

Bilinear forms

A. Eremenko

September 1, 2024

What happens if we drop the positivity condition in the definition of a dot product?

1. A map which assigns a number $B(\mathbf{x}, \mathbf{y})$ to any pair of vectors is called a *bilinear form* if it is linear with respect to each argument:

$$B(\alpha\mathbf{x} + \beta\mathbf{y}, \mathbf{z}) = \alpha B(\mathbf{x}, \mathbf{z}) + \beta B(\mathbf{y}, \mathbf{z}),$$

$$B(\mathbf{x}, \alpha\mathbf{y} + \beta\mathbf{z}) = \alpha B(\mathbf{x}, \mathbf{y}) + \beta B(\mathbf{x}, \mathbf{z}).$$

A bilinear form is called symmetric if

$$B(\mathbf{x}, \mathbf{y}) = B(\mathbf{y}, \mathbf{x})$$

for all $\mathbf{x}, \mathbf{y}$.

For example, any real scalar product is a symmetric bilinear form. If we put $\mathbf{x} = \mathbf{y}$ in a bilinear form we obtain a function $Q(\mathbf{x}) = B(\mathbf{x}, \mathbf{x})$ which maps vectors to numbers. This function is called a *quadratic form* corresponding to B . A symmetric bilinear form can be recovered from its quadratic form by the formula

$$B(\mathbf{x}, \mathbf{y}) = \frac{1}{4} (Q(\mathbf{x} + \mathbf{y}) - Q(\mathbf{x} - \mathbf{y})). \quad (1)$$

Indeed,

$$Q(\mathbf{x} + \mathbf{y}) = B(\mathbf{x} + \mathbf{y}, \mathbf{x} + \mathbf{y}) = B(\mathbf{x}, \mathbf{x}) + 2B(\mathbf{x}, \mathbf{y}) + B(\mathbf{y}, \mathbf{y}), \quad (2)$$

and

$$Q(\mathbf{x} - \mathbf{y}) = B(\mathbf{x} - \mathbf{y}, \mathbf{x} - \mathbf{y}) = B(\mathbf{x}, \mathbf{x}) - 2B(\mathbf{x}, \mathbf{y}) + B(\mathbf{y}, \mathbf{y}). \quad (3)$$

By subtracting and dividing by 4 we obtain (1).

Exercise 1. If we add (2) and (3), we obtain the “parallelogram theorem”, a. k. a. *polarization identity*:

$$Q(\mathbf{x} + \mathbf{y}) + Q(\mathbf{x} - \mathbf{y}) = 2(Q(\mathbf{x}, \mathbf{x}) + Q(\mathbf{y}, \mathbf{y})). \quad (4)$$

Prove that if a continuous function satisfies $Q(\mathbf{x}) > 0$, $\mathbf{x} \neq 0$, and $Q(-\mathbf{x}) = Q(\mathbf{x})$ and (4), then formula (1) defines a bilinear form. Thus all such functions are quadratic forms. This is a somewhat difficult exercise. For hints and a solution, see [1] or [2].

Suppose that we have a finite basis $\mathbf{v}_1, \dots, \mathbf{v}_n$. Then to each bilinear form we can associate a square matrix

$$A = (B(\mathbf{v}_i, \mathbf{v}_j))_{i,j=1}^n.$$

It is called the *Gram matrix* of the bilinear form. The Gram matrix is symmetric if and only if the form is symmetric.

From now on we only consider symmetric bilinear forms on real vector spaces.

This matrix defines the form completely: if

$$\mathbf{x} = \sum x_j \mathbf{v}_j, \quad \mathbf{y} = \sum y_j \mathbf{v}_j,$$

then

$$B(\mathbf{x}, \mathbf{y}) = \sum_{i,j} x_i y_j B(\mathbf{v}_i, \mathbf{v}_j) = \mathbf{x}^T A \mathbf{y},$$

and

$$Q(\mathbf{x}) = \mathbf{x}^T A \mathbf{x}.$$

For example, when $v_1, \dots, v_j$ is the standard basis then x_j are the usual coordinates of a vector $\mathbf{x} \in \mathbf{R}^n$.

