
Introduction

The goal of expressing geometrical relationships through algebraic equations has dominated much of the development of mathematics. This line of thinking goes back to the ancient Greeks, who constructed a set of geometric laws to describe the world as they saw it. Their view of geometry was largely unchallenged until the eighteenth century, when mathematicians discovered new geometries with different properties from the Greeks' *Euclidean* geometry. Each of these new geometries had distinct algebraic properties, and a major preoccupation of nineteenth century mathematicians was to place these geometries within a unified algebraic framework. One of the key insights in this process was made by W.K. Clifford, and this book is concerned with the implications of his discovery.

Before we describe Clifford's discovery (in chapter 2) we have gathered together some introductory material of use throughout this book. This chapter revises basic notions of vector spaces, emphasising pictorial representations of the underlying algebraic rules — a theme which dominates this book. The material is presented in a way which sets the scene for the introduction of Clifford's product, in part by reflecting the state of play when Clifford conducted his research. To this end, much of this chapter is devoted to studying the various products that can be defined between vectors. These include the scalar and vector products familiar from three-dimensional geometry, and the complex and quaternion products. We also introduce the *outer* or *exterior* product, though this is covered in greater depth in later chapters. The material in this chapter is intended to be fairly basic, and those impatient to uncover Clifford's insight may want to jump straight to chapter 2. Readers unfamiliar with the outer product are encouraged to read this chapter, however, as it is crucial to understanding Clifford's discovery.

1.1 Vector (linear) spaces

At the heart of much of geometric algebra lies the idea of vector, or linear spaces. Some properties of these are summarised here and assumed throughout this book. In this section we talk in terms of *vector* spaces, as this is the more common term. For all other occurrences, however, we prefer to use the term *linear* space. This is because the term ‘*vector*’ has a very specific meaning within geometric algebra (as the grade-1 elements of the algebra).

1.1.1 Properties

Vector spaces are defined in terms of two objects. These are the vectors, which can often be visualised as directions in space, and the scalars, which are usually taken to be the real numbers. The vectors have a simple addition operation rule with the following obvious properties:

- (i) Addition is *commutative*:

$$a + b = b + a. \quad (1.1)$$

- (ii) Addition is *associative*:

$$a + (b + c) = (a + b) + c. \quad (1.2)$$

This property enables us to write expressions such as $a + b + c$ without ambiguity.

- (iii) There is an identity element, denoted 0:

$$a + 0 = a. \quad (1.3)$$

- (iv) Every element a has an inverse $-a$:

$$a + (-a) = 0. \quad (1.4)$$

For the case of directed line segments each of these properties has a clear geometric equivalent. These are illustrated in figure 1.1.

Vector spaces also contain a multiplication operation between the scalars and the vectors. This has the property that for any scalar λ and vector a , the product λa is also a member of the vector space. Geometrically, this corresponds to the dilation operation. The following further properties also hold for any scalars λ, μ and vectors a and b :

- (i) $\lambda(a + b) = \lambda a + \lambda b$;
- (ii) $(\lambda + \mu)a = \lambda a + \mu a$;
- (iii) $(\lambda\mu)a = \lambda(\mu a)$;
- (iv) if $1\lambda = \lambda$ for all scalars λ then $1a = a$ for all vectors a .

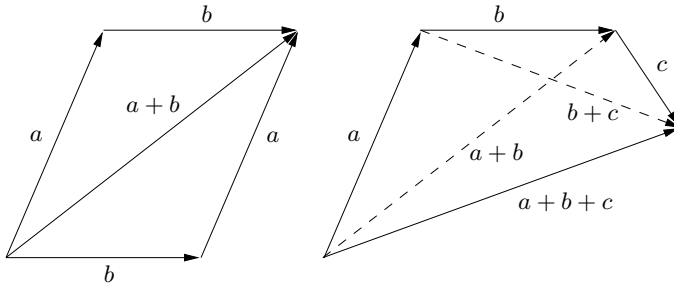


Figure 1.1 A geometric picture of vector addition. The result of $a + b$ is formed by adding the tail of b to the head of a . As is shown, the resultant vector $a + b$ is the same as $b + a$. This finds an algebraic expression in the statement that addition is commutative. In the right-hand diagram the vector $a + b + c$ is constructed two different ways, as $a + (b + c)$ and as $(a + b) + c$. The fact that the results are the same is a geometric expression of the associativity of vector addition.

The preceding set of rules serves to define a vector space completely. Note that the $+$ operation connecting scalars is different from the $+$ operation connecting the vectors. There is no ambiguity, however, in using the same symbol for both.

The following two definitions will be useful later in this book:

- (i) Two vector spaces are said to be *isomorphic* if their elements can be placed in a one-to-one correspondence which preserves sums, and there is a one-to-one correspondence between the scalars which preserves sums and products.
- (ii) If $\mathcal{U}$ and $\mathcal{V}$ are two vector spaces (sharing the same scalars) and all the elements of $\mathcal{U}$ are contained in $\mathcal{V}$, then $\mathcal{U}$ is said to form a *subspace* of $\mathcal{V}$.

