



MAGISTERARBEIT | MASTER'S THESIS

Titel | Title

Der Witt'sche Kürzungs- und Erweiterungssatz

verfasst von | submitted by

Georg Mitterecker BEd

angestrebter akademischer Grad | in partial fulfilment of the requirements for the degree of
Master of Education (MEd)

Wien | Vienna, 2026

Studienkennzahl lt. Studienblatt | Degree
programme code as it appears on the
student record sheet:

UA 199 514 520 02

Studienrichtung lt. Studienblatt | Degree
programme as it appears on the student
record sheet:

Masterstudium Lehramt Sek (AB) Unterrichtsfach
Informatik Unterrichtsfach Mathematik

Betreut von | Supervisor:

ao. Univ.-Prof. Dr. Joachim Mahnkopf

Abstract

Diese Masterarbeit beschäftigt sich mit Struktursätzen und dem Klassifikationsproblem für Bilinearformen. Zu Beginn sollen Grundlagen, die das Arbeiten mit Bilinearformen ermöglichen dargestellt werden. Zentral für die Untersuchung der Struktur einer orthogonalen oder symplektischen Geometrie ist die Konstruktion der nicht-singulären Erweiterung eines singulären Raumes, die selbst ein hyperbolischer Raum ist. Damit lässt sich insbesondere die anisotrope Zerlegung einer symmetrischen Bilinearform herleiten. Die Klassifikation trennt sich auf in symmetrische Bilinearformen und alternierende Bilinearformen. Für symmetrische Bilinearformen hängt die Klassifikation – im Gegensatz zum alternierenden Fall – vom Grundkörper ab. Am Ende der beiden Kapitel steht jeweils der Witt'sche Erweiterungs- und Kürzungssatz, welcher die Eindeutigkeit der anisotropen Zerlegung einer orthogonalen Geometrie beweist.

Abstract

This master thesis is about structure theorems and the classification problem for bilinear forms. First, basic principles for working with bilinear forms are discussed. The central point in studying the structure of an orthogonal or a symplectic geometry is the construction of a non-singular extension of a singular space. The non-singular extension itself is a hyperbolic space. Using non-singular extensions, the anisotropic decomposition of symmetric bilinear forms is derived. The classification is divided in symmetric and alternating bilinear forms. For symmetric bilinear forms the classification depends on the basefield (this is in contrast to the case of alternating bilinear forms). At the end of both chapters Witt's extension and cancellation theorems are obtained; the cancellation theorem in particular yields the uniqueness of the anisotropic decomposition of an orthogonal geometry.

Inhaltsverzeichnis

I	Grundlagen	3
1	Bilinearform	3
2	Matrix einer Bilinearform und Isometrien	6
2.1	Darstellungsmatrix einer Bilinearform	6
2.2	Die Diskriminante einer Form	10
2.3	Isometrien	10
3	Quadratische Formen	13
II	Orthogonalität	15
4	Definitionen und grundlegende Eigenschaften	15
5	Orthogonales Komplement und Orthogonale Summe	20
6	Orthogonale Zerlegung einer Bilinearform	21
7	Hyperbolische Räume	26
III	Klassifikation	29
8	Klassifikationsproblem für metrische Vektorräume	29
9	Klassifikation symmetrischer Bilinearformen	29
9.1	Die Struktur von orthogonalen Geometrien	30
9.2	Orthogonalbasis	32
9.3	Orthogonalbasen: der Fall $\text{char}(F) = 2$	33
9.4	Klassifikation symmetrischer Bilinearform	37
9.4.1	Algebraisch geschlossene Körper	37
9.4.2	Körper der reellen Zahlen	39
9.4.3	Endliche Körper	41

9.5	Rotation, Spiegelung, Symmetrien	45
9.6	Witt'sche Kürzungs- und Erweiterungssatz	46
9.7	Anisotrope Zerlegung einer orthogonalen Geometrie . .	47
10	Symplektische Bilinearform	52
10.1	Nicht-Singuläre Erweiterung eines Untervektorraumes .	52
10.2	Klassifikation symplektischer Bilinearformen	55
10.3	Pfaff'sche Determinante einer alternierenden Matrix . .	58
10.4	Witt'sche Erweiterungs- und Kürzungssatz	59
	Literatur	62

Einleitung

Diese Arbeit befasst sich mit der Struktur und Klassifikation von Bilinearformen. In den Grundlagen wird der Begriff der Bilinearform, sowie die Eigenschaften „schiefsymmetrisch“, aber besonders „symmetrisch“ und „alternierend“ beziehungsweise „orthogonal“ und „symplektisch“ erläutert. Weiters wird die Darstellungsmatrix einer Bilinearform vorgestellt und entsprechend Kongruenz zweier Matrizen und die Diskriminante einer Form definiert. Die Grundlagen werden mit den Isometrien und quadratischen Formen, wobei hier gezeigt wird, dass jede quadratische Form eine symmetrische Bilinearform definiert, abgeschlossen.

Im Teil der Orthogonalität sind wichtige Grundbausteine enthalten, die für die Klassifikation unabdingbar sind. Darunter die Definition der orthogonal direkten Summe, sowie der orthogonalen Zerlegung einer Bilinearform. Auch der Satz von Riesz wird hier hergeleitet. Zudem werden auch die hyperbolischen Räume behandelt, welche eine zentrale Rolle für die Klassifikation symplektischer Bilinearformen spielen.

Die Klassifikation metrischer Vektorräume (sowohl orthogonal, als auch symplektisch) bedeutet, herauszufinden, ob zwei metrische Vektorräume isometrisch sind.

Die Arbeit beschäftigt sich zunächst mit symmetrischen Bilinearformen. Wichtig ist hier natürlich die Erwähnung, dass eine orthogonale Geometrie ebenso auch symplektisch sein kann. Die Frage der Klassifikation wird zunächst in Richtung der Orthonormalbasis gestellt. Für orthogonale Geometrien existiert jedenfalls eine Orthogonalbasis, sofern sie nicht die beiden Eigenschaften: symplektisch und die Charakteristik gleich 2 besitzen, das heißt orthogonale Geometrien sind diagonalisierbar. Die finale Klassifikation hängt vom Grundkörper ab. Verschiedene Grundkörper führen zu verschiedenen Diagonalformen, wodurch die Klassifikation von Grundkörper zu Grundkörper verschieden ist. In der Arbeit wurden algebraisch abgeschlossene Körper, Körper der reellen Zahlen und endliche Körper genauer ausgeführt. Quadratische Formen über $\mathbb{Q}$ und $\mathbb{Q}_p$ können in J.-P. Serre „A course in Arithmetic“ Chapter IV, Paragraph 2 und 3 nachgelesen werden[5][6], da diese Körper den Rahmen dieser Arbeit überschreiten würden.

Zur Klassifikation gehörend ist der Witt'sche Kürzungssatz, welcher auch im Titel der Arbeit widergespiegelt wird. Mit ihm wird die Eindeutigkeit der anisotropen Zerlegung bewiesen, aus welcher zum Beispiel die Wohldefiniert der Summenoperation der Witt Gruppe folgt [8].

Das Klassifikationsproblem für symplektische Bilinearformen gestaltet sich einfacher, da hier nicht zwischen verschiedenen Grundkörpern unterschieden werden muss. Ausschlaggebend ist jedoch die nicht-singuläre Erweiterung, um statt einer Orthogonalbasis eine hyperbolische Basis finden zu können, da symplektische Geometrien nur dann, wenn sie total degeneriert sind eine Orthogonalbasis besitzen. Schlussendlich wird gezeigt, dass die Klassifikation lediglich vom Rang der Form abhängig ist.

Das Kapitel wird mit der Pfaff'schen Determinante einer alternierenden Matrix, welche zusammengefasst aussagt, dass die Determinante einer alternierenden Matrix A mit dem Quadrat der Spezialisierung der Pfaff'schen bei A bestimmt werden kann, und dem

Witt'schen Erweiterungs- und Kürzungssatz für symplektische Formen abgeschlossen.

Ein Ausblick soll an dieser Stelle auch auf die Witt Ringe gegeben werden, welche in M. Kneser „Quadratische Formen“ [8] behandelt werden.

Der Witt'sche Erweiterungs und Kürzungssatz stammt von Ernst Witt und findet sich in E. Witt „Theorie der Quadratischen Formen in beliebigen Körpern“ [9]. In diesem Aufsatz betrachtet Witt Quadratische Formen auch über speziellen Körpern wie Zahlkörper und reelle Funktionenkörper.

Danksagung

Der erste Name, der mir für die Danksagung einfällt, ist mein Betreuer für diese Arbeit, Joachim Mahnkopf, welcher mir nicht nur die Chance gab, in dieses spannende Thema, das mir sonst als Gymnasiumlehrperson verwehrt geblieben wäre, einzusteigen, sondern mich auch kräftig unterstützte, die Arbeit von Beginn bis Ende durchzuziehen, obwohl das für ihn ebenso viel Arbeit bedeutete.

Neben meinem Vollzeitberuf als Lehrperson und allerhand seelischen Kummer und Sorgen hätte ich jedoch bereits aufgegeben, hätte ich nicht meine Familie, speziell meinen älteren Bruder, Thomas Mitterecker und meine Mutter, Ramona Mitterecker, welche immer an mich und meinen Erfolg glaubten und auch weiterhin glauben, gehabt.

Nicht weniger erwähnenswert sind mein Kollegium, welches mich auch in dieser Zeit in meinem Beruf unterstützte und nebenbei auch meine Fortschritte an der Masterarbeit schätzte, sowie meine Therapeutin – welche ich aus Schutzgründen hier namentlich nicht erwähne –, die mir half, den Abschluss meines Studiums und die Fertigstellung dieser Arbeit in greifbare Nähe rücken zu sehen.

Zu guter Letzt möchte ich noch meiner Düse, Nicole Varionova und meiner guten Freundin Cornelia Götz, danken, da sie für mich eine seelische Stütze für jegliche Probleme waren und auch weiterhin in Zukunft sein werden. Danke, dass auch ihr an mich glaubt.

Teil I

Grundlagen

In diesem Abschnitt sollen verschiedene grundlegende Definitionen, wie Bilinearform oder Quadratische Form, sowie deren Eigenschaften behandelt werden.

1 Bilinearform

Alle Definitionen Sätze und Anmerkungen dieses Kapitels, wenn nicht anders genannt, basieren auf Steven Romans Graduate Texts in Mathematics[1].

Eine Bilinearform kann die Eigenschaften symmetrisch, schiefsymmetrisch und alternierend aufweisen.

Definition 1.1 (Bilinearform). Sei V ein Vektorraum über einen beliebigen Körper F . Die Abbildung $\langle \cdot, \cdot \rangle : V \times V \rightarrow F$ heißt Bilinearform, falls sie in jeder Komponente linear ist. Diese Eigenschaft ist genau dann erfüllt, wenn für alle $\alpha, \beta \in F$ und $x, y, z \in V$ gilt:

$$\langle \alpha x + \beta y, z \rangle = \alpha \langle x, z \rangle + \beta \langle y, z \rangle$$

und

$$\langle z, \alpha x + \beta y \rangle = \alpha \langle z, x \rangle + \beta \langle z, y \rangle.$$

Eine Bilinear Form heißt

1) **symmetrisch**, wenn

$$\langle x, y \rangle = \langle y, x \rangle.$$

2) **schiefsymmetrisch**, wenn

$$\langle x, y \rangle = -\langle y, x \rangle.$$

3) **alternierend**, wenn

$$\langle x, x \rangle = 0.$$

Definition 1.2 (Metrischer Vektorraum). Eine Bilinearform $\langle \cdot, \cdot \rangle : V \times V \rightarrow F$, welche symmetrisch, schiefsymmetrisch oder alternierend ist, heißt **Inneres Produkt** auf V und ein Paar $(V, \langle \cdot, \cdot \rangle)$, bei dem V ein Vektorraum und $\langle \cdot, \cdot \rangle$ das innere Produkt von V ist, heißt **metrischer Vektorraum**. Zudem kann unterschieden werden:

1) Ist die Bilinearform $\langle \cdot, \cdot \rangle$ symmetrisch, so spricht man von einer **orthogonalen Geometrie** über F .

2) Ist die Bilinearform $\langle \cdot, \cdot \rangle$ alternierend, so spricht man von einer **symplektische Geometrie** über F .

Anmerkung 1.3. Das Wort symplektisch wurde von Hermann Weyl in seinem Buch *The Classical Groups* eingeführt. Es handelt sich dabei um eine Substitution für das Wort »komplex«.[2]

Beispiel 1.4 (Der Minkowski Raum). M_4 ist eine 4-dimensionale reelle orthogonale Geometrie mit einem inneren Produkt, welches folgendermaßen definiert ist:

$$\langle e_1, e_1 \rangle = \langle e_2, e_2 \rangle = \langle e_3, e_3 \rangle = 1$$

$$\langle e_4, e_4 \rangle = -1$$

$$\langle e_i, e_j \rangle = 0 \text{ für } i \neq j$$

wobei $e_1, \dots, e_4$ die Standardbasis von $\mathbb{R}^4$ bilden.

Anmerkung 1.5. Wenn S ein Untervektorraum eines metrischen Vektorraums $(V, \langle \cdot, \cdot \rangle)$ ist, so erbt dieser die metrische Struktur von V . Wir nennen $(S, \langle \cdot, \cdot \rangle|_{S \times S})$ einen **Unterraum** von V .

Tatsächlich sind die Eigenschaften symmetrisch, schiefsymmetrisch und alternierend nicht voneinander unabhängig. Ihre Relation hängt, von der Charakteristik des Basiskörpers F ab. Der nächste Satz zeigt, dass wir keine Rücksicht auf schiefsymmetrische Formen nehmen müssen, da diese Eigenschaft äquivalent zu symmetrisch oder alternierend ist.

Satz 1.6. Sei $(V, \langle \cdot, \cdot \rangle)$ ein Vektorraum über einem Körper F . Dann gelten die zwei Aussagen:

1) ist $\text{char}(F) = 2$, dann

$$\text{alternierend} \Rightarrow \text{symmetrisch} \Leftrightarrow \text{schiefsymmetrisch}.$$

2) ist $\text{char}(F) \neq 2$, dann

$$\text{alternierend} \Leftrightarrow \text{schiefsymmetrisch}.$$

3) Außerdem ist, falls $\text{char}(F) \neq 2$, die einzige Bilinearform, welche alternierend und symmetrisch ist, die Nullform $\langle x, y \rangle = 0$ für alle $x, y \in V$.

Das heißt unabhängig von der Charakteristik von F gilt: $\text{alternierend} \Rightarrow \text{schiefsymmetrisch}$.

Beweis. Wir zeigen nach der Reihe, dass die Aussagen stimmen. Für die erste Aussage schauen wir uns zunächst an, was durch die Eigenschaft Alternierend per Umformung folgt und schließen dann mit der zusätzlichen Eigenschaft $\text{char}(F) = 2$ auf die Äquivalenzbeziehung. Aussage 2) kann mit einfachen Äquivalenzumformungen gezeigt werden und die letzte Aussage erhalten wir, wenn wir die Eigenschaften gegenüberstellen.

a) Sei also $\langle \cdot, \cdot \rangle$ alternierend, dann folgt

$$\begin{aligned} 0 &= \langle x + y, x + y \rangle \\ &= 1 \cdot \langle x, x + y \rangle + 1 \cdot \langle y, x + y \rangle \\ &= 1 \cdot \langle x, x \rangle + 1 \cdot \langle x, y \rangle + 1 \cdot \langle y, x \rangle + 1 \cdot \langle y, y \rangle \end{aligned}$$

Aus der Eigenschaft alternierend folgt nun $\langle x, x \rangle = \langle y, y \rangle = 0$ und daher:

$$\begin{aligned} 0 &= \langle x, y \rangle + \langle y, x \rangle \\ \implies \langle x, y \rangle &= -\langle y, x \rangle, \text{ das hei\ss t } \langle \cdot, \cdot \rangle \text{ ist schief-symmetrisch.} \end{aligned}$$

b) Es folgt nun wenn $\text{char}(F)=2$, dann ist $-1 = 1$ und die Definitionen von symmetrisch und schief-symmetrisch sind äquivalent.

Damit ist die Aussage 1) bewiesen.

c) Für einen Körper F mit $\text{char}(F) \neq 2$ und einer Bilinearform $\langle \cdot, \cdot \rangle$, welche schief-symmetrisch ist, gilt für jedes $x \in V$:

$$\begin{aligned} \langle x, x \rangle &= -\langle x, x \rangle \\ \iff 2\langle x, x \rangle &= 0 \\ \iff \langle x, x \rangle &= 0, \end{aligned}$$

Das heißt alternierend $\Leftarrow$ schief-symmetrisch und mit der Eigenschaft, dass, unabhängig von der Charakteristik von F gilt, wegen a): alternierend $\Rightarrow$ schief-symmetrisch, folgt die Äquivalenz.

Also können wir zusammenfassen, dass, falls $\text{char}(F) \neq 2$, dann sind die Eigenschaften Alternierend und Schief-symmetrisch äquivalent und Aussage 2) ist bewiesen.

d) Nach Voraussetzung, ist die Bilinearform sowohl alternierend als auch symmetrisch, somit folgt aus 2), dass sie schief-symmetrisch ist und $\langle u, v \rangle = \langle v, u \rangle = -\langle u, v \rangle$. Somit folgt

$$\begin{aligned} 2 \cdot \langle u, v \rangle &= 0 \\ \iff \langle u, v \rangle &= 0 \end{aligned}$$

Damit ist gezeigt, dass eine Bilinearform mit den Eigenschaften Alternierend und Symmetrie nur die Nullform sein kann.

Damit sind alle drei Aussagen bewiesen. $\square$

Beispiel 1.7. Wir untersuchen nun ein Beispiel für eine alternierende Bilinearform. Es wird folgen, wir bereits in Satz 1.6 dass sie ebenfalls schief-symmetrisch ist. Seien

$\vec{v} = \begin{pmatrix} v_1 \\ v_2 \end{pmatrix}$ und $\vec{u} = \begin{pmatrix} u_1 \\ u_2 \end{pmatrix} \in \mathbb{R}^2$. Dann gilt für die Bilinearform $\langle v, w \rangle := \det(v, w)$:

$$\begin{aligned} \det(\vec{v}, \vec{v}) &= \det \begin{pmatrix} v_1 & v_1 \\ v_2 & v_2 \end{pmatrix} \\ &= v_1 \cdot v_2 - v_2 \cdot v_1 \\ &= 0 \end{aligned}$$

und

$$\begin{aligned}
 \det(\vec{v}, \vec{u}) &= \det \begin{pmatrix} v_1 & u_1 \\ v_2 & u_2 \end{pmatrix} \\
 &= v_1 \cdot u_2 - v_2 \cdot u_1 \\
 &= -(u_1 \cdot v_2 - u_2 \cdot v_1) \\
 &= -\det \begin{pmatrix} u_1 & v_1 \\ u_2 & v_2 \end{pmatrix} \\
 &= -\det(\vec{u}, \vec{v})
 \end{aligned}$$

Das heißt $\langle, \rangle$ ist eine alternierende und schiefsymmetrische Bilinearform.

Beispiel 1.8. Im Gegensatz zum vorangegangenen Beispiel schauen wir uns nun eine Bilinearform an, welche zwar schiefsymmetrisch, aber nicht alternierend ist. Es sei $V = F^n$ ein Vektorraum mit dem standardmäßigen inneren Produkt $\langle, \rangle$ welches durch

$$(x_1, \dots, x_n) \cdot (y_1, \dots, y_n) = x_1 y_1 + \dots + x_n y_n$$

definiert ist. Dieses ist zwar eine symmetrische Bilinearform, aber nicht alternierend, wie bereits das einfache Beispiel:

$$(1, 0, \dots, 0) \cdot (1, 0, \dots, 0) = 1 \neq 0$$

zeigt.

Falls $\text{char}(F) = 2$ ist die Bilinearform $\langle, \rangle$ auch schiefsymmetrisch, aber nicht alternierend.

2 Matrix einer Bilinearform und Isometrien

Alle Definitionen Sätze und Anmerkungen dieses Kapitels, wenn nicht anders genannt, basieren auf Steven Romans Graduate Texts in Mathematics[1].

2.1 Darstellungsmatrix einer Bilinearform

Wir können eine Bilinearform mit einer Matrix beschreiben. Dafür benötigen wir eine geordnete Basis $\mathcal{B} = (b_1, \dots, b_n)$ des Körpers V , welcher eine Bilinearform $\langle, \rangle$ besitzt. Die Bilinearform $\langle, \rangle$ ist festgelegt durch die $n \times n$ Matrix mit den Werten

$$M_{\mathcal{B}} = (a_{i,j}) = (\langle b_i, b_j \rangle)_{i,j=1,\dots,n}$$

Wir nennen die Matrix $M_{\mathcal{B}}$ die Matrix der Bilinearform $\langle, \rangle$ oder kurz die Matrix von $\langle, \rangle$ bezüglich der geordneten Basis $\mathcal{B}$. Weiters gilt, jede $n \times n$ Matrix über F ist die Matrix einer Bilinearform über V .

Satz 2.1. Sei $\mathcal{B} = (b_1, \dots, b_n)$ eine geordnete Basis für den metrischen Vektorraum $(V, \langle \cdot, \cdot \rangle)$; dann gilt für die Matrix

$$M_{\mathcal{B}} = (\langle b_i, b_j \rangle)_{i,j=1,\dots,n},$$

dass

$$\langle x, y \rangle = [x]_{\mathcal{B}}^t M_{\mathcal{B}} [y]_{\mathcal{B}}.$$

Beweis. Wir schauen uns zunächst das Produkt $M_{\mathcal{B}} [x]_{\mathcal{B}}$ an, indem wir für die leichtere Ansicht $M_{\mathcal{B}} [x]_{\mathcal{B}}$ genauer unter die Lupe nehmen und dann im Folgeschritt das ganze Produkt $[x]_{\mathcal{B}}^t M_{\mathcal{B}} [y]_{\mathcal{B}}$ umformen. Für einen Vektor $x = \sum r_i b_i$ gilt

$$\begin{aligned} M_{\mathcal{B}} [x]_{\mathcal{B}} &= (\langle b_i, b_j \rangle)_{i,j} (r_j)_{j=1,\dots,n} \\ &= \begin{pmatrix} \sum_{j=1}^n \langle b_1, b_j \rangle \cdot r_j \\ \vdots \\ \sum_{j=1}^n \langle b_n, b_j \rangle \cdot r_j \end{pmatrix} \\ &= \begin{pmatrix} \langle b_1, \sum_j r_j b_j \rangle \\ \vdots \\ \langle b_n, \sum_j r_j b_j \rangle \end{pmatrix} \\ &= \begin{pmatrix} \langle b_1, x \rangle \\ \vdots \\ \langle b_n, x \rangle \end{pmatrix} \end{aligned}$$

Also gilt:

$$M_{\mathcal{B}} [x]_{\mathcal{B}} = \begin{bmatrix} \langle b_1, x \rangle \\ \vdots \\ \langle b_n, x \rangle \end{bmatrix}$$

Analog gilt

$$[x]_{\mathcal{B}}^t M_{\mathcal{B}} = (\langle x, b_1 \rangle \dots \langle x, b_n \rangle).$$

Mit einem weiteren Vektor $y = \sum_i s_i b_i$, folgt

$$\begin{aligned} [x]_{\mathcal{B}}^t M_{\mathcal{B}} [y]_{\mathcal{B}} &= (\langle x, b_1 \rangle \dots \langle x, b_n \rangle) \begin{bmatrix} s_1 \\ \vdots \\ s_n \end{bmatrix} \\ &= \sum_{i=1}^n \langle x, b_i \rangle s_i \\ &= \left\langle x, \sum_{i=1}^n b_i s_i \right\rangle \\ &= \langle x, y \rangle \end{aligned}$$

Somit ist gezeigt, dass die Matrix $M_{\mathcal{B}}$ eindeutig definiert durch den Ausdruck $[x]_{\mathcal{B}}^t A [y]_{\mathcal{B}} = \langle x, y \rangle$ für alle $x, y \in V$ und $M_{\mathcal{B}} = A$. $\square$

Satz 2.2 (Basiswechsel). *Seien $\mathcal{B} = (b_1, \dots, b_n)$ und $\mathcal{C} = (c_1, \dots, c_n)$ geordnete Basen für den metrischen Raum V und $\langle \cdot, \cdot \rangle$ eine Bilinearform auf V , mit der Matrix*

$$M_{\mathcal{B}} = (\langle b_i, b_j \rangle)_{i,j}$$

dann gilt

$$M_{\mathcal{C}} = M_{\mathcal{C},\mathcal{B}}^t M_{\mathcal{B}} M_{\mathcal{C},\mathcal{B}},$$

wobei $M_{\mathcal{C},\mathcal{B}}$ die Übergangsmatrix von $\mathcal{C}$ auf $\mathcal{B}$ repräsentiert. (siehe dazu auch Steven Romans Graduate Texts in Mathematics[1].)