A quadratic form is called *positive definite* if $Q(\mathbf{x}) \geq 0$ for all $\mathbf{x}$, and $= 0$ only for $\mathbf{x} = 0$. A symmetric matrix A is called *positive definite* if the quadratic form $Q(\mathbf{x}) = \mathbf{x}^T A \mathbf{x}$ is positive definite. Bilinear forms corresponding to positive definite quadratic forms are exactly the scalar products, as previously defined.

2. What happens to the matrix of the quadratic or bilinear form when we change the basis? Let $\mathbf{x} = (x_j)$ be the column vector, and $\mathbf{x}'$ be the columns of its coordinates in another basis. Then $\mathbf{x} = C\mathbf{x}'$ with a non-singular matrix

C . Let Q be a quadratic form with matrix A in the standard basis and matrix A' in the other basis. Then

$$Q(\mathbf{x}) = \mathbf{x}^T A \mathbf{x} = (C\mathbf{x}')^T A C\mathbf{x}' = \mathbf{x}'^T C^T A C \mathbf{x}',$$

so the matrix of Q in the new basis is $C^T A C$.

Notice the difference between the transformation rules of the matrix of a linear operator and the matrix of a quadratic form!

How can we simplify a matrix of a symmetric quadratic form by choosing an appropriate basis? The matrix of a quadratic form A is symmetric, so by the spectral theorem for symmetric matrices, there is an orthogonal matrix V such that

$$A = V \Lambda V^{-1}.$$

Orthogonal means $V^{-1} = V^T$. Thus taking $C = V^{-1}$ we obtain

$$\Lambda = C^T A C,$$

which means that every quadratic form can be diagonalized by a an appropriate choice of orthonormal basis, in other words, for every quadratic form there is a basis in which it becomes

$$\sum \lambda_j (x'_j)^2,$$

with some real numbers λ_j . One can further simplify by putting $y_j = \sqrt{|\lambda_j|} x'_j$, and obtain

$$Q(x) = \sum \pm y_j^2, \tag{5}$$

but this last change is in general not orthogonal.

The numbers of pluses, minuses and zeros in this sum cannot be changed by *any* changes of basis, it depends only on Q , and this triple of non-negative integers is called the *signature* of the form Q . By the “number of zeros” we mean the dimension of the space minus the number of terms in (5).

It is easy to see that the signature depends on Q only. Indeed, the number of zeros can be defined as the maximum dimension of the subspace such that the restriction of Q on this subspace is zero. Similarly the number of pluses can be described as the maximum dimension of a subspace such that the restriction of Q on them is non-negative minus the number of zeros.

The form is positive definite if and only if it has n (=dimension of the space) pluses in its signature. Such forms can serve as usual dot products.

3. Indefinite forms are also useful in many applications, including basic physics. As an example we consider the *Minkowski space* which is $\mathbf{R}^4$ equipped with the bilinear form

$$B(\mathbf{u}, \mathbf{v}) = u_0v_0 - u_1v_1 - u_2v_2 - u_3v_3,$$

where we denote vectors in $\mathbf{R}^4$ by $\mathbf{u} = (u_0, u_1, u_2, u_3)^T$. The corresponding quadratic form (which is called the Minkowski metric) is

$$Q(\mathbf{u}) = u_0^2 - u_1^2 - u_2^2 - u_3^2.$$

We can define “orthogonality”, and “length” of a vector as we did for the usual dot product, but with the new definition, the “length” can be imaginary. A vector u is called *time-like* if $Q(u) \geq 0$, in which case the Minkowski “length” is defined as the positive square root

$$s(\mathbf{u}) = \sqrt{u_0^2 - u_1^2 - u_2^2 - u_3^2}.$$

The standard name for this Minkowski “length” s in physics is *interval*. This terminology, “time-like”, “interval” etc. will be explained later.

The theory of the usual dot product begins with the Schwarz inequality. For the Minkowski space we have the

Reverse Schwarz Inequality. *If $Q(\mathbf{u}) > 0$, $Q(\mathbf{v}) > 0$ and $u_0v_0 > 0$. Then*

$$B(\mathbf{u}, \mathbf{v})^2 \geq Q(\mathbf{u})Q(\mathbf{v})$$

with equality only when $\mathbf{u}, \mathbf{v}$ are proportional.