1.1.2 Bases and dimension

The concept of dimension is intuitive for simple vector spaces — lines are one-dimensional, planes are two-dimensional, and so on. Equipped with the axioms of a vector space we can proceed to a formal definition of the dimension of a vector space. First we need to define some terms.

- (i) A vector b is said to be a *linear combination* of the vectors $a_1, \dots, a_n$ if scalars $\lambda_1, \dots, \lambda_n$ can be found such that

$$b = \lambda_1 a_1 + \dots + \lambda_n a_n = \sum_{i=1}^n \lambda_i a_i. \quad (1.5)$$

- (ii) A set of vectors $\{a_1, \dots, a_n\}$ is said to be *linearly dependent* if scalars

$\lambda_1, \dots, \lambda_n$ (not all zero) can be found such that

$$\lambda_1 a_1 + \dots + \lambda_n a_n = 0. \quad (1.6)$$

If such a set of scalars cannot be found, the vectors are said to be *linearly independent*.

- (iii) A set of vectors $\{a_1, \dots, a_n\}$ is said to *span* a vector space $\mathcal{V}$ if every element of $\mathcal{V}$ can be expressed as a linear combination of the set.
- (iv) A set of vectors which are both linearly independent and span the space $\mathcal{V}$ are said to form a *basis* for $\mathcal{V}$.

These definitions all carry an obvious, intuitive picture if one thinks of vectors in a plane or in three-dimensional space. For example, it is clear that two independent vectors in a plane provide a basis for all vectors in that plane, whereas any three vectors in the plane are linearly dependent. These axioms and definitions are sufficient to prove the *basis theorem*, which states that *all bases of a vector space have the same number of elements*. This number is called the *dimension* of the space. Proofs of this statement can be found in any textbook on linear algebra, and a sample proof is left to work through as an exercise. Note that any two vector spaces of the same dimension and over the same field are isomorphic.

The axioms for a vector space define an abstract mathematical entity which is already well equipped for studying problems in geometry. In so doing we are not compelled to interpret the elements of the vector space as displacements. Often different interpretations can be attached to isomorphic spaces, leading to different types of geometry (affine, projective, finite, *etc.*). For most problems in physics, however, we need to be able to do more than just add the elements of a vector space; we need to multiply them in various ways as well. This is necessary to formalise concepts such as angles and lengths and to construct higher-dimensional surfaces from simple vectors.

Constructing suitable products was a major concern of nineteenth century mathematicians, and the concepts they introduced are integral to modern mathematical physics. In the following sections we study some of the basic concepts that were successfully formulated in this period. The culmination of this work, Clifford's *geometric product*, is introduced separately in chapter 2. At various points in this book we will see how the products defined in this section can all be viewed as special cases of Clifford's geometric product.

1.2 The scalar product

Euclidean geometry deals with concepts such as lines, circles and perpendicularity. In order to arrive at Euclidean geometry we need to add two new concepts

to our vector space. These are distances between points, which allow us to define a circle, and angles between vectors so that we can say that two lines are perpendicular. The introduction of a scalar product achieves both of these goals.

Given any two vectors a , b , the scalar product $a \cdot b$ is a rule for obtaining a number with the following properties:

- (i) $a \cdot b = b \cdot a$;
- (ii) $a \cdot (\lambda b) = \lambda(a \cdot b)$;
- (iii) $a \cdot (b + c) = a \cdot b + a \cdot c$;
- (iv) $a \cdot a > 0$, unless $a = 0$.

(When we study relativity, this final property will be relaxed.) The introduction of a scalar product allows us to define the length of a vector, $|a|$, by

$$|a| = \sqrt{a \cdot a}. \quad (1.7)$$

Here, and throughout this book, the positive square root is always implied by the $\sqrt{}$ symbol. The fact that we now have a definition of lengths and distances means that we have specified a *metric space*. Many different types of metric space can be constructed, of which the simplest are the *Euclidean* spaces we have just defined.

The fact that for Euclidean space the inner product is positive-definite means that we have a Schwarz inequality of the form

$$|a \cdot b| \leq |a| |b|. \quad (1.8)$$

The proof is straightforward:

$$\begin{aligned} (a + \lambda b) \cdot (a + \lambda b) &\geq 0 && \forall \lambda \\ \Rightarrow a \cdot a + 2\lambda a \cdot b + \lambda^2 b \cdot b &\geq 0 && \forall \lambda \\ \Rightarrow (a \cdot b)^2 &\leq a \cdot a \, b \cdot b, \end{aligned} \quad (1.9)$$

where the last step follows by taking the discriminant of the quadratic in λ . Since all of the numbers in this inequality are positive we recover (1.8). We can now define the *angle* θ between a and b by

$$a \cdot b = |a| |b| \cos(\theta). \quad (1.10)$$

Two vectors whose scalar product is zero are said to be *orthogonal*. It is usually convenient to work with bases in which all of the vectors are mutually orthogonal. If all of the basis vectors are further normalised to have unit length, they are said to form an *orthonormal* basis. If the set of vectors $\{e_1, \dots, e_n\}$ denote such a basis, the statement that the basis is orthonormal can be summarised as

$$e_i \cdot e_j = \delta_{ij}. \quad (1.11)$$

Here the δ_{ij} is the Kronecker delta function, defined by

$$\delta_{ij} = \begin{cases} 1 & \text{if } i = j, \\ 0 & \text{if } i \neq j. \end{cases} \quad (1.12)$$

