

Therefore,

$$\int_{\mathbf{a}}^{\mathbf{b}} \mathbf{E} \cdot d\mathbf{l} = \frac{1}{4\pi\epsilon_0} \int_{\mathbf{a}}^{\mathbf{b}} \frac{q}{r^2} dr = \frac{-1}{4\pi\epsilon_0} \frac{q}{r} \Big|_{r_a}^{r_b} = \frac{1}{4\pi\epsilon_0} \left(\frac{q}{r_a} - \frac{q}{r_b} \right), \quad (2.18)$$

where r_a is the distance from the origin to the point $\mathbf{a}$ and r_b is the distance to $\mathbf{b}$. The integral around a *closed* path is evidently zero (for then $r_a = r_b$):

$$\oint \mathbf{E} \cdot d\mathbf{l} = 0, \quad (2.19)$$

and hence, applying Stokes' theorem,

$$\nabla \times \mathbf{E} = \mathbf{0}. \quad (2.20)$$

Now, I proved Eqs. 2.19 and 2.20 only for the field of a single point charge at the origin, but these results make no reference to what is, after all, a perfectly arbitrary choice of coordinates; they hold no matter *where* the charge is located. Moreover, if we have many charges, the principle of superposition states that the total field is a vector sum of their individual fields:

$$\mathbf{E} = \mathbf{E}_1 + \mathbf{E}_2 + \dots,$$

so

$$\nabla \times \mathbf{E} = \nabla \times (\mathbf{E}_1 + \mathbf{E}_2 + \dots) = (\nabla \times \mathbf{E}_1) + (\nabla \times \mathbf{E}_2) + \dots = \mathbf{0}.$$

Thus, Eqs. 2.19 and 2.20 hold for *any static charge distribution whatever*.

Problem 2.19 Calculate $\nabla \times \mathbf{E}$ directly from Eq. 2.8, by the method of Sect. 2.2.2. Refer to Prob. 1.63 if you get stuck.

2.3 ■ ELECTRIC POTENTIAL

2.3.1 ■ Introduction to Potential

The electric field $\mathbf{E}$ is not just *any* old vector function. It is a very special *kind* of vector function: one whose curl is zero. $\mathbf{E} = y\hat{\mathbf{x}}$, for example, could not possibly be an electrostatic field; *no* set of charges, regardless of their sizes and positions, could ever produce such a field. We're going to exploit this special property of electric fields to reduce a *vector* problem (finding $\mathbf{E}$) to a much simpler *scalar* problem. The first theorem in Sect. 1.6.2 asserts that any vector whose curl is zero is equal to the gradient of some scalar. What I'm going to do now amounts to a proof of that claim, in the context of electrostatics.

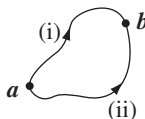


FIGURE 2.30

Because $\nabla \times \mathbf{E} = \mathbf{0}$, the line integral of $\mathbf{E}$ around any closed loop is zero (that follows from Stokes' theorem). Because $\oint \mathbf{E} \cdot d\mathbf{l} = 0$, the line integral of $\mathbf{E}$ from point $\mathbf{a}$ to point $\mathbf{b}$ is the same for all paths (otherwise you could go *out* along path (i) and return along path (ii)—Fig. 2.30—and obtain $\oint \mathbf{E} \cdot d\mathbf{l} \neq 0$). Because the line integral is independent of path, we can define a function⁶

$$V(\mathbf{r}) \equiv - \int_{\mathcal{O}}^{\mathbf{r}} \mathbf{E} \cdot d\mathbf{l}. \quad (2.21)$$

Here $\mathcal{O}$ is some standard reference point on which we have agreed beforehand; V then depends only on the point $\mathbf{r}$. It is called the **electric potential**.