Proof. We denote $\mathbf{u} = (u_0, \mathbf{x})^T$ where $\mathbf{x} = (u_1, u_2, u_3)$, and $v = (v_0, \mathbf{y})$, where $\mathbf{y} = (u_1, u_2, u_3)$, and let $\mathbf{x} \cdot \mathbf{y} = u_1v_1 + u_2v_2 + u_3v_3$ be the usual dot product of 3-vectors. Then

$$B(\mathbf{u}, \mathbf{v}) = u_0v_0 - \mathbf{x} \cdot \mathbf{y} \geq |u_0v_0| - \|\mathbf{x}\|\|\mathbf{y}\| \geq \sqrt{u_0^2 - \|\mathbf{x}\|^2} \sqrt{v_0^2 - \|\mathbf{y}\|^2}, \quad (6)$$

where we used the usual Schwarz inequality $|\mathbf{x} \cdot \mathbf{y}| \leq \|\mathbf{x}\|\|\mathbf{y}\|$, and the elementary inequality (prove it!)

$$ab - cd \geq \sqrt{a^2 - c^2} \sqrt{b^2 - d^2},$$

which holds when $a > c > 0$ and $b > d > 0$. After squaring we obtain our inequality.

From this we obtain the

Reverse Triangle inequality. *If $Q(\mathbf{v}) > 0$ and $Q(\mathbf{u}) > 0$ then*

$$s(\mathbf{u} + \mathbf{v}) \geq s(\mathbf{u}) + s(\mathbf{v}),$$

where $s(\mathbf{u}) = \sqrt{Q(\mathbf{u})}$ is the interval.

This means that the straight line is the longest among all time-like paths between two points!

Proof. With the same notation as before, this is equivalent to

$$\sqrt{(u_0 + v_0)^2 - \|\mathbf{x} + \mathbf{y}\|^2} \geq \sqrt{u_0^2 - \|\mathbf{x}\|^2} + \sqrt{v_0^2 - \|\mathbf{y}\|^2}.$$

Squaring both sides and canceling equal terms on the right and on the left, we obtain an equivalent inequality

$$u_0 v_0 - \mathbf{x} \cdot \mathbf{y} \geq \sqrt{u_0^2 - \|\mathbf{x}\|^2} \sqrt{v_0^2 - \|\mathbf{y}\|^2},$$

and this is the same as the reverse Schwarz inequality (6).

4. Let us find the linear transformations which preserve the interval. They are called *Lorentz transformations*. They are analogous to the orthogonal transformation which preserve the usual dot product. For simplicity we do it in 2 dimensional Minkowski space, which is $\mathbf{R}^2$ consisting of vectors $\mathbf{u} = (t, x)^T$ and equipped with the quadratic form

$$Q(\mathbf{u}) = t^2 - x^2.$$

Let A be a matrix of such a transformation with respect to the standard basis. We look for matrices that satisfy

$$Q(A\mathbf{u}) = Q(\mathbf{u}) \quad \text{for all } \mathbf{u}.$$

Let

$$A = \begin{pmatrix} a & b \\ c & d \end{pmatrix},$$

$$Au = \begin{pmatrix} at + bx \\ ct + dx \end{pmatrix}.$$

So our condition is

$$(at + bx)^2 - (ct + dx)^2 = t^2 - x^2,$$

and this must hold for all t and x . We obtain the system

$$a^2 - c^2 = 1, \quad (7)$$

$$b^2 - d^2 = -1, \quad (8)$$

$$ab - cd = 0. \quad (9)$$

Notice that $ad \neq 0$, this follows from (7), and (8). Dividing (9) by ad we obtain

$$\frac{c}{a} = \frac{b}{d}.$$

Denote this by k . Then

$$c = ak, \quad b = dk.$$

Substitute this to (7) and (8) and obtain

$$a = \frac{\epsilon_1}{\sqrt{1 - k^2}}, \quad d = \frac{\epsilon_2}{\sqrt{1 - k^2}},$$

where $\epsilon_j \in \{\pm 1\}$. and finally

$$b = \frac{\epsilon_2 k}{\sqrt{1 - k^2}}, \quad c = \frac{\epsilon_1 k}{\sqrt{1 - k^2}}.$$

For our matrix to be real, we need $0 \leq k < 1$. So the general form of a real matrix preserving the interval is

$$\frac{1}{\sqrt{1 - k^2}} \begin{pmatrix} \epsilon_1 & \epsilon_2 k \\ \epsilon_1 k & \epsilon_2 \end{pmatrix}, \quad -1 < k < 1. \quad (10)$$

