



de Rham complexes for weak Galerkin finite element spaces

Chunmei Wang^{a,*}, Junping Wang^{b,2}, Xiu Ye^{c,3}, Shangyou Zhang^d^a Department of Mathematics, University of Florida, Gainesville, FL 32611, USA^b Division of Mathematical Sciences, National Science Foundation, Alexandria, VA 22314, USA^c Department of Mathematics, University of Arkansas at Little Rock, Little Rock, AR 72204, USA^d Department of Mathematical Sciences, University of Delaware, Newark, DE 19716, USA

ARTICLE INFO

Article history:

Received 29 April 2020

Received in revised form 23 March 2021

MSC:

primary 65N15

65N30

secondary 35J50

Keywords:

Weak Galerkin

Finite element methods

de Rham complex

Polyhedral elements

ABSTRACT

Two de Rham complex sequences of the finite element spaces are introduced for weak finite element functions and weak derivatives developed in the weak Galerkin (WG) finite element methods on general polyhedral elements. One of the sequences uses polynomials of equal order for all the finite element spaces involved in the sequence and the other one uses polynomials of naturally descending orders. It is shown that the diagrams in both de Rham complexes commute for general polyhedral elements. The exactness of one of the complexes is established for the lowest order element.

© 2021 Elsevier B.V. All rights reserved.

1. Introduction

Finite element exterior calculus is a framework that explores the structural properties of finite element approximating spaces and their connection with basic differential operators such as gradient, curl, and divergence operators in differential calculus. Such differential operators usually form the basis of mathematical models for physical, engineering, and biological problems. The classical conforming finite element exterior calculus has provided good guidance in the design and analysis of various finite element methods for solving partial differential equations (PDE). In particular, the framework has proved to be powerful in analyzing well-posedness of finite element discretizations and revealing new finite elements for solving various PDE modeling problems arising from science and engineering. Finite element exterior calculus has been well developed for conforming finite element methods on simplicial elements, and they often appear in the form of de Rham complexes in the application of numerical solutions for PDEs [1–14].

In the last two decades, an extensive research effort has been made by the computational mathematics community in the development of finite element methods with discontinuous approximating functions. The weak Galerkin (WG) finite element method, first introduced in [15], is one of the recently developed finite element techniques based on discontinuous finite element functions. The essence of weak Galerkin finite element methods is the use of weak finite

* Corresponding author.

E-mail addresses: chunmei.wang@ufl.edu (C. Wang), jwang@nsf.gov (J. Wang), xye@ualr.edu (X. Ye), szhang@udel.edu (S. Zhang).¹ The research of Chunmei Wang was partially supported by National Science Foundation, United States Award DMS-1849483.² The research of Junping Wang was supported by the NSF IR/D program, United States, while working at National Science Foundation. However, any opinion, finding, and conclusions or recommendations expressed in this material are those of the author and do not necessarily reflect the views of the National Science Foundation.³ This research was supported in part by National Science Foundation, United States Grant DMS-1620016.

$$\begin{array}{ccccccc}
\mathbb{R}^1 & \xrightarrow{I} & C^\infty(T) & \xrightarrow{\nabla} & [C^\infty(T)]^3 & \xrightarrow{\nabla \times} & [C^\infty(T)]^3 & \xrightarrow{\nabla \cdot} & C^\infty(T) & \xrightarrow{N} & 0 \\
& & Q_h^{(1)} \downarrow & & Q_h^{(2)} \downarrow & & Q_h^{(3)} \downarrow & & Q_h^{(4)} \downarrow & & \\
\mathbb{R}^m & \xrightarrow{I_w} & V_k^{(1)}(T) & \xrightarrow{\nabla_w} & V_k^{(2)}(T) & \xrightarrow{\nabla_w \times} & V_k^{(3)}(T) & \xrightarrow{\nabla_w \cdot} & V_k^{(4)}(T) & \xrightarrow{N} & 0
\end{array}$$

Fig. 1.1. The WG de Rham Complex 1.

$$\begin{array}{ccccccc}
\mathbb{R}^1 & \xrightarrow{I} & C^\infty(T) & \xrightarrow{\nabla} & [C^\infty(T)]^3 & \xrightarrow{\nabla \times} & [C^\infty(T)]^3 & \xrightarrow{\nabla \cdot} & C^\infty(T) & \xrightarrow{N} & 0 \\
& & R_h^{(1)} \downarrow & & R_h^{(2)} \downarrow & & R_h^{(3)} \downarrow & & R_h^{(4)} \downarrow & & \\
\mathbb{R}^n & \xrightarrow{I_w} & W_k^{(1)}(T) & \xrightarrow{\nabla_w} & W_{k-1}^{(2)}(T) & \xrightarrow{\nabla_w \times} & W_{k-2}^{(3)}(T) & \xrightarrow{\nabla_w \cdot} & W_{k-3}^{(4)}(T) & \xrightarrow{N} & 0
\end{array}$$

Fig. 1.2. The WG de Rham Complex 2.

element functions and their weak derivatives defined as discrete distributions in polynomial subspaces. In general, weak Galerkin finite element formulations for partial differential equations can be derived naturally by replacing usual derivatives by discrete weak derivatives in the corresponding variational forms, with the option of adding stabilization term(s) to enforce a weak continuity of the approximating functions. Like most discontinuous finite element methods, WG method is applicable for finite element partitions with arbitrary shape of polygons/polyhedra.

The purpose of this paper is to present two de Rham diagrams for weak Galerkin finite element spaces with weakly defined differential operators. In particular, two finite element differential complexes are developed in the WG context on polyhedral elements. It is proved that the diagrams in the two de Rham complexes commute with simple, yet powerful L^2 projection operators from the continuous functional spaces to the corresponding WG finite element “subspaces”.

WG DE RHAM COMPLEX 1: POLYNOMIALS OF EQUAL ORDER $k \geq 0$

In the de Rham complex 1 (see Fig. 1.1), the degree of polynomials in all four WG finite element spaces $V_k^{(i)}$, $i = 1, \dots, 4$, is of the same value $k \geq 0$, and $C^\infty(T)$ is the space of infinitely smooth functions in the closed region T . This special feature of using polynomials of equal order in all the finite element spaces is possible only for weakly defined differential operators, as the weak derivative of a k th order polynomial can be polynomials of any degree in the WG context. Polynomials in the finite element spaces in the WG de Rham Complex 2 (see Fig. 1.2) are in a naturally descending order. Details about these finite element spaces and the weak differential operators can be found in forthcoming sections.

WG DE RHAM COMPLEX 2: POLYNOMIALS OF DESCENDING ORDER $k \geq 3$

The above two WG de Rham complexes have some unique features: (1) they both work for discontinuous approximation on general polyhedral elements; (2) all the cochain projections $Q_h^{(i)}$ in Fig. 1.1 and $R_h^{(i)}$ in Fig. 1.2 are the standard L^2 projections; and (3) weakly defined differential operators are employed first time in finite element differential complexes. It should be noted that the weak gradient ∇_w , weak curl $\nabla_w \times$ and weak divergence $\nabla_w \cdot$ in the WG de Rham complexes 1 and 2 have been employed for solving various partial differential equations in many existing literatures, including [16–18]. The WG de Rham complexes 1 and 2 additionally provide discrete weak differential operators defined on surfaces, which is of great interest to the development of numerical methods for PDEs on surfaces or manifolds in general.

2. WG de Rham complexes

Let T be a shape-regular polyhedron in the sense of [18]. Denote by $F(T)$, $E(T)$ and $V(T)$ the set of faces, edges, and vertices of T , respectively. For any face $f \in F(T)$, let $\mathbf{n}$ be a unit normal vector to f and $\mathbf{t}$ be a unit tangential vector on $e \subset \partial f$ which obey the right-hand rule. Throughout this paper, we adopt the notation of $(\cdot, \cdot)_D$ for the L^2 inner product in $L^2(D)$, where D could be a volume, a surface, or a curve.

