

DE ÆMULATIONIBUS REVOLUTIONUM
(MAWF '25/26 1.XIII & 1.XIV [RRS])

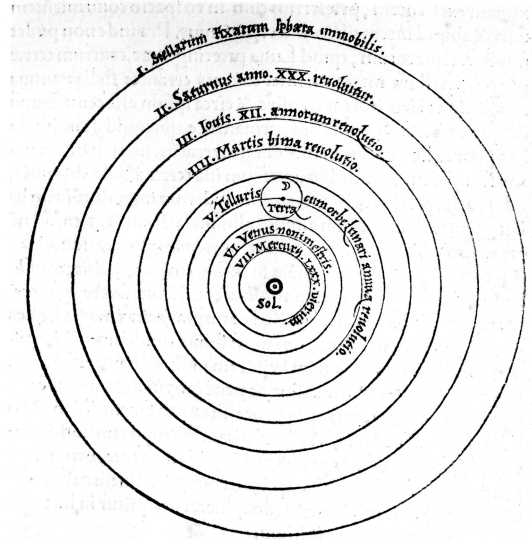


FIG. 1. „To kłamstwo! Kopernik była kobietą!”

SPIS TREŚCI

1. Cliffordowskie realizacje izometrii	1
Dodatek A. Uzupełnienie z teorii przestrzeni kwadratowych	10
Literatura	17

1. CLIFFORDOWSKIE REALIZACJE IZOMETRII

W niniejszym cyklu wykładów zajmiemy się szczegółową rekonstrukcją funktorialnego podniesienia izometrycznych automorfizmów przestrzeni kwadratowej (V, Q) do jej algebry Clifforda $\text{Cliff}(V, Q)$, co doprowadzi nas ostatecznie do identyfikacji – w odwołaniu do prezentacji grupy ortogonalnej w terminach generatorów (i relacji), którą precyzuje Twierdzenie Cartana–Dieudonnégo – podgrupy w algebrze Clifforda implementującej na kanonicznie zanurzonej w $\text{Cliff}(V, Q)$ wyjściowej przestrzeni V (poprzez ograniczenie doń odnośnych automorfizmów) transformacje ortogonalne. Nasze dociekania rozpoczniemy od

Definicja 1. Przyjmijmy zapis dotychczasowy. **Multyplikatywna grupa jedności** algebry Clifforda to zbiór

$$\text{Cliff}(V, Q)^\times := \{ \gamma \in \text{Cliff}(V, Q) \mid \exists \gamma^{-1} \in \text{Cliff}(V, Q) : \gamma \cdot \gamma^{-1} = e^C = \gamma^{-1} \cdot \gamma \}$$

odwracalnych elementów $\text{Cliff}(V, Q)$ ze strukturą grupy indukowaną z tejże algebry unitalnej (tj. w szczególności z operacją binarną daną przez mnożenie w algebrze Clifforda). **Reprezentacja**

dołączona tej grupy to homomorfizm grup

$$\begin{aligned} \text{Ad} &: \text{Cliff}(V, Q)^\times \longrightarrow \text{Inn}(\text{Cliff}(V, Q)) \\ &: u \longmapsto \left(\text{Cliff}(V, Q) \ni \gamma \longmapsto u \cdot \gamma \cdot u^{-1} \equiv \text{Ad}_u(\gamma) \right). \end{aligned}$$

Bywa on także nazywany **reprezentacją wektorową** $\text{Cliff}(V, Q)^\times$.

W bezpośrednim nawiązaniu do uwagi wprowadzającej wysławiamy

Stwierdzenie 1. Przyjmijmy zapis dotychczasowy. Dla dowolnego wektora nieizotropowego $v \in V \setminus Q^{-1}(\{0_{\mathbb{K}}\})$ i dowolnego wektora $w \in V$ spełniona jest tożsamość

$$-\text{Ad}_v(w) = w - 2 \frac{\Phi_Q(w, v)}{Q(v)} \triangleright v,$$

więc też – w szczególności –

$$\text{Ad}_v(V) = V,$$

przy czym utożsamiliśmy – jak zwykle – przestrzeń V z jej monomorficznym obrazem w $\text{Cliff}(V, Q)$.

Dowód: Odwrotność wektora nieizotropowego v w $\text{Cliff}(V, Q)$ to

$$v^{-1} = Q(v)^{-1} \triangleright v,$$

stąd

$$-Q(v) \triangleright \text{Ad}_v(w) = -v \cdot w \cdot v \equiv -v \cdot \{w, v\} + v^2 \cdot w = -2\Phi_Q(w, v) \triangleright v + Q(v) \triangleright w.$$

□

Uwaga 1. Działanie dołączone w ograniczeniu (obustronnym) do $V \subset \text{Cliff}(V, Q)$ realizuje superpozycje przeciwności (odbicia wzgl. wektora zerowego) z odbiciami w płaszczyznach ortogonalnych (wzgl. Q) do wektorów nieizotropowych.

Powyższe obserwacje prowadzą nas wprost do

Stwierdzenie 2. Przyjmijmy zapis dotychczasowy i dla

$$V^\times := V \setminus Q^{-1}(\{0_{\mathbb{K}}\})$$

zdefiniujmy grupę

$$(1) \quad P(V, Q) := \{v \in V^\times\} \subset \text{Cliff}(V, Q)^\times,$$

tj. grupę generowaną mnożeniem (wzgl. mnożenia indukowanego z algebry $\text{Cliff}(V, Q)$) przez wektory nieizotropowe. Reprezentacja dołączona $\text{Cliff}(V, Q)^\times$ na $\text{Cliff}(V, Q)$ ogranicza się do reprezentacji

$$\underline{\text{Ad}} : P(V, Q) \longrightarrow \text{O}(V, Q) : u \longmapsto \text{Ad}_u|_{V \subset \text{Cliff}(V, Q)}.$$

Dowód: Homomorficzność $\underline{\text{Ad}}$ jest oczywista,

$$\forall u_1, u_2 \in P(V, Q) : \underline{\text{Ad}}_{u_1 \cdot u_2} = \underline{\text{Ad}}_{u_1} \circ \underline{\text{Ad}}_{u_2},$$

pozostaje zatem tylko sprawdzić, że endomorfizm $\underline{\text{Ad}}_u$ zachowuje formę kwadratową Q , co czynimy – dla dowolnych $v \in V \setminus Q^{-1}(\{0_{\mathbb{K}}\})$ i $w \in V$ – w bezpośrednim rachunku odwołującym się do Stw. 1,

$$\begin{aligned} Q \circ \underline{\text{Ad}}_v(w) &\equiv Q(v \cdot w \cdot v^{-1}) = Q(-v \cdot w \cdot v^{-1}) = Q\left(w - 2 \frac{\Phi_Q(w, v)}{Q(v)} \triangleright v\right) \\ &\equiv \Phi_Q\left(w - 2 \frac{\Phi_Q(w, v)}{Q(v)} \triangleright v, w - 2 \frac{\Phi_Q(w, v)}{Q(v)} \triangleright v\right) \\ &= \Phi_Q(w, w) - 4 \frac{\Phi_Q(w, v)}{Q(v)} \cdot_{\mathbb{K}} \Phi_Q(w, v) + 4 \frac{\Phi_Q(w, v)^2}{Q(v)^2} \cdot_{\mathbb{K}} \Phi_Q(v, v) \end{aligned}$$

$$= Q(w).$$

□

Reprezentacja dołączona, jakkolwiek przydatna przy wyrabianiu sobie pewnych elementarnych intuicji algebraicznych i geometrycznych, z których skorzystamy w dalszej części wywodu, ma jedną istotną niedoskonałość: Oto jej ograniczenie $\underline{\text{Ad}}$ nie zawiera w ogólności transformacji zmieniających orientację V na przeciwną, czyli – w świetle Tw. Cartana–Dieudonnégo 2 – zbioru generującego grupę $O(V, Q)$, której emulacja jest naszym celem. Istotnie, w przypadku $\dim_{\mathbb{K}} V = 2n + 1 \in 2\mathbb{N} + 1$ endomorfizm Ad_v , $v \in V^\times$ odwzorowuje wektor v w siebie, $\text{Ad}_v(v) = v$, a dowolny wektor $w \in \langle v \rangle_{\mathbb{K}}^{\perp_Q}$ – w wektor przeciwny, $\text{Ad}_v(w) = -w$, co oznacza, że wyznacznik tego endomorfizmu ma wartość $\det \text{Ad}_v = 1 \cdot (-1)^{2n} = 1$. Ażeby naprawić tę niedoskonałość, musimy w naszej rekonstrukcji funktorialnego podniesienia automorfizmów przestrzeni kwadratowej wyjść poza zbiór automorfizmów wewnętrznych, co czynimy w następnej

Definicja 2. Przyjmijmy zapis dotychczasowy. **Reprezentacja dołączona zwichrowana** modyfikatywną grupy jednostki algebry Clifforda, zwana także **reprezentacją J_V -wektorową** modyfikatywną grupy jednostki, to homomorfizm grup

$$\begin{aligned} \widetilde{\text{Ad}} &: \text{Cliff}(V, Q)^\times \longrightarrow \text{Aut}(\text{Cliff}(V, Q)) \\ &: u \longmapsto \left(\text{Cliff}(V, Q) \ni \gamma \longmapsto J_V(u) \cdot \gamma \cdot u^{-1} \equiv \widetilde{\text{Ad}}_u(\gamma) \right). \end{aligned}$$

Reprezentacja ta pozwala wyróżnić grupę

$$\Gamma(V, Q) := \{ u \in \text{Cliff}(V, Q)^\times \mid \widetilde{\text{Ad}}_u(V) = V \} \supset P(V, Q),$$

określaną mianem **grupy Clifforda**.