We see that the set of these matrices consists of 4 components, depending on the signs ϵ_1, ϵ_2 . Transformations with $\epsilon_1 = 1$ are called *orthochronous*, and transformations with $\epsilon_2 = 1$ are called *proper*. Thus a proper orthochronous Lorentz transformation has the form

$$\frac{1}{\sqrt{1 - k^2}} \begin{pmatrix} 1 & k \\ k & 1 \end{pmatrix}, \quad -1 < k < 1. \quad (11)$$

Notice that the determinant is 1, as in the case of rotations. The transformations we obtained are called *boosts*. They are special Lorentz transformations: in the real physical the space-time of dimension 4, Lorentz transformations are compositions of boosts and usual rotations of the 3-dimensional space not affecting the time coordinate.

5. Minkowski space is interpreted in physics as the *space-time*. Its points correspond to *events*. The coordinates (u_1, u_2, u_3) are the usual coordinates of the point where the event occurs, and u_0 is the *time* when it occurs.

(We choose the units where the speed of light equals 1. In other units, the square of the interval is

$$c^2 c u_0^2 - u_1^2 - u_2^2 - u_3^2$$

where c_0 is the speed of light in these units.)

If some point moves in the space, we can record its coordinates at every moment, and obtain a curve in the space time which is called the *world-line* of the point. If the point moves with constant speed, this world line is straight.

For example, if a point does not move, its world line will be parallel to the u_0 axis. If a point moves along u_1 axis with speed k starting at the origin, then its world line is a straight line in the (u_0, u_1) -plane and the equations of this line are $u_1 = k u_0$, $u_2 = \text{const}$, $u_3 = \text{const}$.

So far, there is nothing new.

In what follows we neglect for simplicity the coordinates u_2, u_3 and consider only the motion along the line u_1 with constant speed k . So we are interested only in two coordinates (u_0, u_1) which we denote by (t, x) , so t is the time coordinate and x is the space (position) coordinate.

Now we discuss the changes of coordinates. Consider two observers (coordinate systems) whose origin is the same, one system (t, x) we call fixed, and another one (t', x') is moving along the x -axis to the left with speed v . When the observers pass each other at the origin, they synchronize their clocks. If (t, x) are coordinates of some event with respect to the fixed observer, and (t', x') are coordinates of the same event for the moving observer, then according to classical physics

$$\begin{aligned} t' &= t \\ x' &= x + vt \end{aligned} \tag{12}$$

So the space-time coordinate transformation is described by the matrix

$$\begin{pmatrix} 1 & 0 \\ v & 1 \end{pmatrix}. \tag{13}$$

If some point moves with velocity v , for the fixed observer, then it moves with velocity

$$v' = v + k \tag{14}$$

with respect to the moving observer. This is called the law of addition of velocities.

However, in the very beginning of 20th century it was understood all that this cannot be true in the nature. Two fundamental facts were established:

- a) All laws of physics must be the same in two frames of reference moving with respect to each other with constant velocity. This is called the Relativity Principle, which was formulated by Galileo Galilei, and which is a fundamental principle of all physics.
- b) The speed of light in vacuum is the same in all frames of reference.

The first principle was very well verified by the whole development of physics since Newton. The second principle follows from the theory of light and electromagnetism due to Maxwell, which was also very well verified by the end of 19th century.

But the two principles seem to contradict each other, when we use the addition rule of velocities (14). Direct very precise experiments were made to measure the velocity of light in different directions. The Earth moves in space, so if the law of addition of velocities holds, the speed of light in the direction along the Earth motion would be different from the speed of light across this direction. The experiment of Michelson and Morley showed that these speeds are equal.

Still a great courage was required to say that equation (14) and the rule (12) must be modified.

If the speed of light is constant in all coordinate systems then the interval between any two events must be independent of the coordinate system.

Indeed, suppose that at time 0 we have a flash at the origin. If light will reach x in time t we have

$$ct = x, \quad c^2t^2 - x^2 = 0,$$

where c is the speed of light, and this should be independent of the coordinate system. We can choose the units so that $c = 1$. For example we can measure distances in light-years and time in years. Or instead we can measure time in meters; 1 time-meter is the time required to light to pass the distance 1 meter in vacuum. This is where the Minkowski quadratic form comes from. So we conclude that the correct transformation is not (12) but the Lorentz boost (8). One may ask, why did not people notice this before, after all I

said that the laws of mechanics were very well verified by the end of 19th century.