2.1. WG de Rham complex for polynomials of equal order

For a given non-negative integer $k \geq 0$, we introduce four vector spaces $V_k^{(1)}(T)$, $V_k^{(2)}(T)$, $V_k^{(3)}(T)$ and $V_k^{(4)}(T)$ that allow the operator operations shown as in Fig. 1.1. The first space $V_k^{(1)}(T)$ is defined by

$$\begin{aligned}
V_k^{(1)}(T) = \{v = \{v_0, v_f, v_e, v_n\} : v_0 \in P_k(T), v_f \in P_k(f), v_e \in P_k(e), \\
v_n \in \mathbb{R}, f \in F(T), e \in E(T), n \in V(T)\}.
\end{aligned} \tag{2.1}$$

The second space $V_k^{(2)}(T)$ is given by

$$V_k^{(2)}(T) = \{\mathbf{u} = \{\mathbf{u}_0, \mathbf{u}_f, \mathbf{u}_e\} : \mathbf{u}_0 \in V_{k,0}^{(2)}(T), \mathbf{u}_f \in V_{k,f}^{(2)}(f), \mathbf{u}_e \in V_{k,e}^{(2)}(e), f \in F(T), e \in E(T)\}, \quad (2.2)$$

where

$$V_{k,0}^{(2)}(T) = [P_k(T)]^3, \quad (2.3)$$

$$V_{k,f}^{(2)}(f) = \{\mathbf{u}_f = u_1 \mathbf{t}_{f,1} + u_2 \mathbf{t}_{f,2} : u_1, u_2 \in P_k(f)\}, \quad (2.4)$$

$$V_{k,e}^{(2)}(e) = \{\mathbf{u}_e = u \mathbf{t}_e : u \in P_k(e)\}, \quad (2.5)$$

$\mathbf{t}_{f,1}$ and $\mathbf{t}_{f,2}$ in (2.4) are two orthogonal unit vectors tangential to f , $\mathbf{t}_e$ in (2.5) is a tangent vector on e . The third space $V_k^{(3)}(T)$ is defined as

$$V_k^{(3)}(T) = \{\mathbf{w} = \{\mathbf{w}_0, \mathbf{w}_f\} : \mathbf{w}_0 \in V_{k,0}^{(3)}(T), \mathbf{w}_f \in V_{k,f}^{(3)}(f), f \in F(T)\}, \quad (2.6)$$

where

$$V_{k,0}^{(3)}(T) = [P_k(T)]^3, \quad (2.7)$$

$$V_{k,f}^{(3)}(f) = \{\mathbf{w}_f = w_f \mathbf{n}_f : w_f \in P_k(f)\}, \quad (2.8)$$

with $\mathbf{n}_f$ being a unit normal vector to the surface f . The three unit vectors $\mathbf{t}_{f,1}$, $\mathbf{t}_{f,2}$, and $\mathbf{n}_f$ are assumed to form an orthogonal right-hand system, though this assumption is not necessary. Our fourth space $V_k^{(4)}(T)$ is given by

$$V_k^{(4)}(T) = P_k(T). \quad (2.9)$$

Next we define three weak gradient operators: (1) the regular weak gradient $\nabla_{w,0}$ on volume T , (2) the surface weak gradient $\nabla_{w,f}$ on face f , and (3) the edge weak gradient or directional derivative $\nabla_{w,e}$ on e . More precisely, for $v = \{v_0, v_f, v_e, v_n\} \in V_k^{(1)}(T)$, the weak gradient on volume T , denoted by $\nabla_{w,0}v \in V_{k,0}^{(2)}(T)$, is defined on T by

$$(\nabla_{w,0}v, \varphi)_T = -(v_0, \nabla \cdot \varphi)_T + (v_f, \varphi \cdot \mathbf{n})_{\partial T} \quad \forall \varphi \in V_{k,0}^{(2)}(T). \quad (2.10)$$

The surface weak gradient on the face $f \in F(T)$, denoted by $\nabla_{w,f}v \in V_{k,f}^{(2)}(f)$, is defined as follows:

$$(\nabla_{w,f}v, \theta \times \mathbf{n}_f)_f = -(v_f, \nabla \times \theta \cdot \mathbf{n}_f)_f + (v_e, \theta \cdot \mathbf{t}_{\partial f})_{\partial f} \quad \forall \theta \in V_{k,f}^{(2)}(f), \quad (2.11)$$

where $\mathbf{t}_{\partial f}$ is chosen such that $\mathbf{t}_{\partial f}$ and $\mathbf{n}_f$ obey the right-hand rule. Analogously, the edge weak gradient or directional derivative on edge e , denoted as $\nabla_{w,e}v \in V_{k,e}^{(2)}(e)$, is defined on e such that

$$(\nabla_{w,e}v, \varphi \mathbf{t}_e)_e = -(v_e, \nabla \varphi \cdot \mathbf{t}_e)_e + (v_n, \varphi \mathbf{t}_e \cdot \mathbf{n}_{\partial e})_{\partial e} \quad \forall \varphi \mathbf{t}_e \in V_{k,e}^{(2)}(e), \quad (2.12)$$

where $\mathbf{n}_{\partial e}$ is the unit outward direction at the two end points of e . It is clear that $(v_n, \varphi \mathbf{t}_e \cdot \mathbf{n}_{\partial e})_{\partial e}$ is in fact the difference of the value of $v_n \varphi$ at two end points of e with sign determined by $\mathbf{t}_e$.

With the help of these weak gradient operators, we may define a composite weak gradient operator $\nabla_w : V_k^{(1)}(T) \mapsto V_k^{(2)}(T)$ as follows:

$$\nabla_w v := \{\nabla_{w,0}v, \nabla_{w,f}v, \nabla_{w,e}v\} \in V_k^{(2)}(T) \quad (2.13)$$

for all $v \in V_k^{(1)}(T)$.

Next, we shall introduce two weak curl operators, denoted as $\nabla_{w,0} \times$ and $\nabla_{w,f} \times$, on T and f respectively. For any $\mathbf{u} = \{\mathbf{u}_0, \mathbf{u}_f, \mathbf{u}_e\} \in V_k^{(2)}(T)$, the weak curl $\nabla_{w,0} \times \mathbf{u} \in V_{k,0}^{(3)}(T)$ is defined on T by

$$(\nabla_{w,0} \times \mathbf{u}, \theta)_T = (\mathbf{u}_0, \nabla \times \theta)_T + (\mathbf{u}_f, \theta \times \mathbf{n})_{\partial T} \quad \forall \theta \in V_{k,0}^{(3)}(T). \quad (2.14)$$

The surface weak curl on each face $f \in F(T)$, denoted as $\nabla_{w,f} \times \mathbf{u} \in V_{k,f}^{(3)}(f)$, is defined on f satisfying

$$(\nabla_{w,f} \times \mathbf{u}, \tau \mathbf{n}_f)_f = (\mathbf{u}_f, \nabla \tau \times \mathbf{n}_f)_f + (\mathbf{u}_e, \tau \mathbf{t})_{\partial f} \quad \forall \tau \mathbf{n}_f \in V_{k,f}^{(3)}(f). \quad (2.15)$$

The composite weak curl operator $\nabla_w \times : V_k^{(2)}(T) \mapsto V_k^{(3)}(T)$ is then given by setting

$$\nabla_w \times \mathbf{u} = \{\nabla_{w,0} \times \mathbf{u}, \nabla_{w,f} \times \mathbf{u}\}. \quad (2.16)$$

Finally, for any $\mathbf{w} = \{\mathbf{w}_0, \mathbf{w}_f\} \in V_k^{(3)}(T)$, we define its weak divergence $\nabla_w \cdot \mathbf{w} \in V_k^{(4)}(T)$ by the following equation:

$$(\nabla_w \cdot \mathbf{w}, \tau)_T = -(\mathbf{w}_0, \nabla \tau)_T + (\mathbf{w}_f, \tau \mathbf{n})_{\partial T} \quad \forall \tau \in V_k^{(4)}(T). \quad (2.17)$$

It is clear that the weak divergence operator $\nabla_w \cdot$ maps $V_k^{(3)}(T)$ to $V_k^{(4)}(T)$.

