

CSE439 Fall 2026 Week 2: Linear Algebra, Unitary Operations, and Qubits

First, one item added to review of last lecture:

- $f(n) = O(g(n))$ means there is a constant $C > 0$ such that for all n past a certain point n_0 , $f(n) \leq C \cdot g(n)$.
- $f(n) = \Theta(g(n))$ means there are constants c and C such that beyond some point n_0 , $f(n)$ stays in the range $c \cdot g(n)$ to $C \cdot g(n)$.
- $f(n) \sim g(n)$ is stronger---it says that the limit $f(n)/g(n)$ exists and equals 1.

With exp-log functions, however, the only cases that occur (whiteboard [photo 1](#)) are:

- $f(n) = o(g(n))$, or simply $f = o(g)$, which means that $\lim_{n \rightarrow \infty} \frac{f(n)}{g(n)}$ exists and is 0.
- $f(n) \sim Cg(n)$ for some constant $C > 0$, which we often want to specify. This implies $f(n) = \Theta(g(n))$ as defined above, but is simpler.

Past years' lectures did examples like: $f(n) = 3n^2$ versus $g(n) = 3n^2 + 100n + 500$. The last bullet is the same as $\lim_{n \rightarrow \infty} \frac{g(n)}{f(n)} = 1$. In the example:

$$\frac{g(n)}{f(n)} = \frac{3n^2 + 100n + 500}{3n^2} = 1 + \frac{100n}{3n^2} + \frac{500}{3n^2} = 1 + \frac{100}{3n} + \frac{500}{3n^2}$$

The limit is 1 because the limits of the latter two fractions are 0. In F25, lecture gave a different example where the leading constants were different:

$$\frac{f(n)}{g(n)} = \frac{3n^3 + 25n^2 - 600n + 7}{7n^3 - 40n^2 + 80n - 4} \rightarrow \frac{3}{7} \text{ as } n \rightarrow \infty.$$

Thus we can write $f(n) \sim \frac{3}{7}g(n)$. This is stronger than saying " $f(n) = \Theta(g(n))$ " because it identifies the **principal constant** involved exactly.

One useful tool that helps you not only prove this limit but see-and-justify it at a glance is [L'Hôpital's Rule](#). Here's an *informal* statement: If taking limits of the numerator and denominator of $\frac{f(x)}{g(x)}$ gives the indeterminate form $\frac{0}{0}$ or $\frac{\infty}{\infty}$, then the "limit status" of $\frac{f(x)}{g(x)}$ is the same as that of $\frac{f'(x)}{g'(x)}$.
Fall 2026: this is equivalent to how I wrote it on the whiteboard ([photo 2](#)).

In the above example, we get indeterminate forms until we take the *third* derivatives. By then, the lower-order terms have all been "derived" to zero, the $3n^3$ has become $9n^2$, $18n$, and finally 18, and the $7n^3$ has become 42. This gives $\frac{18}{42}$ which is the same as $\frac{3}{7}$. For Fall 2026 I did an example in

which you don't just iterate the numerator as $f(n), f'(n), f''(n), \dots$ and likewise with $g(n)$; rather, the functions called "f" and "g" can change as one re-works the fraction:

$$\frac{(\ln n)^2}{n^{1/3}} \mapsto \frac{\frac{2 \ln n}{n}}{\frac{1}{3}n^{-2/3}} = \frac{6 \ln n}{n^{1/3}} \mapsto \frac{\frac{6}{n}}{\frac{1}{3}n^{-2/3}} = \frac{18}{n^{1/3}},$$

and now it is clear that the limit is 0, so $(\ln n)^2 = o(n^{1/3})$. Other tricks are to take logs of both sides and/or to expand logarithms in exponents. The latter is exemplified by $2^{3 \log_2 n} = (2^{\log_2 n})^3 = n^3$.

Tensor Products

Fall 2026 wrote a more-detailed version of the following on the whiteboard as [photo 3](#), including vectors treated the same as $1 \times n$ or $m \times 1$ matrices.

