_0), (I - \mathcal{Q}_h^{(k)}) w \rangle_T \\
 &= I_{31} - I_{32}.
 \end{aligned}$$

To estimate I_{31} , we have from the Cauchy-Schwarz inequality, the trace inequality (6.2), (7.5), the estimate (6.5) with $m = s$ that (with $F_e = a \nabla e_0 + \mathbf{b} e_0$)

$$\begin{aligned}
 |I_{31}| &= \sum_{T \in \mathcal{T}_h} \langle F_e \cdot \mathbf{n} - e_n, (I - \mathcal{Q}_h^{(k)}) w \rangle_{\partial T} + \langle (a \nabla_w e_h - a \nabla e_0) \cdot \mathbf{n}, (I - \mathcal{Q}_h^{(k)}) w \rangle_{\partial T} \\
 &\leq \left\{ \left(\sum_{T \in \mathcal{T}_h} \|F_e \cdot \mathbf{n} - e_n\|_{\partial T}^2 \right)^{\frac{1}{2}} + \left(\sum_{T \in \mathcal{T}_h} \|(a \nabla_w e_h - a \nabla e_0) \cdot \mathbf{n}\|_{\partial T}^2 \right)^{\frac{1}{2}} \right\} \\
 &\quad \cdot \left(\sum_{T \in \mathcal{T}_h} \|(I - \mathcal{Q}_h^{(k)}) w\|_{\partial T}^2 \right)^{\frac{1}{2}} \\
 &\leq C \left\{ \left(\sum_{T \in \mathcal{T}_h} \|F_e \cdot \mathbf{n} - e_n\|_{\partial T}^2 \right)^{\frac{1}{2}} + \left(\sum_{T \in \mathcal{T}_h} h_T^{-1} \|(a \nabla_w e_h - a \nabla e_0) \cdot \mathbf{n}\|_T^2 \right)^{\frac{1}{2}} \right\} h^{s+\frac{1}{2}} \|w\|_{1+s} \\
 &\leq C \left\{ h^{\frac{1}{2}} \|e_h\|_{W_h} + \left(\sum_{T \in \mathcal{T}_h} h_T^{-2} \|e_0 - e_b\|_{\partial T}^2 \right)^{\frac{1}{2}} \right\} h^{s+\frac{1}{2}} \|w\|_{1+s} \\
 &\leq Ch^s \|e_h\|_{W_h} \|w\|_{1+s}.
 \end{aligned}$$

Similarly, for the term I_{32} , we have from the Cauchy-Schwarz inequality, the estimate (6.5) with $m = s$, (7.5), the inverse inequality that

$$\begin{aligned}
 |I_{32}| &= \left| \sum_{T \in \mathcal{T}_h} \langle \nabla \cdot F_e, (I - \mathcal{Q}_h^{(k)}) w \rangle_T + \langle \nabla \cdot (a \nabla_w e_h - a \nabla e_0), (I - \mathcal{Q}_h^{(k)}) w \rangle_T \right| \\
 &\leq C \left\{ \left(\sum_{T \in \mathcal{T}_h} \|\nabla \cdot F_e\|_T^2 \right)^{\frac{1}{2}} + \left(\sum_{T \in \mathcal{T}_h} h_T^{-2} \|a \nabla_w e_h - a \nabla e_0\|_T^2 \right)^{\frac{1}{2}} \right\} h^{1+s} \|w\|_{1+s} \\
 &\leq C \left\{ \left(\sum_{T \in \mathcal{T}_h} \|\nabla \cdot F_e\|_T^2 \right)^{\frac{1}{2}} + \left(\sum_{T \in \mathcal{T}_h} h_T^{-3} \|e_0 - e_b\|_{\partial T}^2 \right)^{\frac{1}{2}} \right\} h^{1+s} \|w\|_{1+s} \\
 &\leq C \left\{ \left(\sum_{T \in \mathcal{T}_h} \|\nabla \cdot (a \nabla e_0 + \mathbf{b} e_0)\|_T^2 \right)^{\frac{1}{2}} + \|e_h\|_{W_h} \right\} h^{1+s} \|w\|_{1+s}.
 \end{aligned}$$

Consequently,

$$|I_3| \leq C \left\{ \left(\sum_{T \in \mathcal{T}_h} \|\nabla \cdot (a \nabla e_0 + \mathbf{b} e_0)\|_T^2 \right)^{\frac{1}{2}} + \|e_h\|_{W_h} \right\} h^{1+s} \|w\|_{1+s}.$$

As to the term I_4 , we have from the Cauchy-Schwarz inequality, trace inequality (6.1), the estimate (6.4) with $m = s$, that

$$\begin{aligned}
 |I_4| &\leq \left(\sum_{T \in \mathcal{T}_h} \|e_0 - e_b\|_{\partial T}^2 \right)^{\frac{1}{2}} \left(\sum_{T \in \mathcal{T}_h} \|a(Q_h^{(k-1)} - I) \nabla w \cdot \mathbf{n}\|_{\partial T}^2 \right)^{\frac{1}{2}} \\
 &\leq \left(\sum_{T \in \mathcal{T}_h} h_T^{-3} \|e_0 - e_b\|_{\partial T}^2 \right)^{\frac{1}{2}} \left(\sum_{T \in \mathcal{T}_h} h_T^3 \|a(Q_h^{(k-1)} - I) \nabla w \cdot \mathbf{n}\|_{\partial T}^2 \right)^{\frac{1}{2}} \\
 &\leq C \|e_h\|_{W_h} \left(\sum_{T \in \mathcal{T}_h} h_T^2 \|(Q_h^{(k-1)} - I) \nabla w\|_T^2 + h_T^4 \|(Q_h^{(k-1)} - I) \nabla w\|_{1,T}^2 \right)^{\frac{1}{2}} \\
 &\leq C \|e_h\|_{W_h} h^{s+1} \|w\|_{1+s}.
 \end{aligned} \tag{7.13}$$