Beweis. Wir zeigen den Basiswechsel, indem wir von Satz 2.1 ausgehen und mittels weniger Umformungen auf die gewünschte Form bringen.

$$\begin{aligned} \langle x, y \rangle &= [x]_{\mathcal{B}}^t M_{\mathcal{B}} [y]_{\mathcal{B}} \\ &= \left([x]_{\mathcal{C}}^t M_{\mathcal{C},\mathcal{B}}^t \right) M_{\mathcal{B}} (M_{\mathcal{C},\mathcal{B}} [y]_{\mathcal{C}}) \\ &= [x]_{\mathcal{C}}^t \left(M_{\mathcal{C},\mathcal{B}}^t M_{\mathcal{B}} (M_{\mathcal{C},\mathcal{B}} [y]_{\mathcal{C}}) \right) \end{aligned}$$

Das zweite „ $=$ “ folgt dabei aus der Aussage $[v]_{\mathcal{B}} = M_{\mathcal{C},\mathcal{B}} [v]_{\mathcal{C}}$, beziehungsweise $[v]_{\mathcal{B}}^t = [v]_{\mathcal{C}}^t M_{\mathcal{C},\mathcal{B}}^t$, welche den Basiswechsel des Vektors v von der Basis $\mathcal{C}$ in die Basis $\mathcal{B}$ beschreibt.

Insgesamt folgt, da generell gilt für eine Darstellungsmatrix $M_{\mathcal{C}}$ die Eigenschaft $\langle x, y \rangle = [x]_{\mathcal{C}}^t M_{\mathcal{C}} [y]_{\mathcal{C}}$ gilt, durch Anwenden auf $x = e_i, y = e_j$, dass

$$M_{\mathcal{C}} = M_{\mathcal{C},\mathcal{B}}^t M_{\mathcal{B}} M_{\mathcal{C},\mathcal{B}}.$$

Somit ist gezeigt, dass der Basiswechsel von einer Basis $\mathcal{B}$ auf eine andere $\mathcal{C}$ mit der Übergangsmatrix $M_{\mathcal{C},\mathcal{B}}$ gegeben ist. $\square$

Definition 2.3. Zwei Matrizen $A, B \in \mathcal{M}_n(F)$ heißen **kongruent**, wenn es eine invertierbare Matrix P gibt, sodass

$$A = P^t B P$$

Die Äquivalenzklassen unter der Kongruenz nennen wir **Kongruenzklassen**.

Satz 2.4. *Zwei Matrizen A beziehungsweise B repräsentieren, bezüglich geeigneter Basen $\mathcal{B}$ beziehungsweise $\mathcal{C}$, dieselbe Bilinearform genau dann wenn sie kongruent sind.*

Beweis. Wir schauen uns bei diesem Beweis beide Richtungen getrennt voneinander an und verwenden die Definition 2.3 und den Satz 2.1.

Wenn zwei Matrizen dieselbe Bilinearform auf V repräsentieren, aber bezüglich verschiedener Basen, müssen die Darstellungsmatrizen kongruent sein, das folgt nach Definition

2.2.

Anders herum repräsentiere nun $B = M_{\mathcal{B}}$ eine Bilinearform auf V und es sei

$$A = P^t B P,$$

wobei P invertierbar ist, dann gibt es eine geordnete Basis $\mathcal{C}$ für V , sodass

$$P = M_{\mathcal{C}, \mathcal{B}},$$

und oben eingesetzt erhalten wir

$$A = M_{\mathcal{C}, \mathcal{B}}^t M_{\mathcal{B}} M_{\mathcal{C}, \mathcal{B}}.$$

Es gilt somit nach Definition 2.1 $A = M_{\mathcal{C}}$ und A repräsentiert dieselbe Bilinearform, aber bezüglich $\mathcal{C}$. $\square$

Satz 2.5. Sei $\langle \cdot, \cdot \rangle$ eine Bilinearform auf V und $A = M_B = (a_{i,j})$ die Darstellungsmatrix von $\langle \cdot, \cdot \rangle$ bezüglich der Basis $B = (v_1, \dots, v_n)$. Dann gilt für $\langle (v_1, \dots, v_n) \rangle$:

1) Die Bilinearform $\langle \cdot, \cdot \rangle$ ist orthogonal, genau dann wenn A symmetrisch ist, das heißt $A = A^t$.

2) Die Bilinearform $\langle \cdot, \cdot \rangle$ ist schiefsymmetrisch, genau dann wenn A schiefsymmetrisch ist, das heißt $A = -A^t$.

3) Die Bilinearform $\langle \cdot, \cdot \rangle$ ist symplektisch, genau dann wenn A alternierend ist das heißt A ist schiefsymmetrisch und $a_{i,i} = 0$ für alle $i = 1, \dots, n$. (siehe dazu auch Steven Romans Graduate Texts in Mathematics[1].)

Beweis. Wir beweisen nur 3):

($\Leftarrow$) : Sei $x = \sum_i \alpha_i v_i \in V$ beliebig, dann folgt

$$\begin{aligned} \langle x, x \rangle &= [x]_B^t A [x]_B \\ &= \sum_{i,j=1,\dots,n} \alpha_i a_{i,j} \alpha_j \\ &= \sum_{i>j} \alpha_i a_{i,j} \alpha_j + \sum_{i<j} \alpha_i a_{i,j} \alpha_j + \sum_{i=j} \alpha_i a_{i,j} \alpha_j \\ &= \sum_{i>j} \alpha_i (a_{i,j} + a_{j,i}) \alpha_j + \sum_{i=j} \alpha_i a_{i,j} \alpha_j, \end{aligned}$$

und weil A schiefsymmetrisch ist folgt weiters

$$\begin{aligned} &= \sum_{i>j} \alpha_i \cdot 0 \cdot \alpha_j + \sum_{i=j} \alpha_i \cdot 0 \cdot \alpha_j \\ &= 0. \end{aligned}$$

Also ist die Bilinearform alternierend, das heißt symplektisch.

$(\Rightarrow)$: Aus der Eigenschaft symplektisch beziehungsweise alternierend folgt

$$a_{i,j} = \langle v_i, v_j \rangle = 0, \text{ für } i = j,$$

das heißt, die Einträge der Diagonale sind alle gleich 0.

Für die restlichen Einträge wollen wir zeigen, dass $a_{i,j} + a_{j,i} = 0$ gilt. Dafür formen wir um:

$$\begin{aligned} & \langle v_i + v_j, v_i + v_j \rangle = 0 \\ \Leftrightarrow & \langle v_i, v_i \rangle + \langle v_i, v_j \rangle + \langle v_j, v_i \rangle + \langle v_j, v_j \rangle = 0 \\ \Leftrightarrow & 0 + \langle v_i, v_j \rangle + \langle v_j, v_i \rangle + 0 = 0 \\ \Leftrightarrow & \langle v_i, v_j \rangle + \langle v_j, v_i \rangle = 0 \\ \Leftrightarrow & a_{i,j} + a_{j,i} = 0 \end{aligned}$$

Somit ist gezeigt, dass $A = A^t$, das heißt die Matrix A ist schiefssymmetrisch. $\square$

2.2 Die Diskriminante einer Form

Wenn A und B kongruente Matrizen sind, dann gilt

$$\det(A) = \det(P^t B P) = \det(P)^2 \det(B).$$

Somit unterscheiden sich $\det(A)$ und $\det(B)$ genau um einen quadratischen Faktor, $\det(P)^2 \in (F^\times)^2$.

Definition 2.6. Die **Diskriminante** Δ einer Bilinearform $\langle, \rangle$ ist die Menge von Determinanten von allen Matrizen, welche jene Bilinearform $\langle, \rangle$ darstellen. Daher, wenn $\mathcal{B} = (v_1, \dots, v_n)$ eine geordnete Basis für einen Vektorraum V über einem Grundkörper F ist, gilt

$$\Delta = (F^\times)^2 \det(M_{\mathcal{B}}) = \{r^2 \det(M_{\mathcal{B}}), 0 \neq r \in F\}.$$

2.3 Isometrien

Definition 2.7. Es seien $(V, \langle, \rangle)$ und $(W, \langle, \rangle)$ metrische Vektorräume. Wir nutzen wieder die Notation $\langle, \rangle$ für eine Bilinearform für jeden dieser Räume. Eine bijektive lineare Abbildung $\tau : V \rightarrow W$ wird **Isometrie** genannt, wenn

$$\langle \tau u, \tau v \rangle = \langle u, v \rangle$$

für alle Vektoren $u, v \in V$ gilt. Wenn eine solche Isometrie von V nach W existiert, heißen V und W isometrisch und wir schreiben $V \approx W$. Es ist offenkundig, dass die Menge aller Isometrien von V zu V eine Komposition besitzt.

Anmerkung 2.8. Im Gegensatz zu reellen inneren Produkt-Räumen, müssen wir die Voraussetzung, dass τ bijektiv ist, stellen, weil das nicht automatisch folgt, wie das untere Beispiel zeigt.

Beispiel 2.9. Es sei $\langle (u_1, u_2), (v_1, v_2) \rangle = u_1 v_1$ eine Bilinearform über den Raum $\mathbb{R}^2$ und $\tau : \mathbb{R}^2 \rightarrow \mathbb{R}^2, \tau(u_1, u_2) \mapsto (u_1, 0)$ eine lineare Abbildung.

Die Abbildung τ erfüllt für alle $u, v \in \mathbb{R}^2$ die Eigenschaft $\langle \tau u, \tau v \rangle = \langle u, v \rangle$:

$$\langle \tau(u_1, u_2), \tau(v_1, v_2) \rangle = \langle (u_1, 0), (v_1, 0) \rangle = u_1 v_1 = \langle (u_1, u_2), (v_1, v_2) \rangle,$$

aber ist nicht bijektiv, denn, man sieht sofort, dass das Bild von τ zur Gänze in $\mathbb{R}$ liegt und somit eine echte Untermenge von $\mathbb{R}^2$ ist.

Vormerkung: die folgenden Definitionen verwenden den Begriff der Singularität, beziehungsweise nicht-Singularität, welcher erst im Kapitel der Orthogonalität geklärt wird. Die Definitionen sind trotzdem hier bereits angeführt, um den Definitionsfluss nicht zu stören.

Definition 2.10. Wenn $(V, \langle \cdot, \cdot \rangle)$ eine nicht-singuläre orthogonale Geometrie ist, dann heißt eine Isometrie auf V **orthogonale Transformation**. Die Menge $\mathcal{O}(V)$ aller orthogonalen Transformationen auf V ist eine Gruppe unter einer Komposition, genannt **orthogonale Gruppe** von V .

Definition 2.11. Wenn $(V, \langle \cdot, \cdot \rangle)$ eine nicht-singuläre symplektische Geometrie ist, dann heißt eine Isometrie auf V **symplektische Transformation**. Die Menge $\text{Sp}(V)$ aller symplektischen Transformationen auf V ist eine Gruppe unter einer Komposition, genannt **symplektische Gruppe** von V .

Satz 2.12. Es sei $\tau \in \mathcal{L}(V, W)$ eine Lineare Transformation zwischen den endlichen metrischen Vektorräumen $(V, \langle \cdot, \cdot \rangle)$ und $(W, \langle \cdot, \cdot \rangle)$. Zudem sei $\mathcal{B} = \{v_1, \dots, v_n\}$ eine Basis für V . Dann ist τ genau dann eine Isometrie, wenn τ bijektiv ist und

$$\langle \tau v_i, \tau v_j \rangle = \langle v_i, v_j \rangle$$

für alle i, j .

Beweis. Wir zeigen zunächst die eine Richtung, unter der Annahme, dass τ eine Isometrie zwischen V und W ist und zeigen, dass die Folgerungen gelten. Danach schauen wir uns die Gegenrichtung an und zeigen, dass wenn τ bijektiv ist, so ist τ eine Isometrie.

($\Rightarrow$)

Wenn τ eine Isometrie ist, so ist τ bijektiv und $\langle \tau v_i, \tau v_j \rangle = \langle v_i, v_j \rangle$. Diese Aussagen folgen direkt aus der Definition.

($\Leftarrow$)

Umgekehrt, sei nun τ bijektiv, dann ist zu zeigen

$$\langle v, w \rangle = \langle \tau v, \tau w \rangle \forall v, w \in V$$

Wir können v und w als Linearkombination darstellen:

$$v = \sum_i \alpha v_i \text{ und } w = \sum_j \beta w_j$$

dann

$$\begin{aligned}\langle v, w \rangle &= \left\langle \sum_i \alpha v_i, \sum_j \beta w_j \right\rangle \\ &= \sum_{i,j} \alpha_i \beta_j \langle v_i, w_j \rangle.\end{aligned}$$

Weil nun nach Voraussetzung $\langle v_i, w_j \rangle = \langle \tau v_i, \tau w_j \rangle$ folgt nun weiter

$$\begin{aligned}\langle v, w \rangle &= \sum_{i,j} \alpha_i \beta_j \langle v_i, w_j \rangle \\ &= \sum_{i,j} \alpha_i \beta_j \langle \tau v_i, \tau w_j \rangle \\ &= \left\langle \tau \left(\sum_i \alpha_i v_i \right), \tau \left(\sum_j \beta_j w_j \right) \right\rangle \\ &= \langle \tau v, \tau w \rangle\end{aligned}$$

In Summe ist nun die Äquivalenz gezeigt. □

Satz 2.13. *Es sei wieder $\tau \in \mathcal{L}(V, W)$ eine lineare Transformation zwischen den endlichen metrischen Vektorräumen $(V, \langle \cdot, \cdot \rangle)$ und $(W, \langle \cdot, \cdot \rangle)$. Es gelte nun, dass V orthogonal ist und $\text{char}(F) \neq 2$. Dann ist τ eine Isometrie, genau dann wenn τ bijektiv ist und es gilt*

$$\langle \tau v, \tau v \rangle = \langle v, v \rangle$$

für alle $v \in V$.

Beweis. Die Richtung „ $\Rightarrow$ “ ist offensichtlich.

Für die andere Richtung, „ $\Leftarrow$ “, ist zu zeigen, dass

$$\langle v, w \rangle = \langle \tau v, \tau w \rangle, \forall v, w \in V$$

Es gilt nach Voraussetzung

$$\begin{aligned}\langle v + w, v + w \rangle &= \langle \tau(v + w), \tau(v + w) \rangle \\ &= \langle \tau v, \tau v \rangle + \langle \tau w, \tau w \rangle + 2 \langle \tau v, \tau w \rangle\end{aligned}$$

und außerdem

$$\begin{aligned}\langle v + w, v + w \rangle &= \langle v, v \rangle + \langle w, w \rangle + 2 \langle v, w \rangle \\ &= \langle \tau v, \tau v \rangle + \langle \tau w, \tau w \rangle + 2 \langle v, w \rangle\end{aligned}$$

Daraus folgt

$$\begin{aligned}\langle \tau v, \tau v \rangle + \langle \tau w, \tau w \rangle + 2 \langle \tau v, \tau w \rangle &= \langle \tau v, \tau v \rangle + \langle \tau w, \tau w \rangle + 2 \langle v, w \rangle \\ \Rightarrow 2 \langle \tau v, \tau w \rangle &= 2 \langle v, w \rangle \\ \Rightarrow \langle \tau v, \tau w \rangle &= \langle v, w \rangle\end{aligned}$$

somit ist gezeigt, dass $\langle \tau v, \tau w \rangle = \langle v, w \rangle$ und der Beweis ist erbracht. $\square$

Satz 2.14. *Es sei wieder $\tau \in \mathcal{L}(V, W)$ eine lineare Transformation zwischen den endlichen metrischen Vektorräumen $(V, \langle \cdot, \cdot \rangle)$ und $(W, \langle \cdot, \cdot \rangle)$, Angenommen $\tau : V \approx W$ ist eine Isometrie und*

$$V = S \odot S^\perp \text{ und } W = T \odot T^\perp,$$

dann gilt, wenn $\tau S = T$, so $\tau(S^\perp) = T^\perp$.

Beweis. Es seien $z \in S^\perp$ und $t \in T$, dann können wir, wegen $T = \tau S$, schreiben $t = \tau s$ für ein $s \in S$. Es folgt aus

$$\langle \tau z, t \rangle = \langle \tau z, \tau s \rangle = \langle z, s \rangle = 0,$$

dass $\tau(S^\perp) \subseteq T^\perp$ gilt, aber, da die Dimensionen gleich sind, folgt $\tau(S^\perp) = T^\perp$ $\square$

3 Quadratische Formen

Alle Definitionen Sätze und Anmerkungen dieses Kapitels, wenn nicht anders genannt, basieren auf Steven Romans Graduate Texts in Mathematics[1].

Es gibt eine enge Verbindung zwischen quadratischen Formen auf V und Bilinearformen auf V .

Definition 3.1. Eine **quadratische Form** auf einem Vektorraum V ist eine Abbildung $Q : V \rightarrow F$ mit den Eigenschaften:

1) Für alle $r \in F, v \in V$,

$$Q(rv) = r^2 Q(v)$$

2) Die Abbildung

$$\langle u, v \rangle_Q := Q(u + v) - Q(u) - Q(v)$$

ist eine (symmetrische) Bilinearform

Anmerkung 3.2. Jede quadratische Form Q über V definiert also eine symmetrische Bilinearform $\langle u, v \rangle_Q$ auf V . Wir zeigen nun, umgekehrt, ist $\langle \cdot, \cdot \rangle$ eine symmetrische Bilinearform auf V und gilt $\text{char}(F) \neq 2$, dann ist

$$Q(x) = \frac{1}{2} \langle x, x \rangle$$

eine quadratische Form, denn es gilt für ein $r \in F$:

$$\begin{aligned} Q(rx) &= \frac{1}{2} \langle rx, rx \rangle \\ &= \frac{1}{2} r \langle x, rx \rangle \\ &= \frac{1}{2} rr \langle x, x \rangle \\ &= \frac{1}{2} r^2 \langle x, x \rangle \\ &= r^2 \frac{1}{2} \langle x, x \rangle \\ &= r^2 Q(x) \end{aligned}$$

Somit ist die Form quadratisch. Gleichzeitig erkennt man an

$$\begin{aligned} \langle u, v \rangle_Q &= Q(u+v) - Q(u) - Q(v) \\ &= \frac{1}{2} \langle u+v, u+v \rangle - \frac{1}{2} \langle u, u \rangle - \frac{1}{2} \langle v, v \rangle \\ &= \frac{1}{2} (\langle u, u+v \rangle + \langle v, u+v \rangle) - \frac{1}{2} \langle u, u \rangle - \frac{1}{2} \langle v, v \rangle \\ &= \frac{1}{2} (\langle u, u \rangle + \langle u, v \rangle + \langle v, u \rangle + \langle v, v \rangle) - \frac{1}{2} \langle u, u \rangle - \frac{1}{2} \langle v, v \rangle \\ &= \frac{1}{2} \langle u, v \rangle + \frac{1}{2} \langle v, u \rangle = \langle u, v \rangle, \end{aligned}$$

dass die Abbildung $\langle u, v \rangle_Q$ eine Bilinearform ist.

2) Man erkennt, dass die Abbildungen $\langle, \rangle \rightarrow Q$ und $Q \rightarrow \langle, \rangle_Q$ zueinander invers sind und daher gibt es eine Bijektion zwischen *symmetrischen* Bilinearformen auf V und quadratischen Formen auf V falls $\text{char}(F) \neq 2$. (siehe dazu auch Steven Romans Graduate Texts in Mathematics[1].)

Anmerkung 3.3. Wenn wir außerdem wieder die Annahme treffen, dass $\text{char}(F) \neq 2$ und falls $\mathcal{B} = (v_1, \dots, v_n)$ eine geordnete Basis einer orthogonalen Geometrie $(V, \langle, \rangle)$ ist und die Matrix $M_{\mathcal{B}} = (a_i, a_j)_{i,j=1,\dots,n}$ eine Darstellungsmatrix der symmetrischen Bilinearform $\langle, \rangle$ auf V ist, dann folgt mit Satz 2.1 für $x = \sum x_i v_i$

$$Q(x) = \frac{1}{2} \langle x, x \rangle = \frac{1}{2} [x]_{\mathcal{B}}^t M_{\mathcal{B}} [x]_{\mathcal{B}} = \sum_{i,j} \frac{1}{2} a_{i,j} x_i x_j.$$

Damit ist $Q(x)$ ein homogenes Polynom zweiten Grades in den Koordinaten x_i . (siehe dazu auch Steven Romans Graduate Texts in Mathematics[1].)

Teil II

Orthogonalität

4 Definitionen und grundlegende Eigenschaften

Alle Definitionen Sätze und Anmerkungen dieses Kapitels, wenn nicht anders genannt, basieren auf Steven Romans Graduate Texts in Mathematics[1].

Definition 4.1 (Orthogonalität). Es sei $(V, \langle \cdot, \cdot \rangle)$ ein metrischer Vektorraum. Dann gilt:
1) zwei Vektoren $x, y \in V$ heißen orthogonal, man schreibt $x \perp y$, wenn $\langle x, y \rangle = 0$.

2) ein Vektor $x \in V$ ist orthogonal zur Untermenge S von V , wenn gilt $\langle x, s \rangle = 0$ für alle $s \in S$.

3) Eine Untermenge T von V ist orthogonal zu einer Untermenge S von V , wenn gilt $\langle s, t \rangle = 0$ für alle $t \in T$ und $s \in S$.

Definition 4.2 (Orthogonales Komplement). Das **orthogonale Komplement** $X^\perp$ einer Untermenge X von V ist der Unterraum:

$$X^\perp = \{v \in V \mid v \perp X\}.$$

Anmerkung 4.3. Es kommt nicht darauf an, ob die Form symmetrisch oder alternierend ist, die Orthogonalität ist jedenfalls eine symmetrische Beziehung. Also $x \perp y$ ist äquivalent zu $y \perp x$. Daher werden wir uns genau auf diese beiden Typen von Bilinearformen konzentrieren.

Anmerkung 4.4. Es gibt zwei Arten, wie ein Vektor „degeneriert“ sein kann. Er kann zu sich selbst, oder gar zu jedem anderen Vektor in V orthogonal sein. Dementsprechend definieren wir die nachfolgenden Begriffe.