2.2. WG de Rham complex for polynomials of descending order

For a given integer $k \geq 3$, we introduce four polynomial spaces with descending orders for the diagram in Fig. 1.2: $W_k^{(1)}(T)$, $W_{k-1}^{(2)}(T)$, $W_{k-2}^{(3)}(T)$ and $W_{k-3}^{(4)}(T)$. The first polynomial space $W_k^{(1)}(T)$ is given by

$$W_k^{(1)}(T) = \{v = \{v_0, v_f, v_e, v_n\} : v_0 \in P_k(T), v_f \in P_{k-1}(f), v_e \in P_{k-2}(e), v_n \in \mathbb{R}, f \in F(T), e \in E(T), n \in V(T)\}. \quad (2.18)$$

The second space $W_k^{(2)}(T)$ is given by

$$W_{k-1}^{(2)}(T) = \{\mathbf{u} = \{\mathbf{u}_0, \mathbf{u}_f, \mathbf{u}_e\} : \mathbf{u}_0 \in W_{k-1,0}^{(2)}(T), \mathbf{u}_f \in W_{k-2,f}^{(2)}(f), \mathbf{u}_e \in W_{k-3,e}^{(2)}(e), f \in F(T), e \in E(T)\}, \quad (2.19)$$

where

$$W_{k-1,0}^{(2)}(T) = [P_{k-1}(T)]^3, \quad (2.20)$$

$$W_{k-2,f}^{(2)}(f) = \{\mathbf{u}_f = u_1 \mathbf{t}_{f,1} + u_2 \mathbf{t}_{f,2} : u_1, u_2 \in P_{k-2}(f)\}, \quad (2.21)$$

$$W_{k-3,e}^{(2)}(e) = \{\mathbf{u}_e = u \mathbf{t}_e : u \in P_{k-3}(e)\}. \quad (2.22)$$

The third space $W_k^{(3)}(T)$ is defined as

$$W_k^{(3)}(T) = \{\mathbf{w} = \{\mathbf{w}_0, \mathbf{w}_f\} : \mathbf{w}_0 \in W_{k-2,0}^{(3)}(T), \mathbf{w}_f \in W_{k-3,f}^{(3)}(f)\}, \quad (2.23)$$

where

$$W_{k-2,0}^{(3)}(T) = [P_{k-2}(T)]^3, \quad (2.24)$$

$$W_{k-3,f}^{(3)}(f) = \{\mathbf{w}_f = w_f \mathbf{n}_f : w_f \in P_{k-3}(f)\}, \quad (2.25)$$

with $\mathbf{n}_f$ a unit normal vector to the face f . The fourth polynomial space $W_k^{(4)}(T)$ is defined as follows:

$$W_{k-3}^{(4)}(T) = P_{k-3}(T). \quad (2.26)$$

Analogously, we define three weak gradient operators $\nabla_{w,0}$, $\nabla_{w,f}$ and $\nabla_{w,e}$ on T , f and e accordingly. More precisely, for $v = \{v_0, v_f, v_e, v_n\} \in W_k^{(1)}(T)$, the weak gradient operator on volume T , denoted as $\nabla_{w,0}v \in W_{k-1,0}^{(2)}(T)$, is defined by

$$(\nabla_{w,0}v, \varphi)_T = -(v_0, \nabla \cdot \varphi)_T + (v_f, \varphi \cdot \mathbf{n})_{\partial T} \quad \forall \varphi \in W_{k-1,0}^{(2)}(T). \quad (2.27)$$

The surface weak gradient operator on face $f \in F(T)$, denoted as $\nabla_{w,f}v \in W_{k-2,f}^{(2)}(f)$, is defined on f satisfying

$$(\nabla_{w,f}v, \theta \times \mathbf{n}_f)_f = -(v_f, \nabla \times \theta \cdot \mathbf{n}_f)_f + (v_e, \theta \cdot \mathbf{t}_{\partial f})_{\partial f} \quad \forall \theta \in W_{k-2,f}^{(2)}(f). \quad (2.28)$$

The edge weak gradient or directional derivative on edge $e \in E(T)$, denoted as $\nabla_{w,e}v \in W_{k-3,e}^{(2)}(e)$, is defined on e such that

$$(\nabla_{w,e}v, \varphi \mathbf{t}_e)_e = -(v_e, \nabla \varphi \cdot \mathbf{t}_e)_e + (v_n, \varphi \mathbf{t}_e \cdot \mathbf{n}_{\partial e})_{\partial e} \quad \forall \varphi \mathbf{t}_e \in W_{k-3,e}^{(2)}(e). \quad (2.29)$$

Collectively, we define a weak gradient mapping $\nabla_w : W_k^{(1)}(T) \mapsto W_{k-1}^{(2)}(T)$ as follows:

$$\nabla_w v := \{\nabla_{w,0}v, \nabla_{w,f}v, \nabla_{w,e}v\}. \quad (2.30)$$

Next, we define two weak curl operators $\nabla_{w,0} \times$ and $\nabla_{w,f} \times$ on T and f respectively. For any $\mathbf{u} = \{\mathbf{u}_0, \mathbf{u}_f\} \in W_{k-1}^{(2)}(T)$, the weak curl on volume T , denoted as $\nabla_{w,0} \times \mathbf{u} \in W_{k-2,0}^{(3)}(T)$, is defined on T by

$$(\nabla_{w,0} \times \mathbf{u}, \theta)_T = (\mathbf{u}_0, \nabla \times \theta)_T + (\mathbf{u}_f, \theta \times \mathbf{n})_{\partial T} \quad \forall \theta \in W_{k-2,0}^{(3)}(T). \quad (2.31)$$

The surface weak curl on face $f \in F(T)$, denoted as $\nabla_{w,f} \times \mathbf{u} \in W_{k-3,f}^{(3)}(f)$, is defined on f by

$$(\nabla_{w,f} \times \mathbf{u}, \tau \mathbf{n}_f)_f = (\mathbf{u}_f, \nabla \tau \times \mathbf{n}_f)_f + \langle \mathbf{u}_e, \tau \mathbf{t} \rangle_{\partial f} \quad \forall \tau \mathbf{n} \in W_{k-3,f}^{(3)}(f). \quad (2.32)$$

The composite weak curl mapping $\nabla_w \times : W_{k-1}^{(2)}(T) \mapsto W_{k-2}^{(3)}(T)$ is then defined by

$$\nabla_w \times \mathbf{u} := \{\nabla_{w,0} \times \mathbf{u}, \nabla_{w,f} \times \mathbf{u}\}. \quad (2.33)$$

Finally, the weak divergence $\nabla_w \cdot \mathbf{w} \in W_{k-3}^{(4)}(T)$ for any $\mathbf{w} = \{\mathbf{w}_0, \mathbf{w}_f\} \in W_{k-2}^{(3)}(T)$ is defined on T by the following equation

$$(\nabla_w \cdot \mathbf{w}, \tau)_T = -(\mathbf{w}_0, \nabla \tau)_T + (\mathbf{w}_f, \tau \mathbf{n})_{\partial T} \quad \forall \tau \in W_{k-3}^{(4)}(T). \quad (2.34)$$

3. Complex sequences

The goal of this section is to show that the WG sequences in Figs. 1.1 and 1.2 are complexes in the sense that the composition of any two consecutive operations is zero.

Lemma 3.1. For the weak gradient operator ∇_w and the weak curl operator $\nabla_w \times$ defined in (2.13) and (2.16) or in (2.30) and (2.33), the following identity holds true

$$\nabla_w \times (\nabla_w v) = 0 \quad (3.1)$$

for all v in $V_k^{(1)}(T)$ or $W_k^{(1)}(T)$.

Proof. By the definition of $\nabla_w \times$ in (2.16) or (2.33), we need to show $\nabla_{w,0} \times (\nabla_w v) = 0$ and $\nabla_{w,f} \times (\nabla_w v) = 0$.

For any v in $V_k^{(1)}(T)$ or $W_k^{(1)}(T)$, we have $\nabla_w v = \{\nabla_{w,0} v, \nabla_{w,f} v, \nabla_{w,e} v\}$. By letting $\mathbf{u} = \nabla_w v$ in (2.14) or (2.31) and using (2.10), (2.11), (2.27) and (2.28), we have for any θ in $V_{k,0}^{(3)}(T)$ or $W_{k-2,0}^{(3)}(T)$ that

$$\begin{aligned} & (\nabla_{w,0} \times (\nabla_w v), \theta)_T \\ &= (\nabla_{w,0} v, \nabla \times \theta)_T + (\nabla_{w,f} v, \theta \times \mathbf{n})_{\partial T} \\ &= -(v_0, \nabla \cdot (\nabla \times \theta))_T + (v_f, (\nabla \times \theta) \cdot \mathbf{n})_{\partial T} + (\nabla_{w,f} v, \theta \times \mathbf{n})_{\partial T} \\ &= (v_f, (\nabla \times \theta) \cdot \mathbf{n})_{\partial T} - (v_f, (\nabla \times \theta) \cdot \mathbf{n})_{\partial T} + \sum_{f \in F(T)} \langle v_e, \theta \cdot \mathbf{t} \rangle_{\partial f} \\ &= 0, \end{aligned}$$

where we have used the fact that $\sum_{f \in F(T)} \langle v_e, \theta \cdot \mathbf{t} \rangle_{\partial f} = 0$, as each edge is shared by two adjacent faces with tangential vector $\mathbf{t}$ of opposite directions.