Substituting (7.10)–(7.13) into (7.9) and using the regularity assumption (7.4) with the error estimates (6.6) and (6.15) gives (7.6). This completes the proof of this theorem. $\square$

We shall establish the error estimates for the two boundary components u_b and u_n of the PDWG solution $u_h = \{u_0, u_b, u_n\}$ in the usual L^2 norms defined as follows:

$$\|e_b\| := \|u_b - Q_b^{(k)} u\| = \left(\sum_{T \in \mathcal{T}_h} h_T \|e_b\|_{\partial T}^2 \right)^{\frac{1}{2}}, \tag{7.14}$$

$$\|e_n\| := \|u_n - Q_n^{(l)} ((a \nabla u + \mathbf{b}u) \cdot \mathbf{n})\| = \left(\sum_{T \in \mathcal{T}_h} h_T \|e_n\|_{\partial T}^2 \right)^{\frac{1}{2}}. \tag{7.15}$$

Theorem 7.3. Under the assumptions of Theorem 7.2, there exists a constant C such that

$$\|e_b\| \leq \begin{cases} Ch^{k+s} \|u\|_{k+1}, & l = k, \\ C(1 + \tau_1^{\frac{1}{2}}(1 + \tau_1^{-\frac{1}{2}}))h^k \|u\|_{k+1}, & k = 1, l = k - 1, \\ C(1 + \tau_2^{\frac{1}{2}})(1 + \tau_2^{-\frac{1}{2}})h^{k+s} \|u\|_{k+1}, & k \geq 2, l = k - 1. \end{cases} \tag{7.16}$$

$$\|e_n\| \leq \begin{cases} Ch^{k+s-1} \|u\|_{k+1}, & l = k, \\ C(1 + \tau_1^{\frac{1}{2}}(1 + \tau_1^{-\frac{1}{2}}))h^{k-1} \|u\|_{k+1}, & k = 1, l = k - 1, \\ C(1 + \tau_2^{\frac{1}{2}})(1 + \tau_2^{-\frac{1}{2}})h^{k+s-1} \|u\|_{k+1}, & k \geq 2, l = k - 1, \end{cases} \tag{7.17}$$

provided that the meshsize h is sufficiently small.

Proof. On each element $T \in \mathcal{T}_h$, we have from the triangle inequality that

$$\|e_b\|_{\partial T} \leq \|e_0\|_{\partial T} + \|e_b - e_0\|_{\partial T}.$$

Thus, by (7.14), and the trace inequality (6.2), we obtain

$$\begin{aligned}
 \|e_b\|^2 &= \sum_{T \in \mathcal{T}_h} h_T \|e_b\|_{\partial T}^2 \leq C \sum_{T \in \mathcal{T}_h} h_T \|e_0\|_{\partial T}^2 + Ch^4 \sum_{T \in \mathcal{T}_h} h_T^{-3} \|e_b - e_0\|_{\partial T}^2 \\
 &\leq C(\|e_0\|_0^2 + h^4 \|e_h\|_{W_h}^2),
 \end{aligned}$$

which, together with the error estimates (6.6) and (7.6), gives rise to (7.16).

To derive (7.17), applying the same approach to the error component e_n by using triangle inequality, trace inequality (6.2) and inverse inequality gives

$$\begin{aligned}
\|e_n\|^2 &= \sum_{T \in \mathcal{T}_h} h_T \|e_n\|_{\partial T}^2 \\
&\leq \sum_{T \in \mathcal{T}_h} h_T \|(a \nabla e_0 + \mathbf{b} e_0) \cdot \mathbf{n} - e_n\|_{\partial T}^2 + h_T \|(a \nabla e_0 + \mathbf{b} e_0) \cdot \mathbf{n}\|_{\partial T}^2 \\
&\leq Ch^2 \sum_{T \in \mathcal{T}_h} h_T^{-1} \|(a \nabla e_0 + \mathbf{b} e_0) \cdot \mathbf{n} - e_n\|_{\partial T}^2 + \sum_{T \in \mathcal{T}_h} (h_T \|a \nabla e_0\|_{\partial T}^2 + h_T \|\mathbf{b} e_0\|_{\partial T}^2) \\
&\leq C(h^2 \|e_h\|_{W_h}^2 + \sum_{T \in \mathcal{T}_h} (h_T h_T^{-3} \|e_0\|_T^2 + h_T h_T^{-1} \|e_0\|_T^2)) \\
&\leq C(h^2 \|e_h\|_{W_h}^2 + h^{-2} \|e_0\|^2),
\end{aligned}$$

which, together with the error estimates (6.6) and (7.6), gives rise to the error estimate (7.17). $\square$

8. Numerical results

Two types of domains are considered in the numerical experiment: (1) an unit square domain $\Omega_1 = [0, 1]^2$, and (2) a L-shaped domain Ω_2 with vertices $(0, 0)$, $(0.5, 0)$, $(0.5, 0.5)$, $(1, 0.5)$, $(1, 1)$, $(0, 1)$. In all the computation, the finite element partition $\mathcal{T}_h$ is obtained through a successive uniform refinement of a coarse triangulation of the domain Ω by dividing each coarse element into four congruent sub-elements by connecting the mid-points of the three edges of the triangle. The right-hand side function f , the Dirichlet boundary data g_1 and the Neumann boundary data g_2 are set correspondingly. For simplicity, the parameters in the PDWG numerical scheme (2.6)-(2.7) are chosen as $\tau_1 = \tau_2 = 1$.