Main Point: Matrix multiplication is "wrong"
 Instead, it is Tensor Product: $A \otimes B = \begin{bmatrix} a_{11}B & a_{12}B & \dots & a_{1n}B \\ a_{21}B & a_{22}B & \dots & a_{2n}B \\ \vdots & \vdots & \ddots & \vdots \\ a_{n1}B & a_{n2}B & \dots & a_{nn}B \end{bmatrix}$
 How to visualize it: Say $A = \begin{bmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ a_{n1} & a_{n2} & \dots & a_{nn} \end{bmatrix}$:
 As a matrix, if A is $n \times n$ and B is $r \times r$, then $A \otimes B$ is $(nr) \times (nr)$. If we form $\underbrace{A \otimes A \otimes \dots \otimes A}_{K \text{ items}}$, then the matrix $A^{\otimes K}$ has size n^K . If $K \ll n$, this is exponential size.

If A is $\ell \times m$ and B is $n \times r$ then $A \otimes B$ is $\ell n \times mr$, so the dimensions can be anything. In particular, A and B can both be column vectors with $m = r = 1$, whereupon $A \otimes B$ is a column vector of length ℓn .

Example Hadamard Matrix (without normalizing)

$$H = H_2 = \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix}$$

(normalized: $\begin{bmatrix} 1/\sqrt{2} & 1/\sqrt{2} \\ 1/\sqrt{2} & -1/\sqrt{2} \end{bmatrix}$).

$$H \otimes H = \begin{bmatrix} 1 \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix} & 1 \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix} \\ 1 \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix} & -1 \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix} \end{bmatrix} = \begin{bmatrix} 1 & 1 & 1 & 1 \\ 1 & -1 & 1 & -1 \\ 1 & 1 & -1 & -1 \\ 1 & -1 & -1 & 1 \end{bmatrix} \quad \begin{matrix} 00 & 01 & 10 & 11 \\ 00 & 01 & 10 & 11 \end{matrix} \quad 4 \times 4$$

$H \otimes H \otimes H$ is $2^3 \times 2^3$
ie 8×8

$H^{\otimes n}$ is $N \times N$ where $N = 2^n$.

Rule: +1 if $\langle \text{row} \oplus \text{col} \rangle = 0$
Multiply $\text{row}_i \cdot \text{col}_i$ for all i
then add up mod 2.

Two Quantum Coordinates

In 4-space, the standard basis is given by the vectors:

$$e_0 = (1, 0, 0, 0), e_1 = (0, 1, 0, 0), e_2 = (0, 0, 1, 0), e_3 = (0, 0, 0, 1).$$

The indexing scheme for **quantum coordinates** changes the labels to come from $\{0, 1\}^2$ instead of from $\{1, 2, 3, 4\}$, using the canonical binary order 00, 01, 10, 11. Then we have:

$$e_{00} = (1, 0, 0, 0), e_{01} = (0, 1, 0, 0), e_{10} = (0, 0, 1, 0), e_{11} = (0, 0, 0, 1).$$

The big advantage is that these basis elements are all separable and the labels respect the tensor products involved:

$$\begin{aligned} |00\rangle &= e_{00} = (1, 0, 0, 0) = (1, 0) \otimes (1, 0) = e_0 \otimes e_0 = |0\rangle \otimes |0\rangle = |0\rangle|0\rangle \\ |01\rangle &= e_{01} = (0, 1, 0, 0) = (1, 0) \otimes (0, 1) = e_0 \otimes e_1 = |0\rangle \otimes |1\rangle = |0\rangle|1\rangle \\ |10\rangle &= e_{10} = (0, 0, 1, 0) = (0, 1) \otimes (1, 0) = e_1 \otimes e_0 = |1\rangle \otimes |0\rangle = |1\rangle|0\rangle \\ |11\rangle &= e_{11} = (0, 0, 0, 1) = (0, 1) \otimes (0, 1) = e_1 \otimes e_1 = |1\rangle \otimes |1\rangle = |1\rangle|1\rangle \end{aligned}$$

It is OK to picture the tensoring with row vectors, but because humanity chose to write matrix-vector products as Mv rather than vM , they need to be treated as column vectors. This will lead to cognitive dissonance when we read quantum circuits left-to-right but have to compose matrices right-to-left.