Definition 4.5. 1) Ein $x \in V$, welches nicht dem Nullvektor entspricht heißt **isotrop** (oder **null**), wenn $\langle x, x \rangle = 0$, ansonsten heißt es **nicht-isotrop**.

2) V heißt **isotrop**, wenn es zumindest einen isotropen Vektor enthält, ansonsten heißt V **nicht-isotrop** (oder **anisotrop**).

3) V heißt **total isotrop**, wenn alle Vektoren in V isotrop sind.

Anmerkung 4.6. Wenn $v \in V$ ein isotroper Vektor ist, dann ist auch av für alle $a \in F$ ein isotroper Vektor.

Beweis. Wir zeigen dafür, dass $\langle v, v \rangle = 0 \Rightarrow \langle av, av \rangle = 0$ für ein beliebiges $a \in F$:

$$\begin{aligned} \langle av, av \rangle &= a \langle v, av \rangle \\ &= a (a \langle v, v \rangle) \\ &= a^2 \langle v, v \rangle \\ &= a^2 \cdot 0 \\ &= 0 \end{aligned}$$

□

Um Vektoren v zu beschreiben, die zu Vektoren $w \neq v$ orthogonal sind, definieren wir die folgenden Begriffe:

Definition 4.7. Es sei $(V, \langle \cdot, \cdot \rangle)$ ein metrischer Vektorraum.

1) Ein Vektor $v \in V$ heißt **degeneriert**, wenn $v \perp V$. Die Menge $V^\perp$ aller degenerierten Vektoren heißt **Radikal** von V und wird notiert mit $\text{rad}(V)$. Also:

$$\text{rad}(V) = V^\perp$$

2) V heißt **nicht-singulär** oder **nicht-degeneriert**, wenn $\text{rad}(V) = \{0\}$.

3) V heißt **singulär** oder **degeneriert**, wenn $\text{rad}(V) \neq \{0\}$.

4) V heißt **total-singulär** oder **total-degeneriert**, wenn $\text{rad}(V) = V$.

Satz 4.8. Ein metrischer Vektorraum $(V, \langle \cdot, \cdot \rangle)$ ist nicht-singulär genau dann wenn alle repräsentierenden Matrizen $M_{\mathcal{B}}$ nicht-singulär sind.

Beweis. Der Beweis erfolgt wieder in zwei Richtungen. Die Implikation beweisen wir indirekt, indem wir die Existenz einer singulären Matrix $M_{\mathcal{B}}$ annehmen und auf einen Widerspruch führen. Die andere Richtung des Beweises ergibt sich durch die logische Folge der Voraussetzung.

($\Rightarrow$) :

Nach Voraussetzung ist $\text{rad}(V) = 0$. Angenommen eine der repräsentierenden Matrizen $M_{\mathcal{B}}$ mit der Basis $\mathcal{B} = \{b_1, b_2, \dots, b_n\}$ sei singulär. Dann gibt es ein $w \neq 0$ mit den beiden Eigenschaften

$$w = \lambda_{w_1} b_1 + \dots \lambda_{w_n} b_n \text{ und } M_{\mathcal{B}}[w]_{\mathcal{B}} = 0$$

woraus folgt

$$\begin{aligned} [v]_{\mathcal{B}} M_{\mathcal{B}} [w]_{\mathcal{B}} &= 0 \Rightarrow \langle v, w \rangle = 0 \quad \forall v \in V \\ &\Rightarrow w \perp v, \forall v \in V. \end{aligned}$$

Daher $w \in \text{rad}(V)$; das heißt $(V, \langle \cdot, \cdot \rangle)$ ist singulär, was einen Widerspruch darstellt.

($\Leftarrow$) :

Wenn alle repräsentierenden Matrizen $M_{\mathcal{B}}$ nicht-singulär sind, dann gilt für jedes $v \in V, v \neq 0$: $[M]_{\mathcal{B}}[v]_{\mathcal{B}} \neq 0$. Folgend existiert ein $w \in V$ mit $[w]^t [M]_{\mathcal{B}}[v]_{\mathcal{B}} \neq 0$, woraus weiter folgt $\langle w, v \rangle \neq 0$ und ist somit klar, dass $v \notin \text{rad}(V) \forall v \in V, v \neq 0$. Also ist V nicht singulär.

Somit ist die Äquivalenz der Eigenschaft der nicht-Singularität von einem metrischen Vektorraum $(V, \langle \cdot, \cdot \rangle)$ und dessen repräsentierenden Matrizen $M_{\mathcal{B}}$ gezeigt. □

Anmerkung 4.9. Es gilt anzumerken, wenn S ein Unterraum eines metrischen Vektorraums $(V, \langle \cdot, \cdot \rangle)$ ist, dann ist $S = (S, \langle \cdot, \cdot \rangle_{|_{S \times S}})$ ebenfalls ein metrischer Vektorraum. Unter $\text{rad}(S)$ verstehen wir dann das Radikal des metrischen Vektorraumes $(S, \langle \cdot, \cdot \rangle_{|_{S \times S}})$. Das

heißt $\text{rad}(S)$ ist die Menge der Vektoren von S welche degeneriert *in* S sind. Wenngleich, $S^\perp$ die Menge aller Vektoren *in* V bezeichnet, welche orthogonal zu S sind; daher gilt

$$\text{rad}(S) = S \cap S^\perp.$$

Alternativ gilt $\text{rad}(S) = S^{\perp S}$, wobei $S^{\perp S} := \{v \in S : \langle v, s \rangle = 0 \ \forall s \in S\}$.

Anmerkung 4.10. Da $S \subseteq S^{\perp\perp}$ gilt

$$\text{rad}(S) = S \cap S^\perp \subseteq S^{\perp\perp} \cap S^\perp = \text{rad}(S^{\perp\perp}),$$

wodurch gezeigt ist, wenn S singular ist, so ist es auch $S^{\perp\perp}$.

Beispiel 4.11. Wir betrachten $V(2, 4) := F_2^4$, ein Vektorraum über F_2 , der Dimension 4, mit dem Standard Skalarprodukt, wie aus Beispiel 1.8. Eine Untermenge davon wäre

$$S = \{0000, 1100, 0011, 1111\}.$$

Es gilt $S = S^{\perp S}$, was bedeutet, dass S total singular ist. Gleichzeitig ist $V(4, 2)$ nicht total singular.

Beweis. Wir zeigen zunächst, dass S total singular ist, indem wir zeigen, dass jedes Element orthogonal zu jedem anderem Element in S ist.

$$\langle 0000, 0000 \rangle = 0 \equiv 0 \pmod{2}$$

$$\langle 0000, 1100 \rangle = 0 \equiv 0 \pmod{2}$$

$$\langle 0000, 0011 \rangle = 0 \equiv 0 \pmod{2}$$

$$\langle 0000, 1111 \rangle = 0 \equiv 0 \pmod{2}$$

$$\langle 1100, 1100 \rangle = 2 \equiv 0 \pmod{2}$$

$$\langle 1100, 0011 \rangle = 0 \equiv 0 \pmod{2}$$

$$\langle 1100, 1111 \rangle = 2 \equiv 0 \pmod{2}$$

$$\langle 0011, 0011 \rangle = 2 \equiv 0 \pmod{2}$$

$$\langle 0011, 1111 \rangle = 2 \equiv 0 \pmod{2}$$

$$\langle 1111, 1111 \rangle = 4 \equiv 0 \pmod{2}$$

Dadurch, dass das Skalarprodukt symmetrisch ist, ist nun gezeigt, dass jedes Element von S orthogonal zu jedem anderem Element von S ist, wodurch gilt $S = S^{\perp S}$ und somit heißt S total singular.

Wir zeigen nun, dass V nicht singular ist. Das bedeutet $\text{rad}(V) = \{0\}$.

Für Elemente mit genau einer vorkommenden 1 gilt, dass sie zu sich selbst nicht orthogonal sind:

$$\langle 1000, 1000 \rangle = 1 \equiv 1 \pmod{2}$$

Für Elemente mit genau zwei vorkommenden 1en gilt, dass sie zu Elementen, bei denen genau ein 1 an einer derselben Stellen ist, nicht orthogonal sind:

$$\langle 1100, 0100 \rangle = 1 \equiv 1 \pmod{2}$$

Für Elemente mit genau drei vorkommenden 1 gilt, dass sie zu sich selbst nicht orthogonal sind:

$$\langle 1110, 1110 \rangle = 3 \equiv 1 \pmod{2}$$

Für das Element 1111 gilt, dass es zum Beispiel zu 1000 nicht orthogonal ist:

$$\langle 1111, 1000 \rangle = 1 \equiv 1 \pmod{2}$$

Somit ist gezeigt, dass $\text{rad}(V) = 0$ und somit ist V nicht singulär. □

Satz 4.12. *Es sei V ein Vektorraum mit einer Bilinearform $\langle, \rangle$ und $x, y \in V$. Die folgenden Aussagen sind äquivalent:*

1) *Orthogonalität ist eine symmetrische Beziehung, also*

$$x \perp y \Rightarrow y \perp x$$

2) *Die Form auf V ist symmetrisch oder alternierend, also V ist ein metrischer Vektorraum.*

Beweis. Die Richtung „(1) $\Leftarrow$ (2)“ ist offensichtlich nach Satz 1.6, welcher besagt, dass aus der Eigenschaft Alternierend Symmetrie oder Schiefsymmetrie folgt.

Wir zeigen nun die andere Richtung (1) $\Rightarrow$ (2). Für das weitere Vorgehen führen wir eine neue Symbolik ein. Aus praktischen Gründen bedeute nun $x \bowtie y$, dass $\langle x, y \rangle = \langle y, x \rangle$ und $x \bowtie V$, dass $\langle x, v \rangle = \langle v, x \rangle$ für alle $v \in V$. Falls $x \bowtie V$ für alle $x \in V$, dann ist V orthogonal, das heißt $\langle, \rangle$ ist symmetrisch, und der Beweis ist erfüllt. Daher werfen wir nun einen Blick auf $x \not\bowtie V$, und versuchen

$$x \not\bowtie V \Rightarrow (x \text{ ist isotrop und } (x \bowtie y \Rightarrow x \perp y)) \tag{*}$$

zu zeigen.

Wenn wir die zweite Folge ($x \bowtie y \Rightarrow x \perp y$) gezeigt haben, folgt aus $x \not\bowtie V$, dass x isotrop ist, das heißt dass auch die erste Folge gilt. Sei nun also $x \not\bowtie y$. Da $x \not\bowtie V$, muss es ein $z \in V$ geben, für das gilt $\langle x, z \rangle \neq \langle z, x \rangle$. Damit gilt $\langle x, z \rangle - \langle z, x \rangle \neq 0$ und daher gilt $x \perp y$ genau dann, wenn

$$\langle x, y \rangle \cdot (\langle x, z \rangle - \langle z, x \rangle) = 0.$$

Durch Umformen erhalten wir:

$$\begin{aligned}
\langle x, y \rangle (\langle x, z \rangle - \langle z, x \rangle) &= \langle x, y \rangle \langle x, z \rangle - \langle x, y \rangle \langle z, x \rangle \\
&= \langle y, x \rangle \langle x, z \rangle - \langle x, y \rangle \langle z, x \rangle \\
&= \langle y \langle x, z \rangle - \langle x, y \rangle z, x \rangle \\
&= \langle x, y \langle x, z \rangle - \langle x, y \rangle z \rangle \\
&= \langle x, y \langle x, z \rangle \rangle - \langle x, \langle x, y \rangle z \rangle \\
&= \langle x, z \rangle \langle x, y \rangle - \langle x, y \rangle \langle x, z \rangle \\
&= 0.
\end{aligned}$$

Somit impliziert die Symmetrie der Orthogonalität, dass $\langle x, y \rangle = 0$ und die Aussage $x \rtimes y \Rightarrow x \perp y$ ist bewiesen. Also gilt (*).

Nehmen wir nun an, dass $\langle, \rangle$ nicht symmetrisch ist. Wir zeigen nun, dass alle Vektoren in V isotrop sind, woraus folgt, dass die Bilinearform $\langle, \rangle$ alternierend ist.

Da V nicht orthogonal ist, existieren $u, v \in V$ für welche $u \not\rtimes v$ und somit $u \not\rtimes V$ und $v \not\rtimes V$. Daher folgt aus (*), dass (*) u und v isotrop und für alle $y \in V$ gilt

$$y \rtimes u \Rightarrow y \perp u; \quad y \rtimes v \Rightarrow y \perp v. \quad (+)$$

Da alle Vektoren $w \in V$ für die gilt $w \not\rtimes V$ schon nach (*) isotrop sind, das heißt die Bedingung für alternierend ist erfüllt, soll nun die Annahme gelten, dass $w \rtimes V$. Dann folgt $w \rtimes u$ und $w \rtimes v$ und somit $w \perp u$ und $w \perp v$, wegen (+). Wir schreiben

$$w = (w - u) + u.$$

Mit $w \perp u$ und der Isotropie von u folgt nun

$$\langle w - u, u \rangle = \langle w, u \rangle - \langle u, u \rangle = 0 - 0,$$

und wir sehen $w - u \perp u$. Da die Summe von zwei orthogonal isotropen Vektoren isotrop ist, folgt, dass w isotrop ist, falls $w - u$ isotrop ist. Da nun $u \not\rtimes v$ folgt aber

$$\langle w - u, v \rangle = -\langle u, v \rangle \neq -\langle v, u \rangle = \langle v, w - u \rangle$$

und somit $(w - u) \not\rtimes V$, was heißt, dass $w - u$ isotrop nach (*) ist. Somit ist gezeigt, dass auch w isotrop ist. Also sind alle Vektoren von V isotrop und es folgt $\langle, \rangle$ ist alternierend. $\square$

Anmerkung 4.13 (Orthogonale und Symplektische Geometrien). Wenn ein metrischer Vektorraum $(V, \langle, \rangle)$ sowohl orthogonal als auch symplektisch ist, dann ist die Form symmetrisch und schief-symmetrisch und daher

$$\langle u, v \rangle = \langle v, u \rangle = -\langle u, v \rangle.$$

Folglich, wenn $\text{char}(F) \neq 2$, dann ist V genau dann orthogonal und symplektisch, wenn V total degeneriert ist.

Wenn $\text{char}(F) = 2$, dann gibt es orthogonale symplektische Geometrien, welche nicht total degeneriert sind. Zum Beispiel, es sei $V = \text{span}(u, v)$ ein zwei-dimensionaler Vektorraum, welcher eine Form auf V definiert, deren Matrix mit

$$M = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}$$

gegeben ist. Da M sowohl symmetrisch als auch alternierend ist, so ist es auch die Form.

5 Orthogonales Komplement und Orthogonale Summe

Alle Definitionen Sätze und Anmerkungen dieses Kapitels, wenn nicht anders genannt, basieren auf Steven Romans Graduate Texts in Mathematics[1].

Definition 5.1. Ein metrischer Vektorraum $(V, \langle \cdot, \cdot \rangle)$ ist die **orthogonale direkte Summe** der Unterräume S und T , geschrieben als

$$V = S \odot T,$$

wenn $V = S \oplus T$ und $S \perp T$.

Anmerkung 5.2. 1) Bei allgemeinen metrischen Vektorräumen, das heißt wenn $\langle \cdot, \cdot \rangle$ singular ist, gilt, dass ein orthogonales Komplement kein Vektorraum Komplement sein muss. Tatsächlich zeigt Beispiel 4.11, dass es Fälle gibt, bei denen $S^{\perp S} = S$ gilt. In anderen Fällen, zum Beispiel, wenn $v \in V$, wobei $(V, \langle \cdot, \cdot \rangle)$ ein metrischer Vektorraum ist, degeneriert ist, dann ist $\langle v \rangle^{\perp} = V$. Trotzdem, ist das orthogonale Komplement von S ein Vektorraum Komplement, genau dann wenn entweder die Summe $V = S + S^{\perp}$ oder die Schnittmenge $S \cap S^{\perp S} = \{0\}$, wobei das letztere bedeute, dass S nicht-singular ist.

2) Einige Eigenschaften der Orthogonalität in reellen inneren Produkträumen können auf nicht-singuläre metrische Vektorräumen übertragen werden. Weiters zeigt das Kapitel über das Radikal, dass die Beschränkung auf nicht-singuläre Räume nicht schwerwiegend ist.

Satz 5.3. Es sei $\tau \in \mathcal{L}(V, W)$ eine lineare Transformation zwischen den endlich-dimensionalen metrischen Vektorräumen $(V, \langle \cdot, \cdot \rangle)$ und $(W, \langle \cdot, \cdot \rangle)$. Angenommen $\tau : V \approx W$ ist eine Isometrie und es gilt

$$V = S \odot S^{\perp} \text{ und } W = T \odot T^{\perp};$$

dann folgt aus $\tau S = T$, dass $\tau(S^{\perp}) = T^{\perp}$.

Beweis. Wir zeigen $\tau(S^{\perp}) = T^{\perp}$ indem wir die Orthogonalität der Elemente mit der Bilinearform prüfen.

Es sei $t \in T$, dann gilt $t = \tau s$ für ein $s \in S$, weil $\tau S = T$. Es folgt mit $z \in S^{\perp}$

$$\langle \tau z, t \rangle = \langle \tau z, \tau s \rangle$$

und wegen 2.12 gilt nun

$$\langle \tau z, \tau s \rangle = \langle z, s \rangle.$$

Da nun $z \in S^\perp$ und somit $z \perp s$ gilt, folgt

$$\langle z, s \rangle = 0$$

Also

$$\langle \tau z, t \rangle = 0 \quad \forall t \in T.$$

Somit ist $\tau(S^\perp) \subseteq T^\perp$, aber da $\dim(\tau(S^\perp)) = \dim(T^\perp)$, folgt die Gleichheit $\tau(S^\perp) = T^\perp$.

Die eben genannte Gleichheit der Dimensionen folgt, weil $\tau S = T \Rightarrow \dim(S) = \dim(T)$ und $\dim(V) = \dim(W)$, was gegeben ist, weil τ bijektiv ist und somit folgt mit der Dimensionsformel, $\dim(S^\perp) = \dim(V) - \dim(S) = \dim(W) - \dim(T) = \dim(T^\perp)$ das heißt wir erhalten $\dim(\tau(S^\perp)) = \dim(T^\perp)$ $\square$

6 Orthogonale Zerlegung einer Bilinearform

Alle Definitionen Sätze und Anmerkungen dieses Kapitels, wenn nicht anders genannt, basieren auf Steven Romans Graduate Texts in Mathematics[1].

Der **Riesz Darsellungssatz** besagt, dass jede lineare Funktion f auf einem endlich-dimensionalen reellen oder auch komplexen Innenproduktraum $(V, \langle \cdot, \cdot \rangle)$ mit einem Rieszvektor $R_f \in V$ dargestellt werden kann:

$$f(v) = \langle v, R_f \rangle, \forall v \in V.$$

Ähnliche Ergebnisse halten nicht-singuläre metrische Vektorräume bereit.

Es sei $(V, \langle \cdot, \cdot \rangle)$ ein metrischer Vektorraum über einem Körper F . Es sei zudem $x \in V$ und die Abbildung $\langle \cdot, x \rangle : V \rightarrow F$ definiert als

$$\langle \cdot, x \rangle v = \langle v, x \rangle.$$

Dabei handelt es sich offenkundig um eine lineare Funktion, daher können wir eine lineare Abbildung $\tau : V \rightarrow V^*$

$$\tau x = \langle \cdot, x \rangle$$

definieren. Die Bilinearität der Form ergibt die Linearität von τ und der Kern von τ ist definiert durch

$$\ker(\tau) = \{x \in V \mid \langle V, x \rangle = \{0\}\} = V^\perp = \text{rad}(V).$$

Injektivität bei linearen Funktionen folgt, wenn ausschließlich der Nullvektor auf den Nullvektor abgebildet wird. Daher ist τ injektiv genau dann, wenn V nicht-singulär ist. (siehe dazu auch Steven Romans Graduate Texts in Mathematics[1].)

Satz 6.1 (Der Riesz Darstellungssatz). *Es sei $(V, \langle \cdot, \cdot \rangle)$ ein endlich-dimensionaler nicht-singulärer metrischer Vektorraum. Die Abbildung $\tau : V \rightarrow V^*$ definiert durch*

$$\tau x = \langle \cdot, x \rangle$$

ist ein Isomorphismus. (siehe dazu auch Steven Romans Graduate Texts in Mathematics[1].)

Beweis. Wir zeigen, dass die Abbildung bijektiv ist, wodurch der Rest der Behauptung folgen wird. Es handelt sich bei τ um eine lineare Abbildung, von V zu V^* mit $\ker(\tau) = \text{rad}(V)$. Die Injektivität folgt daher aus der nicht-Singularität von V . Die Surjektivität folgt aus der Injektivität und

$$\dim(V) = \dim(V^*)$$

Somit folgt, dass es für jedes $f \in V^*$ einen eindeutigen Vektor $x \in V$ gibt, für den gilt:

$$fv = \langle v, x \rangle$$

für alle $v \in V$. □

Satz 6.2. *Das gilt für nicht-singuläre metrische Vektorräume $(V, \langle \cdot, \cdot \rangle)$, er kann aber auf singuläre Unterräume $S \subseteq V$ ausgeweitet werden. Die Idee ist, dass jedes lineare Funktional $f \in S^*$ zu einem linearen Funktional $\bar{f}$ auf V erweitert werden kann, wobei der Riesz Vektor definiert ist als*

$$\bar{f}v = \langle v, R_{\bar{f}} \rangle = \langle \cdot, R_{\bar{f}} \rangle v.$$

Damit hat f jene Form mit einem Riesz Vektor $R_f \in V$, aber es gilt nicht unbedingt $R_f \in S$.

Satz 6.3 (Der Riesz Darstellungssatz für Unterräume). *Es sei $S \subseteq V$ ein Unterraum des metrischen Vektorraums $(V, \langle \cdot, \cdot \rangle)$. Wenn entweder V oder S nicht-singulär ist, dann ist die lineare Abbildung $\tau : V \rightarrow S^*$, definiert mit*

$$\tau x = \langle \cdot, x \rangle|_S,$$

surjektiv und hat den Kern $S^{\perp_V}$. Jedes $f \in S^$ ist also von der Form $f(s) = \langle s, v \rangle$ für ein allgemeines, nicht eindeutig bestimmtes $v \in V$. Ist S nicht singulär, dann kann $v \in S$ gewählt werden und dann ist v eindeutig bestimmt. (siehe dazu auch Steven Romans Graduate Texts in Mathematics[1].)*

Beweis. Die Abbildung τ ist, weil entweder V oder S nicht-singulär ist, surjektiv nach Satz 6.1 und hat den Kern $S^{\perp_V}$. Daher gibt es einen (nicht unbedingt eindeutigen) Vektor $x \in V$, sodass $f(s) = \langle s, x \rangle$ für alle $s \in S$. Weiters gilt, ist S nicht-singulär, dann ist nach Satz 6.1 $\tau : S \rightarrow S^*$ bijektiv und es folgt $\exists v \in S$ mit $f(s) = \langle s, v \rangle$, wobei v, s eindeutig bestimmt sind. □

Satz 6.4. *Es sei $(V, \langle \cdot, \cdot \rangle)$ ein endlich dimensionaler metrischer Vektorraum, dann gilt*

$$V = \text{rad}(V) \odot S,$$

wobei S ein nicht-singulär Unterraum und $\text{rad}(V)$ total singulär ist.

Beweis. Wir zeigen zunächst die nicht-Singularität von S und schauen uns dann an, wieso $\text{rad}(V)$ total-singulär ist.