The finite element spaces for the primal variable u_h and the dual variable λ_h are given by

$$\begin{aligned}
W_{k,h} &= \{u_h = \{u_0, u_b, u_n\} : u_0 \in P_k(T), u_b \in P_k(e), u_n \in P_l(e), \forall e \subset \partial T, \forall T \in \mathcal{T}_h\}, \\
M_{k,h} &= \{\lambda_h : \lambda_h|_T \in P_k(T), \forall T \in \mathcal{T}_h\},
\end{aligned}$$

where $l = k$ or $l = k - 1$. The primal-dual weak Galerkin scheme (2.6)-(2.7) is implemented for the case of $k = 1$ and $k = 2$.

Denote by $e_h = \{e_0, e_b, e_n\} = u_h - Q_h u$ the error function. The following L^2 norms are used to measure the errors:

$$\begin{aligned}
\|e_0\| &= \left(\sum_{T \in \mathcal{T}_h} \int_T e_0^2 dT \right)^{\frac{1}{2}}, \quad \|\nabla e_0\| = \left(\sum_{T \in \mathcal{T}_h} \int_T (\nabla e_0)^2 dT \right)^{\frac{1}{2}}, \\
\|e_b\| &= \left(\sum_{T \in \mathcal{T}_h} h_T \int_{\partial T} e_b^2 ds \right)^{\frac{1}{2}}, \quad \|e_n\| = \left(\sum_{T \in \mathcal{T}_h} h_T \int_{\partial T} e_n^2 ds \right)^{\frac{1}{2}}.
\end{aligned}$$

Test Example 1 (Constant diffusion a and convection $\mathbf{b}$). The diffusion tensor $a \in \mathbb{R}^{2 \times 2}$ and the convection tensor $\mathbf{b} \in \mathbb{R}^2$ are taken by constants as follows:

$$a_{11} = 1, \quad a_{12} = a_{21} = 1, \quad a_{22} = 6; \quad b_1 = 1, \quad b_2 = 1.$$

The exact solution is given by $u(x, y) = \sin(\pi x) \sin(\pi y)$. The domain the unit square domain Ω_1 . The Neumann boundary is $\Gamma_N = \{(0, y) : y \in [0, 1]\}$, and the rest of the boundary is of Dirichlet.

Tables 8.1-8.2 demonstrate the approximation errors and the corresponding convergence rates for $k = 1$ and $k = 2$ with $l = k$ and $l = k - 1$, respectively. For the case of $l = k$, we observe from Table 8.1 that the convergence orders for e_0 and e_b in the discrete L^2 norm are both of an optimal order $\mathcal{O}(h^{k+1})$, and the convergence order for e_n in the discrete L^2 norm is of an optimal order $\mathcal{O}(h^k)$, for $k = 1$ and $k = 2$ respectively, which are all consistent with the theoretical results in Theorems 7.2 - 7.3. For the case of $l = k - 1$, we can see from Table 8.2 that the convergence rates for e_0 and e_b in the discrete L^2 norm are of an order $\mathcal{O}(h^{k+1})$, and the convergence rate for e_n in the discrete L^2 norm is of an order $\mathcal{O}(h^k)$ for $k = 1$ and $k = 2$, respectively. Note that for the case of $l = k - 1$ and $k = 2$, the convergence rates for $\|e_0\|$, $\|e_b\|$ and $\|e_n\|$ are consistent with the theoretical results developed in Theorems 7.2 - 7.3; while for the case of $l = k - 1$ and $k = 1$, the convergence rates for $\|e_0\|$, $\|e_b\|$ and $\|e_n\|$ are of an order which is 1 order higher than the expected convergence order given by (7.6) and (7.16)-(7.17), respectively.

Test Example 2 (Continuous diffusion a and convection $\mathbf{b}$). We choose the diffusion tensor $a \in \mathbb{R}^{2 \times 2}$ and the convection tensor $\mathbf{b} \in \mathbb{R}^2$ in the model problem (1.1) as continuous functions as follows:

$$a_{11} = 1 + x, \quad a_{12} = a_{21} = 0, \quad a_{22} = 1 + y; \quad b_1 = e^{1-x}, \quad b_2 = e^{xy}.$$

The Neumann boundary is $\Gamma_N = \{(0, y) : y \in [0, 1]\}$ and the Dirichlet boundary is $\Gamma_D = \partial\Omega \setminus \Gamma_N$. The exact solution is given by $u(x, y) = \sin(x) \cos(y)$. The domain is the unit square Ω_1 .

Table 8.1Various errors and corresponding convergence rates for $k = 1, 2$ with $l = k$ on Ω_1 .

	$1/h$	$\ e_b\ $	rate	$\ e_n\ $	rate	$\ \nabla e_0\ $	rate	$\ e_0\ $	rate
$k = 1$	4	3.48e-1	–	7.88e-0	–	7.53e-1	–	7.97e-2	–
	8	8.72e-2	2.00	3.75e-0	1.07	3.38e-1	1.16	2.10e-2	1.92
	16	2.01e-2	2.11	1.72e-0	1.12	1.56e-1	1.11	4.92e-3	2.09
	32	4.77e-3	2.08	8.31e-1	1.05	7.58e-2	1.04	1.16e-3	2.08
	64	1.17e-3	2.02	4.11e-1	1.02	3.74e-2	1.02	2.85e-4	2.02
$k = 2$	2	1.59e-1	–	6.89e-0	–	7.90e-1	–	9.94e-2	–
	4	2.87e-2	2.47	1.43e-0	2.27	2.28e-1	1.79	1.63e-2	2.61
	8	3.63e-3	2.98	3.48e-1	2.03	6.28e-2	1.86	2.28e-3	2.84
	16	4.60e-4	2.98	8.79e-2	1.99	1.65e-2	1.93	3.01e-4	2.92
	32	5.89e-5	2.96	2.21e-2	1.99	4.22e-3	1.97	3.86e-5	2.96

Table 8.2Various errors and corresponding convergence rates for $k = 1, 2$ with $l = k - 1$ on Ω_1 .