Contrast with Cartesian Products

In a calculus or linear algebra course you have likely encountered the spaces $\mathbb{R}^2$ of points (a, b) in the plane and $\mathbb{R}^3$ of points (x, y, z) or (x_1, x_2, x_3) in 3-dimensional space. Then $\mathbb{R}^n$ means n -dimensional real space, whether you called it a vector space or not. Maybe you also covered the complex vector spaces $\mathbb{C}^n$ or specialized to vectors of rational numbers---which make the vector space $\mathbb{Q}^n$. When we care more about the "space" aspect than the particular kind of numbers allowed, we use the umbrella term "Hilbert space" after the mathematician David Hilbert. That term is often employed by physicists not only to avoid having to specify the dimension n but also to allow it to be infinite. We, however, will stay in finite dimensional spaces and care a lot about what the dimension is.

The usual rule for the *product* of two vector spaces is to add the dimensions. Thus a member of $\mathbb{R}^2 \times \mathbb{R}^3$, which formally is an ordered pair like $((a, b), (x, y, z))$, is considered the same as the 5-tuple (a, b, x, y, z) , which we could re-label as $(x_1, x_2, x_3, x_4, x_5)$. So $\mathbb{R}^2 \times \mathbb{R}^3 = \mathbb{R}^5$.

The F25 lecture gave a concrete example here, doing $\mathbb{R}^3 \times \mathbb{R}^2 = \mathbb{R}^5$ and thus mapping pairs $(x, y, z), (a, b)$ to vectors (x, y, z, a, b) , which you might read as "flattening" the nested tuple $((x, y, z), (a, b))$. In the example, (x, y, z) were the position of a baseball on its way to home plate, a was the spin velocity in revolutions per minute, and b was the axis of spin---which a batter needs to read in order to anticipate how the pitched ball will curve. [The notion of "quantum spin" as one physical implementation of a **qubit** will be roughly similar but less numerical---it is "quantized" after all.] The point is that these are five separate numbers. If you include velocity toward home plate, you get a sixth number: $\mathbb{R}^3 \times \mathbb{R}^2 \times \mathbb{R}^1 = \mathbb{R}^6$. The dimensions add.

The **tensor product**, however, multiplies the dimensions. When defined between our vectors (a, b) and (x, y, z) , it doesn't just ram them together. Instead it combines a copy of the second vector into each part of the first vector. In symbols:

$$(a, b) \otimes (x, y, z) = (a \cdot (x, y, z), b \cdot (x, y, z)) = (ax, ay, az, bx, by, bz).$$

The vectors have dimension 6 not 5. Not every vector in the target space, here $\mathbb{R}^6$, arises as a tensor product. But if we close out $\mathbb{R}^2 \otimes \mathbb{R}^3$ under linear combinations, then we do get all of $\mathbb{R}^6$.

After a couple more examples on the whiteboard ([photo 4](#)), Tuesday's lecture ended here.

Abstract View of Tensor Products (note---from Chapter 14)

What does tensor product *do*? We feel again that good intuition comes by thinking about abstract attributes first, numbers later. Let us make the card suits into the abstract "attribute vector"

$$u = (\clubsuit, \diamondsuit, \heartsuit, \spadesuit).$$

And the ranks of cards becomes the attribute vector

$$v = (2, 3, 4, 5, 6, 7, 8, 9, 10, J, Q, K, A).$$

Then the tensor product---in the order $u \otimes v$ (it's not commutative)---is

$$w = (\clubsuit 2, \clubsuit 3, \dots, \clubsuit K, \clubsuit A, \diamondsuit 2, \diamondsuit 3, \dots, \diamondsuit A, \heartsuit 2, \dots, \heartsuit A, \spadesuit 2, \dots, \spadesuit A).$$

This sorts the deck by suits. If we tensored the other way around, we'd get

$$v \otimes u = (2\clubsuit, 2\diamondsuit, 2\heartsuit, 2\spadesuit, 3\clubsuit, 3\diamondsuit, 3\heartsuit, 3\spadesuit, 4\clubsuit, \dots, 4\spadesuit, 5\clubsuit, \dots, \dots, A\clubsuit, A\diamondsuit, A\heartsuit, A\spadesuit),$$

which sorts the deck by ranks instead. Always the second vector gets "copied inner" while the first vector is "outer." The ordinary product would have just rammed v after u to give a vector of 17 items, four of type `Suit` and thirteen of type `Rank`. This is inhomogeneous mishmash---like "not playing with a full deck" as we say. Whereas, in either order, the tensor product creates a homogeneous length-52 vector of type "Suit and Rank." (Between `Suit` and `Rank`, the order might or might not matter.)