We can expand any vector a in this basis as

$$a = \sum_{i=1}^n a_i \mathbf{e}_i = a_i \mathbf{e}_i, \quad (1.13)$$

where we have started to employ the *Einstein summation convention* that pairs of indices in any expression are summed over. This convention will be assumed throughout this book. The $\{a_i\}$ are the *components* of the vector a in the $\{\mathbf{e}_i\}$ basis. These are found simply by

$$a_i = \mathbf{e}_i \cdot a. \quad (1.14)$$

The scalar product of two vectors $a = a_i \mathbf{e}_i$ and $b = b_i \mathbf{e}_i$ can now be written simply as

$$a \cdot b = (a_i \mathbf{e}_i) \cdot (b_j \mathbf{e}_j) = a_i b_j \mathbf{e}_i \cdot \mathbf{e}_j = a_i b_j \delta_{ij} = a_i b_i. \quad (1.15)$$

In spaces where the inner product is not positive-definite, such as Minkowski spacetime, there is no equivalent version of the Schwarz inequality. In such cases it is often only possible to define an ‘angle’ between vectors by replacing the cosine function with a cosh function. In these cases we can still introduce orthonormal frames and use these to compute scalar products. The main modification is that the Kronecker delta is replaced by η_{ij} which again is zero if $i \neq j$, but can take values ± 1 if $i = j$.

1.3 Complex numbers

The scalar product is the simplest product one can define between vectors, and once such a product is defined one can formulate many of the key concepts of Euclidean geometry. But this is by no means the only product that can be defined between vectors. In two dimensions a new product can be defined via complex arithmetic. A complex number can be viewed as an ordered pair of real numbers which represents a direction in the complex plane, as was realised by Wessel in 1797. Their product enables complex numbers to perform geometric operations, such as rotations and dilations. But suppose that we take the complex number $z = x + iy$ and square it, forming

$$z^2 = (x + iy)^2 = x^2 - y^2 + 2xyi. \quad (1.16)$$

In terms of vector arithmetic, neither the real nor imaginary parts of this expression have any geometric significance. A more geometrically useful product

is defined instead by

$$zz^* = (x + iy)(x - iy) = x^2 + y^2, \quad (1.17)$$

which returns the square of the length of the vector. A product of two vectors in a plane, z and $w = u + vi$, can therefore be constructed as

$$zw^* = (x + iy)(u - iv) = xu + vy + i(uy - vx). \quad (1.18)$$

The real part of the right-hand side recovers the scalar product. To understand the imaginary term consider the polar representation

$$z = |z|e^{i\theta}, \quad w = |w|e^{i\phi} \quad (1.19)$$

so that

$$zw^* = |z||w|e^{i(\theta - \phi)}. \quad (1.20)$$

The imaginary term has magnitude $|z||w|\sin(\theta - \phi)$, where $\theta - \phi$ is the angle between the two vectors. The magnitude of this term is therefore the area of the parallelogram defined by z and w . The sign of the term conveys information about the *handedness* of the area element swept out by the two vectors. This will be defined more carefully in section 1.6.

We thus have a satisfactory interpretation for both the real and imaginary parts of the product zw^* . The surprising feature is that these are still both parts of a complex number. We thus have a second interpretation for complex addition, as a sum between scalar objects and objects representing plane segments. The advantages of adding these together are precisely the advantages of working with complex numbers as opposed to pairs of real numbers. This is a theme to which we shall return regularly in following chapters.

1.4 Quaternions

The fact that complex arithmetic can be viewed as representing a product for vectors in a plane carries with it a further advantage — it allows us to divide by a vector. Generalising this to three dimensions was a major preoccupation of the physicist W.R. Hamilton (see figure 1.2). Since a complex number $x + iy$ can be represented by two rectangular axes on a plane it seemed reasonable to represent directions in space by a triplet consisting of one real and two complex numbers. These can be written as $x + iy + jz$, where the third term jz represents a third axis perpendicular to the other two. The complex numbers i and j have the properties that $i^2 = j^2 = -1$. The norm for such a triplet would then be

$$(x + iy + jz)(x - iy - jz) = (x^2 + y^2 + z^2) - yz(ij + ji). \quad (1.21)$$

The final term is problematic, as one would like to recover the scalar product here. The obvious solution to this problem is to set $ij = -ji$ so that the last term vanishes.