	$1/h$	$\ e_b\ $	rate	$\ e_n\ $	rate	$\ \nabla e_0\ $	rate	$\ e_0\ $	rate
$k = 1$	4	3.77e-1	–	7.57e-0	–	7.19e-1	–	9.30e-2	–
	8	9.77e-2	1.95	3.59e-0	1.07	3.28e-1	1.13	2.49e-2	1.90
	16	2.46e-2	1.99	1.73e-0	1.05	1.56e-1	1.07	6.27e-3	2.00
	32	6.16e-3	2.00	8.54e-1	1.02	7.65e-2	1.03	1.57e-3	2.00
	64	1.54e-3	2.00	4.24e-1	1.01	3.79e-2	1.01	3.91e-4	2.00
$k = 2$	2	2.16e-1	–	6.00e-0	–	8.56e-01	–	1.12e-1	–
	4	3.39e-2	2.67	1.42e-0	2.08	2.44e-01	1.81	1.75e-2	2.67
	8	4.08e-3	3.05	3.54e-1	2.00	6.54e-2	1.90	2.39e-3	2.87
	16	4.88e-4	3.06	8.90e-2	1.99	1.69e-2	1.96	3.08e-4	2.95
	32	6.05e-5	3.01	2.23e-2	2.00	4.26e-3	1.98	3.90e-5	2.98

Table 8.3Various errors and corresponding convergence rates for $k = 1, 2$ with $l = k$ on Ω_1 .

	$1/h$	$\ e_b\ $	rate	$\ e_n\ $	rate	$\ \nabla e_0\ $	rate	$\ e_0\ $	rate
$k = 1$	4	7.65e-3	–	2.0277e-01	–	4.60e-2	–	2.34e-3	–
	8	1.96e-3	1.96	9.7447e-02	1.06	2.22e-2	1.05	5.92e-4	1.98
	16	4.96e-4	1.99	4.7834e-02	1.03	1.10e-2	1.02	1.48e-4	2.00
	32	1.24e-4	2.00	2.3696e-02	1.01	5.46e-3	1.01	3.69e-5	2.00
	64	3.10e-5	2.00	1.1792e-02	1.01	2.72e-3	1.00	9.21e-6	2.00
$k = 2$	2	1.06e-2	–	1.27e-1	–	3.22e-2	–	4.33e-3	–
	4	9.07e-4	3.54	2.57e-2	2.31	7.19e-3	2.16	5.03e-4	3.10
	8	9.04e-5	3.33	6.07e-3	2.08	1.75e-3	2.04	6.16e-5	3.03
	16	9.86e-6	3.20	1.48e-3	2.03	4.33e-4	2.01	7.64e-6	3.01
	32	1.14e-6	3.11	3.67e-4	2.01	1.08e-4	2.00	9.53e-7	3.00

Tables 8.3–8.4 demonstrate the numerical errors and the convergence rates arising from the PDWG scheme (2.6)–(2.7) for the convection-diffusion model problem (1.1). We observe from Tables 8.3–8.4 that the numerical performance is the same as those in Tables 8.1–8.2.

Test Example 3 (Convection-dominated diffusion problem). We consider the diffusion tensor $a \in \mathbb{R}^{2 \times 2}$ and the convection tensor $\mathbf{b} \in \mathbb{R}^2$ given by

$$a_{11} = a_{12} = \epsilon, \quad a_{12} = a_{21} = 0, \quad b_1 = 1, \quad b_2 = 1,$$

where ϵ assumes some small and positive constants. The exact solution is given by $u(x, y) = (x + 0.5)(y + 0.5)e^{1-x}e^y$ with the full Dirichlet boundary condition $\Gamma_D = \partial\Omega_1$.

Tables 8.5–8.6 show the approximation errors and the convergence rates for the convection-dominated diffusion problem with $\epsilon = 10^{-10}$ on the unit square domain Ω_1 . For the case of $l = k$ and $k = 1$, we observe from Table 8.5 that the convergence rates for the errors e_0 , e_b and e_n in the discrete L^2 -norm are of an order $\mathcal{O}(h^2)$, respectively, which are consistent with the theory for both $\|e_0\|$ and $\|e_b\|$; and are of one order higher than the expected order $\mathcal{O}(h)$ for $\|e_n\|$. For the case of $l = k - 1$ with $k = 1$, the convergence rates of e_0 , e_b , and e_n in the discrete L^2 -norm shown in Table 8.6 are all of order

Table 8.4Various errors and corresponding convergence rates for $k = 1, 2$ with $l = k - 1$ on Ω_1 .

	$1/h$	$\ e_b\ $	rate	$\ e_n\ $	rate	$\ \nabla e_0\ $	rate	$\ e_0\ $	rate
$k = 1$	4	2.79e-2	–	4.69e-1	–	8.16e-2	–	9.72e-3	–
	8	6.53e-3	2.10	2.20e-1	1.09	2.87e-2	1.51	1.81e-3	2.42
	16	1.60e-3	2.03	1.07e-1	1.04	1.20e-2	1.26	3.98e-4	2.18
	32	3.97e-4	2.01	5.32e-2	1.01	5.62e-3	1.09	9.60e-5	2.05
	64	9.90e-5	2.00	2.65e-2	1.00	2.75e-3	1.03	2.38e-5	2.01
$k = 2$	2	1.14e-2	–	1.69e-1	–	3.55e-2	–	4.96e-3	–
	4	1.10e-3	3.37	3.89e-2	2.11	7.54e-3	2.24	5.49e-4	3.17
	8	1.18e-4	3.22	9.49e-3	2.04	1.79e-3	2.08	6.59e-5	3.06
	16	1.37e-5	3.11	2.35e-3	2.01	4.39e-4	2.02	8.14e-6	3.02
	32	1.65e-6	3.05	5.85e-4	2.01	1.09e-4	2.01	1.01e-6	3.01

Table 8.5Various errors and corresponding convergence rates for $k = 1, 2$ with $l = k$ and $\epsilon = 10^{-10}$ on Ω_1 .