Next we show that $\nabla_{w,f} \times (\nabla_w v) = 0$ on each face $f \in F(T)$. To this end, on any given face $f \in F(T)$, by letting $\mathbf{u} = \nabla_w v$ in (2.15) or (2.32) and using (2.11), (2.12), (2.28) and (2.29), we have for any τ in $V_{k,f}^{(3)}(f)$ or $W_{k-3,f}^{(3)}(f)$,

$$\begin{aligned} (\nabla_{w,f} \times (\nabla_w v), \tau \mathbf{n})_f &= (\nabla_{w,f} v, \nabla \tau \times \mathbf{n})_f + (\nabla_{w,e} v, \tau \mathbf{t}_{\partial f})_{\partial f} \\ &= -(v_f, \nabla \times (\nabla \tau) \cdot \mathbf{n})_f + (v_e, \nabla \tau \cdot \mathbf{t}_{\partial f})_{\partial f} + (\nabla_{w,e} v, \tau \mathbf{t}_{\partial f})_{\partial f} \\ &= (v_e, \nabla \tau \cdot \mathbf{t}_{\partial f})_{\partial f} - (v_e, \nabla \tau \cdot \mathbf{t}_{\partial f})_{\partial f} + \sum_{e \in \partial f} (v_n, \tau \mathbf{t}_{\partial f} \cdot \mathbf{n}_{\partial(\partial f)})_{\partial e} \\ &= \sum_{e \in \partial f} (v_n, \tau \mathbf{t}_{\partial f} \cdot \mathbf{n}_{\partial(\partial f)})_{\partial e} = 0. \end{aligned}$$

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

Lemma 3.2. For the weak curl and the weak divergence operator $\nabla_w \times$ and $\nabla_w \cdot$ defined in (2.16) and (2.17) or (2.33) and (2.34) respectively, the following identity holds true

$$\nabla_w \cdot (\nabla_w \times \mathbf{u}) = 0 \quad (3.2)$$

for all $\mathbf{u}$ in $V_k^{(2)}(T)$ or $W_{k-1}^{(2)}(T)$.

Proof. Recall that for $\mathbf{u} = \{\mathbf{u}_0, \mathbf{u}_f, \mathbf{u}_e\}$ in $V_k^{(2)}(T)$ or $W_{k-1}^{(2)}(T)$, we have $\nabla_w \times \mathbf{u} = \{\nabla_{w,0} \times \mathbf{u}, \nabla_{w,f} \times \mathbf{u}\}$. By letting $\mathbf{w} = \nabla_w \times \mathbf{u}$ in (2.17) or (2.34) and using (2.14) and (2.15) or (2.31) and (2.32), we have for any τ in $V_k^{(4)}(T)$ or $W_{k-3}^{(4)}(T)$ that

$$\begin{aligned} (\nabla_w \cdot (\nabla_w \times \mathbf{u}), \tau)_T &= -(\nabla_{w,0} \times \mathbf{u}, \nabla \tau)_T + (\nabla_{w,f} \times \mathbf{u}, \tau \mathbf{n})_{\partial T} \\ &= -(\mathbf{u}_0, \nabla \times \nabla \tau)_T - (\mathbf{u}_f, \nabla \tau \times \mathbf{n})_{\partial T} \\ &\quad + (\mathbf{u}_f, \nabla \tau \times \mathbf{n})_{\partial T} + \sum_{f \subset \partial T} (\mathbf{u}_e, \tau \mathbf{t})_{\partial f} \\ &= 0 \end{aligned}$$

which implies (3.2) and thus completes the proof of the lemma. $\square$

4. Commutative properties for the WG de Rham complex 1

Denote by $Q_0^{(1)}$, $Q_f^{(1)}$ and $Q_e^{(1)}$ the L^2 projection operators onto $P_k(T)$, $P_k(f)$, and $P_k(e)$ respectively. For any $v \in H^1(T) \cap C(T)$, define $Q_h^{(1)} v$ by

$$Q_h^{(1)} v = \{Q_0^{(1)} v, Q_f^{(1)} v|_f, Q_e^{(1)} v|_e, v|_n\} \in V_k^{(1)}(T). \quad (4.1)$$

Next, let $Q_0^{(2)}$, $Q_f^{(2)}$ and $Q_e^{(2)}$ be the L^2 projection operators onto $V_{k,0}^{(2)}(T)$, $V_{k,f}^{(2)}(f)$, and $V_{k,e}^{(2)}(e)$ respectively. For $\mathbf{u} \in H(\text{curl}; T) \cap [C(T)]^3$, define $Q_h^{(2)}\mathbf{u}$ by

$$Q_h^{(2)}\mathbf{u} = \{Q_0^{(2)}\mathbf{u}, Q_f^{(2)}(\mathbf{n}_f \times (\mathbf{u}|_f \times \mathbf{n}_f)), Q_e^{(2)}(\mathbf{u}|_e \cdot \mathbf{t}_e)\mathbf{t}_e\} \in V_k^{(2)}(T). \quad (4.2)$$

Analogously, with $Q_0^{(3)}$ and $Q_f^{(3)}$ being the L^2 projection operators onto $V_{k,0}^{(3)}(T)$ and $V_{k,f}^{(3)}(f)$ we define

$$Q_h^{(3)}\mathbf{w} = \{Q_0^{(3)}\mathbf{w}, Q_f^{(3)}(\mathbf{w}|_f \cdot \mathbf{n}_f)\mathbf{n}_f\} \in V_k^{(3)}(T) \quad (4.3)$$

for any $\mathbf{w} \in H(\text{div}; T) \cap [C(T)]^3$. Denote by $Q_h^{(4)}$ the L^2 projection operator from $L^2(T)$ to $V_k^{(4)}(T)$.

Lemma 4.1. For the projection operators $Q_h^{(1)}$ and $Q_h^{(2)}$ defined in (4.1) and (4.2) respectively and the weak gradient ∇_w defined in (2.13), we have

$$\nabla_w(Q_h^{(1)}v) = Q_h^{(2)}\nabla v \quad \forall v \in H^1(T) \cap C^1(T). \quad (4.4)$$

Proof. By (2.13), one has $\nabla_w Q_h^{(1)}v = \{\nabla_{w,0}Q_h^{(1)}v, \nabla_{w,f}Q_h^{(1)}v, \nabla_{w,e}Q_h^{(1)}v\}$. From (4.2), we have $Q_h^{(2)}\nabla v = \{Q_0^{(2)}\nabla v, Q_f^{(2)}(\mathbf{n}_f \times (\nabla v|_f \times \mathbf{n}_f)), Q_e^{(2)}(\nabla v|_e \cdot \mathbf{t}_e)\mathbf{t}_e\}$. Thus it suffices to prove

$$\nabla_{w,0}Q_h^{(1)}v = Q_0^{(2)}\nabla v, \quad \text{in } T, \quad (4.5)$$

$$\nabla_{w,f}Q_h^{(1)}v = Q_f^{(2)}(\mathbf{n}_f \times (\nabla v \times \mathbf{n}_f)), \quad \text{on } f \in F(T), \quad (4.6)$$

$$\nabla_{w,e}Q_h^{(1)}v = Q_e^{(2)}(\nabla v \cdot \mathbf{t}_e)\mathbf{t}_e, \quad \text{on } e \in E(T). \quad (4.7)$$

To prove (4.5), we have from (2.10) that for $v \in H^1(T) \cap C^1(T)$ and any $\varphi \in V_{k,0}^{(2)}(T)$,

$$\begin{aligned} (\nabla_{w,0}Q_h^{(1)}v, \varphi)_T &= -(Q_0^{(1)}v, \nabla \cdot \varphi)_T + (Q_f^{(1)}v, \varphi \cdot \mathbf{n})_{\partial T} \\ &= -(v, \nabla \cdot \varphi)_T + (v, \varphi \cdot \mathbf{n})_{\partial T} \\ &= (\nabla v, \varphi)_T = (Q_0^{(2)}\nabla v, \varphi)_T, \end{aligned}$$

which leads to $\nabla_{w,0}Q_h^{(1)}v = Q_0^{(2)}\nabla v$ in T , and thus proves (4.5).