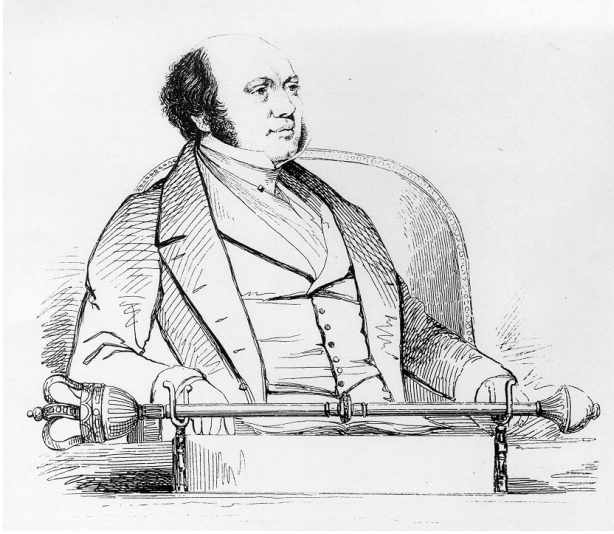


Figure 1.2 *William Rowan Hamilton 1805–1865*. Inventor of quaternions, and one of the key scientific figures of the nineteenth century. He spent many years frustrated at being unable to extend his theory of couples of numbers (complex numbers) to three dimensions. In the autumn of 1843 he returned to this problem, quite possibly prompted by a visit he received from the young German mathematician Eisenberg. Among Eisenberg’s papers was the observation that matrices form the elements of an algebra that was much like ordinary arithmetic except that multiplication was non-commutative. This was the vital step required to find the quaternion algebra. Hamilton arrived at this algebra on 16 October 1843 while out walking with his wife, and carved the equations in stone on Brougham Bridge. His discovery of quaternions is perhaps the best-documented mathematical discovery ever.

The anticommutative law $ij = -ji$ ensures that the norm of a triplet behaves sensibly, and also that multiplication of triplets in a plane behaves in a reasonable manner. The same is not true for the general product of triplets, however. Consider

$$(a + ib + jc)(x + iy + jz) = (ax - by - cz) + i(ay + bx) + j(az + cx) + ij(bz - cy). \quad (1.22)$$

Setting $ij = -ji$ is no longer sufficient to remove the ij term, so the algebra does not close. The only thing for Hamilton to do was to set $ij = k$, where k is some unknown, and see if it could be removed somehow. While walking along the Royal Canal he suddenly realised that if his triplets were instead made up of four terms he would be able to close the algebra in a simple, symmetric way.

To understand his discovery, consider

$$(a + ib + jc + kd)(a - ib - jc - kd) = a^2 + b^2 + c^2 + d^2(-k^2) - bd(ik + ki) - cd(jk + kj), \quad (1.23)$$

where we have assumed that $i^2 = j^2 = -1$ and $ij = -ji$. The expected norm of the above product is $a^2 + b^2 + c^2 + d^2$, which is obtained by setting $k^2 = -1$ and $ik = -ki$ and $jk = -kj$. So what values do we use for jk and ik ? These follow from the fact that $ij = k$, which gives

$$ik = i(ij) = (ii)j = -j \quad (1.24)$$

and

$$kj = (ij)j = -i. \quad (1.25)$$

Thus the multiplication rules for quaternions are

$$i^2 = j^2 = k^2 = -1 \quad (1.26)$$

and

$$ij = -ji = k, \quad jk = -kj = i, \quad ki = -ik = j. \quad (1.27)$$

These can be summarised neatly as $i^2 = j^2 = k^2 = ijk = -1$. It is a simple matter to check that these multiplication laws define a closed algebra.

Hamilton was so excited by his discovery that the very same day he obtained leave to present a paper on the quaternions to the Royal Irish Academy. The subsequent history of the quaternions is a fascinating story which has been described by many authors. Some suggested material for further reading is given at the end of this chapter. In brief, despite the many advantages of working with quaternions, their development was blighted by two major problems.

The first problem was the status of vectors in the algebra. Hamilton identified vectors with *pure quaternions*, which had a null scalar part. On the surface this seems fine — pure quaternions define a three-dimensional vector space. Indeed, Hamilton invented the word ‘*vector*’ precisely for these objects and this is the origin of the now traditional use of i , j and k for a set of orthonormal basis vectors. Furthermore, the full product of two pure quaternions led to the definition of the extremely useful cross product (see section 1.5). The problem is that the product of two pure vectors does not return a new pure vector, so the vector part of the algebra does not close. This means that a number of ideas in complex analysis do not extend easily to three dimensions. Some people felt that this meant that the full quaternion product was of little use, and that the scalar and vector parts of the product should be kept separate. This criticism misses the point that the quaternion product is *invertible*, which does bring many advantages.

The second major difficulty encountered with quaternions was their use in

describing rotations. The irony here is that quaternions offer the clearest way of handling rotations in three dimensions, once one realises that they provide a ‘spin-1/2’ representation of the rotation group. That is, if a is a vector (a pure quaternion) and R is a unit quaternion, a new vector is obtained by the *double-sided* transformation law

$$a' = RaR^*, \quad (1.28)$$

where the $*$ operation reverses the sign of all three ‘imaginary’ components. A consequence of this is that each of the basis quaternions i , j and k generates rotations through π . Hamilton, however, was led astray by the analogy with complex numbers and tried to impose a single-sided transformation of the form $a' = Ra$. This works if the axis of rotation is perpendicular to a , but otherwise does not return a pure quaternion. More damagingly, it forces one to interpret the basis quaternions as generators of rotations through $\pi/2$, which is simply wrong!

Despite the problems with quaternions, it was clear to many that they were a useful mathematical system worthy of study. Tait claimed that quaternions ‘freed the physicist from the constraints of coordinates and allowed thoughts to run in their most natural channels’ — a theme we shall frequently meet in this book. Quaternions also found favour with the physicist James Clerk Maxwell, who employed them in his development of the theory of electromagnetism. Despite these successes, however, quaternions were weighed down by the increasingly dogmatic arguments over their interpretation and were eventually displaced by the hybrid system of vector algebra promoted by Gibbs.