	$1/h$	$\ e_b\ $	rate	$\ e_n\ $	rate	$\ \nabla e_0\ $	rate	$\ e_0\ $	rate
$k = 1$	4	9.52e-2	–	1.11e-1	–	7.76e-01	–	2.78e-2	–
	8	2.64e-2	1.85	2.99e-2	1.90	4.07e-01	0.93	7.53e-3	1.88
	16	6.77e-3	1.96	7.65e-3	1.96	2.06e-01	.098	1.92e-3	1.97
	32	1.70e-3	1.99	1.93e-3	1.99	1.03e-01	1.00	4.83e-4	1.99
	64	4.27e-4	2.00	4.84e-4	2.00	5.18e-02	1.00	1.21e-4	2.00
$k = 2$	2	3.73e-2	–	5.77e-2	–	2.27e-1	–	1.28e-2	–
	4	6.65e-3	2.49	9.50e-3	2.60	6.44e-2	1.82	2.17e-3	2.55
	8	1.36e-3	2.29	1.86e-3	2.35	1.97e-2	1.71	4.64e-4	2.23
	16	3.13e-4	2.11	4.23e-4	2.14	7.38e-3	1.41	1.10e-4	2.08
	32	7.61e-5	2.04	1.03e-4	2.04	3.30e-3	1.16	2.69e-5	2.03

Table 8.6Various errors and corresponding convergence rates for $k = 1, 2$ with $l = k - 1$ and $\epsilon = 10^{-10}$ on Ω_1 .

	$1/h$	$\ e_b\ $	rate	$\ e_n\ $	rate	$\ \nabla e_0\ $	rate	$\ e_0\ $	rate
$k = 1$	4	2.96e-01	–	1.54e-0	–	1.09e-0	–	1.09e-1	–
	8	1.80e-01	0.73	7.93e-1	0.96	5.62e-1	0.95	6.43e-2	0.75
	16	1.06e-01	0.76	4.06e-1	0.97	3.25e-1	0.79	3.77e-2	0.77
	32	5.86e-02	0.85	2.06e-1	0.98	2.05e-1	0.67	2.08e-2	0.86
	64	3.10e-02	0.92	1.04e-1	0.99	1.35e-1	0.66	1.10e-2	0.92
$k = 2$	2	9.82e-2	–	2.77e-1	–	4.25e-1	–	3.98e-2	–
	4	2.81e-2	1.80	7.47e-2	1.89	2.07e-1	1.04	1.04e-2	1.94
	8	7.93e-3	1.82	1.93e-2	1.95	1.14e-1	0.86	2.84e-3	1.86
	16	2.11e-3	1.91	4.92e-3	1.98	6.13e-2	0.89	7.59e-4	1.91
	32	5.44e-4	1.96	1.24e-3	1.99	3.20e-2	0.94	1.97e-4	1.95

$\mathcal{O}(h)$, which are consistent with the expected order for both $\|e_0\|$ and $\|e_b\|$; and is of one order higher than the expected convergence rate given by (7.17) for $\|e_n\|$. For the case of $l = k$ and $k = 2$, we observe from Table 8.5 that the convergence rates for the errors e_0 , e_b and e_n in the discrete L^2 -norm are of order $\mathcal{O}(h^2)$, which are consistent with the theory for $\|e_n\|$; and are of one order lower than the expected optimal order $\mathcal{O}(h^3)$ for both $\|e_0\|$ and $\|e_b\|$. For the case of $l = k - 1$ and $k = 2$, we see from Table 8.6 that the convergence rates for the errors e_0 , e_b and e_n in the discrete L^2 -norm are of order $\mathcal{O}(h^2)$, which are consistent with the theory for $\|e_0\|$ and $\|e_b\|$; and are of one order higher than the expected optimal order $\mathcal{O}(h)$ for $\|e_n\|$.

Table 8.7 shows the approximation errors and the convergence rates for $\epsilon = 10^{-2}$ on the unit square domain Ω_1 when $k = 2$ is employed. For the case of $l = k$, we observe from Table 8.7 that the convergence rates for the errors e_0 and e_b in the discrete L^2 -norm are of optimal order $\mathcal{O}(h^3)$, which is consistent with the theory; and the convergence rate for the error e_n in the discrete L^2 -norm is one order higher than the expected optimal order $\mathcal{O}(h^2)$, which outperforms the theory. For the case of $l = k - 1$, the convergence rates of e_0 , e_b and e_n in the discrete L^2 -norm are of one order higher than the expected optimal order, which are better than the theory. It can be seen that the convergence rates are improved and the theoretical results in Theorems 7.2–7.3 are recovered for the case of $\epsilon = 10^{-2}$. The results indicate that the diffusion coefficient a has an influence on the convergence rate for $k = 2$.

Table 8.7Various errors and corresponding convergence rates for $k = 2$ with $\epsilon = 10^{-2}$ on Ω_1 .

	$1/h$	$\ e_b\ $	rate	$\ e_n\ $	rate	$\ \nabla e_0\ $	rate	$\ e_0\ $	rate
$l = k$	2	4.31e-02	–	6.79e-2	–	2.44e-1	–	1.59e-2	–
	4	6.18e-03	2.80	8.97e-3	2.92	6.47e-2	1.92	2.02e-3	2.98
	8	7.98e-04	2.95	1.25e-3	2.84	1.73e-2	1.90	2.55e-4	2.99
	16	9.42e-05	3.08	1.94e-4	2.69	4.31e-3	2.00	3.12e-5	3.03
	32	1.11e-05	3.08	3.56e-5	2.44	1.03e-3	2.06	3.72e-6	3.07
$l = k - 1$	4	2.71e-2	–	7.51e-2	–	1.96e-1	–	1.02e-2	–
	8	7.01e-3	1.95	1.94e-2	1.96	9.48e-2	1.05	2.64e-3	1.95
	16	1.44e-3	2.29	4.77e-3	2.02	3.63e-2	1.38	5.59e-4	2.24
	32	2.05e-4	2.81	1.13e-3	2.07	9.36e-3	1.96	8.21e-5	2.77
	64	1.97e-5	3.38	2.73e-4	2.05	1.64e-3	2.51	8.01e-6	3.36

Table 8.8Various errors and corresponding convergence rates for $k = 1, 2$ with $l = k$ and $\epsilon = 10^{-2}$ on Ω_2 .