Next, from (2.11) we have for any $\theta \in V_{k,f}^{(2)}(f)$

$$\begin{aligned} (\nabla_{w,f}Q_h^{(1)}v, \theta \times \mathbf{n}_f)_f &= -(Q_f^{(1)}v, \nabla \times \theta \cdot \mathbf{n}_f)_f + (Q_e^{(1)}v, \theta \cdot \mathbf{t})_{\partial f} \\ &= -(v, \nabla \times \theta \cdot \mathbf{n}_f)_f + (v, \theta \cdot \mathbf{t})_{\partial f} \\ &= (\nabla v, \theta \times \mathbf{n}_f)_f = (\mathbf{n}_f \times (\nabla v \times \mathbf{n}_f), \theta \times \mathbf{n}_f)_f \\ &= (Q_f^{(2)}(\mathbf{n}_f \times (\nabla v \times \mathbf{n}_f)), \theta \times \mathbf{n}_f)_f, \end{aligned}$$

which verifies the identity (4.6).

To derive (4.7), from (2.12) we have for $v \in H^1(T) \cap C^1(T)$ and any $\varphi \in V_{k,e}^{(2)}(e)$ that

$$\begin{aligned} (\nabla_{w,e}Q_h^{(1)}v, \varphi \mathbf{t}_e)_e &= -(Q_e^{(1)}v, \nabla \varphi \cdot \mathbf{t}_e)_e + (v, \varphi \mathbf{t}_e \cdot \mathbf{n}_{\partial e})_{\partial e} \\ &= -(v, \nabla \varphi \cdot \mathbf{t}_e)_e + (v, \varphi \mathbf{t}_e \cdot \mathbf{n}_{\partial e})_{\partial e} \\ &= (\nabla v, \varphi \mathbf{t}_e)_e = (Q_e^{(2)}(\nabla v \cdot \mathbf{t}_e)\mathbf{t}_e, \varphi \mathbf{t}_e)_e \end{aligned}$$

which verifies (4.7). This completes the proof of the lemma. $\square$

Lemma 4.2. For the projections $Q_h^{(2)}$ and $Q_h^{(3)}$ defined in (4.2) and (4.3) respectively and the weak curl $\nabla_w \times$ defined in (2.16), we have

$$\nabla_w \times (Q_h^{(2)}\mathbf{u}) = Q_h^{(3)}\nabla \times \mathbf{u} \quad \forall \mathbf{u} \in H(\text{curl}, T) \cap [C^1(T)]^3. \quad (4.8)$$

Proof. By (2.16), one has $\nabla_w \times Q_h^{(2)}\mathbf{u} = \{\nabla_{w,0} \times Q_h^{(2)}\mathbf{u}, \nabla_{w,f} \times Q_h^{(2)}\mathbf{u}\}$. By (4.3), $Q_h^{(3)}\nabla \times \mathbf{u} = \{Q_0^{(3)}\nabla \times \mathbf{u}, Q_f^{(3)}(\nabla \times \mathbf{u} \cdot \mathbf{n}_f)\mathbf{n}_f\}$. Thus, it suffices to prove

$$\nabla_{w,0} \times Q_h^{(2)}\mathbf{u} = Q_0^{(3)}\nabla \times \mathbf{u}, \quad (4.9)$$

$$\nabla_{w,f} \times Q_h^{(2)}\mathbf{u} = Q_f^{(3)}(\nabla \times \mathbf{u} \cdot \mathbf{n}_f)\mathbf{n}_f. \quad (4.10)$$

It follows from (2.14) that for $\mathbf{u} \in H(\text{curl}, T) \cap [C^1(T)]^3$ and any $\theta \in V_{k,0}^{(3)}(T)$,

$$\begin{aligned} (\nabla_{w,0} \times Q_h^{(2)} \mathbf{u}, \theta)_T &= (Q_0^{(2)} \mathbf{u}, \nabla \times \theta)_T + (Q_f^{(2)} (\mathbf{n} \times (\mathbf{u} \times \mathbf{n})), \theta \times \mathbf{n})_{\partial T} \\ &= (\mathbf{u}, \nabla \times \theta)_T + (\mathbf{n} \times (\mathbf{u} \times \mathbf{n}), \theta \times \mathbf{n})_{\partial T} \\ &= (\mathbf{u}, \nabla \times \theta)_T + (\mathbf{u}, \theta \times \mathbf{n})_{\partial T} \\ &= (\nabla \times \mathbf{u}, \theta)_T = (Q_0^{(3)} \nabla \times \mathbf{u}, \theta)_T \end{aligned}$$

which verifies (4.9). Next, from (2.15) we have for any $\tau \mathbf{n}_f \in V_{k,f}^{(3)}(f)$,

$$\begin{aligned} (\nabla_{w,f} \times Q_h^{(2)} \mathbf{u}, \tau \mathbf{n}_f)_f &= (Q_f^{(2)} (\mathbf{n}_f \times (\mathbf{u} \times \mathbf{n}_f)), \nabla \tau \times \mathbf{n}_f)_f + (Q_e^{(2)} (\mathbf{u} \cdot \mathbf{t}) \mathbf{t}, \tau \mathbf{t})_{\partial f} \\ &= (\mathbf{n}_f \times (\mathbf{u} \times \mathbf{n}_f), \nabla \tau \times \mathbf{n}_f)_f + ((\mathbf{u} \cdot \mathbf{t}) \mathbf{t}, \tau \mathbf{t})_{\partial f} \\ &= (\mathbf{u}, \nabla \tau \times \mathbf{n}_f)_f + (\mathbf{u}, \tau \mathbf{t})_{\partial f} \\ &= (\nabla \times \mathbf{u} \cdot \mathbf{n}_f, \tau)_f = (Q_f^{(3)} (\nabla \times \mathbf{u} \cdot \mathbf{n}_f) \mathbf{n}_f, \tau \mathbf{n}_f)_f, \end{aligned}$$

which leads to (4.10). This completes the proof of the lemma. $\square$

Lemma 4.3. For the operator $Q_h^{(3)}$ defined in (4.3), the L^2 projection operator $Q_h^{(4)}$ from $L^2(T)$ to $V_k^{(4)}(T)$, and the weak divergence operator $\nabla_w \cdot$ defined in (2.17), the following identity holds true

$$\nabla_w \cdot (Q_h^{(3)} \mathbf{w}) = Q_h^{(4)} \nabla \cdot \mathbf{w} \quad \forall \mathbf{w} \in H(\text{div}; T) \cap [C(T)]^3. \quad (4.11)$$

Proof. It follows from (2.17) that for $\mathbf{w} \in H(\text{div}, T) \cap [C(T)]^3$ and any $\tau \in V_k^{(4)}(T)$,

$$\begin{aligned} (\nabla_w \cdot Q_h^{(3)} \mathbf{w}, \tau)_T &= -(Q_0^{(3)} \mathbf{w}, \nabla \tau)_T + (Q_f^{(3)} (\mathbf{w} \cdot \mathbf{n}) \mathbf{n}, \tau \mathbf{n})_{\partial T} \\ &= -(\mathbf{w}, \nabla \tau)_T + (\mathbf{w}, \tau \mathbf{n})_{\partial T} \\ &= (\nabla \cdot \mathbf{w}, \tau)_T = (Q_h^{(4)} \nabla \cdot \mathbf{w}, \tau)_T, \end{aligned}$$

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

5. Commutative properties for the WG de Rham complex 2

In this section, we show that the diagram in Fig. 1.2 commutes with properly defined operators $R_h^{(1)}$, $R_h^{(2)}$, $R_h^{(3)}$ and $R_h^{(4)}$. To this end, denote by $R_0^{(1)}$, $R_f^{(1)}$ and $R_e^{(1)}$ the L^2 projection operators onto $P_k(T)$, $P_{k-1}(f)$, and $P_{k-2}(e)$ respectively. For any $v \in H^1(T) \cap C^1(T)$, we define $R_h^{(1)} v$ as follows

$$R_h^{(1)} v = \{R_0^{(1)} v, R_f^{(1)} v|_f, R_e^{(1)} v|_e, v|_n\} \in W_k^{(1)}(T). \quad (5.1)$$