Es sei $S \subseteq V$ ein Vektorraumkomplement von $\text{rad}(V)$, also

$$V = \text{rad}(V) \oplus S.$$

Somit gilt $\text{rad}(V) \cap S = \{0\}$. Wegen $\text{rad}(V) = V^\perp$ gilt $\text{rad}(V) \perp S$, und es folgt:

$$V = \text{rad}(V) \odot S. \quad (*)$$

Sei $s \in S$ mit $\langle s, s' \rangle = 0$ für alle $s' \in S$, dann folgt $\langle s, v \rangle = 0$ für alle $v \in V$, da $V = \text{rad}(V) \oplus S$. Also ist $s \in \text{Rad}(V)$ und es folgt $s \in \text{rad}(V \cap S) = \{0\}$, womit $s = 0$. Das heißt S ist nicht-singulär.

Nun gilt für alle $x \in \text{rad}(V) : \langle x, v \rangle = 0$, für alle $v \in V$, wodurch folgt $\langle x, v \rangle = 0$, für alle $v \in \text{rad}(V)$, was der totalen Singularität von V entspricht. $\square$

Satz 6.5. *Es sei $S \subseteq V$ ein Unterraum eines endlich-dimensionalen metrischen Vektorraums V . Falls V oder S nicht-singulär sind, dann gilt*

$$\dim(S) + \dim(S^\perp) = \dim(V)$$

weitere gilt: die folgenden Aussagen sind äquivalent:

- a) $V = S + S^\perp$
- b) $S \cap S^\perp = \{0\}$
- c) $V = S \odot S^\perp$.

Beweis. Zunächst erläutern wir, wieso $\dim(S) + \dim(S^\perp) = \dim(V)$ gilt. Dann folgen wir daraus, die Äquivalenz der drei Aussagen a) b) und c).

Zunächst folgt, aufgrund der Eigenschaften der Abbildung $\tau : V \rightarrow S^*$, welche im Satz 6.3, gezeigt sind, – es handelt sich um die Surjektivität und, dass die Abbildung den Kern $S^\perp$ besitzt –, mit der Dimensionsformel

$$\dim(S^*) + \dim(S^\perp) = \dim(V).$$

Wegen $\dim(S^*) = \dim(S)$, gilt nun

$$\dim(S) + \dim(S^\perp) = \dim(V).$$

a $\Rightarrow$ b:

Angenommen unter der Bedingung von a) gelte

$$S \cap S^\perp \neq \{0\},$$

dann gilt

$$\dim(S \cap S^\perp) > 0$$

und $S \cap S^\perp$ enthält nicht nur triviale Elemente, welche sowohl in S , als auch in $S^\perp$ enthalten sind und somit für die Dimension von V wieder abgezogen werden müssten, das heißt, es gilt

$$\dim(V) = \dim(S) + \dim(S^\perp) - \dim(S \cap S^\perp).$$

Das wäre aber ein Widerspruch zu a). **b $\Rightarrow$ a:**

Wegen $S \cap S^\perp = \{0\}$ ist die Dimensionsformel mit

$$\dim(S + S^\perp) = \dim(S) + \dim(S^\perp) = \dim(V)$$

erfüllt. Somit folgt

$$V = S + S^\perp,$$

und es ist gezeigt, dass aus der Aussage b) die Aussage a) folgt.

b $\Rightarrow$ c:

Da $S \cap S^\perp = \{0\}$ und $\dim(S) + \dim(S^\perp) = \dim(V)$ folgt

$$V = S \oplus S^\perp$$

und weil $S \perp S^\perp$ ist

$$V = S \odot S^\perp.$$

Somit folgt aus der Aussage b) die Aussage c).

c $\Rightarrow$ b:

Diese Richtung ist offensichtlich. □

Satz 6.6. *Es sei $S \subseteq V$ ein Unterraum eines endlich-dimensionalen metrischen Vektorraums V . Dann gelten, falls V nicht-singulär ist, die folgenden Aussage:*

a) $S^{\perp\perp} = S$

b) $\text{rad}(S) = \text{rad}(S^\perp)$

c) S ist genau dann nicht-singulär, wenn $S^\perp$ nicht-singulär ist.

Beweis. Um die Aussagen a) zu beweisen, müssen wir Bezug auf den vorangegangenen Satz nehmen, die anderen beiden Aussagen folgen aufgrund von Definitionen.

a):

Da V nicht-singulär ist, gilt für $S \subseteq V$, wegen 6.5 wieder die Dimensionsformel

$$\dim(S) + \dim(S^\perp) = \dim(V)$$

Insbesondere, weil $S^\perp \subseteq V$ gilt auch wieder wegen Satz 6.5 ebenso

$$\dim(S^\perp) + \dim(S^{\perp\perp}) = \dim(V)$$

Durch Gleichsetzen folgt

$$\begin{aligned} \dim(S) + \dim(S^\perp) &= \dim(S^\perp) + \dim(S^{\perp\perp}) \\ \Rightarrow \dim(S) &= \dim(S^{\perp\perp}) \end{aligned}$$

Da nun $S \subseteq S^{\perp\perp}$ muss gelten

$$S = S^{\perp\perp}$$

und die Aussage a) ist gezeigt.

b):

Per Definition gilt

$$\text{rad}(S) = S \cap S^\perp$$

und

$$\text{rad}(S^\perp) = S^\perp \cap S^{\perp\perp}.$$

Nun gilt wegen a), dass $S = S^{\perp\perp}$ und daher

$$\text{rad}(S) = S \cap S^\perp = S^{\perp\perp} \cap S^\perp = S^\perp \cap S^{\perp\perp} = \text{rad}(S^\perp).$$

Daher ist nun die Aussage b) bewiesen.

c):

Wenn S nicht-singulär ist gilt

$$\text{rad}(S) = \{0\}$$

und wegen b) gilt, $\text{rad}(S) = \text{rad}(S^\perp)$ und somit folgt aus der nicht-Singularität von S , dass

$$\text{rad}(S^\perp) = \{0\}$$

und somit ist $S^\perp$ nicht-singulär. Analog folgt aus der nicht-Singularität von $S^\perp$ die nicht-Singularität von S , wodurch nun auch c) bewiesen ist.

□

Der vorangegangene Satz kann im allgemeinen nicht strenger gefasst werden. Um das zu zeigen, diene das nachfolgende Beispiel.

Beispiel 6.7. Ist V singulär, dann gilt Satz 6.6 nicht. Wir geben ein Beispiel. Gegeben sei der zwei-dimensionale metrische Vektorraum $(V = \text{span}(u, v), \langle, \rangle)$, mit

$$\langle u, u \rangle = 1 \quad \langle u, v \rangle = 0, \quad \langle v, v \rangle = 0$$

Wenn wir einen Unterraum $S = \text{span}(u)$ betrachten, dann ist $S^\perp = \text{span}(v)$. Nun ist S nicht-singulär, aber $S^\perp$, wegen der Eigenschaft $\langle v, v \rangle = 0$ sehr wohl. Das heißt Aussage c) aus Satz 6.6 gilt nicht, nämlich S ist genau dann nicht-singulär, wenn $S^\perp$ nicht singulär. Außerdem ist $\text{rad}(S) = \{0\}$ und $\text{rad}(S^\perp) = S^\perp$, somit auch Aussage b), $\text{rad}(S) = \text{rad}(S^\perp)$, von Satz 6.6 nicht. Zum Schluss gilt auch noch $S^{\perp\perp} = V \neq S$, womit auch die Aussage a), $S^{\perp\perp} = S$ von Satz 6.6 nicht gilt.

7 Hyperbolische Räume

Alle Definitionen Sätze und Anmerkungen dieses Kapitels, wenn nicht anders genannt, basieren auf Steven Romans Graduate Texts in Mathematics[1].

Ein spezieller Typ von zwei-dimensionalen metrischen Vektorräumen spielt eine wichtige Rolle in der Struktur von metrischen Vektorräumen.

Definition 7.1 (Hyperbolisches Paar). Es sei $(V, \langle, \rangle)$ ein metrischer Vektorraum. Ein **hyperbolisches Paar** ist ein Paar von Vektoren $u, v \in V$ für welche gilt

$$\langle u, u \rangle = \langle v, v \rangle = 0, \quad \langle u, v \rangle = 1$$

Anmerkung 7.2. Wenn der metrische Vektorraum $(V, \langle, \rangle)$ orthogonal ist, gilt, $\langle v, u \rangle = 1$. Ist jedoch der metrische Vektorraum $(V, \langle, \rangle)$ symplektisch, gilt $\langle v, u \rangle = -1$. Außerdem ist im symplektischen Fall stets $\langle u, u \rangle = \langle v, v \rangle = 0$ per Definition, was im orthogonalen Fall nicht gelten muss und daher Teil der Annahmen ist. In beiden Fällen kann jedoch die folgende, Definition formuliert werden.

Definition 7.3. (Hyperbolische Ebene): Es sei $(V, \langle, \rangle)$ ein metrischer Vektorraum und $u, v \in V$ ein hyperbolisches Paar. Dann heißt der Unterraum $H = \text{span}(u, v)$ **hyperbolische Ebene**.

Definition 7.4. (Hyperbolischer Raum): Es seien nun H_i Hyperbolische Ebenen, dann wird ein jeder Raum in der Form

$$\mathcal{H} = H_1 \odot \cdots \odot H_k$$

hyperbolischer Raum genannt.

Definition 7.5. (Hyperbolische Basis): Es seien $u_i, v_i \in (H_i, \langle \cdot, \cdot \rangle)$ hyperbolische Paare für eine jede Hyperbolische Ebene H_i und $\mathcal{H} = H_1 \odot \cdots \odot H_k$ ein hyperbolischer Raum, dann nennen wir die Basis

$$(u_1, v_1, \dots, u_k, v_k)$$

eine **hyperbolische Basis** von $\mathcal{H}$. Im symplektischen Fall nutzen wir den Begriff **symplektische Basis**.

Satz 7.6. Jeder hyperbolischer Raum $\mathcal{H}$ ist nicht-singulär.

Beweis. Angenommen, $\mathcal{H}$ sei singulär, also es gäbe ein $x \neq 0 \in \mathcal{H}$, derart, dass $\langle x, y \rangle = 0$ für alle $y \in \mathcal{H}$ gilt. Wir können x in der Linearkombination mit den Basisvektoren $(u_1, v_1, \dots, u_k, v_k)$ anschreiben.

$$x = \sum_i \alpha_i u_i + \beta_i v_i$$

Also müsste für alle Basisvektoren v_j gelten:

$$\langle x, v_j \rangle = \left\langle \sum_i \alpha_i u_i + \beta_i v_i, v_j \right\rangle = \sum_i (\alpha_i \langle u_i, v_j \rangle + \beta_i \langle v_i, v_j \rangle) = 0$$

Da, wegen $H_1 \perp \dots \perp H_k$, nun für alle i und j $\langle v_i, v_j \rangle = 0$ gilt, sowie für alle $i \neq j$, $\langle u_i, v_j \rangle = 0$ gilt, (für $i = j$ gilt $\langle u_i, v_j \rangle = 1$), folgt nun

$$\langle x, v_j \rangle = \alpha_j = 0,$$

Analog folgt für alle Basisvektoren u_j

$$\langle u_j, x \rangle = \beta_j = 0$$

und es folgt

$$x = \sum_i 0 \cdot u_i + 0 \cdot v_i = 0.$$

Also ist x der Nullvektor, was einen Widerspruch zur Grundannahme darstellt und somit folgt $\mathcal{H}$ ist nicht-singulär. $\square$

Anmerkung 7.7. Im orthogonalen Fall können Hyperbolische Ebenen durch ihren Grad der Isotropie charakterisiert werden. Im symplektischen Fall sind alle Räume total isotrop per Definition.

Beispiel 7.8. Eine zwei-dimensionale, nicht-singuläre, orthogonale Geometrie V ist eine hyperbolische Ebene, genau dann wenn V genau zwei eindimensionale, isotrope Unterräume besitzt. Anders gesagt, der Kegel der isotropen Vektoren ist die Vereinigung zweier eindimensionaler Unterräume von V .

($\Rightarrow$):

Wenn V eine zwei-dimensionale hyperbolische Ebene ist, dann besitzt sie die Basis (u, v) , derart, dass

$$\langle u, u \rangle = \langle v, v \rangle = 0, \langle u, v \rangle = 1.$$

Also sind die beiden Unterräume $\text{span}(u)$ und $\text{span}(v)$ von V isotrop.

Diese sind die einzigen isotropen Unterräume, denn für jeden beliebigen anderen Unterraum $U \subseteq V$ und $x \in U \subseteq V$ gilt

$$x = \alpha u + \beta v$$

und somit

$$\begin{aligned} 0 = \langle x, x \rangle &= \langle \alpha u + \beta v, \alpha u + \beta v \rangle = 2\alpha\beta \langle u, v \rangle = 2\alpha\beta \\ &\Rightarrow \alpha = 0 \text{ oder } \beta = 0 \end{aligned}$$

Somit ist gezeigt, dass $x \in \text{span}(U)$ oder $x \in \text{span}(V)$.

($\Leftarrow$):

Angenommen eine nicht-singuläre, orthogonale Geometrie V besitzt genau zwei ein-dimensionale, isotrope Unterräume. Dann gibt es zwei Vektoren $u', v' \in V$ für die gilt

$$\langle u', u' \rangle = \langle v', v' \rangle = 0$$

Es bleibt $\langle u', v' \rangle$ zu untersuchen. Angenommen

$$\langle u', v' \rangle = 0.$$

Da u', v' linear unabhängig sind, würde für alle $x, y \in V$ gelten

$$x, y \in \text{span}(u', v')$$

und somit

$$\langle x, y \rangle = 0,$$

was einen Widerspruch zur nicht-Singularität darstellen würde. Es muss also gelten

$$\langle u', v' \rangle \neq 0.$$

Wir setzen nun $u := \frac{u'}{\langle u', v' \rangle}$ und $v = v'$ und es folgt

$$\begin{aligned} \langle u, v \rangle &= \left\langle \frac{u'}{\langle u', v' \rangle}, v' \right\rangle \\ &= \frac{\langle u', v' \rangle}{\langle u', v' \rangle} \\ &= 1. \end{aligned}$$

Außerdem gilt wegen $\text{span}(u) = \text{span}(u')$ und $\text{span}(v) = \text{span}(v')$, dass

$$\langle u, u \rangle = \langle v, v \rangle = 0$$

und somit bilden u und v ein hyperbolisches Paar, welches V aufspannt. Also ist V eine hyperbolische Ebene.

Teil III

Klassifikation

8 Klassifikationsproblem für metrische Vektorräume

Beim Klassifikationsproblem für eine Klasse von metrischen Vektorräumen (sowohl orthogonal, als auch symplektisch), geht es darum herauszufinden, ob zwei metrische Vektorräume aus jener Klasse isometrisch sind. Das Problem gilt als gelöst, zumindest theoretisch, wenn eine Menge kanonischer Formen oder eine komplette Menge von Invarianten für Matrizen unter der Bedingung der Kongruenz gefunden ist. Um zu verstehen warum, nehmen wir an, dass $\tau : V \rightarrow W$ eine Isometrie und $\mathcal{B} = (v_1, \dots, v_n)$ eine geordnete Basis für V ist. Dann ist $\mathcal{C} = (\tau v_1, \dots, \tau v_n)$ eine geordnete Basis für W und es gilt

$$M_{\mathcal{B}}(V) = (\langle v_i, v_j \rangle) = (\langle \tau v_i, \tau v_j \rangle) = M_{\mathcal{C}}(W).$$

Daher ist die Kongruenzklasse der Matrizen, welche V darstellt identisch mit der Kongruenzklasse der Matrizen, welche W darstellt.

Auf der anderen Seite, angenommen V und W sind metrische Vektorräume mit derselben Kongruenzklasse von Darstellungsmatrizen. Dann gilt, wenn $\mathcal{B} = (v_1, \dots, v_n)$ eine geordnete Basis von V ist, dann gibt es eine geordnete Basis $\mathcal{C} = (w_1, \dots, w_n)$ für W , sodass

$$(\langle v_i, v_j \rangle) = M_{\mathcal{B}}(V) = M_{\mathcal{C}}(W) = (\langle w_i, w_j \rangle).$$

Daher ist die Abbildung $\tau : V \rightarrow W$, welche mit $\tau v_i = w_i$ definiert ist, eine Isometrie von V zu W .

Somit ist also gezeigt, dass zwei metrische Vektorräume genau dann isometrisch sind, wenn es Basen $\mathcal{B}$ in V und $\mathcal{C}$ in W gibt, sodass die Darstellungsmatrizen $M_{\mathcal{B}}(V)$ und $M_{\mathcal{C}}(W)$ kongruent sind. Daher können wir nun, unter der Verwendung von kanonischen Formen oder Mengen kompletter Invarianten, herausfinden, ob zwei metrische Vektorräume isometrisch sind, indem wir bestimmen, ob die jeweiligen Darstellungsmatrizen zur selben kanonischen Form kongruent sind. (siehe dazu auch Steven Romans Graduate Texts in Mathematics[1].)

9 Klassifikation symmetrischer Bilinearformen

Alle Definitionen Sätze und Anmerkungen dieses Kapitels, wenn nicht anders genannt, basieren auf Steven Romans Graduate Texts in Mathematics[1].

Als Beiwerke für das Kapitel 9 seien jedoch Serre, J.P. A course in Arithmetic Chapter IV, Paragraph 2 und 3 zu nennen. Sie komplettieren die Klassifikation der quadratischen Formen über $\mathbb{Q}_p$ und $\mathbb{Q}$. Da diese Klassifikationssätze derart komplex sind, dass sie den

Rahmen dieser Arbeit sprengen würden, seien die beiden Paragraphen an dieser Stelle lediglich erwähnt und nicht näher behandelt. (siehe [5] und [6]).

9.1 Die Struktur von orthogonalen Geometrien

Wir unterscheiden für die Klassifikation von orthogonalen Geometrien zwischen jenen, die auch symplektisch sind und jenen, die es nicht sind. Wie wir gleich sehen werden, haben nahezu alle orthogonalen Geometrien V eine Orthogonalbasis. Im Gegenzug haben symplektische Geometrien nur im total degenerierten Fall eine Orthogonalbasis.

Orthogonale nicht-symplektische Matrizen haben immer eine Orthogonalbasis, wie wir später sehen werden (Satz 9.3). Orthogonale symplektische Geometrien haben hingegen genau dann eine Orthogonalbasis, wenn sie total degeneriert sind. Weiters, für $\text{char}(F) \neq 2$ sind alle orthogonalen symplektischen Geometrien total degeneriert und haben somit eine Orthogonalbasis. Nur mit der Bedingung $\text{char}(F) = 2$ gibt es orthogonale symplektische Geometrien, welche nicht total degeneriert sind und die somit keine Orthogonalbasis besitzen (siehe dazu Satz 10.8).

Wenn wir also orthogonale symplektische Geometrien mit $\text{char}(F) = 2$ exkludieren, können wir sagen, dass alle orthogonalen Geometrien eine Orthogonalbasis besitzen.

Wenn ein metrischer Vektorraum V eine Orthogonalbasis hat, ist der nächste Schritt nach einer Orthonormalbasis zu suchen. Wenn jedoch V singulär ist, dann gibt es einen Vektor $v \in V^\perp$ mit $v \neq 0$. Ein solcher Vektor kann niemals durch eine Linearkombination von den Vektoren einer Orthonormalbasis $\{u_1, \dots, u_n\}$ dargestellt werden, da für alle Linearkoeffizienten zu den Basisvektoren $u_i \in \{u_1, \dots, u_n\}$ gilt $\langle v, u_i \rangle = 0$.

Jedoch, selbst wenn V nicht-singulär ist, heißt das nicht, dass stets eine Orthonormalbasis existiert. Die Frage, nach einer **Orthonormalbasis** ist abhängig von den Eigenschaften des Grundkörper. Im späteren Verlauf schauen wir uns daher die drei Fälle an: algebraisch abgeschlossene Körper, die reellen Zahlen und endliche Körper.

Auch zu erwähnen sei, dass auch wenn V eine Orthogonalbasis besitzt, führt das Gram-Schmidt'sche Orthogonalisierungsverfahren nicht notwendigerweise zu einer solchen Basis, weil selbst nicht-singuläre orthogonale Geometrien isotrope Vektoren besitzen, wodurch die Division durch $\langle u, u \rangle = 0$ problematisch ist, wie das untenstehende Beispiel zeigt. (siehe dazu auch Steven Romans Graduate Texts in Mathematics[1].)

Beispiel 9.1. Gegeben sei die hyperbolische Ebene $H = \text{span}(u, v)$ und für den Grundkörper F gilt $\text{char}(F) \neq 2$. Wir zeigen, dass das Gram-Schmidt'sche Orthogonalisierungsverfahren nicht verwendet werden kann, um eine orthogonale Basis zu finden, obwohl es eine Orthogonalbasis $B = \{u + v, u - v\}$ gibt.

Da H eine hyperbolische Ebene ist, müssen u, v ein hyperbolisches Paar sein, für welches gilt $\langle u, u \rangle = \langle v, v \rangle = 0$. Das Gram-Schmidt'sche Orthogonalisierungsverfahren sieht nun aber bei einer gegebenen Basis $\{u, v\}$ von H vor, dass die Vektoren einer orthogonalen

Basis w_1 und w_2 mittels

$$w_1 = u, w_2 = v - \frac{\langle u, v \rangle}{\langle v, v \rangle} u$$

gefunden werden können. Da aber nach Voraussetzung $\langle u, u \rangle = \langle v, v \rangle = 0$, liegt eine Division durch 0 vor und das Verfahren versagt.

Jetzt ist noch zu zeigen, dass $B = \{u + v, u - v\}$ tatsächlich eine orthogonale Basis der hyperbolischen Ebene $H = \text{span}(u, v)$ darstellt. Einerseits müssen wir dazu zeigen, dass die beiden Vektoren $u + v$ und $u - v$ linear unabhängig sind und andererseits, dass die Basis $B = \{u + v, u - v\}$ ein erzeugendes System von H ist.

Die beiden Vektoren $u + v$ und $u - v$ sind linear unabhängig weil die Gleichung

$$\begin{aligned} \lambda(u + v) &= (u - v) \\ \Leftrightarrow \lambda u - u + \lambda v + v &= 0 \\ \Leftrightarrow (\lambda - 1)u + (\lambda + 1)v &= 0 \end{aligned}$$

nur dann erfüllt ist, wenn $\lambda = 1$ und $\lambda = -1$, das heißt, es gibt keine Lösung und somit ist die lineare Unabhängigkeit gezeigt.

Die Basis $B = \{u + v, u - v\}$ erzeugt $H = \text{span}(u, v)$. Da $H = \text{span}(u, v)$, erzeugen die Vektoren u und v die Hyperbolische Ebene H und es muss für alle Vektoren $h \in H$ gelten, dass

$$h = \lambda_1 \cdot u + \lambda_2 \cdot v,$$

wobei $\lambda_1, \lambda_2 \in F$. Unser Ziel ist es

$$h = \lambda_3(u + v) + \lambda_4(u - v)$$

zu erreichen. Also betrachten wir umgekehrt die Gleichung

$$\begin{aligned} \lambda_1 u + \lambda_2 v &= \lambda_3(u + v) + \lambda_4(u - v) \\ \Leftrightarrow \lambda_1 u + \lambda_2 v &= \lambda_3 u + \lambda_3 v + \lambda_4 u - \lambda_4 v \\ \Leftrightarrow \lambda_1 u + \lambda_2 v &= (\lambda_3 + \lambda_4)u + (\lambda_3 - \lambda_4)v \\ \Leftrightarrow \lambda_1 &= \lambda_3 + \lambda_4 \text{ und } \lambda_2 = \lambda_3 - \lambda_4 \\ \Leftrightarrow \lambda_3 &= \lambda_1 - \lambda_4 \text{ und } \lambda_4 = \lambda_3 - \lambda_2 \\ \Leftrightarrow \lambda_3 &= \lambda_1 - (\lambda_3 - \lambda_2) \text{ und } \lambda_4 = \lambda_1 - \lambda_4 - \lambda_2 \\ \Leftrightarrow \lambda_3 &= \frac{\lambda_1 + \lambda_2}{2} \text{ und } \lambda_4 = \frac{\lambda_1 - \lambda_2}{2} \end{aligned}$$

Das heißt das Gleichungssystem ist lösbar und $u + v, u - v$ sind ein Erzeugendensystem für H .