	$1/h$	$\ e_b\ $	rate	$\ e_n\ $	rate	$\ \nabla e_0\ $	rate	$\ e_0\ $	rate
$k = 1$	4	1.00e-1	–	1.32e-1	–	7.45e-1	–	3.23e-2	–
	8	3.35e-2	1.58	4.57e-2	1.53	3.97e-1	0.90	1.12e-2	1.53
	16	1.02e-2	1.72	1.47e-2	1.63	1.97e-1	1.01	3.56e-3	1.66
	32	2.58e-3	1.99	4.12e-3	1.84	9.60e-2	1.03	9.18e-4	1.96
	64	6.43e-4	2.00	1.32e-3	1.64	4.76e-2	1.01	2.30e-4	2.00
$k = 2$	4	6.03e-3	–	8.94e-3	–	6.18e-2	–	1.97e-3	–
	8	8.85e-4	2.77	1.39e-3	2.68	1.72e-2	1.85	2.85e-4	2.79
	16	1.02e-4	3.12	2.02e-4	2.78	4.27e-3	2.01	3.23e-5	3.14
	32	1.15e-5	3.15	3.58e-5	2.50	1.02e-3	2.07	3.64e-6	3.15
	64	1.38e-6	3.06	8.00e-6	2.16	2.50e-4	2.03	4.36e-7	3.06

Table 8.9Various errors and corresponding convergence rates for $k = 1, 2$ with $l = k - 1$ and $\epsilon = 10^{-2}$ on Ω_2 .

	$1/h$	$\ e_b\ $	rate	$\ e_n\ $	rate	$\ \nabla e_0\ $	rate	$\ e_0\ $	rate
$k = 1$	4	1.89e-1	–	1.44e-0	–	9.87e-1	–	7.02e-2	–
	8	1.06e-1	0.83	7.33e-1	0.98	4.91e-1	1.00	3.83e-2	0.87
	16	6.53e-2	0.72	3.72e-1	0.98	2.71e-1	0.86	2.33e-2	0.72
	32	3.52e-2	0.89	1.87e-1	0.99	1.54e-1	0.81	1.25e-2	0.90
	64	1.54e-2	1.20	9.28e-2	1.01	7.85e-2	0.98	5.45e-3	1.20
$k = 2$	4	2.49e-2	–	7.00e-2	–	1.76e-1	–	9.52e-3	–
	8	6.56e-3	1.92	1.81e-2	1.95	8.61e-2	1.03	2.49e-3	1.93
	16	1.34e-3	2.29	4.45e-3	2.02	3.32e-2	1.38	5.24e-4	2.25
	32	1.89e-4	2.83	1.05e-3	2.08	8.52e-3	1.96	7.54e-5	2.80
	64	1.80e-5	3.39	2.54e-4	2.05	1.49e-3	2.51	7.25e-6	3.38

Tables 8.8–8.9 show the approximation errors and corresponding convergence rates for $k = 1, 2$ with $l = k$ and $l = k - 1$, respectively, on the L-shaped domain Ω_2 . In both the case of $l = k$ for $k = 1, 2$ and $l = k - 1$ for $k = 2$, we observe a $(k + 1)$ -th order of convergence for $\|e_0\|$ and $\|e_b\|$, and a k -th order of convergence for $\|e_n\|$, which are consistent with our theoretical results established in Theorems 7.2–7.3. As for the case of $l = k - 1$ and $k = 1$, we observe from Table 8.9 that, the convergence rates of $\|e_0\|$, $\|e_b\|$, $\|e_n\|$ are all of order $\mathcal{O}(h)$. Note that the convergence results for $\|e_0\|$ and $\|e_b\|$ for the case of $l = k - 1$ and $k = 1$ are consistent with the theoretical findings in Theorems 7.2–7.3; while for $\|e_n\|$, the convergence rate is 1 order higher than the error estimate given by (7.17).

In what follows of this section, we present the plot of the numerical solution u_h arising from the primal-dual weak Galerkin scheme (2.6)–(2.7) for test problems for which the exact solution is not known.

Test Example 4. The diffusion a is given by $a_{11} = a_{22} = 10^{-4}$, $a_{12} = a_{21} = 0$; the convection is set as $\mathbf{b} = (y, -x)^T$, the domain is an unit square domain Ω_1 ; and the mixed boundary conditions are $g_1 = \sin(3x)$ on the inflow boundary $\Gamma_D = \{(x, y) : \mathbf{b} \cdot \mathbf{n} < 0\}$ and $g_2 = 0$ on $\Gamma_N = \partial\Omega_1 \setminus \Gamma_D$. Fig. 8.1 presents the plots for the numerical solution u_h obtained from the PDWG numerical method (2.6)–(2.7) with $k = 1$ and $l = k - 1$ for the convection-dominated diffusion problem when different load functions $f = 1$ (left) and $f = 0$ (right) are employed, respectively.