Punktu wyjścia do analizy grupy Clifforda dostarcza

Stwierdzenie 3. Przyjmijmy zapis dotychczasowy i niechaj $((V, +_V, P_V, \bullet \mapsto 0_V), \ell_V, Q)$ będzie niezwyrodniałą przestrzenią kwadratową skończonego wymiaru nad ciałem $\mathbb{K}$ o charakterystyce $\text{char}(\mathbb{K}) \neq 2$. Ograniczenie reprezentacji dołączonej zwichrowanej

$$\widetilde{\text{Ad}} : \Gamma(V, Q) \longrightarrow \text{GL}(V; \mathbb{K}) \equiv \text{Aut}_{\mathbb{K}}(V) : u \longmapsto \widetilde{\text{Ad}}_u \upharpoonright_{V \subset \text{Cliff}(V, Q)}.$$

zadaje ciąg dokładny grup

$$(2) \quad \mathbf{1} \longrightarrow \mathbb{K}^\times \longrightarrow \Gamma(V, Q) \xrightarrow{\widetilde{\text{Ad}}} \text{GL}(V; \mathbb{K}),$$

w którym

$$\mathbb{K}^\times \equiv \mathbb{K} \setminus \{0_{\mathbb{K}}\}$$

jest modyfikatywną grupą niezerowych elementów ciała $\mathbb{K}$.

Dowód: Przynależność do jądra $\widetilde{\text{Ad}}$ (obrazów) elementów $\mathbb{K}^\times$ jest rzeczą oczywistą,

$$\forall \lambda \in \mathbb{K}^\times \quad \forall v \in V : \widetilde{\text{Ad}}_{\lambda \triangleright e^C}(v) = (\lambda \triangleright e^C) \cdot v \cdot (\lambda \triangleright e^C)^{-1} \equiv (\lambda \triangleright e^C) \cdot v \cdot (\lambda^{-1} \triangleright e^C) = \lambda \cdot_{\mathbb{K}} \lambda^{-1} \triangleright v = v.$$

Wnioskując odwrotnie, wybierzmy bazę ortogonalną $\{v_i\}_{i \in \overline{1, N}}$, $N \equiv \dim_{\mathbb{K}} V$ w przestrzeni V ,

$$\forall_{i, j \in \overline{1, N}} : \Phi_Q(v_i, v_j) = \lambda_i \delta_{i, j}^{\mathbb{K}}, \quad \lambda_i \in \mathbb{K}^\times,$$

i rozważmy dowolny element $u \in \text{Ker } \widetilde{\text{Ad}}$, który wprost na mocy definicji spełnia warunek

$$\forall_{v \in V} : J_V(u) \cdot v = v \cdot u,$$

tj.

$$\widetilde{\text{Ad}}_u(v) = v \equiv \text{id}_V(v).$$

Rozkładając $u = u_0 + u_1$ na składowe $u_k \in \text{Cliff}(V, Q)^k$, $k \in \{0, 1\}$, przepisujemy powyższe w postaci

$$u_0 \cdot v - u_1 \cdot v = v \cdot u_0 + v \cdot u_1,$$

czyli – wobec Równ. (5.8) – równoważnie w formie układu warunków

$$u_0.v = v.u_0 \quad \wedge \quad u_1.v = -v.u_1.$$

Redukcja stopnia 2 w $\text{Cliff}(V, Q)$ określona przez relacje

$$\forall_{i,j \in \overline{1,N}} : v_i.v_j = -v_j.v_i + 2\lambda_i \delta_{i,j}^{\mathbb{K}}$$

pozwała przy tym rozpisać

$$u_0 = \alpha_0(v_2, v_3, \dots, v_N) + v_1.\alpha_1(v_2, v_3, \dots, v_N),$$

$$u_1 = \beta_1(v_2, v_3, \dots, v_N) + v_1.\beta_0(v_2, v_3, \dots, v_N)$$

w terminach pewnych wielomianów α_k, β_k parzystości k . Kładąc $v = v_1$, otrzymujemy z pierwszego z powyższych warunków równość

$$\alpha_0(v_2, v_3, \dots, v_N).v_1 + v_1.\alpha_1(v_2, v_3, \dots, v_N).v_1 = v_1.\alpha_0(v_2, v_3, \dots, v_N) + v_1^2.\alpha_1(v_2, v_3, \dots, v_N),$$

która sprowadza się ostatecznie do postaci

$$v_1.\alpha_0(v_2, v_3, \dots, v_N) - v_1^2.\alpha_1(v_2, v_3, \dots, v_N) = v_1.\alpha_0(v_2, v_3, \dots, v_N) + v_1^2.\alpha_1(v_2, v_3, \dots, v_N),$$

czyli

$$2\lambda_1 \triangleright \alpha_1(v_2, v_3, \dots, v_N) = 0_{\mathbb{K}}.$$

Jako że Q jest niezwyrodniała, $\lambda_1 \neq 0_{\mathbb{K}}$, przeto $\alpha_1(v_2, v_3, \dots, v_N) = \mathbf{0}_{\text{Cliff}(V, Q)}$, a zatem u_0 nie zależy od v_1 . Powtarzając analogiczne rozumowanie dla v_j , $j \in \overline{2, N}$, stwierdzamy ostatecznie, że

$$u_0 = \lambda \triangleright e^C, \quad \lambda \in \mathbb{K}.$$

W podobny sposób stwierdzamy, że także u_1 nie zależy od v_i , $i \in \overline{1, N}$, co wobec nieparzystości u_1 implikuje jego trywialność, $u_1 = \mathbf{0}_{\text{Cliff}(V, Q)}$. W sumie zatem mamy

$$u = u_0 + u_1 = u_0 = \lambda \triangleright e^C,$$

co z racji odwracalności u daje nam tezę dowodzonego stwierdzenia. $\square$

Uwaga 2. Założenie o niezwyrodnieniu formy kwadratowej Q poczynione w ostatnim stwierdzeniu jest istotne, o czym przekonuje nas analiza przypadku $\text{Cliff}(V, 0) \equiv \wedge^\bullet V$. Oto więc dla dowolnych *niewspółliniowych* $v_1, v_2 \in V \setminus \{0_V\}$ sprawdzamy tożsamość

$$(1_{\mathbb{K}} + v_1 \wedge v_2)^{-1} = 1_{\mathbb{K}} - v_1 \wedge v_2,$$

która oznacza, że $1_{\mathbb{K}} + v_1 \wedge v_2 \in \text{Cliff}(V, Q)^\times$. Przy tym dla dowolnego $v \in V$ zachodzi

$$\widetilde{\text{Ad}}_{1_{\mathbb{K}} + v_1 \wedge v_2}(v) = J_V(1_{\mathbb{K}} + v_1 \wedge v_2) \wedge v \wedge (1_{\mathbb{K}} - v_1 \wedge v_2) = (1_{\mathbb{K}} + v_1 \wedge v_2) \wedge (v - v \wedge v_1 \wedge v_2)$$

$$= v - v \wedge v_1 \wedge v_2 + v_1 \wedge v_2 \wedge v - v_1 \wedge v_2 \wedge v \wedge v_1 \wedge v_2 = v - v \wedge v_1 \wedge v_2 + v \wedge v_1 \wedge v_2 = v,$$

co pokazuje dowodnie, że $\text{Ker } \widetilde{\text{Ad}}$ zawiera elementy nie-skalarne.

Zanim podejmiemy dalszą analizę relacji między grupą Clifforda a grupą automorfizmów przestrzeni kwadratowej (V, Q) , wprowadzimy obecnie nader przydatną konstrukcję pomocniczą:

Definicja 3. Przyjmijmy zapis dotychczasowy. **Norma spinorowa** na $\text{Cliff}(V, Q)$ to odwzorowanie

$$\mathbf{N} : \text{Cliff}(V, Q) \curvearrowright : \gamma \mapsto \gamma.J_V \circ T_V(\gamma),$$

o oczywistej własności

$$\forall_{v \in V} : \mathbf{N}(v) = -Q(v) \triangleright e^C.$$

Kluczową własność określonego powyżej odwzorowania wysławia

Stwierdzenie 4. Przyjmijmy zapis Def. 3 oraz Stw. 3, zakładając przy tym, że przestrzeń kwadratowa (V, Q) jest skończenie wymiarowa i niezwyrodniała. Wówczas norma spinorowa ogranicza się do homomorfizmu grup

$$\mathbf{N}|_{\Gamma(V, Q)} : \Gamma(V, Q) \longrightarrow \mathbb{K}^\times \triangleright e^C \equiv \mathbb{K}^\times.$$

Dowód: Zauważmy na wstępie, że dla dowolnego $u \in \Gamma(V, Q)$ i wszystkich $v \in V$ zachodzi – wprost na mocy definicji grupy Clifforda – $J_V(u).v.u^{-1} \in V$, więc także

$$J_V(u).v.u^{-1} = T_V(J_V(u).v.u^{-1}) = T_V(u)^{-1}.T_V(v).T_V \circ J_V(u) = T_V(u)^{-1}.v.J_V \circ T_V(u),$$

czyli

$$v = (T_V(u).J_V(u)).v.(J_V \circ T_V(u).u)^{-1} = J_V(J_V \circ T_V(u).u).v.(J_V \circ T_V(u).u)^{-1} \equiv \widetilde{\text{Ad}}_{J_V \circ T_V(u).u}(v),$$

co wynika z tego, że $\Gamma(V, Q)$ w reprezentacji $\widetilde{\text{Ad}}$ zachowuje V , a przy tym jest grupą, i oznacza, że

$$J_V \circ T_V(u).u \in \text{Ker } \widetilde{\text{Ad}}|_V.$$

Przy tym wobec pierwszej z powyższych równości zachodzi

$$\begin{aligned} \widetilde{\text{Ad}}_{J_V \circ T_V(u)}(v) &= T_V(u).v.J_V \circ T_V(u)^{-1} \equiv T_V(u^{-1})^{-1}.v.J_V \circ T_V(u^{-1}) = J_V(u^{-1}).v.u \\ &\equiv J_V(u^{-1}).v.(u^{-1})^{-1} \equiv \widetilde{\text{Ad}}_{u^{-1}}(v) \in V, \end{aligned}$$

co pozwala stwierdzić, że $J_V \circ T_V(u) \in \Gamma(V, Q)$, czyli też

$$J_V \circ T_V(u).u \in \Gamma(V, Q) \cap \text{Ker } \widetilde{\text{Ad}}|_V,$$

a w takim razie na mocy Stw. 3 wyprowadzamy wniosek o istnieniu skalaru $\lambda_u \in \mathbb{K}^\times$ spełniającego warunek

$$J_V \circ T_V(u).u = \lambda_u \triangleright e^C.$$

To prowadzi do wniosku, że

$$\mathbf{N}(T_V(u)) \equiv T_V(u).J_V(u) = J_V(J_V \circ T_V(u).u) = J_V(\lambda_u \triangleright e^C) = \lambda_u \triangleright e^C.$$