The potential *difference* between two points $\mathbf{a}$ and $\mathbf{b}$ is

$$\begin{aligned} V(\mathbf{b}) - V(\mathbf{a}) &= - \int_{\mathcal{O}}^{\mathbf{b}} \mathbf{E} \cdot d\mathbf{l} + \int_{\mathcal{O}}^{\mathbf{a}} \mathbf{E} \cdot d\mathbf{l} \\ &= - \int_{\mathcal{O}}^{\mathbf{b}} \mathbf{E} \cdot d\mathbf{l} - \int_{\mathbf{a}}^{\mathcal{O}} \mathbf{E} \cdot d\mathbf{l} = - \int_{\mathbf{a}}^{\mathbf{b}} \mathbf{E} \cdot d\mathbf{l}. \end{aligned} \quad (2.22)$$

Now, the fundamental theorem for gradients states that

$$V(\mathbf{b}) - V(\mathbf{a}) = \int_{\mathbf{a}}^{\mathbf{b}} (\nabla V) \cdot d\mathbf{l},$$

so

$$\int_{\mathbf{a}}^{\mathbf{b}} (\nabla V) \cdot d\mathbf{l} = - \int_{\mathbf{a}}^{\mathbf{b}} \mathbf{E} \cdot d\mathbf{l}.$$

Since, finally, this is true for *any* points $\mathbf{a}$ and $\mathbf{b}$, the integrands must be equal:

$$\mathbf{E} = -\nabla V. \quad (2.23)$$

⁶To avoid any possible ambiguity, I should perhaps put a prime on the integration variable:

$$V(\mathbf{r}) = - \int_{\mathcal{O}}^{\mathbf{r}} \mathbf{E}(\mathbf{r}') \cdot d\mathbf{l}'.$$

But this makes for cumbersome notation, and I prefer whenever possible to reserve the primes for source points. However, when (as in Ex. 2.7) we calculate such integrals explicitly, I will put in the primes.

Equation 2.23 is the differential version of Eq. 2.21; it says that the electric field is the gradient of a scalar potential, which is what we set out to prove.

Notice the subtle but crucial role played by path independence (or, equivalently, the fact that $\nabla \times \mathbf{E} = \mathbf{0}$) in this argument. If the line integral of $\mathbf{E}$ depended on the path taken, then the “definition” of V , Eq. 2.21, would be nonsense. It simply would not define a function, since changing the path would alter the value of $V(\mathbf{r})$. By the way, don’t let the minus sign in Eq. 2.23 distract you; it carries over from Eq. 2.21 and is largely a matter of convention.

Problem 2.20 One of these is an impossible electrostatic field. Which one?

(a) $\mathbf{E} = k[xy \hat{\mathbf{x}} + 2yz \hat{\mathbf{y}} + 3xz \hat{\mathbf{z}}];$

(b) $\mathbf{E} = k[y^2 \hat{\mathbf{x}} + (2xy + z^2) \hat{\mathbf{y}} + 2yz \hat{\mathbf{z}}].$

Here k is a constant with the appropriate units. For the *possible* one, find the potential, using the *origin* as your reference point. Check your answer by computing ∇V . [Hint: You must select a specific path to integrate along. It doesn’t matter *what* path you choose, since the answer is path-independent, but you simply cannot integrate unless you have a definite path in mind.]

2.3.2 ■ Comments on Potential

(i) The name. The word “potential” is a hideous misnomer because it inevitably reminds you of potential *energy*. This is particularly insidious, because there *is* a connection between “potential” and “potential energy,” as you will see in Sect. 2.4. I’m sorry that it is impossible to escape this word. The best I can do is to insist once and for all that “potential” and “potential energy” are completely different terms and should, by all rights, have different names. Incidentally, a surface over which the potential is constant is called an **equipotential**.

(ii) Advantage of the potential formulation. If you know V , you can easily get $\mathbf{E}$ —just take the gradient: $\mathbf{E} = -\nabla V$. This is quite extraordinary when you stop to think about it, for $\mathbf{E}$ is a *vector* quantity (three components), but V is a *scalar* (one component). How can *one* function possibly contain all the information that *three* independent functions carry? The answer is that the three components of $\mathbf{E}$ are not really as independent as they look; in fact, they are explicitly interrelated by the very condition we started with, $\nabla \times \mathbf{E} = \mathbf{0}$. In terms of components,

$$\frac{\partial E_x}{\partial y} = \frac{\partial E_y}{\partial x}, \quad \frac{\partial E_z}{\partial y} = \frac{\partial E_y}{\partial z}, \quad \frac{\partial E_x}{\partial z} = \frac{\partial E_z}{\partial x}.$$

This brings us back to my observation at the beginning of Sect. 2.3.1: $\mathbf{E}$ is a *very special kind of vector*. What the potential formulation does is to exploit this feature to maximum advantage, reducing a vector problem to a scalar one, in which there is no need to fuss with components.