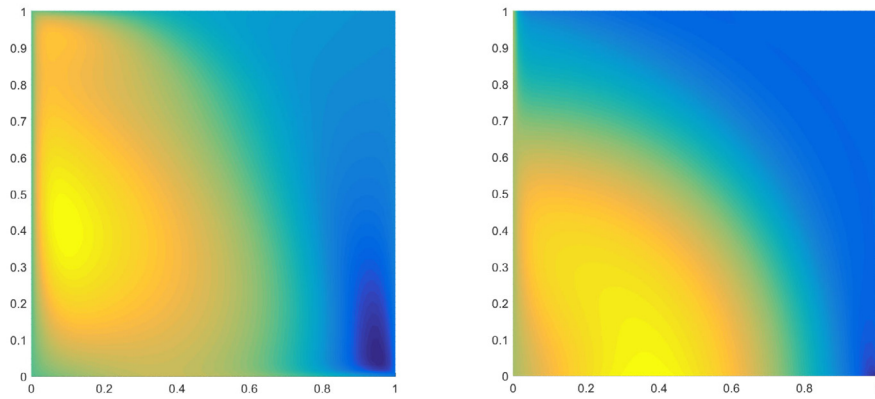


Fig. 8.1. Contour plots of the numerical solution u_h on Ω_1 for the load functions $f = 1$ (left) and $f = 0$ (right).

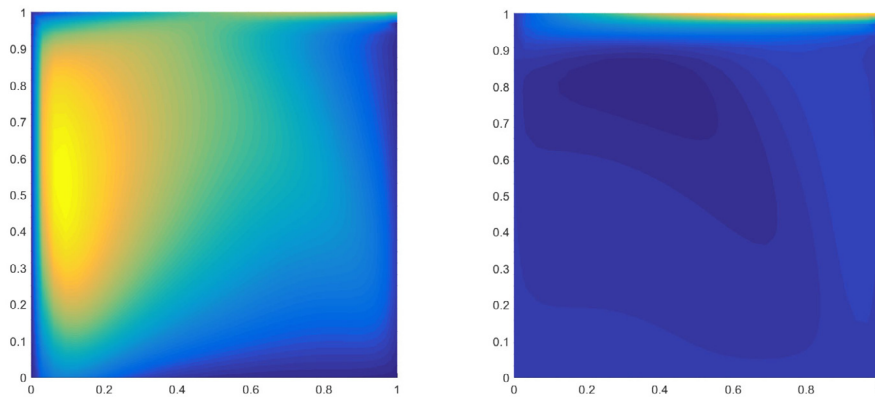


Fig. 8.2. Contour plots of the numerical solution u_h on Ω_1 with the load functions $f = 1$ (left) and $f = 0$ (right).

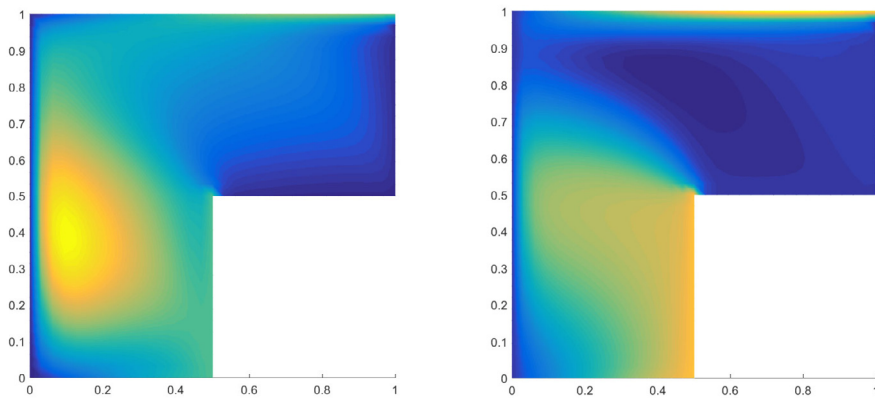


Fig. 8.3. Contour plots of the numerical solution u_h on Ω_2 with the load functions $f = 1$ (left) and $f = 0$ (right).

Test Example 5. The diffusion is given by $a_{11} = a_{22} = 10^{-5}$, $a_{12} = a_{21} = 0$; and the convection vector is set as $\mathbf{b} = (y, -x)^T$. For the unit square domain Ω_1 , $\Gamma_N = \{(x, y) : x = 1 \text{ or } y = 0\}$; and for the L-shaped domain Ω_2 , $\Gamma_N = \{(x, y) : x = 1 \text{ or } y = 0.5\}$. The mixed boundary conditions are $g_1 = \sin(2x)$ on $\Gamma_D = \partial\Omega \setminus \Gamma_N$ and $g_2 = 0$ on Γ_N . We take $l = k - 1$ and $k = 1$. Figs. 8.2-8.3 demonstrate the numerical solutions u_h on the unit square domain Ω_1 and the L-shaped domain Ω_2 when different load functions $f = 0$ (left) and $f = 1$ (right) are employed, respectively.

Test Example 6. We test the positivity-preserving property of the PDWG solution in this example. To this end, we test the problem (1.1) on the unit square domain Ω_1 with the full Dirichlet boundary condition in two cases:

Table 8.10

The minimum value $\bar{u}_b$ of the numerical solution on the edge for piecewise constant and piecewise linear approximations.