Zarazem

$$V = J_V(u).V.u^{-1} \iff V = J_V(u)^{-1}.V.u,$$

przeto

$$V \equiv T_V(V) = T_V(J_V(u)^{-1}.V.u) = T_V(u).T_V(V).T_V \circ J_V(u)^{-1} = T_V(u).V.J_V \circ T_V(u)^{-1},$$

więc ostatecznie

$$\begin{aligned} V \equiv -J_V(V) &= -J_V(T_V(u).V.J_V \circ T_V(u)^{-1}) = -J_V \circ T_V(u).J_V(V).J_V \circ J_V \circ T_V(u)^{-1} \\ &= J_V \circ T_V(u).V.T_V(u)^{-1} \equiv \widetilde{\text{Ad}}_{T_V(u)}(V). \end{aligned}$$

Widzimy zatem, że

$$u \in \Gamma(V, Q) \iff T_V(u) \in \Gamma(V, Q),$$

co pozwala wcześniejszą naszą konstatację przepisać w pożądanej postaci

$$\mathbf{N}(u) = \lambda_{T_V(u)} \triangleright e^C \in \mathbb{K}^\times \triangleright e^C,$$

czyli

$$\mathbf{N}(\Gamma(V, Q)) \subset \mathbb{K}^\times \triangleright e^C.$$

Na koniec bez trudu sprawdzamy, dla dowolnych $u_1, u_2 \in \Gamma(V, Q)$, że

$$\begin{aligned} \mathbf{N}(u_1.u_2) &\equiv u_1.u_2.J_V \circ T_V(u_1.u_2) = u_1.\mathbf{N}(u_2).J_V \circ T_V(u_1) = u_1.\lambda_{T_V(u_2)} \triangleright e^C.J_V \circ T_V(u_1) \\ &= u_1.J_V \circ T_V(u_1).\lambda_{T_V(u_2)} \triangleright e^C \equiv \mathbf{N}(u_1).\mathbf{N}(u_2). \end{aligned}$$

□

Wyposażeni w nowe narzędzie analizy strukturalnej, jakim jest norma spinorowa, możemy przyjrzeć się bliżej ciągowi dokładnemu grup (2). Czynimy to w poniższym

Stwierdzenie 5. Przyjmijmy zapis Stw. 3. Ograniczenie $\widetilde{\text{Ad}}$ reprezentacji dołączonej zwichrowanej zadaje homomorfizm grup

$$\widetilde{\text{Ad}} : \Gamma(V, Q) \longrightarrow \text{O}(V, Q).$$

Dowód: Zaczniemy od wyznaczenia – w odwołaniu do Stw. 4, a dla dowolnego $u \in \Gamma(V, Q)$ – normy

$$\begin{aligned} \text{N}(J_V(u)) &\equiv J_V(u).J_V \circ T_V \circ J_V(u) = J_V(u).T_V(u) = J_V(u.J_V \circ T_V(u)) \equiv J_V(\text{N}(u)) \\ &= J_V(\lambda_{T_V(u)} \triangleright e^C) = \lambda_{T_V(u)} \triangleright e^C \equiv \text{N}(u). \end{aligned}$$

Dla $u \in \Gamma(V, Q)$ mamy tożsamość (w oczywistym zapisie symbolicznym)

$$\begin{aligned} V &\equiv -V = J_V(V) = J_V(\widetilde{\text{Ad}}_u(V)) \equiv J_V(J_V(u).V.u^{-1}) = -J_V(J_V(u)).V.J_V(u^{-1}) \\ &= -J_V(J_V(u)).V.J_V(u)^{-1}, \end{aligned}$$

z której wynika $J_V(u) \in \Gamma(V, Q)$. Uwzględnivszy powyższe oraz $u^{-1} \in \Gamma(V, Q)$, jak również treść Stw. 4, możemy następnie obliczyć, dla dowolnego (anizotropowego) $v \in V^\times \subset \Gamma(V, Q)$,

$$\begin{aligned} \text{N}(\widetilde{\text{Ad}}_u(v)) &\equiv \text{N}(J_V(u).v.u^{-1}) = \text{N}(J_V(u)).\text{N}(v).\text{N}(u^{-1}) = \text{N}(u).\text{N}(v).\text{N}(u)^{-1} \\ &\equiv \lambda_{T_V(u)} \triangleright e^C.\text{N}(v).\lambda_{T_V(u)}^{-1} \triangleright e^C = \text{N}(v) \end{aligned}$$

i na tej podstawie – stwierdzić, że działanie dołączone zwichrowane zachowuje $Q(v)$ dla wszystkich anizotropowych wektorów $v \in V^\times$. Z drugiej strony ilekroć $Q(v) = 0_{\mathbb{K}}$ (a dowolny wektor jest albo anizotropowy, albo spełnia $Q(v) = 0_{\mathbb{K}}$), zachodzi $\widetilde{\text{Ad}}_u(v) \notin V^\times$, gdyż w przeciwnym razie byłoby – w świetle wcześniejszych ustaleń – $v = \widetilde{\text{Ad}}_{u^{-1}}(\widetilde{\text{Ad}}_u(v)) \in \widetilde{\text{Ad}}_u(V^\times) \subset V^\times$, co prowadziłoby do sprzeczności. Także zatem w przypadku izotropowym wartość (zerowa) formy kwadratowej jest zachowywana przez działanie $\widetilde{\text{Ad}}$. □

Ostatnie stwierdzenie w połączeniu z definicją (1) podgrupy $P(V, Q) \subset \Gamma(V, Q)$ oraz tożsamością

$$(3) \quad \widetilde{\text{Ad}}_v \equiv P_v$$

każe postawić pytanie o surjektywny charakter indukowanego przez reprezentację dołączoną zwichrowaną homomorfizmu

$$\widetilde{\text{Ad}} \upharpoonright_{P(V, Q)} : P(V, Q) \longrightarrow \text{O}(V, Q).$$

Odpowiedzi na to pytanie dostarcza

Stwierdzenie 6. Przyjmijmy zapis Def. 2 oraz Stw. 2. Ograniczenie reprezentacji dołączonej zwichrowanej

$$\widetilde{\text{Ad}} \upharpoonright_{P(V, Q)} : P(V, Q) \longrightarrow \text{O}(V, Q)$$

jest epimorfizmem grup.

Dowód: Prosta konsekwencja definicji grupy $P(V, Q)$, tożsamość (3) oraz Tw. 2. □

Zanim postąpimy dalej w naszej eksploracji anatomii grupy Clifforda, wprowadzimy pewne wyróżnione jej podgrupy o absolutnie kluczowym znaczeniu dla naszych rozważań, zorientowanych ostatecznie na konstrukcję spinorów – ich obecność pozwoli wysubtelnić relację (2), dając nam wgląd w strukturę $\text{Ker } \widetilde{\text{Ad}} \upharpoonright_{P(V, Q)}$. Oto więc mamy

Definicja 4. Przyjmijmy zapis dotychczasowy. **Grupa** Pin przestrzeni kwadratowej (V, Q) to podgrupa

$$\text{Pin}(V, Q) := \langle v \in Q^{-1}(\{-1_{\mathbb{K}}, 1_{\mathbb{K}}\}) \rangle \subset P(V, Q).$$

Jej podgrupa

$$\text{Spin}(V, Q) := \text{Pin}(V, Q) \cap \text{Cliff}(V, Q)^0$$

nosi miano **grupy** Spin przestrzeni kwadratowej (V, Q) .

W kontekście poprzedniego stwierdzenia najogólniejsze własności wprowadzonych tu grup opisuje

Stwierdzenie 7. Przyjmijmy zapis Def. 4 oraz Stw. 3. Podgrupy

$$\widetilde{\text{Ad}}(\text{Pin}(V, Q)), \widetilde{\text{Ad}}(\text{Spin}(V, Q)) \subset \text{O}(V, Q)$$

są normalne.

Dowód: Izometryczne działanie definiujące grupy $\text{O}(V, Q)$ na (V, Q)

$$\rho_{\text{def.}} \equiv \text{id}_{\text{O}(V, Q)} : \text{O}(V, Q) \longrightarrow \text{Aut}_{\mathbb{K}}(V, Q)$$

podnosi się funktorialnie do $\text{Cliff}(V, Q)$ na mocy Tw. 7-8.2, przy czym

$$\forall_{\chi \in \text{O}(V, Q)} : [\text{Cliff}(\chi), J_V] \equiv [\text{Cliff}(\chi), \text{Cliff}(P_V)] = \text{Cliff}([\chi, P_V]) = \text{Cliff}(0) = 0,$$

więc też dla dowolnych $v \in V^\times, w \in V$ oraz $\chi \in \text{O}(V, Q)$ obliczamy

$$\begin{aligned} \widetilde{\text{Ad}}_{\chi(v)}(w) &\equiv J_V \circ \chi(v).w.\chi(v)^{-1} = \chi \circ J_V(v).w.\chi(v^{-1}) = \text{Cliff}(\chi)(J_V(v).\chi^{-1}(w).v^{-1}) \\ &= \text{Cliff}(\chi) \circ \widetilde{\text{Ad}}_v(\chi^{-1}(w)), \end{aligned}$$

korzystając przy tym z tożsamości $\text{Cliff}(\chi)(w) \equiv \chi(w)$ słusznej dla dowolnego $w \in V (\subset \text{Cliff}(V, Q))$, czyli – wobec Równ (3) –

$$\widetilde{\text{Ad}}_{\chi(v)}(w) = \chi \circ \widetilde{\text{Ad}}_v \circ \chi^{-1}(w),$$

co pozwala nam zapisać

$$\widetilde{\text{Ad}}_{\chi(v)} \upharpoonright_V = \text{Ad}_\chi(\widetilde{\text{Ad}}_v \upharpoonright_V)$$

i na tej podstawie – w bezpośrednim odwołaniu do Def. 4 (obie grupy są generowane mnożeniowo przez pewne wektory nieizotropowe w stosownej liczbie (parzystej dla $\text{Spin}(V, Q)$), a nadto $Q \circ \chi(v) = Q(v)$) – ustalić pożądane relacje

$$\forall_{H \in \{\widetilde{\text{Ad}}(\text{Pin}(V, Q)), \widetilde{\text{Ad}}(\text{Spin}(V, Q))\}} \forall_{\chi \in \text{O}(V, Q)} : \text{Ad}_\chi(H) \subset H.$$

