

PHYS 480/581: General Relativity

Index Notation and the Metric

Prof. Cyr-Racine
(Dated: August 26, 2026)

In this course, we will be writing a lot of equations relating tensors, vectors, and scalars. We need to introduce a compact notation allowing us to write down these equations without explicitly having to expand the whole calculations into matrices and column vectors. Here and in the future, whenever we say “vector”, we always mean “four-vectors”. Three-dimensional spatial vectors (when discussed) will always be referred to as “three-vectors”.

I. ABSTRACT INDICES

A. Components of vectors and dual vectors

[vector indices] So far, we have denoted four-vectors using italic-bold font (e.g. $\mathbf{p}$). Sometimes, we want to directly refer to the *components* of $\mathbf{p}$ within some basis (we will discuss such basis next time). We will denote these components p^μ , where μ can take the value 0, 1, 2, and 3, with 0 referring to the “time” component, and 1 – 3 refers to the “spatial” components (e.g. $1 = x$, $2 = y$, and $3 = z$ in cartesian coordinates).

[Index placement] Now, the placement of the index is important with four-vector components always having an upper index. The similar quantity with a lower index (i.e. p_μ) is called a dual vector (also referred to as 1-form or covector). Fundamentally, p_μ is a different mathematical object than p^μ , and is therefore very important to keep track of whether the index μ is up or down. Indeed, as we discussed last time, while p^μ is a vector living in the tangent vector space $T_q M$ of some point q of some manifold M , p_μ (which lives in the dual space $T_q^* M$) is a map from that vector space $T_q M$ to the real numbers. Heuristically, you can think of dual vectors as a “machine” that takes in a vector and returns a real number. Of course, we’ve already seen this at play in taking the inner product between a dual vector $\mathbf{p}$ and a vector $\mathbf{q}$

$$\langle \mathbf{p}, \mathbf{q} \rangle = \mathbf{p} \cdot \mathbf{q} = p(\mathbf{q}) = \sum_{\mu=0}^3 p_\mu q^\mu, \quad (1)$$

We see from Eq. (1) that “taking the dot product between two four-vectors” is really about multiplying component-wise the elements of a dual vector with that of an actual vector.

[Greek vs Latin indices] In general, we will use greek letters (e.g. α, β, μ, ν , etc.) to denote spacetime indices, that is, these indices will run from 0 to 3, with 0 denoting the “time” entry. Latin letters (i, j, k , etc.) will exclusively be used to denote spatial coordinates, that is, these indices will run from 1 to 3 (for instance x, y, z in cartesian coordinates).

B. Matrices

[Element of a matrix] Matrices fundamentally carry two indices corresponding to their rows and columns. In general, an element of a matrix $\mathbf{\Lambda}$ from the μ th row and ν th column will be denoted $\Lambda^\mu{}_\nu$. Similar to vectors, the placement of the indices (i.e. whether they are upper or lower indices) is very important. For instance, $\Lambda^\mu{}_\nu$, $\Lambda_{\mu\nu}$, $\Lambda_\mu{}^\nu$ will generally have different numerical values. They are however related by the raising or lowering of indices, which will be covered below.

As you are familiar from linear algebra, matrices can multiply vectors to form other vectors. An example of this we have already encountered is the Lorentz boost along the x -axis between a frame S and S'

$$x'^\mu = \sum_{\nu=0}^3 \Lambda^\mu{}_\nu x^\nu, \quad (2)$$

(note the placement of the ν indices) where the matrix $\Lambda^\mu{}_\nu$ has components given by

$$\Lambda^\mu{}_\nu = \begin{bmatrix} \gamma & -\gamma\beta & 0 & 0 \\ -\gamma\beta & \gamma & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix}. \quad (3)$$

The above is a little bit of an abuse of notation, as $\Lambda^\mu{}_\nu$ is supposed to denote one component of the matrix. You should thus read this as $\Lambda^0{}_0 = \gamma$, $\Lambda^0{}_1 = -\gamma\beta$, $\Lambda^0{}_2 = 0$, etc.

C. Einstein summation convention

To avoid writing the summation symbols for every indices that are summed over, we will use the Einstein summation convention. This says that any pair of repeated indices (with one upper index and one lower index) are always summed over. For instance, the boost Lorentz transformation given in Eq. (2) above can be written more compactly as

$$x'^\mu = \Lambda^\mu{}_\nu x^\nu, \quad (4)$$

where it is now understood that the index ν is summed over from $\nu = 0$ to $\nu = 3$.

[Dummy indices] When an index is summed over, it is usually referred to as a “**bound**” or “**dummy**” index. We are free to relabel any dummy index (that is, switch one greek letter for another). For instance, there is nothing special about using ν in Eq. (4) above, and we could instead have written

$$x'^\mu = \Lambda^\mu{}_\alpha x^\alpha, \quad (5)$$

which has the exact same meaning as Eq. (4).