Next, denote by $R_0^{(2)}$, $R_f^{(2)}$ and $R_e^{(2)}$ the L^2 projection operators onto $W_{k-1,0}^{(2)}(T)$, $W_{k-2,f}^{(2)}(f)$, and $W_{k-3,e}^{(2)}(e)$ respectively. For $\mathbf{u} \in H(\text{curl}; T) \cap [C(T)]^3$, we define $R_h^{(2)} \mathbf{u}$ as follows

$$R_h^{(2)} \mathbf{u} = \{R_0^{(2)} \mathbf{u}, R_f^{(2)} (\mathbf{n}_f \times (\mathbf{u}|_f \times \mathbf{n}_f)), R_e^{(2)} (\mathbf{u}|_e \cdot \mathbf{t}_e) \mathbf{t}_e\} \in W_{k-1}^{(2)}(T). \quad (5.2)$$

Analogously, with $R_0^{(3)}$ and $R_f^{(3)}$ being the L^2 projection operators onto $W_{k-2,0}^{(3)}(T)$ and $W_{k-3,f}^{(3)}(f)$, we may define $R_h^{(3)} \mathbf{w}$ by

$$R_h^{(3)} \mathbf{w} = \{R_0^{(3)} \mathbf{w}, R_f^{(3)} (\mathbf{w} \cdot \mathbf{n}_f) \mathbf{n}_f\} \in W_{k-2}^{(3)}(T) \quad (5.3)$$

for all $\mathbf{w} \in H(\text{div}; T) \cap [C(T)]^3$. Our fourth operator $R_h^{(4)}$ is given as the L^2 projection operator from $L^2(T)$ to $W_{k-3}^{(4)}(T)$.

Lemma 5.1. For the linear operators $R_h^{(1)}$ and $R_h^{(2)}$ defined in (5.1) and (5.2) respectively and the weak gradient ∇_w defined in (2.30), the following identity holds true

$$\nabla_w (R_h^{(1)} v) = R_h^{(2)} \nabla v \quad (5.4)$$

for all $v \in H^1(T) \cap C^1(T)$,

Proof. It follows from (2.30) that $\nabla_w R_h^{(1)} v = \{\nabla_{w,0} R_h^{(1)} v, \nabla_{w,f} R_h^{(1)} v, \nabla_{w,e} R_h^{(1)} v\}$. By (5.2), we have

$$R_h^{(2)} \nabla v = \{R_0^{(2)} \nabla v, R_f^{(2)} (\mathbf{n}_f \times (\nabla v \times \mathbf{n}_f)), R_e^{(2)} (\nabla v \cdot \mathbf{t}_e) \mathbf{t}_e\}.$$

Thus we need to prove

$$\nabla_{w,0} R_h^{(1)} v = R_0^{(2)} \nabla v, \quad \text{in } T, \quad (5.5)$$

$$\nabla_{w,f} R_h^{(1)} v = R_f^{(2)} (\mathbf{n}_f \times (\nabla v \times \mathbf{n}_f)), \quad \text{on } f \in F(T), \quad (5.6)$$

$$\nabla_{w,e} R_h^{(1)} v = R_e^{(2)} (\nabla v \cdot \mathbf{t}_e) \mathbf{t}_e, \quad \text{on } e \in E(T). \quad (5.7)$$

From (2.27), we have for any $\varphi \in W_{k-1,0}^{(2)}(T)$

$$\begin{aligned} (\nabla_{w,0} R_h^{(1)} v, \varphi)_T &= -(R_0^{(1)} v, \nabla \cdot \varphi)_T + (R_f^{(1)} v, \varphi \cdot \mathbf{n})_{\partial T} \\ &= -(v, \nabla \cdot \varphi)_T + (v, \varphi \cdot \mathbf{n})_{\partial T} \\ &= (\nabla v, \varphi)_T = (R_0^{(2)} \nabla v, \varphi)_T, \end{aligned}$$

which proves (5.5).

Next, from (2.28) we have for any $\theta \in W_{k-2,f}^{(2)}(f)$

$$\begin{aligned} (\nabla_{w,f} R_h^{(1)} v, \theta \times \mathbf{n}_f)_f &= -(R_f^{(1)} v, \nabla \times \theta \cdot \mathbf{n}_f)_f + (R_e^{(1)} v, \theta \cdot \mathbf{t})_{\partial f} \\ &= -(v, \nabla \times \theta \cdot \mathbf{n}_f)_f + (v, \theta \cdot \mathbf{t})_{\partial f} \\ &= (\nabla v, \theta \times \mathbf{n}_f)_f \\ &= (\mathbf{n}_f \times (\nabla v \times \mathbf{n}_f), \theta \times \mathbf{n}_f)_f \\ &= (R_f^{(2)}(\mathbf{n}_f \times (\nabla v \times \mathbf{n}_f)), \theta \times \mathbf{n}_f)_f, \end{aligned}$$

which verifies (5.6).

Finally, from (2.29), we have for any $\varphi \in W_{k-3,e}^{(2)}(e)$

$$\begin{aligned} (\nabla_{w,e} R_h^{(1)} v, \varphi \mathbf{t}_e)_e &= -(R_e^{(1)} v, \nabla \varphi \cdot \mathbf{t}_e)_e + (v, \varphi \mathbf{t}_e \cdot \mathbf{n}_{\partial e})_{\partial e} \\ &= -(v, \nabla \varphi \cdot \mathbf{t}_e)_e + (v, \varphi \mathbf{t}_e \cdot \mathbf{n}_{\partial e})_{\partial e} \\ &= (\nabla v \cdot \mathbf{t}_e, \varphi)_e = (R_e^{(2)}(\nabla v \cdot \mathbf{t}_e) \mathbf{t}_e, \varphi \mathbf{t}_e)_e, \end{aligned}$$

which proves (5.7). This completes the proof of the lemma. $\square$

Lemma 5.2. For the linear operators $R_h^{(2)}$ and $R_h^{(3)}$ defined in (5.2) and (5.3) and the weak curl $\nabla_w \times$ defined in (2.33), the following identity holds true:

$$\nabla_w \times (R_h^{(2)} \mathbf{u}) = R_h^{(3)} \nabla \times \mathbf{u} \quad (5.8)$$

for all $\mathbf{u} \in H(\text{curl}; T) \cap [C(T)]^3$.

Proof. By (2.33), one has $\nabla_w \times R_h^{(2)} \mathbf{u} = \{\nabla_{w,0} \times R_h^{(2)} \mathbf{u}, \nabla_{w,f} \times R_h^{(2)} \mathbf{u}\}$. By (5.3), we have $R_h^{(3)} \nabla \times \mathbf{u} = \{R_0^{(3)} \nabla \times \mathbf{u}, R_f^{(3)} (\nabla \times \mathbf{u} \cdot \mathbf{n}_f) \mathbf{n}_f\}$. Thus it suffices to prove

$$\nabla_{w,0} \times R_h^{(2)} \mathbf{u} = R_0^{(3)} \nabla \times \mathbf{u}, \quad \nabla_{w,f} \times R_h^{(2)} \mathbf{u} = R_f^{(3)} (\nabla \times \mathbf{u} \cdot \mathbf{n}_f) \mathbf{n}_f. \quad (5.9)$$

First, it follows from (2.31) that for $\mathbf{u} \in H(\text{curl}; T) \cap [C(T)]^3$ and any $\theta \in W_{k-2,0}^{(3)}(T)$,

$$\begin{aligned} (\nabla_{w,0} \times R_h^{(2)} \mathbf{u}, \theta)_T &= (R_0^{(2)} \mathbf{u}, \nabla \times \theta)_T + (R_f^{(2)}(\mathbf{n} \times (\mathbf{u} \times \mathbf{n})), \theta \times \mathbf{n})_{\partial T} \\ &= (\mathbf{u}, \nabla \times \theta)_T + (\mathbf{n} \times (\mathbf{u} \times \mathbf{n}), \theta \times \mathbf{n})_{\partial T} \\ &= (\mathbf{u}, \nabla \times \theta)_T + (\mathbf{u}, \theta \times \mathbf{n})_{\partial T} \\ &= (\nabla \times \mathbf{u}, \theta)_T = (R_0^{(3)} \nabla \times \mathbf{u}, \theta)_T, \end{aligned}$$

which implies that $\nabla_{w,0} \times R_h^{(2)} \mathbf{u} = R_0^{(3)} \nabla \times \mathbf{u}$.