(iii) **The reference point $\mathcal{O}$.** There is an essential ambiguity in the definition of potential, since the choice of reference point $\mathcal{O}$ was arbitrary. Changing reference points amounts to adding a constant K to the potential:

$$V'(\mathbf{r}) = - \int_{\mathcal{O}'}^{\mathbf{r}} \mathbf{E} \cdot d\mathbf{l} = - \int_{\mathcal{O}'}^{\mathcal{O}} \mathbf{E} \cdot d\mathbf{l} - \int_{\mathcal{O}}^{\mathbf{r}} \mathbf{E} \cdot d\mathbf{l} = K + V(\mathbf{r}),$$

where K is the line integral of $\mathbf{E}$ from the old reference point $\mathcal{O}$ to the new one $\mathcal{O}'$. Of course, adding a constant to V will not affect the potential *difference* between two points:

$$V'(\mathbf{b}) - V'(\mathbf{a}) = V(\mathbf{b}) - V(\mathbf{a}),$$

since the K 's cancel out. (Actually, it was already clear from Eq. 2.22 that the potential difference is independent of $\mathcal{O}$, because it can be written as the line integral of $\mathbf{E}$ from $\mathbf{a}$ to $\mathbf{b}$, with no reference to $\mathcal{O}$.) Nor does the ambiguity affect the gradient of V :

$$\nabla V' = \nabla V,$$

since the derivative of a constant is zero. That's why all such V 's, differing only in their choice of reference point, correspond to the same field $\mathbf{E}$.

Potential as such carries no real physical significance, for at any given point we can adjust its value at will by a suitable relocation of $\mathcal{O}$. In this sense, it is rather like altitude: If I ask you how high Denver is, you will probably tell me its height above sea level, because that is a convenient and traditional reference point. But we could as well agree to measure altitude above Washington, D.C., or Greenwich, or wherever. That would add (or, rather, subtract) a fixed amount from all our sea-level readings, but it wouldn't change anything about the real world. The only quantity of intrinsic interest is the *difference* in altitude between two points, and *that* is the same *whatever* your reference level.

Having said this, however, there *is* a “natural” spot to use for $\mathcal{O}$ in electrostatics—analogueous to sea level for altitude—and that is a point infinitely far from the charge. Ordinarily, then, we “set the zero of potential at infinity.” (Since $V(\mathcal{O}) = 0$, choosing a reference point is equivalent to selecting a place where V is to be zero.) But I must warn you that there is one special circumstance in which this convention fails: when the charge distribution itself extends to infinity. The symptom of trouble, in such cases, is that the potential blows up. For instance, the field of a uniformly charged plane is $(\sigma/2\epsilon_0)\hat{\mathbf{n}}$, as we found in Ex. 2.5; if we naïvely put $\mathcal{O} = \infty$, then the potential at height z above the plane becomes

$$V(z) = - \int_{\infty}^z \frac{1}{2\epsilon_0} \sigma \, dz = - \frac{1}{2\epsilon_0} \sigma (z - \infty).$$

The remedy is simply to choose some other reference point (in this example you might use a point on the plane). Notice that the difficulty occurs only in textbook problems; in “real life” there is no such thing as a charge distribution that goes on forever, and we can *always* use infinity as our reference point.

(iv) Potential obeys the superposition principle. The original superposition principle pertains to the force on a test charge Q . It says that the total force on Q is the vector sum of the forces attributable to the source charges individually:

$$\mathbf{F} = \mathbf{F}_1 + \mathbf{F}_2 + \dots$$

Dividing through by Q , we see that the electric field, too, obeys the superposition principle:

$$\mathbf{E} = \mathbf{E}_1 + \mathbf{E}_2 + \dots$$

Integrating from the common reference point to $\mathbf{r}$, it follows that the potential also satisfies such a principle:

$$V = V_1 + V_2 + \dots$$

That is, the potential at any given point is the sum of the potentials due to all the source charges separately. Only this time it is an *ordinary* sum, not a *vector* sum, which makes it a lot easier to work with.