□

Tytułem ukierunkowania dalszej dyskusji, która pozwoli należycie wyeksponować grupy Pin i Spin , zauważmy, że dla dowolnego składowego $\lambda \in \mathbb{K}^\times$ (i dla każdego wektora $v \in V^\times$) spełniona jest tożsamość

$$P_{\lambda \triangleright v} = P_v,$$

jeśli zatem bylibyśmy w stanie przenormować wszystkie nieizotropowe wektory w V tak, by ich norma (spinorowa, tj. zadawana przez Q) była równa $\pm 1_{\mathbb{K}}$, to wówczas w świetle Tw. 2 reprezentacja dołączona zwichrowana ograniczałaby się w oczywisty sposób do epimorfizmów $\text{Pin}(V, Q) \twoheadrightarrow \text{O}(V, Q)$ oraz $\text{Spin}(V, Q) \twoheadrightarrow \text{SO}(V, Q)$. Kłopot w tym, że tego typu operacja wymaga rozwiązania (w $\mathbb{K}^\times$) równania

$$\{-1_{\mathbb{K}}, 1_{\mathbb{K}}\} \ni Q(\lambda \triangleright v) = \lambda^2 \cdot_{\mathbb{K}} Q(v)$$

dla dowolnej wartości $Q(v)$, co tłumaczy się na warunek

$$\forall_{v \in V^\times} \exists_{\varepsilon_v \in \{-1_{\mathbb{K}}, 1_{\mathbb{K}}\}} : \frac{\varepsilon_v}{Q(v)} \in \{ \lambda^2 \in \mathbb{K}^\times \mid \lambda \in \mathbb{K}^\times \}.$$

Warunek ten jest (szczęśliwie) spełniony dla $\mathbb{K} \in \{\mathbb{R}, \mathbb{C}\}$, w ogólności zaś dostarcza motywacji dla

Definicja 5. Przyjmijmy zapis dotychczasowy i wprowadźmy oznaczenie

$$(\mathbb{K}^\times)^2 \equiv \{ \lambda^2 \in \mathbb{K}^\times \mid \lambda \in \mathbb{K}^\times \}.$$

Ciało $\mathbb{K}$ o charakterystyce $\text{char}(\mathbb{K}) \neq 2$ nazywamy **spinowym**, jeśli spełnia warunek

$$\mathbb{K}^\times = (\mathbb{K}^\times)^2 \cup -(\mathbb{K}^\times)^2,$$

czyli

$$\forall \lambda \in \mathbb{K}^\times \exists \mu \in \mathbb{K}^\times : \mu^2 \in \{-\lambda, \lambda\}.$$

Zwieńczeniem naszych dociekań jest

Twierdzenie 1. Przyjmijmy zapis dotychczasowy. Dla dowolnej skończonej wymiarowej niezwyrodniałej przestrzeni kwadratowej (V, Q) nad ciałem spinowym $\mathbb{K}$ istnieją krótkie ciągi dokładne grup

$$1 \longrightarrow \mathbb{F} \longrightarrow \text{Pin}(V, Q) \xrightarrow{\widetilde{\text{Ad}}} \text{O}(V, Q) \longrightarrow 1,$$

$$1 \longrightarrow \mathbb{F} \longrightarrow \text{Spin}(V, Q) \xrightarrow{\widetilde{\text{Ad}}} \text{SO}(V, Q) \longrightarrow 1,$$

przy czym

$$\mathbb{F} = \begin{cases} \{-1_{\mathbb{K}}, 1_{\mathbb{K}}\} \equiv \mathbb{Z}_2, & \text{gdzie } \sqrt{-1_{\mathbb{K}}} \notin \mathbb{K} \\ \{-\sqrt{-1_{\mathbb{K}}}, \sqrt{-1_{\mathbb{K}}}, -1_{\mathbb{K}}, 1_{\mathbb{K}}\} \equiv \mathbb{Z}_4 & \text{w p.p.} \end{cases},$$

co pokazuje, że grupy Pin i Spin są centralnymi rozszerzeniami grup – odpowiednio – ortogonalnej i specjalnej ortogonalnej przestrzeni kwadratowej. W szczególności dla $\mathbb{K} = \mathbb{R}$ i dowolnej sygnatury $(p, q) \in \mathbb{N}^{*2} \setminus \{(0, 0)\}$, a przy oznaczeniach $G_{\mathbb{R}}(p, q) \equiv G(\mathbb{R}^{p, q})$, $G \in \{\text{O}, \text{SO}, \text{Pin}, \text{Spin}\}$, otrzymujemy krótkie ciągi dokładne

$$1 \longrightarrow \mathbb{Z}_2 \longrightarrow \text{Pin}_{\mathbb{R}}(p, q) \xrightarrow{\widetilde{\text{Ad}}} \text{O}_{\mathbb{R}}(p, q) \longrightarrow 1,$$

$$1 \longrightarrow \mathbb{Z}_2 \longrightarrow \text{Spin}_{\mathbb{R}}(p, q) \xrightarrow{\widetilde{\text{Ad}}} \text{SO}_{\mathbb{R}}(p, q) \longrightarrow 1,$$

w których – o ile tylko $(p, q) \neq (1, 1)$ – podwójne nakrycia grupy ortogonalnej (wzgl. specjalnej ortogonalnej) przez grupę Pin (wzgl. Spin) są topologicznie nietrywialne nad $\text{O}_{\mathbb{R}}(p, q)$ – i tak np. podwójne nakrycie

$$1 \longrightarrow \mathbb{Z}_2 \longrightarrow \text{Spin}_{\mathbb{R}}(n, 0) \xrightarrow{\widetilde{\text{Ad}}} \text{SO}_{\mathbb{R}}(n, 0) \longrightarrow 1$$

jest uniwersalne dla $n \geq 3$.

Dowód: Rozważmy dowolny element $u = v_1.v_2.\dots.v_m \in \text{Pin}(V, Q)$ z jądra homomorfizmu $\widetilde{\text{Ad}}$ (określony przez wektory $v_i \in Q^{-1}(\{-1_{\mathbb{K}}, 1_{\mathbb{K}}\})$, $i \in \overline{1, m}$). Na mocy Stw. 3

$$u \in \mathbb{K}^\times \triangleright e^C,$$

zatem – w świetle Stw. 4 i samej Def. 3 –

$$\begin{aligned} u.u &= u.J_V \circ T_V(u) \equiv \mathbf{N}(u) = \mathbf{N}(v_1).\mathbf{N}(v_2).\dots.\mathbf{N}(v_m) \\ &= (-Q(v_1)) \cdot_{\mathbb{K}} (-Q(v_2)) \cdot_{\mathbb{K}} \dots \cdot_{\mathbb{K}} (-Q(v_m)) \triangleright e^C \in \mathbb{Z}_2 \triangleright e^C. \end{aligned}$$

Mamy tutaj dwie wykluczające się wzajemnie ewentualności:

- albo $\mathbb{K} \not\equiv \sqrt{-1_{\mathbb{K}}}$, a wtedy nieodzownie $u^2 = e^C$, czyli $u \in \{-e^C, e^C\}$, tj.

$$\text{Ker } \widetilde{\text{Ad}} \subseteq \{-e^C, e^C\},$$

ponieważ jednak

$$\widetilde{\text{Ad}}_{\pm e^C} = \text{id}_V,$$

przeto ostatecznie

$$\text{Ker } \widetilde{\text{Ad}} = \{-e^C, e^C\} \cong \mathbb{Z}_2,$$

- albo $\mathbb{K} \ni \sqrt{-1_{\mathbb{K}}}$, a wtedy oba znaki w powyższej tożsamości są dopuszczalne, więc

$$\text{Ker } \widetilde{\text{Ad}} \subseteq \{-\sqrt{-1_{\mathbb{K}}} \triangleright e^C, \sqrt{-1_{\mathbb{K}}} \triangleright e^C, -e^C, e^C\},$$

a że

$$\forall_{x \in \{-\sqrt{-1_{\mathbb{K}}} \triangleright e^C, \sqrt{-1_{\mathbb{K}}} \triangleright e^C, -e^C, e^C\}} : \widetilde{\text{Ad}}_x = \text{id}_V,$$

toteż

$$\text{Ker } \widetilde{\text{Ad}} = \{-\sqrt{-1_{\mathbb{K}}} \triangleright e^C, \sqrt{-1_{\mathbb{K}}} \triangleright e^C, -e^C, e^C\} \cong \mathbb{Z}_4.$$

Topologia podwójnych nakryć ma swój aspekt elementarny, który odróżnia nakrycia $\mathbb{Z}_2 \rightarrow \widehat{B} \rightarrow B$ postaci $\widehat{B} = B \times \mathbb{Z}_2$ (zwane trywialnymi) od nakryć przyjmujących postać trywialną jedynie lokalnie (nad bazą B), ale też – zwłaszcza w obecnym kontekście – aspekt wyższy, który odróżnia nakrycia n -spójne (dla $n \geq 1$) od nie- n -spójnych. Zanalizowanie tego ostatniego aspektu wymaga znajomości szczegółów topologii wszystkich grup uwikłanych w treść twierdzenia na poziomie wykraczającym dalece poza zakres niniejszego kursu – zainteresowanego Słuchacza odsyłamy do literatury źródłowej, jak choćby do doskonałej monografii S. Helgasona [Hel01]. W niniejszym dowodzie skupimy się natomiast na aspekcie elementarnym, wykazując, że podwójne nakrycia wymienione w drugiej części twierdzenia nie mają *globalnej* struktury podwójnej kopii bazy. W tym celu wystarczy połączyć krzywą ciągłą¹ punkty w przestrzeni totalnej nakrycia, które rzutują się na ten sam punkt bazy, a więc np. $-e^C$ i e^C w $\text{Spin}_{\mathbb{R}}(p, q)$. Jako że $(p, q) \neq (1, 1)$ (na podstawie poczynionego tu założenia), zawsze znajdziemy wektory $e_1, e_2 \in \mathbb{R}^{p+q}$ spełniające układ warunków

$$\delta_E^{(p,q)}(e_1) = \delta_E^{(p,q)}(e_2) =: \varepsilon \in \{-1, 1\} \quad \wedge \quad e_1 \perp_{\delta_E^{(p,q)}} e_2.$$