Next, from (2.32) we have for any $\tau \mathbf{n} \in W_{k-3,f}^{(3)}(f)$

$$\begin{aligned} (\nabla_{w,f} \times R_h^{(2)} \mathbf{u}, \tau \mathbf{n}_f)_f &= (R_f^{(2)}(\mathbf{n}_f \times (\mathbf{u} \times \mathbf{n}_f)), \nabla \tau \times \mathbf{n}_f)_f + (R_e^{(2)}(\mathbf{u} \cdot \mathbf{t}) \mathbf{t}, \tau \mathbf{t})_{\partial f} \\ &= (\mathbf{n}_f \times (\mathbf{u} \times \mathbf{n}_f), \nabla \tau \times \mathbf{n}_f)_f + ((\mathbf{u} \cdot \mathbf{t}) \mathbf{t}, \tau \mathbf{t})_{\partial f} \\ &= (\mathbf{u}, \nabla \tau \times \mathbf{n}_f)_f + (\mathbf{u}, \tau \mathbf{t})_{\partial f} \\ &= (\nabla \times \mathbf{u} \cdot \mathbf{n}_f, \tau)_f = (R_f^{(3)}(\nabla \times \mathbf{u} \cdot \mathbf{n}_f) \mathbf{n}_f, \tau \mathbf{n}_f)_f, \end{aligned}$$

which implies $\nabla_{w,f} \times R_h^{(2)} \mathbf{u} = R_f^{(3)} (\nabla \times \mathbf{u} \cdot \mathbf{n}_f) \mathbf{n}_f$ on face $f \in F(T)$. This completes the proof of the lemma. $\square$

Lemma 5.3. For the linear operator $R_h^{(3)}$ defined in (5.3), the L^2 projection $R_h^{(4)}$ from $L^2(T)$ to $W_{k-3}^{(4)}(T)$, and the weak divergence $\nabla_w \cdot$ defined in (2.34), the following identity holds true

$$\nabla_w \cdot (R_h^{(3)} \mathbf{w}) = R_h^{(4)} \nabla \cdot \mathbf{w} \quad (5.10)$$

for all $\mathbf{w} \in H(\text{div}, T) \cap [C(T)]^3$.

Proof. It follows from (2.34) that for $\mathbf{w} \in H(\text{div}, T) \cap [C(T)]^3$ and any $\tau \in W_{k-3}^{(4)}(T)$, we have

$$\begin{aligned} (\nabla_w \cdot R_h^{(3)} \mathbf{w}, \tau)_T &= -(R_0^{(3)} \mathbf{w}, \nabla \tau)_T + (R_f^{(3)}(\mathbf{w} \cdot \mathbf{n})\mathbf{n}, \tau \mathbf{n})_{\partial T} \\ &= -(\mathbf{w}, \nabla \tau)_T + (\mathbf{w} \cdot \mathbf{n}, \tau)_{\partial T} \\ &= (\nabla \cdot \mathbf{w}, \tau)_T = (R_h^{(4)} \nabla \cdot \mathbf{w}, \tau)_T, \end{aligned}$$

which proves the lemma. $\square$

6. Exactness for de Rham complex 1

We show that the weak Galerkin de Rham complex 1 is exact for $k = 0$ on tetrahedra and hexahedra.

Theorem 6.1. Let T be a tetrahedron. The following de Rham complex is exact for $k = 0$

$$\begin{array}{ccccccccc} \mathbb{R}^1 & \xrightarrow{I} & C^\infty(T) & \xrightarrow{\nabla} & [C^\infty(T)]^3 & \xrightarrow{\nabla \times} & [C^\infty(T)]^3 & \xrightarrow{\nabla \cdot} & C^\infty(T) & \xrightarrow{N} & 0 \\ & & \downarrow Q_h^{(1)} & & \downarrow Q_h^{(2)} & & \downarrow Q_h^{(3)} & & \downarrow Q_h^{(4)} & & \\ \mathbb{R}^4 & \xrightarrow{I_w} & V_k^{(1)}(T) & \xrightarrow{\nabla_w} & V_k^{(2)}(T) & \xrightarrow{\nabla_w \times} & V_k^{(3)}(T) & \xrightarrow{\nabla_w \cdot} & V_k^{(4)}(T) & \xrightarrow{N} & 0 \end{array} \quad (6.1)$$

Here I_w is the inclusion map that assigns constant value to v_0, v_f, v_e , and v_n ; N stands for the null operator.

Proof. A necessary condition for exactness is zero difference of dimensions; i.e.,

$$\begin{aligned} & \sum_{i=0}^4 (-1)^i \dim V_k^{(i)}(T) \\ &= 4 \\ & \quad - \frac{(k+1)(k+2)(k+3)}{6} - 4 \frac{(k+1)(k+2)}{2} - 6(k+1) - 4 \\ & \quad + 3 \frac{(k+1)(k+2)(k+3)}{6} + 2 \cdot 4 \frac{(k+1)(k+2)}{2} + 6(k+1) \\ & \quad - 3 \frac{(k+1)(k+2)(k+3)}{6} - 4 \frac{(k+1)(k+2)}{2} \\ & \quad + \frac{(k+1)(k+2)(k+3)}{6} \\ &= 0, \end{aligned} \quad (6.2)$$

where we have set $V_k^{(0)}(T) = \mathbb{R}^4$.

We claim that the kernel of the operator ∇_w has dimension 4. To this end, let $v \in V_0^{(1)}(T) \in \text{Ker}(\nabla_w)$; i.e.

$$0 = \nabla_w v = \{\nabla_{w,0} v, \nabla_{w,f} v, \nabla_{w,e} v\}.$$

From the definition (2.12) for $\nabla_{w,e} v$, we have $v_n = \alpha_4$ with a constant α_4 at all the vertices of T . On each face (triangle) $f \in F(T)$, one may construct a linear function $P_{1,f} v$ so that $P_{1,f} v = v_e$ at the center of each edge $e \in \partial f$. From the definition (2.11) we have

$$\nabla_{w,f} v = \nabla_f P_{1,f} v \quad \forall f \in F(T),$$

where ∇_f is the surface gradient operator on the face f . It follows from $\nabla_{w,f} v = 0$ that $\nabla_f P_{1,f} v = 0$ so that $P_{1,f} v = \alpha_3$ for a constant α_3 . This shows that $v_e = \alpha_3$ on all edges. Analogously, we may obtain $v_f = \alpha_2$ on all faces for a constant α_2 . Finally, $v_0 = \alpha_1$ is already a constant that does not enter into the calculation of the weak derivatives for the case of $k = 0$. This shows that the dimension of $\text{Ker}(\nabla_w)$ is 4 so that

$$\text{Range}(I_w) = \text{Ker}(\nabla_w).$$

Since the dimension of $V_0^{(1)}(T) = 15$, thus the dimension of $\text{Range}(\nabla_w) = 15 - 4 = 11$.

Next, we claim that

$$\dim(\text{Range}(\nabla_w \times)) = 6. \quad (6.3)$$

Since $\dim(V_0^{(4)}) = 1$ and $\dim(V_0^{(3)}) = 7$, we have $\dim(\text{Range}(\nabla_w \cdot)) = 1$ and $\dim(\text{Ker}(\nabla_w \cdot)) = 6$. It follows that

$$\text{Range}(\nabla_w \times) = \text{Ker}(\nabla_w \cdot)$$

provided that (6.3) holds true.