[Free indices] Indices that are not repeated are called **free** indices. These must match and balance on either side of every equal sign. For instance, I cannot make the substitution $\mu \rightarrow \beta$ in the right-hand side of Eq. (4), without also making the substitution also on the left-hand side. That is, this equation is nonsense

$$x'^\mu = \Lambda^\beta{}_\alpha x^\alpha, \quad (\text{wrong!!}), \quad (6)$$

while this one is fine:

$$x'^\beta = \Lambda^\beta{}_\alpha x^\alpha. \quad (7)$$

See **boxes 4.4 and 4.5** for further examples.

II. THE METRIC

The most important mathematical object we will encounter in this course is the metric tensor (the term “tensor” will be defined later) $g_{\mu\nu}$. The metric is usually defined with its two indices down, and is symmetric in its two indices, $g_{\mu\nu} = g_{\nu\mu}$. Mathematically, the metric at point p of some manifold M is a map $g_{(p)} : T_p M \times T_p M \rightarrow \mathbb{R}$, which takes two four-vectors from the tangent space at point p and returns a real number. Given four-vectors $\mathbf{p}$ and $\mathbf{q}$ at point p , we for example have:

$$g_{(p)}(\mathbf{p}, \mathbf{q}) = g_{(p)}(\mathbf{q}, \mathbf{p}) = g_{\mu\nu} p^\mu q^\nu \in \mathbb{R}. \quad (8)$$

For notational convenience, we generally won’t write the subscript (p)

A. Kronecker delta and inverse metric

The Kronecker delta is defined as

$$\delta^\mu{}_\nu = \begin{cases} 1 & \text{if } \mu = \nu, \\ 0 & \text{otherwise.} \end{cases} \quad (9)$$

This means that whenever the Kronecker delta appears in (Einstein convention) sums, its role is to replace the index being summed over with the other index of the delta, that is,

$$\delta^\mu{}_\nu p^\nu = p^\mu. \quad (10)$$

This is just the identity map. In matrix form, $\delta^\mu{}_\nu$ is the identity matrix with ones on the diagonal and zeros everywhere else.

[Inverse metric] Importantly, the Kronecker delta allows us to define the inverse metric $g^{\mu\nu}$ (that with the upper indices) via the relation

$$g^{\mu\nu} g_{\nu\sigma} = g_{\lambda\sigma} g^{\lambda\mu} = \delta^\mu{}_\sigma. \quad (11)$$

Mathematically, the inverse metric at point p of some manifold M is a map $g^{(p)} : T_p^*M \times T_p^*M \rightarrow \mathbb{R}$, which takes two dual vectors (or one-forms) and returns a real number.

B. Raising and Lowering indices

The metric can be also be used to raise and lower indices on vectors, dual vectors, and other tensors. For example, say you have a vector with components p^μ . The components of the related dual vector p_μ are given by

$$p_\mu = g_{\mu\nu} p^\nu, \quad (12)$$

where of course Einstein summation convention is implied. The metric at point p of some manifold M thus defines a map (called an isomorphism) between the tangent space T_pM and its dual T_p^*M . The inverse metric $g^{\mu\nu}$ can also be used to raise indices. For example,

$$p^\mu = g^{\mu\nu} p_\nu, \quad (13)$$

In this case, the inverse metric at point p of some manifold M defines an isomorphism between T_p^*M and T_pM , that is, it takes a dual vector (one-form) and returns the corresponding vector. Since we will exclusively consider manifolds with metric in this course, the tangent space at any spacetime point will always related to its dual via the metric.

[3D Euclidean space] In 3D Euclidean space, the metric is simply $g_{ij} = \delta_{ij}$ (i.e. the identity map), which tells us that the tangent space and its dual at any point are the exact same vector space, which is why we never worry about the distinction between three-vectors and dual three-vector in Euclidean space.

[Higher-order tensors] The same philosophy applies to matrices (or other tensorial objects), for which we can use the metric to raise or lower indices. For example,

$$\Lambda_{\alpha\nu} = g_{\alpha\mu} \Lambda^\mu{}_\nu. \quad (14)$$

C. Relation fo inner product

Given everything we already said above, you probably guessed that the action of the metric is related to the inner product between a dual vector and a vector. From Eq. (1) above,

$$\langle \mathbf{p}, \mathbf{q} \rangle = \mathbf{p} \cdot \mathbf{q} = p_\mu q^\mu = g_{\mu\nu} p^\nu q^\mu = p^\nu q_\nu, \quad (15)$$

where we used Eq. (12) in the third equality. Thus, the metric defines the norm of vectors living in the tangent space of any spacetime point.