Dla dowolnej takiej pary wektorów definiujemy krzywą (ciągłą)

$$\gamma : [0, \frac{\pi}{2}] \rightarrow \text{Cl}_{p,q}^{\mathbb{R}0} : t \mapsto (\cos t \triangleright e_1 + \sin t \triangleright e_2) \cdot (\cos t \triangleright e_1 - \sin t \triangleright e_2).$$

Bez trudu wyznaczamy

$$\delta_E^{(p,q)}(\cos t \triangleright e_1 + \sin t \triangleright e_2) = \cos^2 t \cdot \delta_E^{(p,q)}(e_1) + \sin^2 t \cdot \delta_E^{(p,q)}(e_2) + 2 \sin t \cdot \cos t \cdot \Phi_{\delta_E^{(p,q)}}(e_1, e_2) = \varepsilon,$$

$$\delta_E^{(p,q)}(\cos t \triangleright e_1 - \sin t \triangleright e_2) = \cos^2 t \cdot \delta_E^{(p,q)}(e_1) + \sin^2 t \cdot \delta_E^{(p,q)}(e_2) - 2 \sin t \cdot \cos t \cdot \Phi_{\delta_E^{(p,q)}}(e_1, e_2) = \varepsilon$$

i na tej podstawie stwierdzamy, że krzywa leży w grupie Spin ,

$$\gamma([0, \frac{\pi}{2}]) \subset \text{Spin}_{\mathbb{R}}(p, q),$$

będącej przestrzenią totalną rozpatrywanego nakrycia. Przy tym krzywa ta łączy ze sobą oba punkty w $\widetilde{\text{Ad}}^{-1}(\text{id}_V)$,

$$\gamma(0) = e_1 \cdot e_1 = \varepsilon \triangleright e^C, \quad \gamma(\frac{\pi}{2}) = -e_2 \cdot e_2 = -\varepsilon \triangleright e^C,$$

co nie byłoby możliwe w topologii $\text{O}_{\mathbb{R}}(p, q) \times \mathbb{Z}_2$. □

¹Przy naturalnych założeniach dotyczących topologii rozpatrywanych grup, znajdujących potwierdzenie w szczegółowych dociekaniach, np. na podstawie lektury rzeczonyj monografii S. Helgasona.

DODATEK A. UZUPEŁNIENIE Z TEORII PRZESTRZENI KWADRATOWYCH

Dyskusja cliffordowskiej realizacji grupy izometrii bazuje na Twierdzeniu Cartana–Dieudonnégo, zawierającym identyfikację prostego układu generującego tejże grupy. Ażeby je wysławić, będziemy potrzebowali nieco elementarnego substratu z zakresu algebry liniowej.

Zacznijmy *adagio*...

Stwierdzenie 8. Niechaj (V, Q) będzie skończenie wymiarową niezwyrodniałą przestrzenią kwadratową nad ciałem $\mathbb{K}$, a $W \subset V$ – jej dowolną podprzestrzenią o dopełnieniu ortogonalnym $W^{\perp_Q}$. Wówczas zachodzi równość

$$(4) \quad \dim_{\mathbb{K}} V = \dim_{\mathbb{K}} W + \dim_{\mathbb{K}} W^{\perp_Q}$$

oraz tożsamość

$$W = (W^{\perp_Q})^{\perp_Q}.$$

Dowód: Izomorfizm $\mathbb{K}$ -liniowy

$$l_{\Phi_Q} : V \xrightarrow{\cong} V^* : v \mapsto \Phi_Q(v, \cdot),$$

o którego istnieniu przesądza niezwyrodnienie Q , indukuje epimorfizm

$$W|l_{\Phi_Q} : V \rightarrow W^* : v \mapsto \Phi_Q(v, \cdot)|_W,$$

dla którego możemy wypisać standardowy bilans wymiarów:

$$\dim_{\mathbb{K}} V = \dim_{\mathbb{K}} \text{Ker}_W|l_{\Phi_Q} + \dim_{\mathbb{K}} \text{Im}_W|l_{\Phi_Q} = \dim_{\mathbb{K}} W^{\perp_Q} + \dim_{\mathbb{K}} W^* \equiv \dim_{\mathbb{K}} W + \dim_{\mathbb{K}} W^{\perp_Q},$$

dowodzący słuszności postulowanej równości dla $\dim_{\mathbb{K}} W \leq \dim_{\mathbb{K}} V < \infty$. Skoro zaś

$$W \subset (W^{\perp_Q})^{\perp_Q},$$

co oczywiste, to ta sama równość odniesiona do podprzestrzeni $W^{\perp_Q}$ daje

$$\dim_{\mathbb{K}} V = \dim_{\mathbb{K}} W^{\perp_Q} + \dim_{\mathbb{K}} (W^{\perp_Q})^{\perp_Q},$$

więc też po obustronnym skróceniu (skończonego) składnika $\dim_{\mathbb{K}} W^{\perp_Q}$ w obu reprezentacjach $\dim_{\mathbb{K}} V$,

$$\dim_{\mathbb{K}} (W^{\perp_Q})^{\perp_Q} = \dim_{\mathbb{K}} W,$$

a to już implikuje postulowaną tożsamość. □

W następnej kolejności wprowadzimy

Definicja 6. Przyjmijmy zapis dotychczasowy. Przestrzeń wektorowa V nad ciałem $\mathbb{K}$ wyposażona w formę symetryczną γ jest nazywana **płaszczyzną hiperboliczną**, jeżeli jest niezwyrodniała (względem γ), $\dim_{\mathbb{K}} V = 2$ i istnieje w niej niezerowy wektor izotropowy $v \in V$,

$$v \neq 0_V \quad \wedge \quad \gamma(v, v) = 0_{\mathbb{K}}.$$

Przestrzeń hiperboliczna to suma ortogonalna dowolnej rodziny płaszczyzn hiperbolicznych. Forma symetryczna dwuliniowa (wzgl. kwadratowa) na przestrzeni hiperbolicznej jest określana mianem **formy hiperbolicznej**.

Bazę $\{v, w\}$ płaszczyzny hiperbolicznej V spełniającą układ warunków

$$(5) \quad \gamma(v, v) = 0_{\mathbb{K}} = \gamma(w, w) \quad \wedge \quad \gamma(v, w) = 1_{\mathbb{K}}$$

nazywamy **parą hiperboliczną**.

Elementarnego opisu przestrzeni hiperbolicznych dostarcza

Stwierdzenie 9. Przyjmijmy zapis dotychczasowy. Każda płaszczyzna hiperboliczna nad ciałem $\mathbb{K}$ o charakterystyce $\text{char } \mathbb{K} \neq 2$ ma parę hiperboliczną. I odwrotnie, dowolna dwuwymiarowa przestrzeń symetryczna o bazie spełniającej warunki (5) jest płaszczyzną hiperboliczną.

Dowód: Niech V będzie płaszczyzną hiperboliczną o wektorze $v \neq 0_V$ izotropowym względem formy symetrycznej γ . Wobec $\text{Ker } \gamma = \{0_V\}$ musi istnieć wektor $w \notin \langle v \rangle_{\mathbb{K}}$ spełniający warunek

$$\gamma(v, w) \neq 0_{\mathbb{K}},$$

gdyż w przeciwnym razie $v \in \text{Ker } \gamma$. Położywszy

$$u := \lambda \triangleright_V v +_V w, \quad \lambda := -2_{\mathbb{K}}^{-1} \cdot_{\mathbb{K}} \gamma(v, w)^{-1} \cdot_{\mathbb{K}} \gamma(w, w),$$

bez trudu sprawdzamy, że układ $\{\gamma(v, w)^{-1} \triangleright_{\mathbb{K}} u, v\}$ jest parą hiperboliczną dla V , oto bowiem $u \notin \langle v \rangle_{\mathbb{K}}$, a ponadto

$$\begin{aligned} \gamma(\gamma(v, w)^{-1} \triangleright_{\mathbb{K}} u, \gamma(v, w)^{-1} \triangleright_{\mathbb{K}} u) &\equiv (\gamma(v, w)^{-1})^2 \cdot_{\mathbb{K}} \gamma(\lambda \triangleright_V v +_V w, \lambda \triangleright_V v +_V w) \\ &= (\gamma(v, w)^{-1})^2 \cdot_{\mathbb{K}} (\lambda^2 \cdot_{\mathbb{K}} \gamma(v, v) +_{\mathbb{K}} 2\lambda \cdot_{\mathbb{K}} \gamma(v, w) +_{\mathbb{K}} \gamma(w, w)) \\ &= (\gamma(v, w)^{-1})^2 \cdot_{\mathbb{K}} (2\lambda \cdot_{\mathbb{K}} \gamma(v, w) +_{\mathbb{K}} \gamma(w, w)) = 0_{\mathbb{K}} \end{aligned}$$

oraz

$$\begin{aligned} \gamma(\gamma(v, w)^{-1} \triangleright_{\mathbb{K}} u, v) &\equiv \gamma(v, w)^{-1} \cdot_{\mathbb{K}} \gamma(\lambda \triangleright_V v +_V w, v) \\ &= \gamma(v, w)^{-1} \cdot_{\mathbb{K}} (\lambda \cdot_{\mathbb{K}} \gamma(v, v) +_{\mathbb{K}} \gamma(w, v)) = \gamma(v, w)^{-1} \cdot_{\mathbb{K}} \gamma(v, w) = 1_{\mathbb{K}}. \end{aligned}$$

I odwrotnie, jeżeli elementy bazy $\{v, w\}$ przestrzeni V spełniają relacje (5), to γ ma w tej bazie macierz

$$[\gamma]_{\{v, w\}} = \begin{pmatrix} 0_{\mathbb{K}} & 1_{\mathbb{K}} \\ 1_{\mathbb{K}} & 0_{\mathbb{K}} \end{pmatrix},$$

która jest odwracalna ($\det_{(2)}([\gamma]_{\{v, w\}}) = -1_{\mathbb{K}} \neq 0_{\mathbb{K}}$). To pokazuje, że forma γ jest niezwyrodniała (wszak zerowość wyznacznika formy kwadratowej jest *niezmiennikiem* wyboru bazy), a zatem V jest w istocie płaszczyzną hiperboliczną. $\square$