	$1/h$	$\bar{u}_b$ for Case 1			$\bar{u}_b$ for Case 2		
		$\epsilon = 1$	$\epsilon = 10^{-2}$	$\epsilon = 10^{-4}$	$\epsilon = 1$	$\epsilon = 10^{-2}$	$\epsilon = 10^{-4}$
$k = 0$	4	1.0000	0.0100	0.0001	0.9530	-0.0370	-0.0469
	8	1.0000	0.0100	0.0001	0.9511	-0.0389	-0.0488
	16	1.0000	0.0100	0.0001	0.9505	-0.0395	-0.0494
	32	1.0000	0.0100	0.0001	0.9503	-0.0397	-0.0496
	64	1.0000	0.0100	0.0001	0.9503	-0.0397	-0.0469
$k = 1$	4	1.0000	0.0100	0.0001	1.0000	0.0100	0.0001
	8	0.8497	-0.1403	-0.1502	0.8021	-0.1879	-0.1978
	16	0.9893	-0.0007	-0.0106	0.9894	-0.0006	-0.0105
	32	1.0000	0.0100	0.0001	1.0000	0.0100	0.0001
	64	0.9941	0.0041	-0.0058	0.9929	0.0029	-0.0070

- Case 1: $a_{11} = a_{22} = 1$, $a_{12} = a_{21} = 0$, $\mathbf{b} = (0, 0)^T$;
- Case 2: $a_{11} = a_{22} = 1$, $a_{12} = a_{21} = 0$, $\mathbf{b} = (1, 1)^T$.

We take the right-hand side function $f = 1$, and the boundary function $g_1 = \epsilon > 0$ with $\epsilon = 1, 10^{-2}, 10^{-4}$. In our numerical experiment, we test the positivity of the minimum value of the numerical solution on each edge $e \in \mathcal{E}_h$ for both the piecewise constant (i.e., $k = 0, l = 0$) approximation and piecewise linear ($k = 1, l = k - 1$) approximation. We denote by $\bar{u}_b$ this minimum value, that is

$$\bar{u}_b = \min_{e \in \mathcal{E}_h} u_b(z_0),$$

where $z_0 \in e$ is the middle point of the edge e .

Table 8.10 demonstrates the minimum value $\bar{u}_b$ of the numerical solution for Case 1 and Case 2. As we may observe, the numerical solution $\bar{u}_b$ is positive-preserving in Case 1 for piecewise constant approximation. While for Case 2 (i.e., the convection-diffusion term $\mathbf{b}$ is not zero), $\bar{u}_b$ changes its sign and negative values appear for some small ϵ . In this sense, the coefficient $\mathbf{b}$ has effect on the positive-preserving property of our numerical scheme. On the other hand, for piecewise linear approximation, we observe that even for Poisson equations (i.e., Case 1), $\bar{u}_b$ could be negative for small $\epsilon = 10^{-2}, 10^{-4}$. However, as ϵ increases to 1, $\bar{u}_b$ changes to positive. In other words, the numerical scheme is positive-preserving for large ϵ . From the point of view of positive-preserving property, it seems that the piecewise constant approximation behaviors better than the piecewise linear approximation. The analysis of the positive-preserving property for the PDWG solution is a challenging and interesting work, which deserves a comprehensive study.

Acknowledgement

The authors would like to thank Dr. Junping Wang for his invaluable discussion and suggestions.

References

- [1] L. Mu, J. Wang, X. Ye, Weak Galerkin finite element methods on polytopal meshes, *Int. J. Numer. Anal. Model.* 12 (2015) 31–53, arXiv:1204.3655v2, 2015.
- [2] C. Wang, A new primal-dual weak Galerkin finite element method for ill-posed elliptic Cauchy problems, *J. Comput. Appl. Math.* 371 (2020) 112629.
- [3] C. Wang, J. Wang, A primal-dual weak Galerkin finite element method for second order elliptic equations in non-divergence form, *Math. Comput.* 87 (2018) 515–545.
- [4] C. Wang, J. Wang, A primal-dual weak Galerkin finite element method for Fokker-Planck type equations, *SIAM Numer. Anal.* 58 (5) (2020) 2632–2661.
- [5] C. Wang, J. Wang, Primal-dual weak Galerkin finite element methods for elliptic Cauchy problems, *Comput. Math. Appl.* 78 (2019) 905–928.
- [6] C. Wang, J. Wang, A primal-dual finite element method for first-order transport problems, arXiv:1906.07336.
- [7] C. Wang, New discretization schemes for time-harmonic Maxwell equations by weak Galerkin finite element methods, *J. Comput. Appl. Math.* 341 (2018) 127–143.
- [8] C. Wang, J. Wang, Discretization of div-curl systems by weak Galerkin finite element methods on polyhedral partitions, *J. Sci. Comput.* 68 (2016) 1144–1171.
- [9] C. Wang, J. Wang, A hybridized formulation for weak Galerkin finite element methods for biharmonic equation on polygonal or polyhedral meshes, *Int. J. Numer. Anal. Model.* 12 (2015) 302–317.
- [10] C. Wang, J. Wang, An efficient numerical scheme for the biharmonic equation by weak Galerkin finite element methods on polygonal or polyhedral meshes, *Comput. Math. Appl.* 68 (12) (2014) 2314–2330.
- [11] C. Wang, J. Wang, R. Wang, R. Zhang, A locking-free weak Galerkin finite element method for elasticity problems in the primal formulation, *J. Comput. Appl. Math.* 307 (2016) 346–366.
- [12] J. Wang, C. Wang, Weak Galerkin finite element methods for elliptic PDEs, *Sci. China* 45 (2015) 1061–1092.

- [13] J. Wang, X. Ye, A weak Galerkin mixed finite element method for second-order elliptic problems, *Math. Comput.* 83 (2014) 2101–2126.
- [14] J. Wang, X. Ye, A weak Galerkin finite element method for second-order elliptic problems, *J. Comput. Appl. Math.* 241 (2013) 103–115.
- [15] J. Wang, X. Ye, A weak Galerkin finite element method for the Stokes equations, *Adv. Comput. Math.* 42 (2016) 155–174, <https://doi.org/10.1007/s10444-015-9415-2>.
- [16] C. Wang, L. Zikatanov, Low regularity primal-dual weak Galerkin finite element methods for convection-diffusion equations, submitted for publication, arXiv:1901.06743.
- [17] C. Wang, H. Zhou, A weak Galerkin finite element method for a type of fourth order problem arising from fluorescence tomography, *J. Sci. Comput.* 71 (3) (2017) 897–918.