The F25 lecture did a different example using only the suits $\spadesuit, \heartsuit, \diamondsuit$ and grouping 2--9 as "spot cards" **S** and 10,J,Q,K,A as "honor cards" **H**. This yielded just six combined basic attributes:

$$\spadesuit S, \spadesuit H, \heartsuit S, \heartsuit H, \diamondsuit S, \diamondsuit H.$$

A vector such as $[2, 3, 2, 0, 4, 1]$ means that you hold two spade spot cards, three spade honors, two little hearts but no heart honors, and four little diamonds with one diamond honor, (We're ignoring clubs.) The point of contrast is that it makes sense to combine *attributes*, such as "spade" and "honor card" and talk about numerical values for the combination. Back in the baseball example, it did not make sense to say that the ball had "x+spin" as a combined attribute---the "x" and "spin" are separate components. Another analogy that may be useful is that attributes are like *classes* in OOP, while individual vectors are *instances*. Nevertheless, Paul Dirac's great insight was:

Any attribute **A** can be represented by a standard basis vector $|A\rangle$ in some vector space big enough to house all the independent attributes we want to consider.

The $| \rangle$ by itself is called a **ket**, and presumes that we are writing the vector as a column vector.

Because column vectors are inconvenient in paragraph form, I will often write them in rows but use square brackets and superscript T for "transpose" to say they are really column vectors. (But sometimes the T will be omitted when the intent is clear.)

Now let's see how this works numerically. The vector $\mathbf{u} = [0, 1, 0, 0]^T$ is the standard basis vector corresponding to "diamonds" in our scheme for suits. For ranks, the seven is indexed by

$$\mathbf{v} = [0, 0, 0, 0, 0, 1, 0, 0, 0, 0, 0, 0]^T.$$

Then

$$\begin{aligned}\mathbf{u} \otimes \mathbf{v} &= [0 \cdot [0, 0, 0, 0, 0, 1, 0, 0, 0, 0, 0, 0], 1 \cdot [0, 0, 0, 0, 0, 1, 0, 0, 0, 0, 0, 0], \dots, 0]^T \\ &= [0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 1, 0, 0, 0, 0, 0, 0, 0, \dots, 0]^T.\end{aligned}$$

The single 1 corresponds to the position of the $\diamond 7$ in the abstract indexing scheme. So we write:

$$\mathbf{u} \otimes \mathbf{v} = |\diamond 7\rangle = |\diamond\rangle|7\rangle.$$

Thus the tensor product of two standard basis vectors gives us a standard basis vector in the larger space. Indeed, we can get the entire standard basis of 52 vectors this way. We've also started writing the "invisible dot" product of kets---which are in quantum coordinates---to stand for the tensor product in the underlying coordinates.

Suppose we next take the tensor product of w with itself. Doing this with the attribute vectors, we get

$$w \otimes w = (\clubsuit 2 \clubsuit 2, \clubsuit 2 \clubsuit 3, \clubsuit 2 \clubsuit 4, \dots, \clubsuit 2 \clubsuit A, \clubsuit 2 \diamond 2, \dots, \clubsuit 2 \spadesuit A, \clubsuit 3 \clubsuit 2, \clubsuit 3 \clubsuit 3, \dots, \dots, \spadesuit A \spadesuit A).$$

What does this represent? It conveys the idea of playing two cards in sequence. This allows them to be the same card---casinos usually play blackjack with eight decks shuffled together, for instance---but we will encounter algorithms whose point is either *amplifying* or eliminating such particular possibilities. The point is that tensor product is the underlying way of Nature simply doing a *sequence*, which means *concatenating* the symbolic representations.

Tensor Product = Simple Concatenation = Flow of Events.

This kind of representation is not just useful in quantum. It underlies the idea of the "[TensorFlow](#)" API and library in machine learning.

The common example that matters most is when both spaces have dimension 2. This case is innately confusing because $2 + 2 = 2 \cdot 2 = 2^2$. But hopefully the above will help us avoid confusion.

Another important note: We could encode the suits as $\clubsuit = 00$, $\diamond = 01$, $\heartsuit = 10$, and $\spadesuit = 11$. Then the concept of "black suit" could look like " $|00\rangle + |11\rangle$ " (divided by $\sqrt{2}$ to normalize). Below we will call this our "basic entangled state"---but that is when we are modeling two separate physical entities,

namely two qubits. Here we have only one entity: the *suit* of a *card*, so we don't have separate objects to entangle. We don't say "clubs are entangled with spades."