Przestrzenie hiperboliczne są naturalnym elementem opisu zwyrodniałych ograniczeń niezwyrodniałych form symetrycznych (lub – równoważnie – niezwyrodniałych rozszerzeń zwyrodniałych form symetrycznych) i jako takie pozwalają lepiej zrozumieć geometrię zanurzeń podprzestrzeni izotropowych w niezwyrodniałych przestrzeniach symetrycznych. Doskonałej ilustracji tej tezy dostarcza

Stwierdzenie 10 (O rozszerzeniu hiperbolicznym przestrzeni izotropowej). Przyjmijmy zapis dotychczasowy. Niechaj V będzie skończenie wymiarową przestrzenią wektorową nad ciałem $\mathbb{K}$ o charakterystyce $\text{char } \mathbb{K} \neq 2$ wyposażoną w niezwyrodniałą formę symetryczną γ i niech $W \subset V$ będzie podprzestrzenią V . Oznaczmy $\gamma_W := \gamma|_{W \times W}$ i $W_0 := \text{Ker } \gamma_W$. Wybierzmy w W_0 bazę $\{w_i\}_{i \in \overline{1, K}}$, $K \equiv \dim_{\mathbb{K}} W_0$ i niech D będzie dopełnieniem ortogonalnym W_0 w W ,

$$W = W_0 \oplus D.$$

Wówczas istnieją wektory $\{v_i\}_{i \in \overline{1, K}} \subset D^{\perp} \subset V$ o tej własności, że dla każdego indeksu $i \in \overline{1, K}$ układ $\{v_i, w_i\}$ jest parą hiperboliczną rozpinającą płaszczyznę hiperboliczną $H_i := \langle v_i, w_i \rangle_{\mathbb{K}}$. Pary te zadają rozkład ortogonalny podprzestrzeni

$$\widetilde{W} := \langle v_1, v_2, \dots, v_K, w_1, w_2, \dots, w_K \rangle_{\mathbb{K}} +_V D \subset V$$

dany wzorem

$$\widetilde{W} = H_1 \oplus H_2 \oplus \dots \oplus H_K \oplus D.$$

Dowód: Oznaczmy

$$D_1 := \langle w_2, w_3, \dots, w_K \rangle_{\mathbb{K}} \oplus D.$$

Z inkluzji właściwej podprzestrzeni

$$D_1 \not\subset W_0 \oplus D$$

wynika inkluzja właściwa ich anihilatorów,

$$D_1^{\perp\gamma} \not\subset (W_0 \oplus D)^{\perp\gamma},$$

a stąd dalej – istnienie wektora $v_1 \in D_1^{\perp\gamma} \setminus (W_0 \oplus D)^{\perp\gamma}$, tj. takiego, który spełnia warunki

$$\forall_{j \in \overline{2, K}} : \gamma(v_1, w_j) = 0_{\mathbb{K}}, \quad \forall_{v \in D} : \gamma(v_1, v) = 0_{\mathbb{K}}$$

oraz

$$\gamma(v_1, w_1) \neq 0_{\mathbb{K}}.$$

Wobec izotropowości wektora w_1 i niezwyrodnienia $\gamma|_{\langle v_1, w_1 \rangle_{\mathbb{K}}}$ podprzestrzeń $H_1 := \langle v_1, w_1 \rangle_{\mathbb{K}}$ jest zatem płaszczyzną hiperboliczną, co w świetle Stw. 9 oznacza istnienie takiego wektora $u_1 \in H_1$, który tworzy wraz z w_1 parę hiperboliczną. Przy tym oczywiście

$$\widetilde{W}_1 := \langle v_1, w_1, w_2, w_3, \dots, w_K \rangle_{\mathbb{K}} +_V D = H_1 \oplus \langle w_2, w_3, \dots, w_K \rangle_{\mathbb{K}} \oplus D.$$

Jądrem formy symetrycznej $\gamma|_{\widetilde{W}_1 \times \widetilde{W}_1}$ jest

$$\text{Ker } \gamma|_{\widetilde{W}_1} = \langle w_2, w_3, \dots, w_K \rangle_{\mathbb{K}},$$

wobec czego całą opisaną procedurę możemy powtórzyć w odniesieniu do podprzestrzeni izotropowej

$$\widetilde{W}_1 = \text{Ker } \gamma|_{\widetilde{W}_1} \oplus \widetilde{D}_1, \quad \widetilde{D}_1 := H_1 \oplus D.$$

□

Oczywistą konsekwencją powyższego jest

Corollarium 1. Przyjmijmy zapis dotychczasowy, przy czym zakładamy, że V jest niezwyrodniałą przestrzenią kwadratową wymiaru $n \in \mathbb{N}$ nad ciałem $\mathbb{K}$ o charakterystyce $\text{char } \mathbb{K} \neq 2$, $W \subset V$ zaś – jej podprzestrzenią Q -zerową. Wówczas

$$\dim_{\mathbb{K}} W \leq E\left(\frac{n}{2}\right).$$

Podprzestrzeń, której wymiar wysyca powyższą nierówność, określamy mianem **maksymalnej podprzestrzeni Q -zerowej**.

Mamy także wygodne

Stwierdzenie 11. Przyjmijmy zapis dotychczasowy, przy czym zakładamy, że (V, Q) jest (skończenie wymiarową) **przestrzenią hiperboliczną**, tj. sumą (Q) -ortogonalną skończonej liczby płaszczyzn hiperbolicznych. Niechaj $W_0 \subset V$ będzie dowolną maksymalną podprzestrzenią Q -zerową. Każda izometria $\chi \in O(V, Q)$ o własności

$$\chi|_{W_0} = \text{id}_{W_0}$$

jest obrotem,

$$\chi \in \text{SO}(V, Q).$$

Dowód: Utożsamiając podprzestrzeń W_0 z treści Stw. 10 z maksymalną podprzestrzenią Q -zerową, o której mowa w treści twierdzenia dowodzonego, stwierdzamy istnienie (maksymalnie) Q -zerowego dopełnienia *prostego* $\Delta \subset V$ tejże podprzestrzeni,

$$W_0 \oplus \Delta = V,$$

rozpiętego na wektorach v_i , $i \in \overline{1, 2n}$, $2n \equiv \dim_{\mathbb{K}} V$ dopełniających elementy (dowolnej) bazy $\{w_i\}_{i \in \overline{1, n}}$ do odnośnych par hiperbolicznych $\{w_i, v_i\}$. Istotnie, wektory v_i są parami ortogonalne (wszak $H_i \perp H_j$, ilekroć $i \neq j$), a do tego – izotropowe. Wprost na mocy założenia zachodzi przy tym

$$\forall_{i \in \overline{1, n}} : \chi(w_i) = w_i,$$

a zatem także – dla dowolnych $(w, x) \in W_0 \times \Delta$ –

$$\begin{aligned} \Phi_Q(w, \chi(x) - x) &= \Phi_Q(w, \chi(x)) - \Phi_Q(w, x) \equiv \Phi_Q(\chi(w), \chi(x)) - \Phi_Q(w, x) = \Phi_Q(w, x) - \Phi_Q(w, x) \\ &= \mathbf{0}_{\mathbb{K}}, \end{aligned}$$

skąd wniosek:

$$\forall_{x \in \Delta} : \chi(x) - x \in W_0^{\perp Q}.$$

Q -zerowość W_0 implikuje inkluzję

$$W_0 \subseteq W_0^{\perp Q},$$

która w konsekwencji równości (4), zapisanej dla W_0 :

$$\dim_{\mathbb{K}} V = \dim_{\mathbb{K}} W_0 + \dim_{\mathbb{K}} W_0^{\perp Q},$$

a słusznej wobec niezwyrodnienia V , oraz w konsekwencji maksymalności W_0 , pociągających za sobą równość

$$\dim_{\mathbb{K}} W_0^{\perp Q} = \dim_{\mathbb{K}} V - \dim_{\mathbb{K}} W_0 = 2n - n = n \equiv \dim_{\mathbb{K}} W_0,$$

sprowadza się do tożsamość

$$W_0 = W_0^{\perp Q}.$$

Możemy zatem przepisać wcześniejszy wniosek w postaci

$$\forall_{x \in \Delta} : \chi(x) - x \in W_0,$$

otrzymując tym sposobem w szczególności relacje

$$\forall_{i \in \overline{1, n}} \exists_{\mu_i^1, \mu_i^2, \dots, \mu_i^n \in \mathbb{K}} : \chi(v_i) = v_i +_V \mu_i^j \triangleright w_j,$$

które przesądzą o górnotrójkątnej postaci macierzy endomorfizmu χ względem bazy $\mathcal{B} := \{w_1, w_2, \dots, w_n, v_1, v_2, \dots, v_n\}$

$$[\chi]_{\mathcal{B}}^{\mathcal{B}} = \begin{pmatrix} \mathbf{1}_n & (\mu_i^j)_{i \in \overline{1, n}}^{j \in \overline{1, n}} \\ \mathbf{0}_n & \mathbf{1}_n \end{pmatrix}.$$

Licząc wyznacznik χ w tej właśnie bazie, otrzymujemy pożądaný wynik

$$\det \chi = \det_{(2n)} [\chi]_{\mathcal{B}}^{\mathcal{B}} = \mathbf{1}_{\mathbb{K}}.$$

□

Konsekwencją dużo mniej oczywistą, a fundamentalną dla teorii przestrzeni kwadratowych i teorii spinorów, jest²

Twierdzenie 2 (Cartana–Dieudonnégo). Przyjmijmy zapis dotychczasowy, zakładając dodatkowo, że przestrzeń kwadratowa (V, Q) nad ciałem $\mathbb{K}$ o charakterystyce $\text{char } \mathbb{K} \neq 2$ jest skończonego wymiaru, $N := \dim_{\mathbb{K}} V < \infty$, i jest niezwyrodniała. Wówczas

$$\forall_{\chi \in \text{O}(V, Q)} \exists_{n \in \overline{0, N}} \exists_{v_1, v_2, \dots, v_n \in V^\times} : \chi = P_{v_1} \circ P_{v_2} \circ \dots \circ P_{v_n},$$

przy czym

$$\chi \in \text{SO}(V, Q) \iff n \in 2\mathbb{N}.$$

²W swojej wersji pierwotnej, sformułowanej i udowodnionej dla $\mathbb{K} \in \{\mathbb{R}, \mathbb{C}\}$, twierdzenie to pochodzi od É.J. Cartana [Car38a, Car38b]. Jego wersję ogólną oraz takż dowód podał następnie J.A. Dieudonné w monografii [Die55]. Dowód przedstawiony w niniejszym skrypcie pochodzi zasadniczo od E. Artina [Art61].