1.5 The cross product

Two of the lasting legacies of the quaternion story are the introduction of the idea of a vector, and the cross product between two vectors. Suppose we form the product of two pure quaternions a and b , where

$$a = a_1i + a_2j + a_3k, \quad b = b_1i + b_2j + b_3k. \quad (1.29)$$

Their product can be written

$$ab = -a_ib_i + c, \quad (1.30)$$

where c is the pure quaternion

$$c = (a_2b_3 - a_3b_2)i + (a_3b_1 - a_1b_3)j + (a_1b_2 - a_2b_1)k. \quad (1.31)$$

Writing $c = c_1i + c_2j + c_3k$ the component relation can be written as

$$c_i = \epsilon_{ijk}a_jb_k, \quad (1.32)$$

where the alternating tensor ϵ_{ijk} is defined by

$$\epsilon_{ijk} = \begin{cases} 1 & \text{if } ijk \text{ is a cyclic permutation of } 123, \\ -1 & \text{if } ijk \text{ is an anticyclic permutation of } 123, \\ 0 & \text{otherwise.} \end{cases} \quad (1.33)$$

We recognise the preceding as defining the cross product of two vectors, $a \times b$. This has the following properties:

- (i) $a \times b$ is perpendicular to the plane defined by a and b ;
- (ii) $a \times b$ has magnitude $|a||b|\sin(\theta)$;
- (iii) the vectors a , b and $a \times b$ form a right-handed set.

These properties can alternatively be viewed as defining the cross product, and from them the algebraic definition can be recovered. This is achieved by starting with a right-handed orthonormal frame $\{\mathbf{e}_i\}$. For these we must have

$$\mathbf{e}_1 \times \mathbf{e}_2 = \mathbf{e}_3 \quad \text{etc.} \quad (1.34)$$

so that we can write

$$\mathbf{e}_i \times \mathbf{e}_j = \epsilon_{ijk} \mathbf{e}_k. \quad (1.35)$$

Expanding out a vector in terms of this basis recovers the formula

$$\begin{aligned} a \times b &= (a_i \mathbf{e}_i) \times (b_j \mathbf{e}_j) \\ &= a_i b_j (\mathbf{e}_i \times \mathbf{e}_j) \\ &= (\epsilon_{ijk} a_i b_j) \mathbf{e}_k. \end{aligned} \quad (1.36)$$

Hence the geometric definition recovers the algebraic one.

The cross product quickly proved itself to be invaluable to physicists, dramatically simplifying equations in dynamics and electromagnetism. In the latter part of the nineteenth century many physicists, most notably Gibbs, advocated abandoning quaternions altogether and just working with the individual scalar and cross products. We shall see in later chapters that Gibbs was misguided in some of his objections to the quaternion product, but his considerable reputation carried the day and by the 1900s quaternions had all but disappeared from mainstream physics.

1.6 The outer product

The cross product has one major failing — it only exists in three dimensions. In two dimensions there is nowhere else to go, whereas in four dimensions the concept of a vector orthogonal to a pair of vectors is not unique. To see this, consider four orthonormal vectors $\mathbf{e}_1, \dots, \mathbf{e}_4$. If we take the pair $\mathbf{e}_1$ and $\mathbf{e}_2$ and attempt



Figure 1.3 *Hermann Gunther Grassmann (1809–1877)*, born in Stettin, Germany (now Szczecin, Poland). A German mathematician and school-teacher, Grassmann was the third of his parents' twelve children and was born into a family of scholars. His father studied theology and became a minister, before switching to teaching mathematics and physics at the Stettin Gymnasium. Hermann followed in his father's footsteps, first studying theology, classical languages and literature at Berlin. After returning to Stettin in 1830 he turned his attention to mathematics and physics. Grassmann passed the qualifying examination to win a teaching certificate in 1839. This exam included a written assignment on the tides, for which he gave a simplified treatment of Laplace's work based upon a new geometric calculus that he had developed. By 1840 he had decided to concentrate on mathematics research. He published the first edition of his geometric calculus, the 300 page *Lineale Ausdehnungslehre* in 1844, the same year that Hamilton announced the discovery of the quaternions. His work did not achieve the same impact as the quaternions, however, and it was many years before his ideas were understood and appreciated by other mathematicians. Disappointed by this lack of interest, Grassmann turned his attention to linguistics and comparative philology, with greater immediate impact. He was an expert in Sanskrit and translated the *Rig-Veda* (1876–1877). He also formulated the linguistic law (named after him) stating that in Indo-European bases, successive syllables may not begin with aspirates. He died before he could see his ideas on geometry being adopted into mainstream mathematics.

to find a vector perpendicular to both of these, we see that any combination of $\mathbf{e}_3$ and $\mathbf{e}_4$ will do.