Now Back to Qubits

We can also take tensor products of non-basis vectors of length 2. Let's us try

$$\mathbf{u} = \frac{e_0 + e_1}{\sqrt{2}} = \sqrt{0.5} [1, 1]^T.$$

When we do $\mathbf{u} \otimes \mathbf{u}$, the first thing that happens is that the scalars in front multiply to get $\sqrt{0.5} \cdot \sqrt{0.5} = \frac{1}{2}$ as the multiplier on the whole thing. The vector bodies combine as

$$[1 \cdot [1, 1], 1 \cdot [1, 1]] = [1, 1, 1, 1].$$

(Strictly speaking, we should do this as column vectors---maybe we'll show on the blackboard---but it's always fine to do as row vectors and remember to transpose when needed at the end.) So

$$\mathbf{u} \otimes \mathbf{u} = \left[\frac{1}{2}, \frac{1}{2}, \frac{1}{2}, \frac{1}{2} \right]^T.$$

This is a unit vector. We can do the same with

$$\mathbf{v} = \frac{e_0 - e_1}{\sqrt{2}} = \sqrt{0.5} [1, -1]^T$$

We get $\mathbf{v} \otimes \mathbf{v} = \frac{1}{2}$ times $[1, -1] \otimes [1, -1] = \frac{1}{2} [1, -1, -1, 1]$. OK, $\frac{1}{2} [1, -1, -1, 1]^T$ to be strict. We also get:

$$\mathbf{u} \otimes \mathbf{v} = \frac{1}{2} [1, -1, 1, -1]^T.$$

$$\mathbf{v} \otimes \mathbf{u} = \frac{1}{2} [1, 1, -1, -1]^T.$$

Suppose I instead said

$$\mathbf{u} \otimes \mathbf{v} = \frac{1}{2}[-1, 1, -1, 1]^T?$$

Would I be wrong? Numerically yes: $\mathbf{u} \otimes \mathbf{v}$ does not equal $-\mathbf{u} \otimes \mathbf{v}$. But Nature's "Quantum Rose" does not care. The principle involved will be highlighted in chapter 5 but let's note it now:

Multiplying a vector (one that represents a quantum state) by -1 , or by i or $-i$, or by any other *scalar* of unit magnitude, does not change the quantum state that it represents.

Nature *may*, however, care about the difference between $\mathbf{u} \otimes \mathbf{v}$ and $\mathbf{v} \otimes \mathbf{u}$.

Let's ignore the normalizing multipliers out front for the moment, since they do not matter to the ability to combine the vector bodies. How about the simpler vector

$$\mathbf{w} = [1, 0, 0, 1]?$$

This equals $e_{00} + e_{11}$, i.e., $|00\rangle + |11\rangle$, ignoring the normalizing constant $\sqrt{0.5}$. Can we get this as a tensor product of two vectors of length 2?

The answer is no. We can prove it by representing the general 2-by-2 tensor product as

$$[a, b] \otimes [c, d] = [ac, ad, bc, bd].$$

To get $[1, 0, 0, 1]$ as the result, we need to solve the equations

$$ac = 1, ad = 0, bc = 0, \text{ and } bd = 1.$$

But $ad = 0$ entails that either a is 0 or d is 0. If $a = 0$, then $ac = 1$ is impossible. But if $d = 0$, then $bd = 1$ is impossible. So there is no solution.

Definition. A vector is **separable** if it can be written as the tensor product of two smaller vectors. Otherwise---and especially when the vector represents a quantum state---we call it **entangled**.

To introduce some more quantum terminology, when a unit vector is not a basis vector, it is necessarily a linear combination of two or more basis vectors. Then it is a **superposition**. One of the amazing verities of physics is that we really can put particle-level sustems into superpositions and interact with them. The math of how those interactions behave involves the kind of vectors we are already seeing. The vectors $\mathbf{u}$ and $\mathbf{v}$ above are superpositions. When we re-interpret the attributes $|0\rangle$ and $|1\rangle$ as $|dead\rangle$ and $|alive\rangle$, then $\mathbf{u}$ becomes the superposition

$$\frac{|dead\rangle + |alive\rangle}{\sqrt{2}}$$

This is said to be the state of **Schrödinger's Cat**. The philosophical issue is whether a macro-level being, not a particle, can be put into superposition. (We will later argue the answer is "yes...but...") For now we prefer to take the simple realist view that a particle can have the state $\mathbf{u}$. It is not a cat, but it has a pet-name, indeed a ket-name: $|+\rangle$. The vector $\mathbf{v}$, with its prominent minus sign, is called $|-\rangle$.