To compare the classical laws with relativistic ones, we rewrite the Lorentz transformation in the usual units, in which the speed of light is approximately $c = 3 \cdot 10^8 m/sec$. The interval in these units has the expression

$$\sqrt{c^2 t^2 - x_1^2 - x_2^2 - x_3^2},$$

that is

$$\sqrt{c^2 t^2 - x^2} \tag{15}$$

in the one-dimensional case that we consider. So we have to perform the change of coordinates with the matrix

$$\begin{pmatrix} c & 0 \\ 0 & 1 \end{pmatrix}.$$

If A is the matrix (11), then the matrix preserving the form (15) will be

$$\frac{1}{\sqrt{1-k^2}} \begin{pmatrix} 1 & k/c \\ ck & 1 \end{pmatrix},$$

where we have to put $k = v/c$, where v is the speed in m/sec , so we obtain the Lorentz transformation matrix in the usual units:

$$\frac{1}{\sqrt{1-(v/c)^2}} \begin{pmatrix} 1 & v/c^2 \\ v & 1 \end{pmatrix}.$$

When v is much smaller than c , this is approximately the same as (13). In Michelson–Morley experiment v was the speed of the Earth rotation around the Sun, which is approximately $3 \cdot 10^5 m/sec$ so $v/c \approx 10^{-3}$, and $\sqrt{1-(v/c)^2} \approx 1$ with accuracy 2×10^{-6} .

To obtain the relativistic law of “addition of velocities”, we just compose two transformations (8), that is multiply their matrices.

Appendix 1. More rigorous derivation of Lorentz transformations. Let us look for non-singular linear transformations

$$\begin{pmatrix} t' \\ x' \end{pmatrix} = \begin{pmatrix} a & b \\ c & d \end{pmatrix} \begin{pmatrix} t \\ x \end{pmatrix} = \begin{pmatrix} at + bx \\ ct + dx \end{pmatrix}$$

which preserve the zero set of the quadratic form

$$t^2 - x^2 = 0. \quad (16)$$

In fact we want to impose the stronger condition that the part of the cone

$$t^2 - x^2 > 0$$

where $t > 0$ is mapped into itself. This means *causality* in physics: if an event at $(0, 0)$ causes some event at (t, x) then we must have $t > 0$, and this fact should be independent of the coordinate system.

This means that whenever (16) holds we must have

$$(t')^2 - (x')^2 = 0.$$

Equation (16) is equivalent to $t = x$ or $t = -x$. After a simple computation we obtain

$$ab = cd, \quad (17)$$

$$a^2 - c^2 = d^2 - b^2. \quad (18)$$

Dividing (17) on ad we obtain

$$\frac{b}{d} = \frac{c}{a}.$$

(Think why we cannot have $ad = 0$). Denote this common value by k , then we have $b = kd$, $c = ka$. Introducing these to (18) we get

$$a^2(1 - k^2) = d^2(1 - k^2).$$

If $k = \pm 1$ our matrix will be singular, so this case must be rejected, and we are left with

$$a^2 = d^2.$$

If $a = d$, our matrix is

$$a \begin{pmatrix} 1 & k \\ k & 1 \end{pmatrix}. \quad (19)$$

If $a = -d$ then it is

$$a \begin{pmatrix} 1 & -k \\ k & -1 \end{pmatrix}. \quad (20)$$

This matrix (20) is a product of (19) with the matrix of reflection in $x = 0$:

$$R = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}, \quad (21)$$

so it is enough to consider only (19).

Now we make the following physical assumption: our matrices depend on one parameter, the 1-dimensional vector which is interpreted as velocity, or more precisely the ratio of velocity to the speed of light. we associate velocity with parameter k . Let our matrix be $A(k)$. Then the condition of isotropy of space means

$$A(k)R = RA(-k).$$

from this we obtain that $a(k) = a(-k)$ must be satisfied. And finally the transformations $A(k)$ and $A(-k)$ must be inverse to each other, this gives $a(k) = \pm 1/\sqrt{1 - k^2}$. Finally we have to choose $+$ in this formula to preserve *causality*: a vector with positive t must be mapped to a vector with positive t . So the final form is

$$\frac{1}{\sqrt{1 - k^2}} \begin{pmatrix} 1 & k \\ k & 1 \end{pmatrix},$$

a proper orthochronous Lorentz transformation.