A suitable generalisation of the idea of the cross product was constructed by

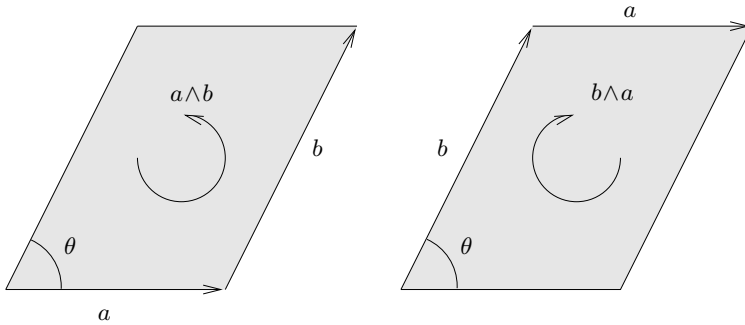


Figure 1.4 *The outer product.* The outer or wedge product of a and b returns a directed area element of area $|a||b|\sin(\theta)$. The orientation of the parallelogram is defined by whether the circuit $a, b, -a, -b$ is right-handed (anticlockwise) or left-handed (clockwise). Interchanging the order of the vectors reverses the orientation and introduces a minus sign in the product.

the remarkable German mathematician H.G. Grassmann (see figure 1.3). His work had its origin in the *Barycentrischer Calcul* of Möbius. There the author introduced expressions like AB for the line connecting the points A and B and ABC for the triangle defined by A , B and C . Möbius also introduced the crucial idea that the sign of the quantity should change if any two points are interchanged. (These *oriented* segments are now referred to as *simplices*.) It was Grassmann's leap of genius to realise that expressions like AB could actually be viewed as a product between vectors. He thus introduced the *outer* or *exterior product* which, in modern notation, we write as $a \wedge b$, or 'a wedge b'.

The outer product can be defined on any vector space and, geometrically, we are not forced to picture these vectors as displacements. Indeed, Grassmann was motivated by a *projective* viewpoint, where the elements of the vector space are interpreted as points, and the outer product of two points defines the line through the points. For our purposes, however, it is simplest to adopt a picture in which vectors represent directed line segments. The outer product then provides a means of encoding a plane, without relying on the notion of a vector perpendicular to it. The result of the outer product is therefore neither a scalar nor a vector. It is a new mathematical entity encoding an oriented plane and is called a *bivector*. It can be visualised as the parallelogram obtained by sweeping one vector along the other (figure 1.4). Changing the order of the vectors reverses the orientation of the plane. The magnitude of $a \wedge b$ is $|a||b|\sin(\theta)$, the same as the area of the plane segment swept out by the vectors.

The outer product of two vectors has the following algebraic properties:

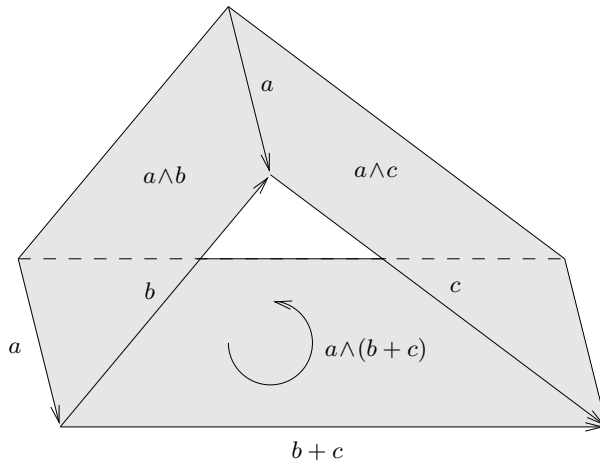


Figure 1.5 *A geometric picture of bivector addition.* In three dimensions any two non-parallel planes share a common line. If this line is denoted a , the two planes can be represented by $a \wedge b$ and $a \wedge c$. Bivector addition proceeds much like vector addition. The planes are combined at a common boundary and the resulting plane is defined by the initial and final edges, as opposed to the initial and final points for vector addition. The mathematical statement of this addition rule is the distributivity of the outer product over addition.

- (i) The product is *antisymmetric*:

$$a \wedge b = -b \wedge a. \quad (1.37)$$

This has the geometric interpretation of reversing the orientation of the surface defined by a and b . It follows immediately that

$$a \wedge a = 0, \quad \text{for all vectors } a. \quad (1.38)$$

- (ii) Bivectors form a linear space, the same way that vectors do. In two and three dimensions the addition of bivectors is easy to visualise. In higher dimensions this addition is not always so easy to visualise, because two planes need not share a common line.
- (iii) The outer product is distributive over addition:

$$a \wedge (b + c) = a \wedge b + a \wedge c. \quad (1.39)$$

This helps to visualise the addition of bivectors which share a common line (see figure 1.5).

While it is convenient to visualise the outer product as a parallelogram, the

actual shape of the object is not conveyed by the result of the product. This can be seen easily by defining $a' = a + \lambda b$ and forming

$$a' \wedge b = a \wedge b + \lambda b \wedge b = a \wedge b. \quad (1.40)$$

The same bivector can therefore be generated by many different pairs of vectors. In many ways it is better to replace the picture of a directed parallelogram with that of a directed circle. The circle defines both the plane and a handedness, and its area is equal to the magnitude of the bivector. This therefore conveys all of the information one has about the bivector, though it does make bivector addition harder to visualise.