Tensoring the cat with another cat, as above with $\mathbf{u} \otimes \mathbf{u}$, gave us $\frac{e_{00} + e_{01} + e_{10} + e_{11}}{2}$. Are the cats entangled? No---they are separable, because we got this as a tensor product to begin with. But the vector $\frac{e_{00} + e_{11}}{\sqrt{2}}$ does represent the state of two cats that are either both dead or both alive.

Those cats---in the same Schrödinger box---are entangled in ways more subtle than if they were both playing with the same ball of yarn.

Inner Products

The **dot-product** of two equal-length vectors $\mathbf{a} = (a_1, \dots, a_n)$ and $\mathbf{b} = (b_1, \dots, b_n)$ is always defined by

$$\mathbf{a} \cdot \mathbf{b} = a_1 b_1 + \dots + a_n b_n.$$

(I actually want to write a bigger solid dot, and our textbook does so, but MathCha doesn't have it.) We have already seen the example of the dot product of binary strings treated as vectors with entries modulo 2, which is called the "Boolean inner product" on page 12 of Chapter 2.

When $\mathbf{a}$ and $\mathbf{b}$ are vectors of real-number entries, the **real inner product** is the same as the dot-product. Besides $\mathbf{a} \cdot \mathbf{b}$, the most common notation for it is $\langle \mathbf{a}, \mathbf{b} \rangle$.

When $\mathbf{a}$ and $\mathbf{b}$ are vectors over the *complex* numbers, however, the entries of the *first* vector are complex-conjugated. The common angle-bracket notation remains the same:

$$\langle \mathbf{a}, \mathbf{b} \rangle = \bar{a}_1 b_1 + \dots + \bar{a}_n b_n.$$

Why do we have to conjugate? Well, the inner product of a vector $\mathbf{a} = (a_1, \dots, a_n)$ with itself is supposed to be the square of its Euclidean length, which is the sum of the squared magnitudes of its components. If a_1 is a real number, then a_1^2 gives its squared magnitude even when a_1 is negative. But if a_1 is complex, we need in general to do $\bar{a}_1 \cdot a_1 = |a_1|^2$.

This should not cause confusion with regard to the real-number case: if the entries a_i all happen to be real numbers, their conjugates don't change, so we get $a_1 b_1 + \dots + a_n b_n$ anyway. But don't forget to

conjugate when they really are complex numbers! In some cases it will also matter that we conjugate the "a" part---not the "b" part. Dirac notation helps keep this straight because if you "break apart the bracket" to make $\langle \mathbf{a} | \cdot | \mathbf{b} \rangle$, the **bra** $\langle \mathbf{a} |$ denotes a conjugated row vector.

The "twist" of conjugating the a_i entries is thus more fundamental. Any, this defines the "standard" inner product on a finite-dimensional real or complex vector space. (If you're curious about the general definition of when a vector space V becomes a **Hilbert Space**, it is when there is an abstract product $\cdot$ such that whenever you have a sequence $\mathbf{v}_1, \mathbf{v}_2, \mathbf{v}_3, \dots$ of vectors in the space and a vector $\mathbf{w}$ (that could be "floating outside") such that the scalar values $\mathbf{v}_1 \cdot \mathbf{w}, \mathbf{v}_2 \cdot \mathbf{w}, \mathbf{v}_3 \cdot \mathbf{w}, \dots$ go to zero, the vector $\mathbf{w}$ in fact belongs to V .)