Dowód: Zauważmy na wstępie, że dowolne odbicie elementarne P_v przyjmuje względem rozkładu Q -ortogonalnego

$$V = \langle v \rangle_{\mathbb{K}} \oplus_Q \langle v \rangle_{\mathbb{K}}^{\perp_Q}$$

postać

$$P_v = (-\text{id}_{\langle v \rangle_{\mathbb{K}}}) \oplus_Q \text{id}_{\langle v \rangle_{\mathbb{K}}^{\perp_Q}},$$

zatem spełnia warunek

$$\det P_v = -\mathbf{1}_{\mathbb{K}}.$$

Przy założeniu słuszności pierwszej części tezy dowodzonego twierdzenia obserwacja ta implikuje natychmiast drugą jej część.

Dowód części pierwszej przeprowadzimy metodą indukcji silnej względem N , zauważając natychmiast oczywistość tezy w przypadku $N = 1$ – istotnie, jeśli $V = \langle v \rangle_{\mathbb{K}}$, to mamy koniecznie $\chi(v) = \lambda \triangleright v$ dla pewnego $\lambda \in \mathbb{K}^\times$, a przy tym $\mathbf{0}_{\mathbb{K}} \neq Q(v) = Q \circ \chi(v) = \lambda^2 \cdot_{\mathbb{K}} Q(v)$ implikuje $\lambda \in \{-\mathbf{1}_{\mathbb{K}}, \mathbf{1}_{\mathbb{K}}\}$, więc albo $\chi = \text{id}_V \in \text{SO}(V, Q)$, albo $\chi = P_v \in \text{O}(V, Q) \setminus \text{SO}(V, Q)$. Poczyniwszy założenie indukcyjne o słuszności dowodzonej tezy dla $N \in \overline{1, N_0 - 1}$, $N_0 > 1$, rozbijemy następnie dowód jej słuszności dla $N = N_0$ na składowe odpowiadające wykluczającym się nawzajem ewentualnościom:

- (i) $\exists_{v \in V^\times} : \chi(v) = v$;
- (ii) $\forall_{v \in V^\times} : \chi(v) - v \neq \mathbf{0}_V \quad \wedge \quad \exists_{w \in V^\times} : \chi(w) - w \in V^\times$;
- (iii) $\forall_{v \in V^\times} : \chi(v) - v \in Q^{-1}(\{0_{\mathbb{K}}\}) \setminus \{0_V\}$.

W przypadku (i) przywołujemy tezę twierdzenia o istnieniu dopełnienia ortogonalnego dowolnej skończonej wymiarowej podprzestrzeni niezwyrodniałej w przestrzeni kwadratowej (na co pozwala nieizotropowość v) i dokonujemy rozkładu

$$(6) \quad V = \langle v \rangle_{\mathbb{K}} \oplus_Q \langle v \rangle_{\mathbb{K}}^{\perp_Q},$$

odnotowując przy tym rozkład rozważanego endomorfizmu

$$\chi = \text{id}_{\langle v \rangle_{\mathbb{K}}} \oplus \chi|_{\langle v \rangle_{\mathbb{K}}^{\perp_Q}}.$$

Wobec równości

$$\dim_{\mathbb{K}} \langle v \rangle_{\mathbb{K}}^{\perp_Q} = N_0 - 1$$

założenie indukcyjne pozwala nam rozłożyć $\chi|_{\langle v \rangle_{\mathbb{K}}^{\perp_Q}}$ na co najwyżej $N_0 - 1$ odbić elementarnych $P_x^{(N_0-1)}$ w hiperpłaszczyznach $Q|_{\langle v \rangle_{\mathbb{K}}^{\perp_Q}} \equiv Q_{\perp v}$ -ortogonalnych do wyróżnionych wektorów nieizotropowych $x \in \langle v \rangle_{\mathbb{K}}^{\perp_Q}$. Odbicie $P_x^{(N_0-1)}$ przyjmuje względem odnośnego rozkładu

$$\langle v \rangle_{\mathbb{K}}^{\perp_Q} = \langle x \rangle_{\mathbb{K}} \oplus_{Q_{\perp v}} \langle x \rangle_{\mathbb{K}}^{\perp_{Q_{\perp v}}},$$

postać blokowo-diagonalną

$$P_x^{(N_0-1)} = (-\text{id}_{\langle x \rangle_{\mathbb{K}}}) \oplus \text{id}_{\langle x \rangle_{\mathbb{K}}^{\perp_{Q_{\perp v}}}},$$

a jego trywialne rozszerzenie do całej przestrzeni V zapisuje się względem rozkładu (6) jako

$$\text{id}_{\langle v \rangle_{\mathbb{K}}} \oplus P_x^{(N_0-1)} \equiv \text{id}_{\langle v \rangle_{\mathbb{K}}} \oplus (-\text{id}_{\langle x \rangle_{\mathbb{K}}}) \oplus \text{id}_{\langle x \rangle_{\mathbb{K}}^{\perp_{Q_{\perp v}}}},$$

czyli sprowadza się do odbicia elementarnego w hiperpłaszczyźnie

$$\langle x \rangle_{\mathbb{K}}^{\perp_Q} = \langle v \rangle_{\mathbb{K}} \oplus_Q \langle x \rangle_{\mathbb{K}}^{\perp_{Q_{\perp v}}},$$

tj. spełnia tożsamość

$$\text{id}_{\langle v \rangle_{\mathbb{K}}} \oplus P_x^{(N_0-1)} \equiv P_x.$$

To rozumowanie przekonuje, że w przypadku (i) izometria χ rozkłada się na co najwyżej $N_0 - 1 < N_0$ odbić elementarnych.

W przypadku (ii) nieizotropowość $\chi(w) - w$ pozwala rozpatrzyć odbicie elementarne $P_{\chi(w)-w}$, które spełnia tożsamość

$$\begin{aligned} P_{\chi(w)-w}(w) &= w - \frac{2\Phi_Q(w, \chi(w)-w)}{Q(\chi(w)-w)} \triangleright (\chi(w) - w) \\ &= w - \frac{2(\Phi_Q(w, \chi(w)) - Q(w))}{2Q(w) - 2\Phi_Q(w, \chi(w))} \triangleright (\chi(w) - w) = \chi(w). \end{aligned}$$

Na jej podstawie stwierdzamy, że izometria $\chi_w := P_{\chi(w)-w} \circ \chi$ zachowuje nieizotropowy wektor w , co w świetle poprzedniej części naszego dowodu oznacza, że χ_w jest superpozycją co najwyżej $N_0 - 1$ odbić elementarnych. Co za tym idzie, wyjściowa izometria χ jest superpozycją co najwyżej $N_0 - 1 + 1 = N_0$ odbić elementarnych, zgodnie z tezą indukcyjną.

W przypadku (iii) przekonujemy się, że $N_0 = 2k$, $k \in \mathbb{N}_{\geq 2}$ oraz że $\chi \in \text{SO}(V, Q)$, co pozwala przeprowadzić następujące proste rozumowanie. Nieizotropowość wektora $v \in V^\times$ (wybranego dowolnie, a wybór taki istnieje z racji niezwyrodnienia V i założenia dotyczącego charakterystyki $\mathbb{K}$ – gdyby było $Q(v) = 0$ dla dowolnego $v \in V$, to mielibyśmy $\Phi_Q \equiv 0$ (formuła polaryzacyjna), więc sprzeczność) oznacza, że przynajmniej jeden z wektorów $\chi(v) \pm v$ jest nieizotropowy, gdyż

$$\begin{aligned} Q(\chi(v) + v) &= \mathbf{0}_{\mathbb{K}} = Q(\chi(v) - v) \\ \iff Q(v) + \Phi_Q(v, \chi(v)) &= \mathbf{0}_{\mathbb{K}} = Q(v) - \Phi_Q(v, \chi(v)) \implies Q(v) = \mathbf{0}_{\mathbb{K}}. \end{aligned}$$

W rozważanym przypadku (iii) oznacza to niechybnie $\chi(v) +_V v \in V^\times$ przy jednoczesnym $\chi(v) \neq v$, ponieważ zaś

$$P_{\chi(v)+_V v}(v) = v - \frac{2\Phi_Q(v, \chi(v)+_V v)}{Q(\chi(v)+_V v)} \triangleright (\chi(v) +_V v) = -\chi(v),$$

przeto $\tilde{\chi}_v := P_v \circ P_{\chi(v)+_V v} \circ \chi$ zachowuje wektor nieizotropowy v . Rozumując jak poprzednio, stwierdzamy, że $\tilde{\chi}_v$ jest superpozycją co najwyżej $N_0 - 1$ odbić elementarnych, przy czym jeśli przyjąć, że – jak pokażemy lada chwila – N_0 jest liczbą parzystą, a χ jest obrotem, to mamy do czynienia z sytuacją, w której *obrót* $\tilde{\chi}_v$ rozkłada się na co najwyżej $N_0 - 1 \in 2\mathbb{N} + 1$ odbić (elementarnych), czyli koniecznie w rozkładzie tym jest co najwyżej $N_0 - 2 \in 2\mathbb{N}$ czynników, to zaś – koniec końców – prowadzi do wniosku, że χ jest superpozycją co najwyżej $N_0 - 2 + 2 = N_0$ odbić elementarnych, w zgodzie z tezą indukcyjną. Pozostaje przeto dowieść, że endomorfizm ten przy spełnionych warunkach z punktu (iii) jest w istocie obrotem w parzystowymiarowej (niezwyrodniałej) przestrzeni kwadratowej. Jasno widać, że $N_0 \neq 1$, oto bowiem w przypadku $N_0 = 1$ jest – jak pokazaliśmy wcześniej – $\chi(v) = -v$ (wszak $\chi(v) \neq v$), więc też $\chi(v) - v = -2_{\mathbb{K}} \triangleright v \in V^\times$, przeciwnie do założenia (iii). Gdyby natomiast było $N_0 = 2$, to wówczas mielibyśmy $\chi(v) \notin \langle v \rangle_{\mathbb{K}}$, gdyż w przeciwnym przypadku byłoby albo $\chi(v) = v$, niezgodnie z założeniem (iii), albo $\chi(v) = -v$ ($\{-1_{\mathbb{K}}, 1_{\mathbb{K}}\}$ to jedyne dopuszczalne wartości własne izometrii χ), a wtedy $\chi(v) - v = -2_{\mathbb{K}} \triangleright v \in V^\times$, również w sprzeczności z założeniem (iii). Skoro jednak $\chi(v) \notin \langle v \rangle_{\mathbb{K}}$, to $\{v, \chi(v)\}$ jest bazą V , w której wprost na mocy założenia (iii) znika gramian