1.6.1 Two dimensions

The outer product of any two vectors defines a plane, so one has to go to at least two dimensions to form an interesting product. Suppose then that $\{\mathbf{e}_1, \mathbf{e}_2\}$ are an orthonormal basis for the plane, and introduce the vectors

$$a = a_1 \mathbf{e}_1 + a_2 \mathbf{e}_2, \quad b = b_1 \mathbf{e}_1 + b_2 \mathbf{e}_2. \quad (1.41)$$

The outer product $a \wedge b$ contains

$$\begin{aligned} a \wedge b &= a_1 b_1 \mathbf{e}_1 \wedge \mathbf{e}_1 + a_1 b_2 \mathbf{e}_1 \wedge \mathbf{e}_2 + a_2 b_1 \mathbf{e}_2 \wedge \mathbf{e}_1 + a_2 b_2 \mathbf{e}_2 \wedge \mathbf{e}_2 \\ &= (a_1 b_2 - a_2 b_1) \mathbf{e}_1 \wedge \mathbf{e}_2, \end{aligned} \quad (1.42)$$

which recovers the imaginary part of the product of (1.18). The term therefore immediately has the expected magnitude $|a| |b| \sin(\theta)$. The coefficient of $\mathbf{e}_1 \wedge \mathbf{e}_2$ is positive if a and b have the same orientation as $\mathbf{e}_1$ and $\mathbf{e}_2$. The orientation is defined by traversing the boundary of the parallelogram defined by the vectors a , b , $-a$, $-b$ (see figure 1.4). By convention, we usually work with a right-handed set of reference axes (viewed from above). In this case the coefficient $a_1 b_2 - a_2 b_1$ will be positive if a and b also form a right-handed pair.

1.6.2 Three dimensions

In three dimensions the space of bivectors is also three-dimensional, because each bivector can be placed in a one-to-one correspondence with the vector perpendicular to it. Suppose that $\{\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3\}$ form a right-handed basis (see comments below), and the two vectors a and b are expanded in this basis as $a = a_i \mathbf{e}_i$ and $b = b_i \mathbf{e}_i$. The bivector $a \wedge b$ can then be decomposed in terms of an orthonormal frame of bivectors by

$$\begin{aligned} a \wedge b &= (a_i \mathbf{e}_i) \wedge (b_j \mathbf{e}_j) \\ &= (a_2 b_3 - b_3 a_2) \mathbf{e}_2 \wedge \mathbf{e}_3 + (a_3 b_1 - a_1 b_3) \mathbf{e}_3 \wedge \mathbf{e}_1 \\ &\quad + (a_1 b_2 - a_2 b_1) \mathbf{e}_1 \wedge \mathbf{e}_2. \end{aligned} \quad (1.43)$$

The components in this frame are therefore the same as those of the cross product. But instead of being the components of a vector perpendicular to a and b , they are the components of the bivector $a \wedge b$. It is this distinction which enables the outer product to be defined in any dimension.

1.6.3 *Handedness*

We have started to employ the idea of *handedness* without giving a satisfactory definition of it. The only space in which there is an unambiguous definition of handedness is three dimensions, as this is the space we inhabit and most of us can distinguish our left and right hands. This concept of ‘left’ and ‘right’ is a man-made convention adopted to make our life easier, and it extends to the concept of a frame in a straightforward way. Suppose that we are presented with three orthogonal vectors $\{\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3\}$. We align the 3 axis with the thumb of our right hand and then close our fist. If the direction in which our fist closes is the same as that formed by rotating from the 1 to the 2 axis, the frame is right-handed. If not, it is left-handed.

Swapping any pair of vectors swaps the handedness of a frame. Performing two such swaps returns us to the original handedness. In three dimensions this corresponds to a cyclic reordering, and ensures that the frames $\{\mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3\}$, $\{\mathbf{e}_3, \mathbf{e}_1, \mathbf{e}_2\}$ and $\{\mathbf{e}_2, \mathbf{e}_3, \mathbf{e}_1\}$ all have the same orientation.

There is no agreed definition of a ‘right-handed’ orientation in spaces of dimensions other than three. All one can do is to make sure that any convention used is adopted consistently. In all dimensions the orientation of a set of vectors is changed if any two vectors are swapped. In two dimensions one does still tend to talk about right-handed axes, though the definition is dependent on the idea of looking down on the plane *from above*. The idea of above and below is not a feature of the plane itself, but depends on how we embed it in our three-dimensional world. There is no definition of left or right-handed which is intrinsic to the plane.

1.6.4 *Extending the outer product*

The preceding examples demonstrate that in arbitrary dimensions the components of $a \wedge b$ are given by

$$(a \wedge b)_{ij} = a_{[i} b_{j]} \quad (1.44)$$

where the $[]$ denotes antisymmetrisation. Grassmann was able to take this idea further by defining an outer product for any number of vectors. The idea is a simple extension of the preceding formula. Expressed in an orthonormal frame, the components of the outer product on n vectors are the totally antisymmetrised

products of the components of each vector. This definition has the useful property that the outer product is *associative*,

$$a \wedge (b \wedge c) = (a \wedge b) \wedge c. \quad (1.45)$$

For example, in three dimensions we have

$$a \wedge b \wedge c = (a_i \mathbf{e}_i) \wedge (b_j \mathbf{e}_j) \wedge (c_k \mathbf{e}_k) = \epsilon_{ijk} a_i b_j c_k \mathbf{e}_1 \wedge \mathbf{e}_2 \wedge \mathbf{e}_3, \quad (1.46)$$

which represents a *directed volume* (see section 2.4).