Appendix 2. Addition of velocities. We have

$$\frac{1}{\sqrt{1 - k_1^2}} \begin{pmatrix} 1 & k_1 \\ k_1 & 1 \end{pmatrix} \frac{1}{\sqrt{1 - k_2^2}} \begin{pmatrix} 1 & k_2 \\ k_2 & 1 \end{pmatrix} = \frac{1}{\sqrt{1 - k_3^2}} \begin{pmatrix} 1 & k_3 \\ k_3 & 1 \end{pmatrix},$$

where

$$k_3 = \frac{k_1 + k_2}{1 + k_1 k_2}. \quad (22)$$

Verify this. This is the relativistic law of addition of velocities. We recall that $k = v/c$ which is very small for the objects that we usually observe. In this case the denominator of (22) is almost 1, and we obtain the familiar law of addition of velocities.

It is interesting to verify that this law is associative and commutative. To do this, we recall the familiar rule for the hyperbolic tangent of a sum:

$$\tanh(\theta_1 + \theta_2) = \frac{\tanh \theta_1 + \tanh \theta_2}{1 + \tanh \theta_1 \tanh \theta_2}.$$

This suggests that we should define a real number θ by the formula

$$k = \tanh \theta.$$

This number is called the *rapidity* to distinguish it from speed. Recalling the formula

$$\cosh^2 \theta - \sinh^2 \theta = 1,$$

we obtain $\sqrt{1 - k^2} = 1/\cosh^2 \theta$, so our matrix (10) of Lorentz transformation becomes

$$\begin{pmatrix} \cosh \theta & \sinh \theta \\ \sinh \theta & \cosh \theta \end{pmatrix},$$

which is very similar to the rotation matrix.

Appendix 3. Physical interpretation of a time-like interval and the twins paradox.

Suppose we have two events separated by a time-like interval. Then there is a coordinate system in which these two events appear in the same place. Indeed, start with some coordinate system where the events occur at $(0, 0)$ and (t, x) . As the interval between them is time like we have $t^2 - x^2 > 0$, so if we choose $k = -x/t$ in (10), then the coordinates in the new system will be $(0, 0)$ and $(t', x') = (\sqrt{t^2 - x^2}, 0)$. So in the new coordinate system the events take occur in the same place, and the time between is exactly what we called the interval.

So in general, the interval is the “proper time” between the events: the time which passes between them in the coordinate system where they occur in the same place.

For the system which does not move with uniform speed, the world line is a curve but the proper time of the system can be defined analogously to the length of the curve in geometry:

$$\int_{\gamma} ds. \tag{23}$$

In Euclidean geometry, $ds = \sqrt{dx^2 + dy^2}$ is the length element while in Minkowski geometry $ds = \sqrt{st^2 - dx^2}$ is the element of the interval. A curve is called time-like if this interval is real everywhere on a curve. The world lines of the real objects can be only time-like. Then we obtain from the reverse Schwarz inequality that the Minkowski length of a time-like curve

is *less than or equal to* the Minkowski distance between its end points! And the *longest* curve between two points is the straight line.

What does this really mean? Suppose that Alice and Bob are twins. At some point Alice decides to travel. She departs and returns, while Bob stays in his place. Departure and return of course occur at the same place from Bob's point of view, so the interval between them is the time which Bob waits for Alice's return. On the other hand for Alice, her proper time of travel is the integral (23) along her world line, which as we have seen is in general *smaller* than the time Bob waits for her. Therefore when Alice returns she is *younger* than Bob! This is called the twins paradox.

For example, suppose that Alice's destination is 4 light years away, and she travels with the speed $k = 1/2$. From Bob's point of view the round trip takes $2 \times 4/(1/2) = 16$ years. From Alice's point of view it takes $2 \times \sqrt{8^2 - 4^2} = 13.8$ years. So Alice is 2.2 years younger than Bob when she returns. I neglected the time Alice needs for acceleration to and deceleration from the speed of $c/2 = 150,000$ kilometers per second.

References

- [1] P. Jordan and J. von Neumann, On inner products in linear metric spaces, Ann. Math., 36,3 (1935) 719–723.
- [2] Mathematics stack exchange,
<https://math.stackexchange.com/questions/21792/>