D. Notion of distance

Since it defines the norm of vectors, the metric allows for an unambiguous notion of distance between two nearby spacetime points. This is perhaps most apparent when we write the spacetime interval ds^2 as

$$ds^2 = g_{\mu\nu} dx^\mu dx^\nu. \quad (16)$$

In flat spacetime, we already saw that the spacetime interval is $ds^2 = -dt^2 + dx^2 + dy^2 + dz^2$ using spatial cartesian coordinates, and we will use the notation

$$ds^2 = \eta_{\mu\nu} dx^\mu dx^\nu, \quad (17)$$

where

$$\eta_{\mu\nu} = \begin{bmatrix} -1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix} \quad (18)$$

is referred to as the *Minkowski* metric. In this course, we will use $\eta_{\mu\nu}$ exclusively to denote the Minkowski metric, while $g_{\mu\nu}$ will be used to denote the more general metric of (usually curved) spacetime. Note that because of the intimate relationship between ds^2 and $g_{\mu\nu}$ given in Eq. (16), we sometime refer to ds^2 as the metric itself, which is a slight abuse of language.

[Inner product on flat spacetime] With the help of the Minkowski metric $\eta_{\mu\nu}$, the inner product admits the familiar form in cartesian coordinates

$$\mathbf{p} \cdot \mathbf{q} = \eta_{\mu\nu} p^\mu q^\nu = -p^t q^t + p^x q^x + p^y q^y + p^z q^z. \quad (19)$$

E. Numerical Examples

Let's say I have a vector with cartesian components $p^\mu = (1, 1, 1, 1)$ in flat spacetime (i.e. the metric is Minkowski). Then, the corresponding dual vector has components

$$p_\mu = (-1, 1, 1, 1). \quad (20)$$

As a matrix example, if $\Lambda^\mu{}_\nu$ has components given in Eq. (3) above, than in flat space time where $g_{\mu\nu} = \eta_{\mu\nu}$, $\Lambda_{\alpha\nu}$ has components

$$\Lambda_{\alpha\nu} = \eta_{\alpha\mu} \Lambda^\mu{}_\nu = \begin{bmatrix} -\gamma & \gamma\beta & 0 & 0 \\ -\gamma\beta & \gamma & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix}. \quad (21)$$

III. INDEX NOTATION VERSUS MATRIX MULTIPLICATION

Consider three 4×4 matrices $\mathbf{A}$, $\mathbf{B}$, and $\mathbf{C}$ such that

$$\mathbf{C} = \mathbf{AB}. \quad (22)$$

In terms of components, this matrix product could be written as

$$C_{\mu\nu} = A_{\mu\alpha} B^\alpha{}_\nu = A_\mu{}^\alpha B_{\alpha\nu}. \quad (23)$$

Now, since the components $A_{\mu\alpha}$ and $B^\alpha{}_\nu$ are just numbers, I'm free to change the order to $C_{\mu\nu} = B^\alpha{}_\nu A_{\mu\alpha}$. Now the trick question is what happens if you are given an equation in terms of its components and want to write it down as a matrix product. Say you are given the actual numerical value for the all the components $A_{\mu\alpha}$ and $B^\alpha{}_\nu$ and ask to compute $C_{\mu\nu} = B^\alpha{}_\nu A_{\mu\alpha}$. Do you write down $\mathbf{C} = \mathbf{BA}$? Since in general $\mathbf{BA} \neq \mathbf{AB}$, it's important to get the order right. The trick is to make sure that repeated indices are next to each other. Notice how the two α indices in Eq. (23) are consecutive, that is the order of indices from left to right is $\mu\alpha\alpha\nu$. This tells you that the right order for multiplying the matrix is $\mathbf{C} = \mathbf{AB}$.

[Matrix notation for upper and lower indices] This kind of situation arise for instance when raising or lowering indices with the metric η . For instance, imagine you are given the components of a matrix $\mathbf{C}$ as $C^{\mu\nu}$. Say you want instead to compute the components $C^\mu{}_\nu$. You use the metric to lower the second index

$$C^\mu{}_\nu = C^{\mu\alpha} \eta_{\alpha\nu} = [\mathbf{C}\boldsymbol{\eta}]^\mu{}_\nu, \quad (24)$$

where in the last equality we have written the matrix product explicitly in the right order. Thus to get the components of $C^\mu{}_\nu$, you need to perform the matrix multiplication $\mathbf{C}\boldsymbol{\eta}$. If instead you wanted to lower the first index and get the components of $C_\mu{}^\nu$, you would instead do

$$C_\mu{}^\nu = \eta_{\mu\alpha} C^{\alpha\nu} = [\boldsymbol{\eta}\mathbf{C}]_\mu{}^\nu. \quad (25)$$

Thus, to get the components of $C_\mu{}^\nu$, you need to perform the matrix multiplication $\boldsymbol{\eta}\mathbf{C}$. If you would like to lower both indices to get $C_{\mu\nu}$, you would then do

$$C_{\mu\nu} = \eta_{\mu\alpha} C^{\alpha\beta} \eta_{\beta\nu} = [\boldsymbol{\eta}\mathbf{C}\boldsymbol{\eta}]_{\mu\nu}. \quad (26)$$

Of course, as we start to introduce higher-order tensors and objects, thinking of indices contraction as matrix multiplication will become much less useful. But for 2-tensors, thinking of matrix multiplication could save you a lot of algebra, if you are careful with the ordering as shown above (see **box 4.6** for another example of such calculation).