A further feature of the antisymmetry of the product is that the outer product of any set of linearly dependent vectors vanishes. This means that statements like ‘this vector lies on a given plane’, or ‘these two hypersurfaces share a common line’ can be encoded algebraically in a simple manner. Equipped with these ideas, Grassmann was able to construct a system capable of handling geometric concepts in arbitrary dimensions.

Despite Grassmann’s considerable achievement, the book describing his ideas, his *Lineale Ausdehnungslehre*, did not have any immediate impact. This was no doubt due largely to his relative lack of reputation (he was still a German schoolteacher when he wrote this work). It was over twenty years before anyone of note referred to Grassmann’s work, and during this time Grassmann produced a second, extended version of the *Ausdehnungslehre*. In the latter part of the nineteenth century Grassmann’s work started to influence leading figures like Gibbs and Clifford. Gibbs wrote a number of papers praising Grassmann’s work and contrasting it favourably with the quaternion algebra. Clifford used Grassmann’s work as the starting point for the development of his geometric algebra, the subject of this book.

Today, Grassmann’s ideas are recognised as the first presentation of the abstract theory of vector spaces over the field of real numbers. Since his death, his work has given rise to the influential and fashionable areas of *differential forms* and *Grassmann variables*. The latter are anticommuting variables and are fundamental to the foundations of much of modern supersymmetry and superstring theory.

1.7 Notes

Descriptions of linear algebra and vector spaces can be found in most introductory textbooks of mathematics, as can discussions of the scalar and cross products and complex arithmetic. Quaternions, on the other hand, are much less likely to be mentioned. There is a large specialised literature on the quaternions, and a good starting point are the works of Altmann (1986, 1989). Altmann’s paper on ‘Hamilton, Rodrigues and the quaternion scandal’ (1989) is also a good introduction to the history of the subject.

The outer product is covered in most modern textbooks on geometry and

physics, such as those by Nakahara (1990), Schutz (1980), and Gockeler & Schucker (1987). In most of these works, however, the exterior product is only treated in the context of differential forms. Applications to wider topics in geometry have been discussed by Hestenes (1991) and others. A useful summary is provided in the proceedings of the conference *Hermann Gunther Grassmann (1809–1877)*, edited by Schubring (1996). Grassmann's *Lineale Ausdehnungslehre* is also finally available in English translation due to Kannenberg (1995).

For those with a deeper interest in the history of mathematics and the development of vector algebra a good starting point is the set of books by Kline (1972). There are also biographies available of many of the key protagonists. Perhaps even more interesting is to return to their original papers and experience first hand the robust and often humorous language employed at the time. The collected works of J.W. Gibbs (1906) are particularly entertaining and enlightening, and contain a good deal of valuable historical information.

1.8 Exercises

- 1.1 Suppose that the two sets $\{a_1, \dots, a_m\}$ and $\{b_1, \dots, b_n\}$ form bases for the same vector space, and suppose initially that $m > n$. By establishing a contradiction, prove the *basis theorem* that all bases of a vector space have the same number of elements.
- 1.2 Demonstrate that the following define vector spaces:
 - (a) the set of all polynomials of degree less than or equal to n ;
 - (b) all solutions of a given linear homogeneous ordinary differential equation;
 - (c) the set of all $n \times m$ matrices.
- 1.3 Prove that in Euclidean space $|a + b| \leq |a| + |b|$. When does equality hold?
- 1.4 Show that the unit quaternions $\{\pm 1, \pm i, \pm j, \pm k\}$ form a discrete group.
- 1.5 The unit quaternions i, j, k are generators of rotations about their respective axes. Are rotations through either π or $\pi/2$ consistent with the equation $ijk = -1$?
- 1.6 Prove the following:
 - (a) $a \cdot (b \times c) = b \cdot (c \times a) = c \cdot (a \times b)$;
 - (b) $a \times (b \times c) = a \cdot c b - a \cdot b c$;
 - (c) $|a \times b| = |a| |b| \sin(\theta)$, where $a \cdot b = |a| |b| \cos(\theta)$.
- 1.7 Prove that the dimension of the space formed by the exterior product of m vectors drawn from a space of dimension n is

$$\frac{n(n-1) \cdots (n-m+1)}{1 \cdot 2 \cdots m} = \frac{n!}{(n-m)!m!}.$$

- 1.8 Prove that the n -fold exterior product of a set of n *dependent* vectors is zero.
- 1.9 A convex polygon in a plane is specified by the ordered set of points $\{x_0, x_1, \dots, x_n\}$. Prove that the directed area of the polygon is given by

$$A = \frac{1}{2}(x_0 \wedge x_1 + x_1 \wedge x_2 + \cdots + x_n \wedge x_0).$$

What is the significance of the sign? Can you extend the idea to a triangulated surface in three dimensions?