Rather than messily write the conjugate transpose of a vector $\mathbf{a}$ as $\bar{\mathbf{a}}^T$, we write it as $\mathbf{a}^*$, where the superscript star can be read as "star", "adjoint", or "dual" depending on the context. (Word-to-the-wise, however: newer sources are using a superscript "dagger" instead: $\mathbf{a}^\dagger$. I avoid the dagger in handwriting because it can look either like a plus sign or like a superscript T for transpose without conjugating.) Now Dirac's inspiration was that we can break the angled "bracket" into two pieces at the vertical bar:

$$\langle \mathbf{a} | \mathbf{b} \rangle = \langle \mathbf{a} | \cdot | \mathbf{b} \rangle = \mathbf{a}^* \cdot \mathbf{b}$$

where the dot is the same as our dot product and ordinary "matrix" multiplication. He called the two parts the "bra" and the "ket". Thus the bra $\langle \mathbf{a} |$ is another way of writing $\mathbf{a}^*$. (And when a is a scalar, a^* is another way of writing $\bar{a}$ since transpose is immaterial.) It comes IMHO with a different philosophy: instead of the "pristine" vector $\mathbf{a}$ being changed by conjugation, it is "moved into dual position." It is considered OK to use " $\mathbf{a}$ " and " $|\mathbf{a}\rangle$ " as synonymous notations, although we've already used the ket notation as wrapper around an *attribute*, to signify the basis state indexing that attribute.

Well, keeping things simple is why Part I of the text avoids the Dirac notation, but I think conveying Dirac's insight will help tie our various "product" ideas together: We can write bras and kets together in products any which way. If we write

$$|\mathbf{a}\rangle \cdot |\mathbf{b}\rangle$$

then this is the tensor product of the two vectors. We've already been doing this with our standard basis states of binary-string attributes under big-endian notation, e.g. (making the dot $\cdot$ invisible):

$$\begin{aligned} |0\rangle|1\rangle &= |0\rangle \otimes |1\rangle = e_0 \otimes e_1 = e_{01} = |01\rangle \\ |10\rangle|11\rangle &= |10\rangle \otimes |11\rangle = e_{10} \otimes e_{11} = e_{1011} = |1011\rangle \end{aligned}$$

And then $\langle \mathbf{a} | \cdot \langle \mathbf{b} |$ is the tensor product of the corresponding row vectors and works out the same as $(|\mathbf{a}\rangle \cdot |\mathbf{b}\rangle)^*$ owing to the scalar conjugation rule $\bar{\bar{a}} \cdot \bar{\bar{b}} = \overline{ab}$ applied to each entry. Finally,

$$|\mathbf{a}\rangle \cdot \langle \mathbf{b}|$$

can be figured out as multiplying an $\ell \times 1$ column vector by a (conjugated) $1 \times n$ row vector. You get an $\ell \times n$ matrix. Here is a sketch picture:



It actually has the same entries $a_i b_j^*$ as $\mathbf{a} \otimes \mathbf{b}^*$, but rolled into a matrix rather than left as a longer vector. Especially when $\ell = n$, this is called the **outer-product** matrix. And the final secret is that when $\mathbf{b} = \mathbf{a}$, the outer-product matrix

$$\mathbf{A} = |\mathbf{a}\rangle \cdot \langle \mathbf{a}|$$

captures all the ways that the quantum state $|\mathbf{a}\rangle$ can interact with the outside world. Most in particular, when we take the ordinary matrix product of $\mathbf{A}$ with another quantum state $|\mathbf{c}\rangle$ of compatible size, the associativity of basic dot multiplication gives us:

$$\mathbf{A} \cdot \mathbf{c} = (|\mathbf{a}\rangle \cdot \langle \mathbf{a}|) \cdot |\mathbf{c}\rangle = |\mathbf{a}\rangle \cdot (\langle \mathbf{a}|\mathbf{c}\rangle),$$

which is just the original vector $\mathbf{a}$ multiplied by the scalar value $\langle \mathbf{a}|\mathbf{c}\rangle$ from the inner product. The essence $|\mathbf{a}\rangle \cdot \langle \mathbf{a}|$ has a neat symmetry that maybe requites the concern about making a left-handed versus right-handed choice with conjugation. Maybe Nature is "evenhanded" after all. Moreover, this shows that all of our products:

- "ordinary" matrix product
- inner product
- outer product, and
- tensor product

are really **fungible**---and belong to a reality where we can climb up a "stairway to heaven" of higher dimensions and back down again, so that maybe dimensionality is not fundamental. But the high-volume ThePhysicsMemes Twitter site had [the last word](#).