Also ist $B = \{u + v, u - v\}$ eine Orthogonalbasis von $H = \text{span}(u, v)$, denn es gilt $\langle u + v, u - v \rangle = \langle u, u \rangle - \langle v, v \rangle = 0 - 0 = 0$.

Alternativ lässt sich auch wie folgt argumentieren: da $u + v, u - v$ linear unabhängig sind und $\dim(H) = (u, v) = 2$ folgt auch direkt, dass $(u + v, u - v)$ den Vektorraum H erzeugt.

9.2 Orthogonalbasis

In diesem Unterkapitel schauen wir uns an, wie für eine orthogonale Geometrie, sei sie nun auch symplektisch oder nicht, eine orthogonale Basis hergeleitet werden kann.

Es sei V eine orthogonale Geometrie. Wie wir später sehen werden in (10.8), wenn V auch symplektisch ist, dann hat V eine orthogonale Basis, genau dann wenn, V total degeneriert ist. Weiters, wenn $\text{char}(F) \neq 2$ gilt, dann ist V ohnehin, sofern orthogonal und auch symplektisch, total degeneriert. Dementsprechend haben alle orthogonal symplektischen Geometrien eine Orthogonalbasis, falls $\text{char}(F) \neq 2$.

Wenn V nun orthogonal, aber nicht symplektisch ist, dann enthält V einen nicht-isotropen Vektor u_1 . Der Unterraum $\text{span}(u_1)$ ist somit nicht-singulär und

$$V = \text{span}(u_1) \odot V_1,$$

wobei $V_1 = \text{span}(u_1)^\perp$. Wenn V_1 nicht symplektisch ist, dann können wir dieses weiter zerlegen und wir erhalten

$$V = \text{span}(u_1) \odot \text{span}(u_2) \odot V_2,$$

wobei, wie zu erahnen, $V_2 = \text{span}(u_2)^\perp$.

Dieser Prozess kann nun solange wiederholt werden, bis eine Zerlegung

$$V = \text{span}(u_1) \odot \dots \odot \text{span}(u_k) \odot U$$

erreicht ist, bei der U symplektisch und orthogonal ist. Es könnte natürlich sein, dass gilt $U = \{0\}$. Wir setzen nun $\mathcal{B} = \{u_1, \dots, u_k\}$.

Da U symplektisch ist, können wir nun die folgende Gleichung betrachten:

$$0 = \langle x + y, x + y \rangle = \langle x, x \rangle + \langle x, y \rangle + \langle y, x \rangle + \langle y, y \rangle = 2 \cdot \langle x, y \rangle \quad \text{mit } x, y \in U. \quad (*)$$

- ist nun $\text{char}(F) \neq 2$, so folgt $\langle x, y \rangle = 0$ für alle $x, y \in U$, das heißt, U ist total degeneriert. Daher, wenn $\mathcal{C}$ eine Basis für U ist, dann, ist die Vereinigung $\mathcal{B} \cup \mathcal{C}$ eine orthogonale Basis für V . Das heißt im Fall $\text{char}(F) \neq 2$ besitzt jede orthogonale Geometrie eine Orthogonalbasis.

- wenn hingegen $\text{char}(F) = 2$ gilt, fällt wegen $2 = 0$ in $(*)$ alles weg, und so gibt es trotz isotroper Vektoren x, y auch die Möglichkeit des Ergebnisses $\langle x, y \rangle = 1$. Also existieren hyperbolische Ebenen in U und es gilt $U = \mathcal{H} \odot \text{rad}(U)$, wobei $\mathcal{H}$ hyperbolisch von der Dimension $2m$ ist und somit

$$V = \text{span}(u_1) \odot \dots \odot \text{span}(u_k) \odot \mathcal{H} \odot \text{rad}(U),$$

wobei $\text{rad}(U)$ total degeneriert ist und alle u_i nicht-isotrop sind. Wenn $\mathcal{C} = (x_1, y_1, \dots, x_m, y_m)$ eine hyperbolische Basis für $\mathcal{H}$ ist und $\mathcal{D} = (z_1, \dots, z_l)$ eine geordnete Basis für $\text{rad}(U)$, dann ist die Vereinigung

$$\mathcal{E} = \mathcal{B} \cup \mathcal{C} \cup \mathcal{D} = (u_1, \dots, u_k, x_1, y_1, \dots, x_m, y_m, z_1, \dots, z_l) \quad (**)$$

keine Orthogonalbasis von V , denn $\langle x_i, y_i \rangle = 1 \neq 0$, aber „nahe daran“ eine zu sein. In der Tat, im nächsten Kapitel werden wir sehen, dass V auch in diesem Fall (das heißt $\text{char}(F) = 2$) eine Orthogonalbasis besitzt.

9.3 Orthogonalbasen: der Fall $\text{char}(F) = 2$

Lemma 9.2. *Angenommen $\text{char}(F) = 2$. Es sei W eine drei-dimensionale orthogonale Geometrie der Form*

$$W = \text{span}(v) \odot \text{span}(x, y),$$

wobei v nicht-isotrop und $H_1 = \text{span}(x, y)$ eine hyperbolische Ebene ist, das heißt $\langle x, x \rangle = 0$ und $\langle y, y \rangle = 0, \langle x, y \rangle = 1$. Dann gilt

$$W = \text{span}(v_1) \odot \text{span}(v_2) \odot \text{span}(v_3),$$

wobei jedes $v_i \in W$ nicht-isotrop ist.

Beweis. Wir finden zunächst die Existenz linear unabhängiger Vektoren v_1, v_2 und v_3 und zeigen im weiteren Schritt, dass sie nicht-isotrop sind. Folgend gilt für sie auch

$$W = \text{span}(v_1) \odot \text{span}(v_2) \odot \text{span}(v_3).$$

Zunächst, weil v nicht-isotrop ist, gilt $\langle v, v \rangle = a$, wobei $a \neq 0$. Dann definieren wir die Vektoren

$$v_1 = v + x + y$$

$$v_2 = v + ax$$

$$v_3 = v + (1 - a)x + y$$

Diese Vektoren sind linear unabhängig, denn, seien $\lambda_1, \lambda_2, \lambda_3 \in F$ und $v \in W$ derart gewählt, dass $W = \text{span}(v) \odot \text{span}(x, y)$, dann ist

$$\lambda_1 v_1 + \lambda_2 v_2 + \lambda_3 v_3 = 0$$

$$\lambda_1(v + x + y) + \lambda_2(v + ax) + \lambda_3(v + (1 - a)x + y) = 0$$

$$\lambda_1 u + \lambda_1 x + \lambda_1 y + \lambda_2 u + \lambda_2 ax + \lambda_3 u + \lambda_3 x - \lambda_3 ax + \lambda_3 y = 0$$

$$v(\lambda_1 + \lambda_2 + \lambda_3) + x(\lambda_1 + \lambda_2 a + \lambda_3 - \lambda_3 a) + y(\lambda_1 + \lambda_3) = 0.$$

Da u, x und y linear unabhängig sind, bleibt das folgende Gleichungssystem unter die Lupe zu nehmen:

$$\lambda_1 + \lambda_2 + \lambda_3 = 0 \tag{1}$$

$$\lambda_1 + \lambda_2 a + \lambda_3 - \lambda_3 a = 0 \tag{2}$$

$$\lambda_1 + \lambda_3 = 0 \tag{3}$$

Wegen (3) gilt $\lambda_1 = -\lambda_3$, es folgt somit aus (2) $\lambda_2 = \lambda_3$. Und somit folgt für (1):

$$\lambda_1 + \lambda_2 + \lambda_3 = 0$$

$$-\lambda_3 + \lambda_2 + \lambda_3 = 0$$

$$\lambda_2 = 0$$

woraus nun folgt, dass $0 = \lambda_2 = \lambda_3 = -\lambda_1$ und somit gilt auch $\lambda_1 = 0$. Also sind die Vektoren v_1, v_2 und v_3 linear unabhängig.

Die drei Vektoren sind weiter paarweise orthogonal. Wir schauen uns zunächst v_1 und v_2 an:

$$\begin{aligned}
\langle v_1, v_2 \rangle &= \langle v + x + y, v + ax \rangle \\
&= \langle v, v + ax \rangle + \langle x, v + ax \rangle + \langle y, v + ax \rangle \\
&= \langle v, v \rangle + a \langle v, x \rangle + \langle x, v \rangle + a \langle x, x \rangle + \langle y, v \rangle + a \langle y, x \rangle \\
&= \langle v, v \rangle + (1 + a) \langle v, x \rangle + a \langle x, x \rangle + \langle y, v \rangle + a \langle y, x \rangle \\
&= a + (1 + a) \cdot 0 + 0 + 0 + a \\
&= 2a = 0.
\end{aligned}$$

Wir fahren nun fort mit den Vektoren v_1 und v_3 :

$$\begin{aligned}
\langle v_1, v_3 \rangle &= \langle v + x + y, v + (1 - a)x + y \rangle \\
&= \langle v, v + (1 - a)x + y \rangle + \langle x, v + (1 - a)x + y \rangle + \langle y, v + (1 - a)x + y \rangle \\
&= \langle v, v \rangle + \langle v, (1 - a)x \rangle + \langle v, y \rangle + \langle x, v \rangle + \langle x, (1 - a)x \rangle + \langle x, y \rangle + \langle y, v \rangle + \\
&\quad + \langle y, (1 - a)x \rangle + \langle y, y \rangle
\end{aligned}$$

analog, durch Verwenden der Entitäten, erhalten wir

$$\begin{aligned}
&a + \langle v, (1 - a)x \rangle + 0 + 0 + \langle x, (1 - a)x \rangle + 1 + 0 + \langle y, (1 - a)x \rangle + 0 \\
&= a + \langle v, (1 - a)x \rangle \langle x, (1 - a)x \rangle + 1 + \langle y, (1 - a)x \rangle.
\end{aligned}$$

Wenn wir jetzt noch weiter umformen:

$$\begin{aligned}
a + \langle v, x \rangle (1 - a) \langle x, x \rangle (1 - a) + 1 + \langle y, x \rangle (1 - a) &= a + 0 + 0 + 1 + (1 - a) \\
&= a + 1 + 1 - a \\
&= 2 = 0
\end{aligned}$$

sehen wir, dass auch v_1 und v_3 orthogonal zueinander sind.

Es bleiben noch v_2 und v_3 :

$$\begin{aligned}
\langle v_2, v_3 \rangle &= \langle v + ax, v + (1 - a)x + y \rangle \\
&= \langle v, v + (1 - a)x + y \rangle + \langle ax, v + (1 - a)x + y \rangle \\
&= \langle v, v \rangle + \langle v, (1 - a)x \rangle + \langle v, y \rangle + \langle ax, v \rangle + \langle ax, (1 - a)x \rangle + \langle ax, y \rangle \\
&= a + 0 + 0 + 0 + \langle ax, (1 - a)x \rangle + a \\
&= a + a(1 - a) \langle x, x \rangle + a \\
&= a + 0 + a = 2a = 0,
\end{aligned}$$

womit auch für v_2 und v_3 die Orthogonalität gezeigt ist.

Es bleibt noch zu zeigen, dass die Vektoren v_i nicht-isotrop sind. Wir beginnen mit

Vektor v_1 :

$$\begin{aligned}
\langle v_1, v_1 \rangle &= \langle v + x + y, v + x + y \rangle \\
&= \langle v, v \rangle + \langle v, x \rangle + \langle v, y \rangle + \langle x, v \rangle + \langle x, x \rangle + \langle x, y \rangle + \langle y, v \rangle + \langle y, x \rangle + \langle y, y \rangle \\
&= a + 0 + 0 + 0 + 0 + 1 + 0 + 1 + 0 \\
&= a
\end{aligned}$$

Somit ist der Vektor v_1 nicht isotrop. Weiter geht es mit v_2 :

$$\begin{aligned}
\langle v_2, v_2 \rangle &= \langle v + ax, v + ax \rangle \\
&= \langle v, v \rangle + a \langle v, x \rangle + a \langle x, v \rangle + a^2 \langle x, x \rangle \\
&= a + 0 + 0 + 0 = a
\end{aligned}$$

Somit ist auch der Vektor v_2 nicht isotrop. Es bleibt nur noch der Vektor v_3 :

$$\begin{aligned}
\langle v_3, v_3 \rangle &= \langle v + (1-a)x + y, v + (1-a)x + y \rangle \\
&= \langle v, v \rangle + \langle v, (1-a)x \rangle + \langle v, y \rangle \\
&\quad + \langle (1-a)x, v \rangle + \langle (1-a)x, (1-a)x \rangle + \langle (1-a)x, y \rangle \\
&\quad + \langle y, v \rangle + \langle y, (1-a)x \rangle + \langle y, y \rangle \\
&= a + 0 + 0 + 0 + (1-a)^2 \langle x, x \rangle + (1-a) \langle x, y \rangle \\
&\quad + 0 + (1-a) \langle y, x \rangle + 0 \\
&= a + 0 + 0 + 0 + 0 + (1-a) \cdot 1 + (1-a) \cdot 1 + 0 \\
&= a + (1-a) + (1-a) \\
&= a + 2(1-a) \\
&= a + 0 = a
\end{aligned}$$

und somit ist für alle v_i gezeigt, dass sie nicht-isotrop sind.

Mit der zuvor gezeigten linearen Unabhängigkeit ergibt sich, $W = \text{span}(v_1) \odot \text{span}(v_2) \odot \text{span}(v_3)$ und der Beweis ist erbracht. □

Mit dem vorangegangenen Lemma, können wir in $\text{char}(F) = 2$ Vektoren, wie $\{u_k, x_1, y_1\}$ in Gleichung (**) durch nicht-isotrope Vektoren $\{v_k, v_{k+1}, v_{k+2}\}$ ersetzen, die eine Orthogonalbasis bilden. Außerdem kann der Prozess solange wiederholt werden, bis die isotropen Vektoren absorbiert sind und eine orthogonale Basis von nicht-isotropen Vektoren übrig bleibt.

Zusammenfassend erhalten wir:

Satz 9.3. *Es sei V eine orthogonale Geometrie.*

1) *Wenn V auch symplektisch ist, das heißt alle $v \in V$ sind isotrop, dann hat V eine orthogonale Basis, genau dann wenn V total degeneriert ist (Anmerkung 10.7). Wenn*

$\text{char}(F) \neq 2$, so haben also alle symplektischen orthogonalen Geometrien eine orthogonale Basis, weil sie total degeneriert sind, aber das ist nicht der Fall, wenn $\text{char}(F) = 2$.

2) Wenn V nicht symplektisch ist, das heißt es existieren $v \in V$, welche nicht isotrop sind, dann hat V für jede Charakteristik von F eine geordnete orthogonale Basis $\mathcal{B} = (u_1, \dots, v_k, z_1, \dots, z_m)$, wobei $\langle u_i, u_i \rangle = a_i \neq 0$ und $\langle z_i, z_i \rangle = 0$. Somit hat die Darstellungsmatrix $M_{\mathcal{B}}$ die diagonale Form:

$$M_{\mathcal{B}} = \begin{bmatrix} a_1 & & & & & \\ & \ddots & & & & \\ & & a_k & & & \\ & & & 0 & & \\ & & & & \ddots & \\ & & & & & 0 \end{bmatrix}$$

mit $k = \text{rang}(M_{\mathcal{B}})$ Elementen, die nicht 0 sind.

Beweis. Wir beweisen an dieser Stelle den zweiten Teil des Satzes (da der erste im Kapitel 10.2 für symplektische Geometrien geführt wird). Dafür suchen wir uns einen nicht-isotropen Vektor u_1 aus V heraus, welcher nach Voraussetzung (V ist nicht-symplektisch) existiert.

Ist $\text{char}(F) \neq 2$, so haben wir in Kapitel 9.1 gesehen, dass V eine Orthogonalbasis besitzt.

Ist $\text{char}(F) = 2$, dann besitzt V eine Basis $V = (u_1, \dots, u_k, x_1, y_1, \dots, x_m, y_m, z_1, \dots, z_l)$, wobei x_i, y_i eine hyperbolische Basis ist, das heißt $H_i = \langle x_i, y_i \rangle$ ist eine hyperbolische Ebene – siehe dazu Gleichung (**) in Kapitel 9.1 –

Wir können nun für u_1 und H_1 das Lemma 9.2 anwenden und erhalten v_1, v_2, v_3 , welche zum einem orthogonal aufeinander stehen und zum anderen nicht-isotrop sind. Für jede hyperbolische Ebene $H_i \perp v_1$ können wir mit v_1 (oder einem anderen orthogonalen Vektor v_i) das Lemma 9.2 anwenden bis eine Zerlegung

$$V = \text{span}(v_1) \odot \dots \odot \text{span}(v_2) \odot \text{span}(z_1) \odot \dots \odot \text{span}(z_m)$$

erreicht ist, bei der v_i nicht-isotrop und z_i isotrop sind, das heißt es gilt $\langle v_i, v_i \rangle = a_i$ und $\langle z_i, z_i \rangle = 0$. Somit erhält man die Basis $\mathcal{B} = (u_1, \dots, u_k, z_1, \dots, z_m)$ und die Darstellungsmatrix

$$M_{\mathcal{B}} = \begin{bmatrix} a_1 & & & & & \\ & \ddots & & & & \\ & & a_k & & & \\ & & & 0 & & \\ & & & & \ddots & \\ & & & & & 0 \end{bmatrix}$$

mit $k = \text{rang}(M_{\mathcal{B}})$ Elementen, die nicht 0 sind. □

Korollar 9.4. *Es sei M eine symmetrische Matrix, die im Fall $\text{char}(F) = 2$ nicht alternierend sein soll, dann ist M kongruent zu einer Diagonalmatrix.*

Beweis. Das Korollar folgt direkt aus Satz 9.3. Die Ausnahme, dass M nur dann alternierend beziehungsweise symplektisch sein darf, wenn $\text{char}(F) \neq 2$, ist im ersten Teil des Satzes klargestellt. □

Im vorangegangenen Satz 9.3 bilden die Diagonalmatrizen keine Mengen kanonischer Vertreter bezüglich der Kongruenz von Matrizen. Wir erhielten aber auch nicht unbedingt eine *Orthonormalbasis*, sondern nur eine Orthogonalbasis.

Wir können jedoch eine zu M_B kongruente Matrix M_C finden, indem wir die Skalare $r_1, \dots, r_k$ ($r_i \neq 0$) und somit die Basis $\mathcal{C} = (r_1 u_1, \dots, r_k u_n, z_1, \dots, z_m)$ betrachten. Wir erhalten also die zu M_B kongruente Matrix

$$M_C = \begin{bmatrix} r_1^2 a_1 & & & & \\ & \ddots & & & \\ & & r_k^2 a_k & & \\ & & & 0 & \\ & & & & \ddots \\ & & & & & 0 \end{bmatrix}. \quad (9.3)$$

Es ist daher möglich durch die Wahl von $r_i = \frac{1}{\sqrt{a_i}}$ eine **Orthonormalbasis** zu finden, falls $\sqrt{a_i}$ im Körper F existiert.

Tatsächlich hängt die Bestimmung einer Menge an kanonischen Vertreter von symmetrischen Matrizen, bezüglich Kongruenz von den Eigenschaften des Basiskörpers ab. Dazu werden wir im nächsten Kapitel drei verschiedene Typen von Basiskörpern unterscheiden: Algebraisch geschlossene Körper, die reellen Zahlen und endliche Körper.

9.4 Klassifikation symmetrischer Bilinearform

9.4.1 Algebraisch geschlossene Körper

Satz 9.5. *Es sei V eine orthogonale Geometrie über einen algebraisch abgeschlossenen Körper F . Im Fall $\text{char}(F) = 2$ soll V nicht symplektisch sein. Dann hat V eine geordnete orthogonale Basis $\mathcal{B} = (u_1, \dots, u_k, z_1, \dots, z_m)$, sodass $\langle u_i, u_i \rangle = 1$ und $\langle z_i, z_i \rangle = 0$ für alle u_i beziehungsweise z_i gilt. Daher hat M_B die diagonale Form*

$$M_{\mathcal{B}} = Z_{k,m} = \begin{bmatrix} 1 & & & & \\ & \ddots & & & \\ & & 1 & & \\ & & & 0 & \\ & & & & \ddots \\ & & & & & 0 \end{bmatrix},$$

wobei der Eintrag 1 auf der Diagonalen k -mal erscheint und der Eintrag 0 m -mal. Speziell, wenn V nicht-singular ist, so hat V eine Orthonormalebasis.

Beweis. Wir verwenden für den Beweis die Matrix aus (9.3) und setzen $r_i = \frac{1}{\sqrt{a_i}}$ (das dürfen wir, weil für jedes $r \in F$ hat das Polynom $x^2 - r$ eine Wurzel in F). Da somit $r_i^2 = \frac{1}{a_i}$, ist erhalten wir

$$r_i^2 a_i = \frac{a_i}{a_i} = 1$$

und somit

$$M_{\mathcal{B}} = Z_{k,m} = \begin{bmatrix} 1 & & & & \\ & \ddots & & & \\ & & 1 & & \\ & & & 0 & \\ & & & & \ddots \\ & & & & & 0 \end{bmatrix}.$$

□

Zusammenfassend erhalten wir:

Satz 9.6. *Es sei die Menge $\mathcal{S}_n$ die Menge aller symmetrischen $n \times n$ Matrizen über einem algebraisch geschlossenen Körper F . Wenn $\text{char}(F) = 2$, schränken wir $\mathcal{S}_n$ auf die Menge aller symmetrischer Matrizen ein, welche zumindest einen Eintrag in der Diagonale haben, welcher nicht gleich 0 ist. (das heißt wir vermeiden alternierende Matrizen im Fall $\text{char}(F) = 2$.) Dann gilt:*

- 1) Jede Matrix M in $\mathcal{S}_n$ ist kongruent zu einer eindeutigen Matrix der Form $Z_{k,m}$. Tatsächlich gilt $k = \text{rang}(M)$ und $m = n - \text{rang}(M)$.
- 2) Die Menge aller Matrizen $Z_{k,m}$ mit $k+m = n$ ist eine Menge kanonischer Vertreter, für $\mathcal{S}_n$ modulo Kongruenz von Matrizen.
- 3) Der Rang einer Matrix ist eine vollständige Invariante bezüglich Kongruenz von Matrizen in $\mathcal{S}_n$. (siehe dazu auch Steven Romans Graduate Texts in Mathematics[1].)