$$\begin{aligned} \det_{(2)} \begin{pmatrix} \Phi_Q(v, v) & \Phi_Q(v, \chi(v)) \\ \Phi_Q(v, \chi(v)) & \Phi_Q(\chi(v), \chi(v)) \end{pmatrix} &= \det_{(2)} \begin{pmatrix} Q(v) & \Phi_Q(v, \chi(v)) \\ \Phi_Q(v, \chi(v)) & Q(v) \end{pmatrix} \\ &= Q(v)^2 - \Phi_Q(v, \chi(v))^2 = 2_{\mathbb{K}}^{-1} \cdot_{\mathbb{K}} Q(\chi(v) - v) \cdot_{\mathbb{K}} 2_{\mathbb{K}}^{-1} \cdot_{\mathbb{K}} Q(\chi(v) +_V v) = \mathbf{0}_{\mathbb{K}}, \end{aligned}$$

co oznacza, że V jest zwyrodniała, wbrew założeniu. Ostatecznie więc $N_0 \geq 3$. Rozważmy następnie izotropowy wektor $w \in V \setminus \{0_V\}$ (taki istnieje, gdyż istnieją $v \in V^\times$ (wszak Q jest niezwyrodniała), a dla takich $\chi(v) - v \in Q^{-1}(\{0_{\mathbb{K}}\})$ jest niezerowy). W świetle Stw. 10 istnieje niepuste (wszak $\dim_{\mathbb{K}} V > 2$) dopełnienie Q -ortogonalne płaszczyzny hiperbolicznej zawierającej w , czyli też – wektor nieizotropowy $v \in V^\times$ o własności $\Phi_Q(v, w) = \mathbf{0}_{\mathbb{K}}$. Istotnie, $W_0 \equiv W := \langle w \rangle_{\mathbb{K}}$ uzupełnia się – wedle schematu ze Stw. 10 – do płaszczyzny hiperbolicznej H , której niezwyrodnienie jako podprzestrzeni V przesądza o istnieniu rozkładu ortogonalnego $V = H \oplus_Q H^{\perp_Q}$ z $H^{\perp_Q} \neq \{0_V\}$, a to z racji nierówności $\dim_{\mathbb{K}} V - \dim_{\mathbb{K}} H = N_0 - 2 \geq 1$. Przy tym $H^{\perp_Q}$ jest koniecznie niezwyrodniała, gdyby bowiem taką nie była, to wobec ortogonalności rozkładu $V = H \oplus_Q H^{\perp_Q}$ sama nadprzestrzeń V byłaby zwyrodniała, niezgodnie z założeniem. Nieizotropowy wektor v

wyberamy z $H^{\perp Q}$ właśnie, a wówczas dla dowolnego $\varepsilon \in \mathbb{K}^\times$ otrzymujemy wektor nieizotropowy $w +_V \varepsilon \triangleright v \in V$,

$$\begin{aligned} Q(w +_V \varepsilon \triangleright v) &\equiv \Phi_Q(w +_V \varepsilon \triangleright v, w +_V \varepsilon \triangleright v) = Q(w) + \varepsilon^2 \cdot_{\mathbb{K}} Q(v) + 2_{\mathbb{K}} \cdot_{\mathbb{K}} \varepsilon \cdot_{\mathbb{K}} \Phi_Q(v, w) \\ &= \varepsilon^2 \cdot_{\mathbb{K}} Q(v) \neq \mathbf{0}_{\mathbb{K}}, \end{aligned}$$

a zatem także – wprost na mocy (iii) – niezerowy wektor izotropowy $\chi(w +_V \varepsilon \triangleright v) - (w +_V \varepsilon \triangleright v)$, przy czym warunek izotropowości tego ostatniego daje nam – w połączeniu z warunkiem izotropowości $\chi(v) - v$ – tożsamość

$$\begin{aligned} \mathbf{0}_{\mathbb{K}} &= Q(\chi(w +_V \varepsilon \triangleright v) - (w +_V \varepsilon \triangleright v)) \equiv Q(\chi(w) - w +_V \varepsilon \triangleright (\chi(v) - v)) \\ &= Q(\chi(w) - w) +_{\mathbb{K}} \varepsilon^2 \cdot_{\mathbb{K}} Q(\chi(v) - v) +_{\mathbb{K}} 2_{\mathbb{K}} \cdot_{\mathbb{K}} \varepsilon \cdot_{\mathbb{K}} \Phi_Q(\chi(w) - w, \chi(v) - v) \\ &= Q(\chi(w) - w) +_{\mathbb{K}} 2_{\mathbb{K}} \cdot_{\mathbb{K}} \varepsilon \cdot_{\mathbb{K}} \Phi_Q(\chi(w) - w, \chi(v) - v). \end{aligned}$$

Dodając do siebie obie strony powyższej równości dla $\varepsilon = -\mathbf{1}_{\mathbb{K}}$ i $\varepsilon = \mathbf{1}_{\mathbb{K}}$, otrzymujemy równość

$$Q(\chi(w) - w) = \mathbf{0}_{\mathbb{K}},$$

śluszną dla *dowolnego* izotropowego wektora w . W konkluzji możemy zapisać, w rozpatrywanym tu przypadku (iii),

$$V_1 := \text{Image}(\chi - \text{id}_V) \subset Q^{-1}(\{\mathbf{0}_{\mathbb{K}}\})$$

(zawieranie to jest implikowane wprost przez (iii) dla wektorów nieizotropowych, a dla izotropowych zostało właśnie udowodnione). Zauważmy przy tym, że

$$(7) \quad \forall_{x, y \in V_1} : \Phi_Q(x, y) = 2_{\mathbb{K}}^{-1} \cdot_{\mathbb{K}} (Q(x +_V y) - Q(x) - Q(y)) = \mathbf{0}_{\mathbb{K}}$$

(wszak V_1 jest podgrupą), więc V_1 jest podprzestrzenią Q -zerową,

$$V_1 \subseteq V_1^{\perp Q}.$$

Wybermy $x \in V$ oraz $y \in V_1^{\perp Q}$, a wtedy – w świetle powyższego –

$$\begin{aligned} \Phi_Q(x, \chi(y) - y) &= \Phi_Q(\chi(x), \chi(y) - y) - \Phi_Q(\chi(x) - x, \chi(y) - y) = \Phi_Q(\chi(x), \chi(y) - y) \\ &= \Phi_Q(x, y) - \Phi_Q(\chi(x), y) = -\Phi_Q(\chi(x) - x, y) = \mathbf{0}_{\mathbb{K}}, \end{aligned}$$

czyli – wobec dowolności x i niezwyrodnienia Q –

$$\forall_{y \in V_1^{\perp Q}} : \chi(y) - y = 0_V,$$

to zaś oznacza, że

$$(8) \quad \chi|_{V_1^{\perp Q}} = \text{id}_{V_1^{\perp Q}}.$$

Przywołując raz jeszcze warunek (iii), konkludujemy, że

$$V_1^{\perp Q} \subset Q^{-1}(\{\mathbf{0}_{\mathbb{K}}\})$$

(gdyby $v \in V_1^{\perp Q}$ miał $Q(v) \neq 0$, czyli $v \in V^\times$, to wówczas byłoby także $v - v \equiv \chi(v) - v \neq 0_V$, więc sprzeczność), a że $V_1^{\perp Q} \subset V$ jest podgrupą, przeto jest też automatycznie podprzestrzenią Q -zerową (por. (7)), a zatem

$$V_1^{\perp Q} \subseteq (V_1^{\perp Q})^{\perp Q}.$$

Ostatecznie otrzymujemy ciąg inkluzji

$$V_1 \subseteq V_1^{\perp Q} \subseteq (V_1^{\perp Q})^{\perp Q} = V_1,$$

w którym ostatnia równość wynika wprost ze Stw. 8. Powyższe implikuje już wprost równość

$$V_1 = V_1^{\perp Q},$$

a z niej – na mocy Równ. (4) – parzystość $N_0 \equiv \dim_{\mathbb{K}} V$. Obecność w niej podprzestrzeni Q -zerowej V_1 wymiaru $\dim_{\mathbb{K}} V_1 = \frac{N_0}{2}$, czyli maksymalnego, przesądza – w świetle Stw. 10 – o hiperboliczności

V i tym samym pozwala nam odnieść tezę Stw. 11 do izometrii χ o własności (8), tj. ograniczającej się trywialnie do maksymalnej podprzestrzeni Q -zerowej $V_1^{\perp Q} = V_1$. Tym sposobem wnioskujemy, że χ jest w istocie obrotem, co kończy dowód. $\square$

LITERATURA

- [Art61] E. Artin, *Geometric algebra*, Interscience Tracts in Pure and Applied Mathematics, vol. 3, Interscience Publishers, 1961.
- [Car38a] É. J. Cartan, *Leçons sur la théorie des spineurs. I : les spineurs de l'espace à trois dimensions*, Actualités scientifiques et industrielles: Exposés de géométrie, vol. 9, Hermann & cie., 1938.
- [Car38b] ———, *Leçons sur la théorie des spineurs. II : les spineurs de l'espace à $n > 3$ dimensions, les spineurs en géométrie riemannienne*, Actualités scientifiques et industrielles: Exposés de géométrie, vol. 11, Hermann & cie., 1938.
- [Die55] J.A. Dieudonné, *La géométrie des groupes classiques*, Ergebnisse der Mathematik und ihrer Grenzgebiete, vol. 5, Springer, 1955.
- [Hel01] S. Helgason, *Differential Geometry, Lie Groups, and Symmetric Spaces*, Graduate studies in mathematics, vol. 34, American Mathematical Society, 2001.