(v) Units of Potential. In our units, force is measured in newtons and charge in coulombs, so electric fields are in newtons per coulomb. Accordingly, potential is newton-meters per coulomb, or joules per coulomb. A joule per coulomb is a **volt**.

Example 2.7. Find the potential inside and outside a spherical shell of radius R (Fig. 2.31) that carries a uniform surface charge. Set the reference point at infinity.

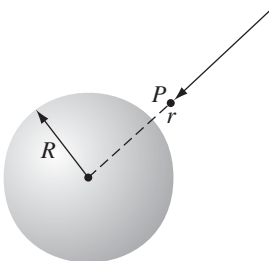


FIGURE 2.31

Solution

From Gauss's law, the field outside is

$$\mathbf{E} = \frac{1}{4\pi\epsilon_0} \frac{q}{r^2} \hat{\mathbf{r}},$$

where q is the total charge on the sphere. The field inside is zero. For points outside the sphere ($r > R$),

$$V(r) = - \int_{\infty}^r \mathbf{E} \cdot d\mathbf{l} = \frac{-1}{4\pi\epsilon_0} \int_{\infty}^r \frac{q}{r'^2} dr' = \frac{1}{4\pi\epsilon_0} \frac{q}{r'} \Big|_{\infty}^r = \frac{1}{4\pi\epsilon_0} \frac{q}{r}.$$

To find the potential inside the sphere ($r < R$), we must break the integral into two pieces, using in each region the field that prevails there:

$$V(r) = \frac{-1}{4\pi\epsilon_0} \int_{\infty}^R \frac{q}{r'^2} dr' - \int_R^r (0) dr' = \frac{1}{4\pi\epsilon_0} \frac{q}{r'} \Big|_{\infty}^R + 0 = \frac{1}{4\pi\epsilon_0} \frac{q}{R}.$$

Notice that the potential is *not* zero inside the shell, even though the field is. V is a *constant* in this region, to be sure, so that $\nabla V = \mathbf{0}$ —that’s what matters. In problems of this type, you must always *work your way in from the reference point*; that’s where the potential is “nailed down.” It is tempting to suppose that you could figure out the potential inside the sphere on the basis of the field there alone, but this is false: The potential inside the sphere is sensitive to what’s going on outside the sphere as well. If I placed a second uniformly charged shell out at radius $R' > R$, the potential inside R would change, even though the field would still be zero. Gauss’s law guarantees that charge exterior to a given point (that is, at larger r) produces no net *field* at that point, provided it is spherically or cylindrically symmetric, but there is no such rule for *potential*, when infinity is used as the reference point.

Problem 2.21 Find the potential inside and outside a uniformly charged solid sphere whose radius is R and whose total charge is q . Use infinity as your reference point. Compute the gradient of V in each region, and check that it yields the correct field. Sketch $V(r)$.

Problem 2.22 Find the potential a distance s from an infinitely long straight wire that carries a uniform line charge λ . Compute the gradient of your potential, and check that it yields the correct field.

Problem 2.23 For the charge configuration of Prob. 2.15, find the potential at the center, using infinity as your reference point.

Problem 2.24 For the configuration of Prob. 2.16, find the potential difference between a point on the axis and a point on the outer cylinder. Note that it is not necessary to commit yourself to a particular reference point, if you use Eq. 2.22.

2.3.3 ■ Poisson’s Equation and Laplace’s Equation

We found in Sect. 2.3.1 that the electric field can be written as the gradient of a scalar potential.

$$\mathbf{E} = -\nabla V.$$

The question arises: What do the divergence and curl of $\mathbf{E}$,

$$\nabla \cdot \mathbf{E} = \frac{\rho}{\epsilon_0} \quad \text{and} \quad \nabla \times \mathbf{E} = \mathbf{0},$$