Beweis. 1) Nach Satz 9.5 existiert zu jeder Matrix M in $\mathcal{S}_n$ eine kongruente Matrix der Form $Z_{k,m}$, wobei der Eintrag 1 auf der Diagonalen k -mal erscheint und der Eintrag 0 m -mal. Es folgt, dass $k = \text{rang}(M)$ und $m = n - \text{rang}(M)$. Sind $Z_{k,m}$ und $Z_{k',m'}$ kongruent, dann folgt $k = \text{rang}(Z_{k,m}) = \text{rang}(Z_{k',m'}) = k'$. Es folgt weiters $m = n - k =$

$n - k' = m' \Rightarrow Z_{k,m} = Z_{k',m'}$. Eine Matrix $M \in \mathcal{S}_n$ ist also kongruent zu genau einer Matrix $Z_{k,m}$.

2) folgt aus 1).

3) folgt aus 1).

□

9.4.2 Körper der reellen Zahlen

Wenn wir den Grundkörper $F = \mathbb{R}$ wählen, so können wir, in Anbetracht der Matrix aus (9.3), $r_i = \frac{1}{\sqrt{|a_i|}}$ wählen. Dadurch werden alle Einträge zu 0, 1 oder -1.

Satz 9.7 (Trägheitsgesetz von Sylvester). *Jede reelle orthogonale Geometrie V hat eine geordnete Basis*

$$\mathcal{B} = (u_1, \dots, u_p, v_1, \dots, v_m, z_1, \dots, z_k)$$

für welche gilt $\langle u_i, u_i \rangle = 1$, $\langle v_i, v_i \rangle = -1$ und $\langle z_i, z_i \rangle = 0$. Daher hat die Matrix $M_{\mathcal{B}}$ die Diagonalform

$$M_{\mathcal{B}} = Z_{p,m,k} = \begin{bmatrix} 1 & & & & & & & \\ & \ddots & & & & & & \\ & & 1 & & & & & \\ & & & -1 & & & & \\ & & & & \ddots & & & \\ & & & & & -1 & & \\ & & & & & & 0 & \\ & & & & & & & \ddots \\ & & & & & & & & 0 \end{bmatrix}$$

wobei der Eintrag 1 auf der Diagonalen p -mal erscheint, der Eintrag -1 m -mal und der Eintrag 0 k -mal.

Beweis. Der Beweis erfolgt durch setzen von $r_i = \frac{1}{\sqrt{|a_i|}}$ in (9.3); wir erhalten

$$r_i^2 a_i = \frac{a_i}{|a_i|} = \begin{pmatrix} 1 & , \text{ wenn } a_i > 0 \\ -1 & , \text{ wenn } a_i < 0 \\ 0 & , \text{ wenn } a_i = 0 \end{pmatrix}.$$

Es folgt, dass die Einträge der Diagonale der Matrix lediglich aus 1er, -1er und 0er besteht. □

Satz 9.8. *Es sei S_n die Menge aller symmetrischen $n \times n$ Matrizen über den Körper der reellen Zahlen $\mathbb{R}$.*

Dann gilt:

1) *Jede Matrix in S_n ist kongruent zu einer eindeutigen Matrix der Gestalt $Z_{p,m,k}$, für*

p, m und $k = n - p - m$.

2) Die Menge aller Matrizen der Form $Z_{p,m,k}$ mit $p + m + k = n$ ist eine Menge kanonischer Vertreter, für $\mathcal{S}_n$ modulo Kongruenz.

3) Sei $M \in \mathcal{S}_n$ und M kongruent zu $Z_{p,m,k}$. Dann ist $p + m$ der Rang von M und $p - m$ heißt die **Signatur** von M und das Triple (p, m, k) heißt die **Trägheit** von M . Das Paar (p, m) oder äquivalent dazu, das Paar $(p + m, p - m)$ ist eine vollständige Invariante bezüglich Kongruenz in $\mathcal{S}_n$. (siehe dazu auch Steven Romans Graduate Texts in Mathematics[1].)

Beweis. Für jede Matrix M existiert eine kongruente Matrix $M_{\mathcal{B}}$ der Gestalt $Z_{p,m,k}$, wegen Satz 9.7. Es bleibt also die Eindeutigkeit zu zeigen. Angenommen es gäbe noch eine zweite zu M kongruente Matrix $M_{\mathcal{C}}$, der Gestalt $Z_{p',m',k'}$. Dann sehen die geordneten Basen für $M_{\mathcal{B}}$ beziehungsweise $M_{\mathcal{C}}$ so aus:

$$\mathcal{B} = (u_1, \dots, u_p, v_1, \dots, v_m, z_1, \dots, z_k)$$

und

$$\mathcal{C} = (u'_1, \dots, u'_{p'}, v'_1, \dots, v'_{m'}, z'_1, \dots, z'_{k'})$$

Der Rang der beiden Matrizen muss äquivalent sein, daher muss gelten $p + m = p' + m'$ und $k = k'$.

Wenn $x \in \text{span}(u_1, \dots, u_p)$ und $x \neq 0$, dann gilt

$$\langle x, x \rangle = \left\langle \sum_i r_i u_i, \sum_j r_j u_j \right\rangle = \sum_{i,j} r_i r_j \langle u_i, u_j \rangle = \sum_{i,j} r_i r_j \delta_{i,j} = \sum_i r_i^2 > 0.$$

Genauso betrachten wir $y \in \text{span}(v'_1, \dots, v'_{m'})$, wobei $y \neq 0$, dann gilt

$$\langle y, y \rangle = \left\langle \sum_i s_i v'_i, \sum_j s_j v'_j \right\rangle = \sum_{i,j} s_i s_j \langle v'_i, v'_j \rangle = \sum_{i,j} s_i s_j \delta_{i,j} = - \sum_i s_i^2 < 0.$$

Daher, sollte $y \in \text{span}(v'_1, \dots, v'_{m'}, z'_1, \dots, z'_{k'})$, dann gilt $\langle y, y \rangle \leq 0$. Es folgt

$$\text{span}(u_1, \dots, u_p) \cap \text{span}(v'_1, \dots, v'_{m'}, z'_1, \dots, z'_{k'}) = \{0\}$$

Also gilt

$$\dim(\text{span}(u_1, \dots, u_p)) + \dim(\text{span}(v'_1, \dots, v'_{m'}, z'_1, \dots, z'_{k'})) \leq n.$$

und somit $p + m' + k' \leq n$, beziehungsweise $p + (n - p') \leq n$.

Daher ist $p \leq p'$. Aufgrund der Symmetrie gilt aber anders herum auch $p' \leq p$, daher $p = p'$. Wegen $k = k'$ folgt $m = m'$.

Somit ist auch die Eindeutigkeit gezeigt und der erste Teil des Satzes ist bewiesen.

Der dritte Teil des Satzes folgt direkt aus den Definitionen.

Somit folgt auch der zweite Teil des Satzes und der Satz ist komplett bewiesen. $\square$

9.4.3 Endliche Körper

Um mit endlichen Körpern klar zu kommen, müssen wir erst die Verteilung von Quadraten in endlichen Körpern, sowie mögliche Werte des Skalars $\langle v, v \rangle$ betrachten.

Satz 9.9. *Es sei F_q ein endlicher Körper mit q Elementen.*

- (1) *Wenn $\text{char}(F_q) = 2$ dann ist jedes Element von F_q ein Quadrat.*
- (2) *Wenn $\text{char}(F_q) \neq 2$, dann sind exakt die Hälfte aller Elemente, ungleich 0, in F_q , Quadrate. Daher sind gesamt $(q-1)/2$ Quadrate, ungleich 0, in F_q vorhanden. Zudem, sei x kein Quadrat in F_q , dann haben alle Elemente, die kein Quadrat sind, die Gestalt r^2x , für ein $r \in F_q$.*

Beweis. Wir schreiben $F = F_q$ und es sei F^* die Untergruppe aller Elemente, welche nicht 0 sind, in F . Es sei außerdem

$$(F^*)^2 = \{a^2 | a \in F^*\}$$

die Untergruppe, aller Quadrate, ungleich 0, in F . Die *Frobenius Abbildung* $\phi : F^* \rightarrow (F^*)^2$ definiert durch $\phi(a) = a^2$ ist ein surjektiver Gruppenhomomorphismus mit dem Kern

$$\ker(\phi) = \{a \in F | a^2 = 1\} = \{-1, 1\}$$

Wenn nun $\text{char}(F) = 2$, dann $\ker(\phi) = 1$ und daher ist ϕ bijektiv und $|F^*| = |(F^*)^2|$, also ist jedes Element von F ein Quadrat, womit (1) bewiesen ist.

Wenn $\text{char}(F) \neq 2$, dann $|\ker(\phi)| = 2$ und es erzeugen immer zwei verschiedene Elemente dasselbe Quadrat (nämlich $-a$ und $a \in F^*$). Somit $|F^*| = 2|(F^*)^2|$, was beweist, dass die Anzahl der Quadrate, ungleich 0, in F gleich $(q-1)/2$ ist.

Nun sei $x \in F^*$ kein Quadrat. Wenn wir x mit einem Quadrat $r^2 \in F^*$ multiplizieren, erhalten wir ein Element $r^2x \in F^*$, welches kein Quadrat ist. Da es $\frac{q-1}{2}$ Quadrate gibt, gibt es auch $\frac{q-1}{2}$ Elemente der Form r^2x . Entscheidend ist nun, dass wir nun alle Elemente von F^* vor uns haben denn

$$|F^*| = q - 1 = \frac{q-1}{2} + \frac{q-1}{2} = |(F^*)^2| + |x(F^*)^2|.$$

Also hat jedes Element in F_q , welches kein Quadrat ist, die Form r^2x , wobei $r \in F_q$ und somit ist auch (2) des Satzes bewiesen. $\square$

Definition 9.10. Eine Bilinearform $\langle , \rangle$ auf V heißt **universell**, wenn für ein beliebiges $c \in F \setminus \{0\}$ ein Vektor $v \in V$ existiert, für welchen $\langle v, v \rangle = c$.

Satz 9.11. *Es sei V eine orthogonale Geometrie über einem endlichen Körper F mit $\text{char}(F) \neq 2$. Wir nehmen an, dass V einen nicht-singulären Unterraum U besitzt für den gilt $\dim(U) \geq 2$. Dann ist die Bilinearform von V universell. (siehe dazu auch Steven Romans Graduate Texts in Mathematics[1].)*

Beweis. Der Satz 9.3 impliziert, dass es eine orthogonale Basis mit $k = \text{rang}(M_{\mathcal{B}})$ vielen Einträgen in der Diagonale gibt, welche ungleich 0 sind. Da $\dim(U) \geq 2$, folgt $k \geq 2$ und es muss also zwei linear unabhängige Vektoren in V geben, für welche

$$\langle u, u \rangle = a \neq 0, \langle v, v \rangle = b \neq 0, \langle u, v \rangle = 0$$

gilt.

Wir wollen nun für jedes $c \in F$ Skalare α und β finden, sodass

$$c = \langle \alpha u + \beta v, \alpha u + \beta v \rangle = a\alpha^2 + b\beta^2$$

oder

$$a\alpha^2 = c - b\beta^2$$

gilt. Wir betrachten nun die beiden Mengen $A = \{a\alpha^2 | \alpha \in F\}$ und $B = \{c - b\beta^2 | \beta \in F\}$. Die Gleichung $a\alpha^2 = c - b\beta^2$ hat eine Lösung, genau dann wenn die Schnittmenge $A \cap B$ nicht leer ist.

Wenn nun $|F| = q$ ist, so ist die Anzahl seiner Quadrate gleich $\frac{q+1}{2}$. Die Menge A entspricht genau diesen Elementen multipliziert mit einer Konstante a , also $|A| = \frac{q+1}{2}$. Ähnlich sieht es mit der Menge B aus, nur dass hier auch noch eine Verschiebung stattfindet, also $|B| = \frac{q+1}{2}$. Wegen

$$|A| + |B| = \frac{q+1}{2} + \frac{q+1}{2} = q+1 > q = |F|$$

muss gelten, dass $A \cap B$ nicht leer ist, also existieren α und β , sodass $a\alpha^2 = c - b\beta^2$ erfüllt ist. Woraus folgt, dass V universal ist. $\square$

Satz 9.12. *Es sei V eine orthogonale Geometrie über einem endlichen Körper F . Im Fall, dass $\text{char}(F) = 2$ nehmen wir an, dass V nicht symplektisch ist. Wenn $\text{char}(F) \neq 2$ dann sei d ein fixiertes Skalar (kein Quadrat) in F^* . Für $a \in F$, wobei $a \neq 0$, setzen wir:*

$$X_k(a) = \begin{bmatrix} 1 & & & & & \\ & \ddots & & & & \\ & & 1 & & & \\ & & & a & & \\ & & & & 0 & \\ & & & & & \ddots \\ & & & & & & 0 \end{bmatrix}$$

wobei $\text{rang}(X_k(a)) = k$. Dann gelten:

- 1) Wenn $\text{char}(F) = 2$, dann gibt es eine geordnete Basis $\mathcal{B}$ von V für welche $M_{\mathcal{B}} = X_k(1)$.
- 2) Wenn $\text{char}(F) \neq 2$, dann gibt es eine geordnete Basis $\mathcal{B}$ von V für welche $M_{\mathcal{B}}$ gleich $X_k(1)$ oder $X_k(d)$ ist. (siehe dazu auch Steven Romans Graduate Texts in Mathematics[1].)

Beweis. 1) Im Fall $\text{char}(F) = 2$ setzen wir für die Matrix (9.3) $r_i = (\sqrt{a_i})^{-1}$. Das ist möglich, weil jedes Element in F eine Quadratwurzel besitzt. Folglich erhalten wir:

$$M_{\mathcal{B}} = \begin{bmatrix} 1 & & & & & \\ & \ddots & & & & \\ & & 1 & & & \\ & & & 1 & & \\ & & & & 0 & \\ & & & & & \ddots \\ & & & & & & 0 \end{bmatrix}$$

und der Beweis ist erbracht.

2) Wenn $\text{char}(F) \neq 2$, dann impliziert der Satz 9.3, dass es eine geordnete orthogonale Basis

$$\mathcal{B} = (u_1, \dots, u_k, z_1, \dots, z_m)$$

gibt, für welche $\langle u_i, u_i \rangle = a_i \neq 0$ und $\langle z_i, z_i \rangle = 0$ ist. Daher hat $M_{\mathcal{B}}$ die diagonale Form

$$M_{\mathcal{B}} = \begin{bmatrix} a_i & & & & & \\ & \ddots & & & & \\ & & a_k & & & \\ & & & 0 & & \\ & & & & \ddots & \\ & & & & & 0 \end{bmatrix}$$

Wir betrachten nun die nicht-singuläre Geometrie $V_1 = \text{span}(u_1, u_2)$, welche einen Unterraum von V bildet. Nach Satz 9.11 ist die Form universell, wenn sie auf V_1 eingeschränkt ist. Daher existiert insbesondere ein $v_1 \in V_1$ für das $\langle v_1, v_1 \rangle = 1$.

Als nächstes setzen wir $v_1 = ru_1 + su_2$ mit $r, s \in F$, wobei entweder r oder s ungleich 0 ist. Daher ist

$$\mathcal{B}_1 = (v_1, u_2, \dots, u_k, z_1, \dots, z_m)$$

eine geordnete Basis für V für welche die Matrix $M_{\mathcal{B}_1}$ eine Diagonalmatrix ist, bei der der erste Eintrag gleich 1 ist. Wir können diesen Vorgang wiederholen, indem wir nun den Unterraum $V_2 = \text{span}(v_2, v_3)$ betrachten. Das weitergedacht führt uns letzten Endes zu einer geordneten Basis

$$\mathcal{C} = (v_1, v_2, \dots, u_{k-1}, u_k, z_1, \dots, z_m)$$

für welche $M_{\mathcal{C}} = X_k(a)$, wobei $a \in F$ und $a \neq 0$.

Wenn a ein Quadrat in F ist, dann können wir v_k durch $\frac{v_k}{\sqrt{a}}$ ersetzen, um eine Basis $\mathcal{D}$ zu bekommen, für welche $M_{\mathcal{D}} = X_k(1)$ gilt.

Wenn a kein Quadrat in F ist, dann gilt nach Satz 9.9: $a = r^2 d$ für ein $r \in F$. Es folgt, dass durch das Ersetzen von v_k durch $\frac{v_k}{r}$ eine Basis $\mathcal{D}$ resultiert, wobei $M_{\mathcal{D}} = X_k(d)$. $\square$

Wir fassen die Ergebnisse im Satz 9.13 zusammen:

Satz 9.13. *Es sei $\mathcal{S}_n$ die Menge aller symmetrischen $n \times n$ Matrizen über einen endlichen Körper F . Im Fall $\text{char}(F) = 2$ beschränken wir $\mathcal{S}_n$ auf die Menge aller symmetrischen Matrizen mit zumindest einem Eintrag in der Hauptdiagonalen ungleich 0.*

1) *Wenn $\text{char}(F) = 2$, dann ist jede Matrix in $\mathcal{S}_n$ kongruent zu einer eindeutigen Matrix der Gestalt $X_k(1)$ und die Menge der Matrizen $\{X_k(1) | k = 0, \dots, n\}$ ist eine Menge kanonischer Vertreter, für $\mathcal{S}_n$ modulo Kongruenz.*

2) *Wenn $\text{char}(F) \neq 2$, so sei $d \in F^*$ fixiert, das kein Quadrat ist. Dann ist jede Matrix in $\mathcal{S}_n$ kongruent zu einer eindeutigen Matrix der Gestalt $X_k(1)$ oder $X_k(d)$. Die Menge $\{X_k(1), X_k(d) | k = 0, \dots, n\}$ ist eine Menge kanonischer Vertreter, für $\mathcal{S}_n$ modulo Kongruenz. (Das bedeutet, dass es genau zwei Kongruenzklassen für jeden Rang gibt.) (siehe dazu auch Steven Romans Graduate Texts in Mathematics[1].)*

Beweis. 1) Nach Satz 9.12(1) existiert für jede Matrix $M \in \mathcal{S}_n$ eine kongruente Matrix $M_{\mathcal{B}}$ der Form $\{X_k(1) | k = 0, \dots, n\}$. Es bleibt wieder die Eindeutigkeit zu zeigen.

Ist $X_k(1)$ kongruent zu $X_{k'}(1)$, so folgt $k = \text{rang}(X_k(1)) = \text{rang}(X_{k'}(1)) = k'$, daher $X_n(1) = X_n'(1)$.

2) Ist $X_k(x)$, mit $x = 1$ oder $x = d$ kongruent zu $X_{k'}(x')$, mit $x' = 1$ oder $x' = d$, dann folgt wieder durch Betrachten des Rangs, dass $k = k'$. Als nächstes werfen wir einen Blick auf die beiden Matrizen selbst:

$$\begin{pmatrix} 1 & & & & \\ & \ddots & & & \\ & & x & & \\ & & & 0 & \\ & & & & \ddots \\ & & & & & 0 \end{pmatrix} \equiv \begin{pmatrix} 1 & & & & \\ & \ddots & & & \\ & & x' & & \\ & & & 0 & \\ & & & & \ddots \\ & & & & & 0 \end{pmatrix}$$

Wir konzentrieren uns lediglich auf die Teile, in denen die Diagonaleinträge ungleich 0 sind, also

$$\begin{pmatrix} 1 & & \\ & \ddots & \\ & & x \end{pmatrix} \equiv \begin{pmatrix} 1 & & \\ & \ddots & \\ & & x' \end{pmatrix}$$

Wir bilden nun die Determinante und, da die beiden Matrizen kongruent sind folgt die Gleichheit

$$\det(x') = \det(x) \cdot (F^*)^2.$$

Es folgt

$$x = x' \cdot (F^*)^2.$$

Da sich die beiden Seiten kongruent um ein Quadrat unterscheiden folgt auch $x = x'$. $\square$

9.5 Rotation, Spiegelung, Symmetrien

Definition 9.14 (Symmetrie). Die Symmetrie H_u $u \in V$ ist definiert als der Operator für den gilt

$$H_u u = -u, H_u w = w \text{ für alle } w \in \langle u \rangle^\perp$$

Anmerkung 9.15. Für ein $x \in V$ gilt

$$H_u x = x - \frac{2 \langle x, u \rangle}{\langle u, u \rangle} u$$

Satz 9.16. Es sei V eine nicht-singuläre orthogonale Geometrie über einem Körper F , für den $\text{char}(F) \neq 2$ gilt. Wenn $u, v \in V$ nicht-isotrope Vektoren derselben Länge (ungleich 0) sind, dann existiert eine Symmetrie σ für welche

$$\sigma u = v \text{ oder } \sigma u = -v$$

gilt.

Beweis. Da u und v nicht-isotrop sind und die gleiche Länge haben, muss mindestens $u - v$ oder $u + v$ nicht-isotrop sein, denn $u - v$ und $u + v$ sind zueinander orthogonal und deren Summe, $2u$, wäre deshalb isotrop, wie folgende Rechnung zeigt:

$$\begin{aligned} \langle 2u, 2u \rangle &= \langle u + v + u - v, u + v + u - v \rangle \\ &\text{durch Substituieren mit } x = u + v \text{ und } y = u - v \text{ erhalten wir} \\ &= \langle x, x \rangle + \langle x, y \rangle + \langle y, x \rangle + \langle y, y \rangle \\ &= 0 + 0 + 0 = 0 \end{aligned}$$

Wenn $u + v$ nicht-isotrop ist, dann gelten

$$\sigma_{u+v}(u + v) = -(u + v)$$

und

$$\sigma_{u+v}(u - v) = u - v,$$

woraus folgt $\sigma_{u+v} u = -v$.

Wenn $u - v$ nicht-isotrop ist, dann gelten

$$\sigma_{u-v}(u - v) = -(u - v)$$

und

$$\sigma_{u-v}(u + v) = u + v$$

woraus folgt $\sigma_{u-v} u = v$. □

9.6 Witt'sche Kürzungs- und Erweiterungssatz

Satz 9.17 (Witt'sche Kürzungssatz). *Es seien V und W isometrische nicht-singuläre orthogonale Geometrien über einem Körper F mit $\text{char}(F) \neq 2$. Weiter sei*

$$V = S \odot S^\perp \text{ und } W = T \odot T^\perp;$$

dann gilt

$$S \approx T \Rightarrow S^\perp \approx T^\perp.$$

(siehe dazu auch Steven Romans *Graduate Texts in Mathematics*[1].)

Beweis. Zunächst zeigen wir, dass es genügt den Satz im Fall $V = W$ zu beweisen. Angenommen der Satz gilt, wenn $V = W$ und $\mu : V \rightarrow W$ eine Isometrie ist, dann folgt

$$\mu(S) \odot \mu(S^\perp) = \mu(S \odot S^\perp) = \mu V = W = T \odot T^\perp.$$

Weiters gilt $\mu S \approx S \approx T$. Wir können daher den Satz anwenden und erhalten

$$S^\perp \approx \mu(S^\perp) \approx T^\perp.$$

Um den Beweis nun für $V = W$ zu erbringen, nehmen wir an

$$V = S \odot S^\perp = T \odot T^\perp,$$

wobei S und T nicht-singulär sind (denn V ist nicht-singulär) und $S \approx T$ gilt. Es sei dann $\tau : S \rightarrow T$ eine Isometrie.

Wir verwenden wir eine Induktion auf $\dim(S)$.

Induktionsanfang: Angenommen $\dim(S) = 1$, wobei $S = \text{span}(s)$. Weil

$$\langle \tau s, \tau s \rangle = \langle s, s \rangle \neq 0$$

gilt, impliziert der Satz 9.16, dass es eine Symmetrie σ gibt, für welche $\sigma s = \epsilon \tau s$, mit $\epsilon = \pm 1$, gilt. Daher ist σ eine Isometrie auf V für welche $T = \sigma S$ gilt und mit Satz 2.14 folgt $T^\perp = \sigma(S^\perp)$. Daher ist $\sigma|_{S^\perp}$ die gesuchte Isometrie.