From $\dim(V_0^{(2)}) = \dim(\text{Range}(\nabla_w \times)) + \dim(\text{Ker}(\nabla_w \times))$ and (6.3), we have

$$\dim(\text{Ker}(\nabla_w \times)) = \dim(V_0^{(2)}) - 6 = 17 - 6 = 11 = \dim(\text{Range}(\nabla_w)).$$

Hence, we have from (3.1) that

$$\text{Ker}(\nabla_w \times) = \text{Range}(\nabla_w).$$

It remains to prove (6.3). Note that by (3.2), we have $\dim(\text{Range}(\nabla_w \times)) \leq \dim(\text{Ker}(\nabla_w \cdot)) = 6$. Consider the following orthogonal decomposition of $V_0^{(3)}(T)$:

$$V_0^{(3)}(T) = \text{Range}(\nabla_w \times) \oplus W.$$

It is clear that (6.3) is equivalent to $\dim(W) = 1$. For any $w = \{w_0, w_f n_f\} \in W$, we have

$$(w, \nabla_w \times v)_T = 0 \quad \forall v \in V_0^{(2)},$$

which implies

$$(w_0, \nabla_{w,0} \times v)_T = 0, \quad (6.4)$$

$$(w_f n_f, \nabla_{w,f} \times v)_f = 0 \quad \text{on each face } f. \quad (6.5)$$

Using the definition of $\nabla_{w,0}$ in (2.14) and (6.4), we have

$$\sum_f (w_0, n_f \times v_f)_f = 0 \quad \forall v = \{v_0, v_f\} \in V_{0,0}^{(2)} \times V_{0,f}^{(2)},$$

which implies $w_0 \times n_f = 0$ on each face f so that $w_0 = 0$. Next, it follows from (6.5) that

$$\sum_f (w_f n_f, \nabla_{w,f} \times v)_f = \sum_e ([w_f]_e, v_e)_e = 0,$$

which implies the jump $[w_f]_e = 0$ on all edges. Thus, w_f assumes a constant value at all faces $f \in F(T)$. This shows that the function $w \in W$ have the form $w = \{0, c n_f\}$ with $c = \text{const}$ so that $\dim(W) = 1$. This verifies the claim (6.3). $\square$

Theorem 6.2. For $k = 0$, the following de Rham complex is exact on cubic element T :

$$\begin{array}{ccccccccc} \mathbb{R}^1 & \xrightarrow{I} & C^\infty(T) & \xrightarrow{\nabla} & [C^\infty(T)]^3 & \xrightarrow{\nabla \times} & [C^\infty(T)]^3 & \xrightarrow{\nabla \cdot} & C^\infty(T) & \xrightarrow{N} & 0 \\ & & \downarrow Q_h^{(1)} & & \downarrow Q_h^{(2)} & & \downarrow Q_h^{(3)} & & \downarrow Q_h^{(4)} & & \\ \mathbb{R}^8 & \xrightarrow{I_w} & V_k^{(1)}(T) & \xrightarrow{\nabla_w} & V_k^{(2)}(T) & \xrightarrow{\nabla_w \times} & V_k^{(3)}(T) & \xrightarrow{\nabla_w \cdot} & V_k^{(4)}(T) & \xrightarrow{N} & 0 \end{array} \quad (6.6)$$

Proof. Again, the necessary condition of zero difference of dimensions for exactness holds true, as it can be easily seen that

$$\begin{aligned} & \sum_{i=0}^4 (-1)^{i-1} \dim V_k^{(i)}(T) \\ &= -8 \\ &+ \frac{(k+1)(k+2)(k+3)}{6} + 6 \frac{(k+1)(k+2)}{2} + 12(k+1) + 8 \\ &- 3 \frac{(k+1)(k+2)(k+3)}{6} - 2 \cdot 6 \frac{(k+1)(k+2)}{2} - 12(k+1) \\ &+ 3 \frac{(k+1)(k+2)(k+3)}{6} + 6 \frac{(k+1)(k+2)}{2} \\ &- \frac{(k+1)(k+2)(k+3)}{6} = 0, \end{aligned}$$

where we have set $V_k^{(0)}(T) = \mathbb{R}^8$.

Let us show that the kernel of the operator ∇_w has dimension 8. In fact, for any $v \in V_0^{(1)}(T) \in \text{Ker}(\nabla_w)$, we have

$$0 = \nabla_w v = \{\nabla_{w,0} v, \nabla_{w,f} v, \nabla_{w,e} v\}.$$

From the definition (2.12) for $\nabla_{w,e} v$, we see that $v_n = \alpha_8$ with a constant α_8 at all eight vertices of T . On each face (rectangle) $f \in F(T)$, the condition of $\nabla_{w,f} v = 0$ implies v_e has the same value on any parallel edges so that v_e has a

total of 3 independent unknowns. Analogously, the condition of $\nabla_{w,0}v = 0$ implies that v_f has the same value on any parallel faces so that v_f has a total of 3 independent unknowns. With the one free unknown for v_0 , we have a total of 8 independent unknowns for functions in $\text{Ker}(\nabla_w)$ so that

$$\dim(\text{Ker}(\nabla_w)) = 8, \quad (6.7)$$

which leads to

$$\text{Range}(I_w) = \text{Ker}(\nabla_w).$$

Since the dimension of $V_0^{(1)}(T) = 27$, thus the dimension of $\text{Range}(\nabla_w) = 27 - 8 = 19$.

Observe that the proof of (6.3) can be adopted without any modification to yield the following result:

$$\dim(\text{Range}(\nabla_w \times)) = 8. \quad (6.8)$$

Since $\dim(V_0^{(4)}) = 1$ and $\dim(V_0^{(3)}) = 9$, we then have $\dim(\text{Range}(\nabla_w \cdot)) = 1$ and $\dim(\text{Ker}(\nabla_w \cdot)) = 8$. It follows from (6.8) that

$$\text{Range}(\nabla_w \times) = \text{Ker}(\nabla_w \cdot).$$

Next, from $\dim(V_0^{(2)}) = \dim(\text{Range}(\nabla_w \times)) + \dim(\text{Ker}(\nabla_w \times))$ and (6.8), we have

$$\dim(\text{Ker}(\nabla_w \times)) = \dim(V_0^{(2)}) - 8 = 27 - 8 = 19 = \dim(\text{Range}(\nabla_w)).$$

Hence, we have from (3.1) that

$$\text{Ker}(\nabla_w \times) = \text{Range}(\nabla_w). \quad \square$$

References

- [1] A. Angoshtari, A. Yavari, Differential complexes in continuum mechanics, *Arch. Ration. Mech. Anal.* 216 (2015) 193–220.
- [2] D. Arnold, R. Falk, R. Winther, Finite element exterior calculus, homological techniques, and applications, *Acta Numer.* (2006) 1–55.
- [3] D. Arnold, R. Falk, R. Winther, Finite element exterior calculus: from hodge theory to numerical stability, *Bull. Amer. Math. Soc.* 47 (2010) 281–354.
- [4] D. Boffi, A note on the discrete compactness property and the de rham complex, *Appl. Math. Lett.* 14 (2001) 33–38.
- [5] A. Bossavit, Whitney forms: A class of finite elements for three-dimensional computations in electromagnetism, *IEEE Proc. A* 135 (1988) 493–500.
- [6] S. Christiansen, J. Hu, K. Hu, Nodal finite element de rham complexes, *Numer. Math.* 139 (2018) 411–446.
- [7] L. Demkowicz, P. Monk, L. Vardapetyan, W. Rachowicz, De rham diagram for hp finite element spaces, *Comput. Math. Appl.* 39 (2000) 29–38.
- [8] L. Demkowicz, Polynomial exact sequences and projection-based interpolation with application to maxwell equations, in: D. Boffi, L. Gastaldi (Eds.), *Mixed Finite Elements, Compatibility Conditions, and Applications*, in: *Lecture Notes in Mathematics*, vol. 1939, Springer, Berlin, Heidelberg, 2008.
- [9] P. Hauret, F. Hecht, A discrete differential sequence for elasticity based upon continuous displacements, *SIAM J. Sci. Comput.* 35 (2013) B291–B314.
- [10] R. Hiptmair, Canonical construction of finite elements, *Math. Comp.* 68 (1999) 1325–1346.
- [11] A. Gillette, K. Hu, S. Zhang, Nonstandard finite element de rham complexes on cubical meshes, *BIT Numer. Math.* 60 (2020) 373–4093.
- [12] M. Neilan, Discrete and conforming smooth de rham complexes in three dimensions, *Math. Comp.* 84 (2015) 2059–2081.
- [13] X. Tai, R. Winther, A discrete de rham complex with enhanced smoothness, *Calcolo* 43 (2006) 287–306.
- [14] H. Whitney, *Geometric Integration Theory*, Princeton University Press, 1957.
- [15] J. Wang, X. Ye, A weak Galerkin finite element method for second-order elliptic problems, *J. Comput. Appl. Math.* 241 (2013) 103–115.
- [16] L. Mu, J. Wang, X. Ye, S. Zhang, A weak Galerkin finite element method for the maxwell equations, *J. Sci. Comput.* 65 (2015) 363–386.
- [17] 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.
- [18] J. Wang, X. Ye, A weak Galerkin mixed finite element method for second-order elliptic problems, *Math. Comp.* 83 (2014) 2101–2126.