Unit Vectors

Back on ground level, note that for a complex scalar c , c^*c and cc^* both equal its squared magnitude $|c|^2$. To wit, if we write $c = a + bi$ with a and b real, then $c^* = a - bi$, and we get

$$cc^* = (a + bi)(a - bi) = a^2 - abi + abi + b^2 = a^2 + b^2 = |c|^2.$$

And for a vector $\mathbf{a}$,

$$\langle \mathbf{a} | \mathbf{a} \rangle = \sum_{i=1}^n a_i^* a_i = \sum_{i=1}^n |a_i|^2$$

is its **squared magnitude**. The **norm** of the vector is the positive square root of this:

$$|\mathbf{a}| = \sqrt{\langle \mathbf{a} | \mathbf{a} \rangle}.$$

Definition: A **unit vector** has norm 1, equivalently, has squared magnitude 1. Any unit vector $\mathbf{v}$ in a Hilbert space gives rise to a "**legal quantum state**", which we can write as $|\mathbf{v}\rangle$ for pretentious emphasis (or decide not to do so).

In the two-dimensional real space $\mathbb{R}^2$, the unit vectors (x, y) correspond to the points on the unit circle in the plane, as defined by $x^2 + y^2 = 1$. The simplest quantum states are unit vectors of two *complex* numbers, i.e., members of $\mathbb{C}^2$. But in many cases we can ignore the distinction between complex and real entries and pretend-visualize them as points on the unit circle in the plane anyway. (There is a non-fudged, non-lossy visualization in 3D called the **Bloch Sphere**---we will come to it later.)

Now for operations, we begin by extending some notation we just used for vectors:

Definition. The **conjugate transpose** of an $m \times n$ matrix A is obtained by first transposing A to make the $n \times m$ matrix A^T , and then taking the complex conjugate of every element. We write A^* for the resulting $n \times m$ matrix. (Many other sources write $A^\dagger$ instead.)

Examples:

- If $A = \begin{bmatrix} 1+i & 1-i \\ 2 & 0 \end{bmatrix}$, then $A^T = \begin{bmatrix} 1+i & 2 \\ 1-i & 0 \end{bmatrix}$ and $A^* = \begin{bmatrix} 1-i & 2 \\ 1+i & 0 \end{bmatrix}$.
- If $\mathbf{Y} = \begin{bmatrix} 0 & -i \\ i & 0 \end{bmatrix}$, then $\mathbf{Y}^T = \begin{bmatrix} 0 & i \\ -i & 0 \end{bmatrix}$, and weirdly, $\mathbf{Y}^* = \begin{bmatrix} 0 & -i \\ i & 0 \end{bmatrix} = \mathbf{Y}$ back again.
- For our old friend the Hadamard matrix $\mathbf{H} = \frac{1}{\sqrt{2}} \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix}$, $\mathbf{H}^T = \mathbf{H}^* = \mathbf{H}$ because $\mathbf{H}$ is symmetric and has all-real entries.

Note that

$$\mathbf{H}e_0 = \mathbf{H}|0\rangle = \frac{1}{\sqrt{2}} \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix} \begin{bmatrix} 1 \\ 0 \end{bmatrix} = \frac{1}{\sqrt{2}} \begin{bmatrix} 1 \\ 1 \end{bmatrix} = \frac{e_0 + e_1}{\sqrt{2}},$$

a vector we saw in the previous lecture, and ditto for

$$\mathbf{H}e_1 = \mathbf{H}|1\rangle = \frac{1}{\sqrt{2}} \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix} \begin{bmatrix} 0 \\ 1 \end{bmatrix} = \frac{1}{\sqrt{2}} \begin{bmatrix} 1 \\ -1 \end{bmatrix} = \frac{e_0 - e_1}{\sqrt{2}}.$$

The two resulting vectors are also unit vectors, and they are linearly independent. Indeed, they are **orthogonal**, because

$$\langle [1, 1], [1, -1] \rangle = 1 \cdot 1 + 1 \cdot (-1) = 0.$$

Thus, like e_0 and e_1 , they form an **orthonormal basis** for $\mathbb{C}^2$ (not to mention also for $\mathbb{R}^2$). Focusing on the plus and minus sign between e_0 and e_1 , the "Dirac names" for these states are $|+\rangle$ and $|-\rangle$, respectively.