Induktionsannahme: Wir nehmen nun an, dass der Satz auch stimmt für $\dim(S) < k$.

Induktionsschritt: Es sei $\dim(S) = k$ und $\tau : S \rightarrow T$ eine Isometrie. Weil S nicht-singulär ist, wählen wir einen nicht-isotropen Vektor $s \in S$ und schreiben $S = \text{span}(s) \odot U$, wobei U nicht-singulär ist. Es folgen

$$V = S \odot S^\perp = \text{span}(s) \odot U \odot S^\perp$$

und

$$V = T \odot T^\perp = \tau(\text{span}(s)) \odot \tau U \odot T^\perp$$

Aus dem Induktionsanfang (S ist eindimensional) können wir nun schließen, dass

$$U \odot S^\perp \approx \tau U \odot T^\perp$$

gilt. Da $U \approx \tau U$ und $\dim(U) = \dim(U) < k$ gilt. Folgt aus der Induktionsannahme $S^\perp \approx \sigma(S^\perp) \approx T^\perp$ womit der Beweis erbracht ist. $\square$

Satz 9.18 (Witt'sche Erweiterungssatz). *Es seien V und V' isometrische nicht-singuläre orthogonale Geometrien auf einem Körper F mit $\text{char}(F) \neq 2$. Angenommen U ist ein Unterraum von V und*

$$\tau : U \rightarrow \tau U \subseteq V'$$

ist eine Isometrie, dann kann τ zu einer Isometrie von V nach V' erweitert werden. (siehe dazu auch Steven Romans Graduate Texts in Mathematics[1].)

Beweis. Laut Satz 10.13 können wir τ auf eine nicht-singuläre Ergänzung von U erweitern. Daraus folgt, dass wir ohne Probleme annehmen können, dass sowohl U , als auch τU nicht-singulär sind. Daher gilt

$$V = U \odot U^\perp$$

und

$$V' = \tau U \odot (\tau U)^\perp.$$

Um die Erweiterung von τ auf V abzuschließen, müssen wir lediglich eine hyperbolische Basis $B_{U^\perp}$ für $U^\perp$,

$$B_{U^\perp} = (e_1, f_1, \dots, e_p, f_p),$$

und eine hyperbolische Basis $B_{(\tau U)^\perp}$ für $(\tau U)^\perp$,

$$B_{(\tau U)^\perp} = (e'_1, f'_1, \dots, e'_p, f'_p)$$

wählen.

Wir definieren nun die Erweiterung durch Setzen von $\tau e_i = e'_i$ und $\tau f_i = f'_i$. □

9.7 Anisotrope Zerlegung einer orthogonalen Geometrie

Wir haben gesehen, dass eine jede orthogonale Geometrie V als

$$V = U \odot \text{rad}(V),$$

wobei U nicht-singulär ist, geschrieben werden kann. Nicht-singuläre Räume mögen jedoch weiterhin isotrope Vektoren beinhalten.

Für ein $u \in U$ können wir, wie wir später in Satz 10.3 sehen werden, $\text{span}(u)$ zu einer hyperbolischen Ebene $H = \text{span}(u, x)$ zu erweitern. Wir können nun

$$V = H \odot H^{\perp_U} \odot \text{rad}(V)$$

schreiben, wobei $H^{\perp_U}$ das orthogonale Komplement von H und U ist und einen isotropen Vektor weniger beinhaltet. Um diesen Prozess zu verallgemeinern schauen wir uns zunächst maximalen total degenerierte Unterräume an und definieren den **Witt Index**. (siehe dazu auch Steven Romans Graduate Texts in Mathematics[1].)

Satz 9.19. *Es sei V eine nicht-singuläre orthogonale Geometrie über einem Körper F mit $\text{char} \neq 2$. Dann haben alle maximal total degenerierten Unterräume von V dieselbe Dimension.*

Beweis. Es seien U und U' maximalen total degenerierte Unterräume von V . Wenn $\dim(U) \leq \dim(U')$ gilt, dann gibt es einen Isomorphismus $\tau : U \rightarrow \tau U \subseteq U'$, welcher eine Isometrie ist, weil sowohl U , als auch U' total degeneriert sind. Der Witt'sche Erweiterungssatz 9.18 impliziert, dass τ auf eine Isometrie $\bar{\tau} : V \rightarrow V$ erweitert werden kann. Tatsächlich ist $\bar{\tau}^{-1}(U')$ ein total degenerierter Raum welcher U enthält, woraus folgt, dass $\bar{\tau}^{-1}(U') = U$ und somit $\dim(U) = \dim(U')$ gilt. $\square$

Definition 9.20 (Witt Index). Die Dimension eines maximalen total degenerierten Unterräumen von V wird als **Witt Index** bezeichnet und mit $w(V)$ notiert.

Satz 9.21. *Es sei V eine nicht-singuläre orthogonale Geometrie über einem Körper F mit $\text{char} \neq 2$. Jeder total degenerierte Unterraum von V mit Dimension $w(V)$ ist maximal.*

Beweis. Angenommen U ist nicht maximal, dann gibt es einen total degenerierten Unterraum U' mit $U' \supsetneq U$, dann folgt $\dim(U') > \dim(U)$, was allerdings einen Widerspruch zu Satz 9.19 darstellt. $\square$

Satz 9.22. *Es sei V eine nicht-singuläre orthogonale Geometrie auf einem Körper F mit $\text{char} \neq 2$, dann gelten*

- 1) *Alle maximal hyperbolischen Unterräume von V haben die Dimension $2w(V)$.*
- 2) *Jeder hyperbolische Unterraum der Dimension $2w(V)$ muss maximal sein.*
- 3) *Der Witt Index eines hyperbolischen Unterraums $\mathcal{H}_{2k}$ ist k .*

Beweis. 1) Der Beweis verwendet dieselbe Idee wie der Beweis aus Satz 9.19. Es seien

$$\mathcal{H}_{2k} = H_1 \odot \dots \odot H_k$$

und

$$\mathcal{K}_{2m} = K_1 \odot \dots \odot K_m$$

zwei maximal hyperbolische Unterräume von V , außerdem soll gelten $H_i = \text{span}(u_i, v_i)$ und $K_i = \text{span}(x_i, y_i)$.

Unter der Annahme, dass $\dim(\mathcal{H}) \leq \dim(\mathcal{K})$ ist die lineare Abbildung $\tau : \mathcal{H} \rightarrow \mathcal{K}$, welche durch

$$\tau u_i = x_i, \tau v_i = y_i$$

definiert ist, eine Isometrie von $\mathcal{H}$ nach $\tau\mathcal{H}$. Der Witt'sche Erweiterungssatz 9.18 impliziert die Existenz einer Isometrie $\bar{\tau} : V \rightarrow V$, welche τ erweitert. Tatsächlich ist $\bar{\tau}^{-1}(\mathcal{K})$ ein hyperbolischer Raum, welcher $\mathcal{H}$ enthält, woraus folgt, dass $\bar{\tau}^{-1}(\mathcal{K}) = \mathcal{H}$ und $\dim(\mathcal{K}) = \dim(\mathcal{H})$ gilt. Das heißt alle maximalen hyperbolischen Räume haben dieselbe Dimension.

Um zu sehen, dass die maximale Dimension $h(V)$ eines hyperbolischen Unterraums von V gleich $2w(V)$, wobei $w(V)$ der Witt Index von V ist, grenzen wir die Dimension beidseitig ein. Die nicht-singuläre Erweiterung eines maximal total degenerierten Unterraum U_w von V ist ein hyperbolischer Raum mit der Dimension $2w(V)$. Woraus folgt, dass

$h(V) \geq 2w(V)$ gelten muss. Auf der anderen Seite gibt es einen total degenerierten Unterraum U_k , welcher in jede hyperbolische Ebene $\mathcal{H}_{2k}$ enthalten ist, woraus $k \leq w(V)$ folgt. Das heißt $\dim(h(V)) \leq 2w(V)$ und $h(V) \leq 2w(V)$. Schlussendlich folgt $h(V) = 2w(V)$.

2) folgt aus 1)

3) Ein hyperbolischer Unterraum $\mathcal{H}_{2k}$ kann als

$$\mathcal{H}_{2k} = H_1 \odot \dots \odot H_k$$

Generell mit $H_i = \langle u_i, v_i \rangle$ geschrieben werden. Es folgt, dass $\langle u_1, \dots, u_k \rangle$ ein maximaler total degenerierter Unterraum ist, daher ist der Witt Index von $\mathcal{H}_{2k}$ gleich $\dim(u_1, \dots, u_k) = k$. $\square$

Satz 9.23 (Die anisotrope Zerlegung einer orthogonalen Geometrie). *Es sei $V = U \odot \text{rad}(V)$ eine orthogonale Geometrie über einem Körper F , wobei $\text{char}(F) \neq 2$ gilt. Es sei zudem $\mathcal{H}$ ein maximaler hyperbolischer Unterraum von U , wobei $\mathcal{H} = \{0\}$, falls U keine isotropen Vektoren beinhaltet. Dann gilt*

$$V = S \odot \mathcal{H} \odot \text{rad}(V),$$

wobei S anisotrop, $\mathcal{H}$ (maximal) hyperbolisch mit Dimension $2w(V)$ und $\text{rad}(V)$ total degeneriert ist. (siehe dazu auch Steven Romans Graduate Texts in Mathematics[1].) Diese Zerlegung ist eindeutig bestimmt bis auf Isomorphie, das heißt $\text{rad}(V)$ ist eindeutig bestimmt und S und $\mathcal{H}$ sind bis auf Isomorphie eindeutig bestimmt. (siehe dazu auch Jens Casten Jantzen und Joachim Schwermer, Algebra, [7].)

Beweis. 1) *Existenz:*

Wenn U bereits keine isotropen Vektoren enthält setzen wir $S = U$ und $\mathcal{H} = \{0\}$ und der Satz ist bewiesen.

Sollte U isotrope Vektoren enthalten, dann finden wir einen maximal hyperbolischen Unterraum $\mathcal{H}$ von U und es gilt

$$U = \mathcal{H} \odot \mathcal{H}^\perp.$$

Weil $\mathcal{H}$ maximal ist, ist $\mathcal{H}^\perp$ anisotrop, denn wäre $u \in \mathcal{H}^\perp$ isotrop, so wäre die nicht-singuläre Erweiterung $\mathcal{H} \odot \text{span}(u)$ ein hyperbolischer Raum, welcher größer als $\mathcal{H}$ wäre, was ein Widerspruch zur Maximalität von $\mathcal{H}$ wäre.

2) *Eindeutigkeit:*

Angenommen es gibt zwei solche Zerlegungen, dann gilt

$$V = S \odot \mathcal{H}(K^r) \odot \text{rad}(V) = S' \odot \mathcal{H}(K^r) \odot \text{rad}(V),$$

wobei S und S' anisotrop sind und $\mathcal{H}(K^r) = \mathcal{H}$ mit $\dim(\mathcal{H}) = 2r$ und $\mathcal{H}(K^s) = \mathcal{H}'$ mit $\dim(\mathcal{H}') = 2s$ gilt.

Da das Radikal bei beiden Zerlegungen ident ist, folgt aus dem Witt'schen Kürzungssatz, dass

$$S \odot \mathcal{H}(K^r) \text{ und } S' \odot \mathcal{H}(K^s)$$

isomorph sind.

Wir können ohne weiteres $r \geq s$ annehmen; dann folgt mit dem Witt'schen Kürzungssatz (10.14) aus

$$S \odot H(K^r) = S \odot H(K^{r-s}) \odot H(K^s) = S' \odot H(K^s),$$

dass

$$S' = S \odot H(K^{s-r})$$

gilt. Da S anisotrop ist, muss $r = s$ gelten, und damit $S = S'$.

(siehe dazu auch Jens Carsten Jantzen und Joachim Schwermer, Algebra, [7].) $\square$

Anmerkung 9.24. Aus der Eindeutigkeit der anisotropen Zerlegung folgt, dass S bis auf Isomorphie eindeutig bestimmt ist; S heißt daher der anisotrope Kern der Bilinearform. Aus der Eindeutigkeit der anisotropen Zerlegung folgt auch die Wohldefiniertheit der Summenoperation in der Witt Gruppe. (siehe dazu auch Martin Kneser, Quadratische Formen, [8].)

In dem nachfolgenden Beispiel wird die Anwendung der anisotropen Zerlegung auf eine orthogonale Geometrie mit dem Grundkörper $\mathbb{R}^4$ gezeigt.

Beispiel 9.25. Es sei $(\mathbb{R}^4, \langle, \rangle)$ eine orthogonale Geometrie, dann gibt es nach Satz 9.3 eine Orthogonalbasis $B = (b_1, b_2, b_3, b_4)$ von $\mathbb{R}^4$, sodass $M_B(\langle, \rangle)$ diagonal ist.

Es sei

$$M_B(\langle, \rangle) = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & -1 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix}$$

die Darstellungsmatrix der Bilinearform.

Wir bestimmen die anisotrope Zerlegung der Geometrie $(\mathbb{R}^4, \langle, \rangle)$, indem wir wie gefolgt vorgehen:

- 1) Das Radikal von $\text{rad}(\langle, \rangle)$ bestimmen
- 2) Alle isotropen Vektoren, welche außerhalb des Radikals liegen und eine passende hyperbolische Ebene bestimmen
- 3) Anisotrope Zerlegung bestimmen

1) Es gilt $\text{rad}(\langle, \rangle) = \langle b_4 \rangle$, denn es gilt $\text{rad}(\langle, \rangle) = \ker(M_B(\langle, \rangle)) = \langle b_4 \rangle$. Wir erhalten also $\mathbb{R}^4 = \langle b_1, b_2, b_3 \rangle \odot \langle b_4 \rangle$ mit $\langle b_4 \rangle = \text{rad}(\langle, \rangle)$.

2) Um die isotropen Vektoren in $\langle b_1, b_2, b_3 \rangle$ zu bestimmen, setzen wir die Linearkombination eines isotropen Vektors $v = \sum_{i=1}^3 \alpha_i \beta_i$ in die Eigenschaft $0 = \langle v, v \rangle$ ein.

$$\begin{aligned} 0 &= \left\langle \sum_i \alpha_i \beta_i, \sum_j \alpha_j \beta_j \right\rangle \\ &= \sum_{i,j} \alpha_i \alpha_j \langle b_i, b_j \rangle \end{aligned}$$

Weil $\langle b_i, b_j \rangle = 0$ für den Fall $i \neq j$ fallen alle Summenglieder für die gilt $i \neq j$ weg und es folgt.

$$\begin{aligned}\sum_{i,j} \alpha_i \alpha_j \langle b_i, b_j \rangle &= \sum_{i=1}^3 \alpha_i^2 \langle b_i, b_i \rangle \\ &= \alpha_1^2 + \alpha_2^2 - \alpha_3^2\end{aligned}$$

Die Gleichung $0 = \alpha_1^2 + \alpha_2^2 - \alpha_3^2$ hat die Lösungen $(1, 0, 1)$ und $(0, -1, -1)$, woraus folgt, dass $b_1 + b_3$ und $-b_2 - b_3$ isotrope Vektoren in $\langle b_1, b_2, b_3 \rangle$ sind.

Diese Vektoren bilden eine hyperbolische Ebene, denn

$$\begin{aligned}\langle b_1 + b_3, -b_2 - b_3 \rangle &= \langle b_1, -b_2 \rangle + \langle b_1, -b_3 \rangle + \langle b_3, -b_2 \rangle + \langle b_3, -b_3 \rangle \\ &= 0 + 0 + 0 + 1 \\ &= 1\end{aligned}$$

Somit dürfen wir $b'_1 = b_1 + b_3$ und $b'_2 = -b_2 - b_3$ setzen und erhalten die hyperbolische Eben $\mathcal{H} = \langle b'_1, b'_2 \rangle$.

An dieser Stelle sei angemerkt, wäre der isotrope Vektor $b_2 + b_3$ statt $-b_2 - b_3$ gewählt worden, wäre bei der obigen Rechnung -1 herausgekommen, woraus ebenso folgt, dass $b_1 + b_3$ und $-(b_2 + b_3) = -b_2 - b_3$ eine hyperbolische Ebene aufspannen.

3) Um die anisotrope Zerlegung zu komplettieren, suchen wir das orthogonale Komplement $\mathcal{H} \odot \text{rad}(V)$.

Dies ist identisch mit dem orthogonalen Komplement von H in $\langle b_1, b_2, b_3 \rangle$. Wir schreiben $b'_3 = \alpha_1 b_1 + \alpha_2 b_2 + \alpha_3 b_3$ und erhalten somit zwei Gleichungen:

$$\begin{aligned}\langle b_1 + b_3, \alpha_1 b_1 + \alpha_2 b_2 + \alpha_3 b_3 \rangle &= 0 \\ \Leftrightarrow \langle b_1, \alpha_1 b_1 + \alpha_2 b_2 + \alpha_3 b_3 \rangle + \langle b_3, \alpha_1 b_1 + \alpha_2 b_2 + \alpha_3 b_3 \rangle &= 0 \\ \Leftrightarrow \alpha_1 \langle b_1, b_1 \rangle + \alpha_2 \langle b_1, b_2 \rangle + \alpha_3 \langle b_1, b_3 \rangle + \alpha_1 \langle b_3, b_1 \rangle + \alpha_2 \langle b_3, b_2 \rangle + \alpha_3 \langle b_3, b_3 \rangle &= 0 \\ \Leftrightarrow \alpha_1 \langle b_1, b_1 \rangle + \alpha_3 \langle b_3, b_3 \rangle &= 0 \\ \Leftrightarrow \alpha_1 - \alpha_3 &= 0 \\ \Leftrightarrow \alpha_1 &= \alpha_3\end{aligned}$$

und

$$\begin{aligned}\langle -b_2 - b_3, \alpha_1 b_1 + \alpha_2 b_2 + \alpha_3 b_3 \rangle &= 0 \\ \Leftrightarrow \langle -b_2, \alpha_1 b_1 + \alpha_2 b_2 + \alpha_3 b_3 \rangle + \langle -b_3, \alpha_1 b_1 + \alpha_2 b_2 + \alpha_3 b_3 \rangle &= 0 \\ \Leftrightarrow \alpha_1 \langle -b_2, b_1 \rangle + \alpha_2 \langle -b_2, b_2 \rangle + \alpha_3 \langle -b_2, b_3 \rangle + \alpha_1 \langle -b_3, b_1 \rangle + \alpha_2 \langle -b_3, b_2 \rangle + \alpha_3 \langle -b_3, b_3 \rangle &= 0 \\ \Leftrightarrow \alpha_2 \langle -b_2, b_2 \rangle + \alpha_3 \langle -b_3, b_3 \rangle &= 0 \\ \Leftrightarrow -\alpha_2 + \alpha_3 &= 0 \\ \Leftrightarrow \alpha_2 &= \alpha_3\end{aligned}$$

Um eine Lösung für das Gleichungssystem zu finden, wählen wir $\alpha_3 = 1$ und erhalten $b'_3 = b_1 + b_2 + b_3$. Also folgt für die anisotrope Zerlegung

$$\mathbb{R}^4 = \mathcal{H} \odot \langle b'_3 \rangle \odot \langle b'_4 \rangle,$$

wobei $b'_4 = b_4$ gilt; dabei ist $\langle b'_4 \rangle = \text{rad}(V)$, $H = \langle b'_1, b'_2 \rangle$ eine hyperbolische Ebene und $\langle b'_3 \rangle$ ein anisotroper Teilraum, denn $\langle b'_3, b'_3 \rangle = \langle b_1 + b_2 + b_3, b_1 + b_2 + b_3 \rangle = \langle b_1, b_1 \rangle + \langle b_2, b_2 \rangle + \langle b_3, b_3 \rangle = 1 + 1 - 1 = 1$. Für die Darstellungsmatrix gilt also

$$M_{B'}(\langle, \rangle) = \begin{bmatrix} 0 & 1 & 0 & 0 \\ -1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix}.$$

Man erkennt in der Teilmatrix $\begin{bmatrix} 0 & 1 \\ -1 & 0 \end{bmatrix}$ die hyperbolische Ebene wieder. Die 1 in der Diagonale von $M_{B'}(\langle, \rangle)$ folgt aus $\langle b'_3, b'_3 \rangle = 1$ und die 0 im letzten Eintrag folgt aus dem Radikal.

10 Symplektische Bilinearform

Alle Definitionen Sätze und Anmerkungen dieses Kapitels, wenn nicht anders genannt, basieren auf Steven Romans Graduate Texts in Mathematics[1].

10.1 Nicht-Singuläre Erweiterung eines Untervektorraumes

Es sei $(V, \langle, \rangle)$ ein beliebiger metrischer Vektorraum (das heißt nicht notwendigerweise orthogonal oder symplektisch) und $U \subseteq V$ ein Unterraum von V . Wenn U singulär ist, ist es von Interesse einen *minimalen* nicht-singulären Unterraum von V zu finden, welcher U enthält. Wir werden am Ende dieses Kapitels herausarbeiten, dass es reicht eine hyperbolische Erweiterung zu finden.

Definition 10.1 (Erweiterung). Es sei $(V, \langle, \rangle)$ ein nicht-singulärer metrischer Vektorraum und $U \subseteq V$ ein Unterraum. Ein Unterraum $\bar{U} \subseteq V$ für den gilt $U \leq \bar{U}$ heißt **Erweiterung** von U

Definition 10.2 (Nicht-singuläre Ergänzung). Es sei wieder $(V, \langle, \rangle)$ ein nicht-singulärer metrischer Vektorraum. Eine **nicht-singuläre Ergänzung** von einem Unterraum $U \subseteq V$, ist eine Erweiterung von U , welche unter der Menge aller nicht-singulären Erweiterungen von U minimal ist.

Satz 10.3. *Es sei $(V, \langle, \rangle)$ ein nicht-singulärer, endlich-dimensionaler metrischer Vektorraum über einem Körper F ; wir nehmen an $\text{char}(F) \neq 2$ wenn V orthogonal ist. Sei S ein Unterraum von V ; wenn $v \in V$ isotrop ist und*

$$\text{span}(v) \odot S$$

gilt (das heißt $\text{span}(v) \odot S$ ist eine orthogonale direkte Summe), dann gibt es eine hyperbolische Ebene $H = \text{span}(v, z)$, sodass

$$H \odot S$$

gilt. Falls v isotrop ist, dann gibt es also eine hyperbolische Ebene H , welche v enthält und orthogonal zu s ist. (siehe dazu auch Satz 11.10 aus [1])

Beweis. Wir schauen uns zunächst den Fall an, dass V symplektisch ist und danach den Fall, dass V orthogonal ist. Der Abschluss des Beweises erfolgt nach einer Zusammenführung der Fälle.

Da V nicht-singulär ist, ist das Radikal trivial und es folgt $S^{\perp\perp} = S$. Da $v \notin S = S^{\perp\perp}$, gibt es ein $x \in S^\perp$, für welches $\langle v, x \rangle \neq 0$ gilt.

Wenn V symplektisch ist, dann sind alle Vektoren isotrop und wir können $z = \frac{1}{\langle v, x \rangle} \cdot x$ setzen. Und z und v bilden ein hyperbolisches Paar:

$$\langle v, z \rangle = \left\langle v, \frac{1}{\langle v, x \rangle} x \right\rangle = \langle v, x \rangle \cdot \frac{1}{\langle v, x \rangle} = 1$$

Außerdem gilt $\langle z, z \rangle = \langle v, v \rangle = 0$, weil V symplektisch ist und somit gibt es für den symplektischen Fall eine hyperbolische Ebene $H = \text{span}(v, z)$.

Wenn V orthogonal ist, setzen wir $z = rv + sx$ für ein $x \in S^\perp$. Da v, x orthogonal zu S sind, ist auch z orthogonal zu S . Weil v isotrop ist, gilt wieder bereits $\langle v, v \rangle = 0$, also bleibt zu zeigen:

$$1 = \langle v, z \rangle = \langle v, rv + sx \rangle = r \langle v, v \rangle + s \langle v, x \rangle = s \langle v, x \rangle$$

und

$$0 = \langle z, z \rangle = \langle rv + sx, rv + sx \rangle = 2rs \langle v, x \rangle + s^2 \langle x, x \rangle = 2r + s^2 \langle x, x \rangle$$

Weil $\langle v, x \rangle \neq 0$ kann die erste Gleichung mit $s = \frac{1}{\langle v, x \rangle}$ gelöst werden.

Da $\text{char}(F) \neq 2$, kann die zweite Gleichung mit $r = \frac{-s^2 \langle x, x \rangle}{2}$ gelöst werden.

Folgend sind auch im orthogonalen Fall v und z ein hyperbolisches Paar.

Also gibt es in beiden Fällen eine hyperbolische Ebene $H = \text{span}(v, z) \subseteq S^\perp$. Daher $S \subset S^{\perp\perp} \subseteq H^\perp$. Und weil H nicht-singulär ist, gilt auch $H \cap H^\perp = \{0\}$ und somit $H \cap S = \{0\}$. Also existiert die orthogonale Summe $H \odot S$. $\square$

Satz 10.4. *Es sei $(V, \langle, \rangle)$ ein nicht-singulärer, endlich-dimensionaler metrischer Vektorraum über einem Körper F ; wir nehmen an $\text{char}(F) \neq 2$ wenn V orthogonal ist.*

Es sei nun U ein Unterraum von V und

$$U = \text{span}(v_1, \dots, v_k) \odot W,$$

wobei W nicht-singulär und die Vektoren $\{v_1, \dots, v_k\}$ eine Basis von $\text{rad}(U)$ sind. Dann gibt es einen hyperbolischen Raum $\mathcal{H} = H_1 \odot \dots \odot H_k$ mit einer hyperbolischen Basis $(v_1, z_1, \dots, v_k, z_k)$, sodass

$$\overline{U} := \mathcal{H}_k \odot W$$

eine echte nicht-singuläre Erweiterung von U ist. Insbesondere gilt

$$\dim(\overline{U}) = \dim(U) + \dim(\text{rad}(U)).$$

Beweis. Der Beweis erfolgt per Induktion von k . Zunächst ist anzumerken, dass alle Vektoren v_i , weil sie in $\text{rad}(U)$ liegen, isotrop sind.

Induktionsbeginn: Es sei $k = 1$, dann existiert $\text{span}(v_1) \odot W$ und Satz 10.3 impliziert die Existenz einer hyperbolischen Ebene $H = \text{span}(v_1, z)$, für welche $H \odot W$ existiert.

Induktionsannahme: wir nehmen nun an, dass das Ergebnis auch für $k \geq 2$ richtig ist.

Induktionsschritt: Da nach Annahme

$$\text{span}(v_1) \odot (\text{span}(v_2, \dots, v_k) \odot W)$$

gilt, impliziert Satz 10.3, dass eine hyperbolische Ebene $H_1 = \text{span}(v_1, z_1)$ existiert, für welche

$$H_1 \odot (\text{span}(v_2, \dots, v_k) \odot W)$$

gilt. Da $v_2, \dots, v_k \in \text{rad}(\text{span}(v_2, \dots, v_k) \odot W)$, impliziert die Induktionsannahme, dass es einen hyperbolischen Raum $H_2 \odot \dots \odot H_k$ mit der hyperbolischen Basis $(v_2, z_2, \dots, v_k, z_k)$ gibt, für welche die direkte orthogonale Summe gleich

$$H_2 \odot \dots \odot H_k \odot W$$

existiert. Daher existiert auch $H_1 \odot \dots \odot H_k \odot W = \mathcal{H}_k \odot W$, das heißt, es gilt $\overline{U} = \mathcal{H}_k \odot W$, wobei

$$\begin{aligned} \dim(\overline{U}) &= \dim(H_k) + \dim(W) \\ &= 2\dim(\text{rad}(U)) + \dim(W) \\ &= \dim(\text{rad}(U)) + \dim(\text{rad}(U)) + \dim(W) \\ &= \dim(\text{rad}(U)) + \dim(U), \end{aligned}$$

und der Beweis ist erbracht. □

Definition 10.5 (Hyperbolische Erweiterung). Wir nennen jedes $\overline{U}$, wie in Satz 10.4 eine **hyperbolische Erweiterung** von U . Wenn U nicht-singulär ist, sagen wir, dass U eine hyperbolische Erweiterung von sich selbst ist.

Wir können nun beweisen, dass die hyperbolische Erweiterung von U genau die minimalen, nicht-singulären Erweiterungen von U sind.

Satz 10.6. *Es sei U ein Unterraum eines nicht-singulären, endlich-dimensionalen metrischen Vektorraums V . Dann sind die drei Aussagen äquivalent:*

- 1) $\mathcal{T} = \mathcal{H} \odot W$ ist eine hyperbolische Erweiterung von U
- 2) $\mathcal{T}$ ist eine minimale nicht-singuläre Erweiterung von U
- 3) $\mathcal{T}$ ist eine nicht-singuläre Erweiterung von U und

$$\dim(\mathcal{T}) = \dim(U) + \dim(\text{rad}(U))$$

Insbesondere, sind zwei nicht-singuläre Erweiterungen von U isometrisch.

Beweis. 2 $\Rightarrow$ 1): Wenn $U \leq X \leq V$, wobei X nicht-singulär ist, dann können wir den Satz 10.4 auf U als Unterraum von X anwenden und erhalten eine hyperbolische Erweiterung $\mathcal{K} \odot W$ von U für welchen gilt

$$U \subseteq \mathcal{K} \odot W \subseteq X.$$

Daher, enthält jede nicht-singuläre Erweiterung X von U eine hyperbolische Erweiterung $\mathcal{K} \odot W$ von U . Weiters, alle hyperbolischen Erweiterungen von U haben dieselbe Dimension:

$$\dim(\mathcal{H} \odot W) = \dim(U) + \dim(\text{rad}(U))$$

und somit gibt es keine echte hyperbolische Erweiterung von U , welche in einer weiteren hyperbolischen Erweiterung von U enthalten ist. Also ist $\mathcal{K} \odot W$ eine minimale nicht-singuläre Erweiterung von U . Das heißt, ist X eine minimale nicht-singuläre Erweiterung von U , dann ist $X = \mathcal{K} \odot W$ und somit ist X eine hyperbolische Erweiterung von U .

1 $\Rightarrow$ 3): Folgt aus Satz 10.4.

3 $\Rightarrow$ 2): Sei $\mathcal{T}$ eine nicht-singuläre Erweiterung von U mit $\dim(\mathcal{T}) = \dim(U) + \dim(\text{rad}(U))$. Angenommen $\mathcal{T}$ ist nicht minimal, dann folgt $\mathcal{T}$ enthält eine echte hyperbolische Erweiterung $\mathcal{K} \odot U$ von U , aber nach Satz 10.4 gilt $\dim(\mathcal{K} \odot U) = \dim(U) + \dim(\text{rad}(U))$. Folglich gilt $\mathcal{T} = \mathcal{K} \odot U$, was bedeutet, dass $\mathcal{T}$ keine echte hyperbolische Erweiterung enthält und somit ein Widerspruch vorliegt. $\square$

10.2 Klassifikation symplektischer Bilinearformen

Anmerkung 10.7. 1) Wir zeigen, dass eine symplektische Geometrie $(V, \langle \cdot, \cdot \rangle)$ nur dann eine Orthogonalbasis besitzt, wenn sie total degeneriert ist.

Beweis. Ist $v_1, \dots, v_n$ eine Orthogonalbasis, so folgt, für alle

$$x = \sum_i \alpha_i v_i \in V, \text{ und } y = \sum_j \beta_j v_j \in V \ (\alpha_i, \beta_i \in F)$$

dass

$$\langle x, y \rangle = \sum_{i,j} \alpha_i \beta_j \langle v_i, v_j \rangle = 0,$$

denn $\langle v_i, v_j \rangle = 0$, falls $i \neq j$, weil $v_1, \dots, v_n$ eine Orthogonalbasis bilden, und $\langle v_i, v_j \rangle = 0$, falls $i = j$, da die Form symplektisch und somit alternierend ist.

Umgekehrt ist für total degenerierte Form jede Basis eine Orthogonalbasis.

2) Wir sehen, weil $\langle v, v \rangle = 1$ für symplektische Formen unmöglich ist, kann es keine Orthonormalbasis für symplektische Geometrien geben. $\square$

Als „Ersatz“ für die Nichtexistenz von Orthogonalbasen haben symplektische Geometrien aber stets hyperbolische Basen:

Satz 10.8. 1) Eine symplektische Geometrie hat eine Orthogonalbasis genau dann wenn sie total degeneriert ist.

2) Jede nicht-singuläre symplektische Geometrie V ist ein hyperbolischer Raum, daher gilt

$$V = H_1 \odot H_2 \odot \dots \odot H_k,$$

wobei jedes H_i eine hyperbolische Ebene ist. Daher existiert eine hyperbolische Basis für V , welche eine Basis B bildet, für welche die Matrix die Gestalt

$$Y_{2k} = \begin{bmatrix} 0 & 1 & & & & \\ -1 & 0 & & & & \\ & & 0 & 1 & & \\ & & -1 & 0 & & \\ & & & & \ddots & \\ & & & & & 0 & 1 \\ & & & & & -1 & 0 \end{bmatrix}$$

besitzt. Insbesondere ist die Dimension von V gerade.

3) Jede symplektische Geometrie $(V, \langle, \rangle)$ hat die Gestalt

$$V = \text{rad}(V) \odot \mathcal{H},$$

wobei $\mathcal{H}$ ein hyperbolischer Raum und $\text{rad}(V)$ ein total degenerierter Raum ist. Der Rang der Form $\langle, \rangle$ ist $\dim(\mathcal{H})$ und V ist eindeutig, bis auf Isometrie, durch seinen Rang und seine Dimension festgelegt. (siehe dazu auch Steven Romans Graduate Texts in Mathematics[1].)

Beweis. 3) Sei $\mathcal{F}$ die Familie aller hyperbolischen Unterräume von V . $\mathcal{F}$ ist nicht leer, weil der Unterraum $\{0\}$ singulär ist. Folglich hat V eine hyperbolische Erweiterung ungleich 0. Weil nun V endlich-dimensional ist, gibt es ein maximal Element (bezüglich der Inklusion „ $\subseteq$ “) in $\mathcal{F}$, welches wir als $\mathcal{H}$ bezeichnen. Weil $\mathcal{H}$ nicht-singulär ist, gilt: wenn $\mathcal{H} \neq F$, so ist

$$V = \mathcal{H} \odot \mathcal{H}^\perp.$$

Da $\mathcal{H}$ nicht singulär ist, gilt $\text{rad}(V) \subseteq \mathcal{H}^\perp$. Wir nehmen an, es existiert ein $x \in \mathcal{H}^\perp$ mit $x \notin \text{rad}(V)$; dann gibt es ein $v \in V$, sodass

$$\langle x, v \rangle \neq 0.$$

Wegen $V = \mathcal{H} \odot \mathcal{H}^\perp$, gibt es $y \in \mathcal{H}^\perp$ und $z \in \mathcal{H}$, für welche gilt

$$v = z + y,$$

woraus

$$\langle x, v \rangle = \langle x, z \rangle + \langle x, y \rangle$$

folgt. Wegen $\langle x, z \rangle = 0$, muss also $\langle x, v \rangle = \langle x, y \rangle \neq 0$ sein.

Es existieren also $x, y \in \mathcal{H}^\perp$, die eine hyperbolische Ebene aufspannen, welche orthogonal zu $\mathcal{H}$ ist (vergleiche auch Beweis von 10.3. Es folgt, dass $\mathcal{H} \odot \langle x, y \rangle \in \mathcal{F}$ existiert. Das widerspricht jedoch der Maximalität von $\mathcal{H}$ und somit muss gelten $\mathcal{H}^\perp = \text{rad}(V)$ und daher hat jede symplektische Geometrie V die Gestalt

$$V = \text{rad}(V) \odot \mathcal{H}$$

Es muss weiters gelten $\text{rang}(V) = \dim(\mathcal{H})$, weil der degenerierte Teil nicht zum Rang beiträgt. Somit bestimmen Rang und Dimension die Geometrie eindeutig.

2) Für nicht-singuläre Geometrien fällt das Radikal von V aus 1) weg und es gilt

$$V = \mathcal{H} = H_1 \odot H_2 \odot \dots \odot H_k,$$

wobei jedes H_i eine hyperbolische Ebene ist. Es folgt die Existenz einer hyperbolischen Basis $\mathcal{B}$ für welche die Matrix Y_{2k} die genannte Gestalt besitzt. Die Dimension ist gerade, weil auf jedes Element H_i ein Block

$$\begin{bmatrix} 0 & 1 \\ -1 & 0 \end{bmatrix}$$

kommt. 1) folgt aus Anmerkung 10.7. □

Im nächsten Satz betrachten die Blockform der Matrix aus Satz 10.8:

Satz 10.9. *Die Menge aller symmetrischen $n \times n$ Matrizen der Gestalt*

$$X_{2k, n-2k} = \begin{bmatrix} Y_{2k} & 0 \\ 0 & 0_{n-2k} \end{bmatrix}$$

ist eine Menge von Vertretern, der Menge der alternierender Matrizen bezüglich Kongruenz. (siehe dazu auch Steven Romans Graduate Texts in Mathematics[1].)

Beweis. Symplektische Geometrien werden durch alternierende Matrizen repräsentiert. Das entspricht schiefsymmetrischen Matrizen mit den Diagonaleinträgen 0. Weiters existiert nach Satz 10.8 für jede symmetrische $n \times n$ Matrix, welche alternierend ist, eine kongruente Matrix der Gestalt

$$X_{2k, n-2k} = \begin{bmatrix} Y_{2k} & 0 \\ 0 & 0_{n-2k} \end{bmatrix}.$$

Der Rang von $X_{2k, n-2k}$ ist gleich $2k$. Zwei Matrizen dieser Form können nur kongruent sein, wenn ihre Ränge übereinstimmen. □

10.3 Pfaff'sche Determinante einer alternierenden Matrix

Satz 10.10. *Der Rang einer alternierend Matrix A über einem Körper V ist immer gerade und die Determinante ist immer ein Quadrat. (siehe dazu auch „Algebra“ und „Basic Algebra I“ von Lang S. und Jacobson N. [3], [4])*

Beweis. Da die Matrix A alternierend ist, können wir schreiben:

$$A = Q \begin{bmatrix} 0 & 1 & & & & \\ -1 & 0 & & & & \\ & & 0 & 1 & & \\ & & -1 & 0 & & \\ & & & \ddots & & \\ & & & & 0 & 1 \\ & & & & -1 & 0 \\ & & & & & 0 \\ & & & & & \ddots & \\ & & & & & & 0 \end{bmatrix} Q^T$$

Da die Matrix aus $\begin{bmatrix} 0 & 1 \\ -1 & 0 \end{bmatrix}$ -Blöcken vom Rang 2 besteht, ist der Rang gerade. Berechnen der Determinante ergibt

$$\det(A) = \det(Q) \cdot \begin{bmatrix} 0 & 1 & & & & \\ -1 & 0 & & & & \\ & & 0 & 1 & & \\ & & -1 & 0 & & \\ & & & \ddots & & \\ & & & & 0 & 1 \\ & & & & -1 & 0 \\ & & & & & 0 \\ & & & & & \ddots & \\ & & & & & & 0 \end{bmatrix} \cdot \det(Q^T) = \det(Q^2) \cdot \begin{Bmatrix} 1 \\ 0 \end{Bmatrix}$$

Da die Determinante von Q^2 , aber auch 0 und 1, ein Quadrat ist, ist der Beweis vollständig erbracht. $\square$

Im nachfolgenden Satz verschärfen wir den Satz 10.10 und erhalten Formeln für Wurzel der Determinante.

Satz 10.11. *Es sei n gerade und A eine alternierende $n \times n$ Matrix, deren (i,j) -Eintrag, $i < j$, die Unbestimmte $a_{i,j}$ ist. (Die Einträge unterhalb der Diagonalen sind dann durch $a_{ij} = -a_{ji}$ gegeben und $a_{ii} = 0$ für alle i .) Dann gibt es genau ein Polynom $Pf(A)$ in den Variablen $a_{i,j}$, $i < j$ sodass*

$$(Pf(A))^2 = \det(A)$$

beziehungsweise

$$„\sqrt{Pf(A)} = \det(A)“$$

und

$$Pf(s) = 1$$

gilt. (siehe dazu auch „Algebra“ und „Basic Algebra I“ von Leng S. und Jacobson N. [3], [4])

Beweis. Da die Matrix

$$A = \begin{bmatrix} a_{11} & \dots & a_{1n} \\ \vdots & & \vdots \\ a_{n1} & \dots & a_{nn} \end{bmatrix}$$

alternierend ist, kann nach Satz 10.10 die Determinante von A als Quadrat in $\mathbb{Q}(a_{i,j}, i, j = 1, \dots, n)$ geschrieben werden.

Wir verwenden dazu die beiden teilerfremden Polynome $P(a_{1,1}, \dots, a_{nn})$ und $Q(a_{11}, \dots, a_{nn})$, welche beide im gleichen Polynomring $\mathbb{Z}[a_{ij}, i, j = 1, \dots, n]$ liegen.

Es gilt also

$$\det(A) = \left(\frac{P(a_{11}, \dots, a_{nn})}{Q(a_{11}, \dots, a_{nn})} \right)^2,$$

woraus nun folgt

$$Q^2 \det(A) = P^2.$$

Da aber Q und P teilerfremd sind muss nun $Q = 1$ gelten und es folgt

$$\det(A) = P(a_{11}, \dots, a_{nn})^2$$

womit der Satz bewiesen ist. □

Anmerkung 10.12. Ist $A = (x_{ij})$ eine alternierende Matrix mit Einträgen aus dem Ring R , dann gilt also $Pf(x_{ij})^2 = \det(A)$. Das heißt um die „Wurzel „aus der Determinante von A zu berechnen genügt es die Pfaff'sche bei $A = (a_{ij})$ zu spezialisieren.

10.4 Witt'sche Erweiterungs- und Kürzungssatz

Satz 10.13. *Es seien V und V' isometrische nicht-singuläre metrische Vektorräume und $U \subseteq V$ ein Unterraum von V , mit einer nicht-singulären Erweiterung $\bar{U}$. Jede Isometrie $\tau : U \rightarrow \tau U$ kann auf eine Isometrie von $\bar{U}$, welche auf eine nicht-singuläre Ergänzung von τU abbildet, erweitert werden.*

Beweis. Wir schreiben $U = \text{rad}(U) \odot W$, wobei $(u_1, \dots, u_k)$ eine Basis für $\text{rad}(U)$ ist, und $\bar{U} = \mathcal{H} \odot W$ sei eine hyperbolische Erweiterung, wobei $(u_1, z_1, \dots, u_k, z_k)$ eine hyperbolische Basis für $\mathcal{H}$ ist (vergleiche 10.5 und 10.4).

Es sei $\tau : U \rightarrow \tau U$ eine Isometrie. Weil $(u_1, \dots, u_k)$ eine Basis für $\text{rad}(U)$ ist, folgt, dass $(\tau u_1, \dots, \tau u_k)$ eine Basis für $\text{rad}(\tau U)$.

Daher können wir eine hyperbolische Erweiterung für $\tau U = \text{rad}(\tau W) \odot \tau W$ finden,

$$\overline{\tau U} = \mathcal{H}' \odot \tau W,$$

wobei $\mathcal{H}'$ die hyperbolische Basis $(\tau u_1, x_1, \dots, \tau u_k, x_k)$ hat (siehe Satz 10.4). Um τ zu erweitern, setzen wir $\bar{\tau} z_i = x_i$ für alle $i = 1, \dots, k$. $\square$

Satz 10.14 (Witt'sche Erweiterungssatz). *Es seien V und V' isometrische nicht-singuläre symplektische Geometrien über einem Körper F . Dann gilt für jede Isometrie*

$$\tau : S \rightarrow \tau S \subseteq V',$$

wobei S ein Unterraum von V ist, dass diese auf eine Isometrie von V nach V' erweitert werden kann. (siehe dazu auch Steven Romans Graduate Texts in Mathematics[1].)

Beweis. Laut Satz 10.13 können wir τ auf eine nicht-singuläre Ergänzung von S erweitern. Daraus folgt, dass wir ohne Probleme annehmen können, dass sowohl S , als auch τS nicht-singulär sind. Daher gilt

$$V = S \odot S^\perp$$

und

$$V' = \tau S \odot (\tau S)^\perp.$$

Um die Erweiterung von τ auf V abzuschließen, müssen wir lediglich eine hyperbolische Basis $B_{S^\perp}$ für $S^\perp$,

$$B_{S^\perp} = (e_1, f_1, \dots, e_p, f_p),$$

und eine hyperbolische Basis $B_{(\tau S)^\perp}$ für $(\tau S)^\perp$,

$$B_{(\tau S)^\perp} = (e'_1, f'_1, \dots, e'_p, f'_p)$$

wählen.

Wir definieren nun die Erweiterung durch Setzen von $\tau e_i = e'_i$ und $\tau f_i = f'_i$. $\square$

Der Witt'sche Kürzungssatz folgt direkt aus dem Witt'schen Erweiterungssatz:

Satz 10.15 (Witt'sche Kürzungssatz). *Es seien V und V' isometrische nicht-singuläre symplektische Geometrien über einem Körper F . Wenn*

$$V = S \odot S^\perp \text{ und } V' = T \odot T^\perp$$

dann folgt

$$S \approx T \Rightarrow S^\perp \approx T^\perp.$$

(siehe dazu auch Steven Romans Graduate Texts in Mathematics[1].)

Beweis. wenn $S \approx T$ gilt, gibt es eine Isometrie

$$\tau : S \rightarrow T.$$

Satz 10.14 besagt, dass die Isometrie τ auf eine Isometrie, welche V auf V' abbildet, erweitert werden kann. Somit muss weiters eine Isometrie

$$\phi : S^\perp \rightarrow T^\perp$$

existieren und es folgt $S^\perp \approx T^\perp$.

□

Literatur

- [1] Roman, S. *Advanced Linear Algebra* 3. Edition, Springer
- [2] Weyl, H. *The Classical Groups*, Princeton University Press
- [3] Lang, S. *Algebra*, Springer, Kapitel XV, Paragraph 9
- [4] Jacobson, N. *Basic Algebra I*, Dover Publications, Kapitel 6.2
- [5] Serre, J.-P. *A course in Arithmetic*, Springer, Kapitel IV, Paragraph 2
- [6] Serre, J.-P. *A course in Arithmetic*, Springer, Kapitel IV, Paragraph 3
- [7] Jantzen, J. C., Schwermer, J. *Algebra*, 2. Auflage, Springer
- [8] Kneser, M. *Quadratische Formen*, Springer
- [9] Witt, E. *Theorie der quadratischen Formen in beliebigen Körpern*, J. f. d. reine u. angewandte Mathematik **176** (1936)