

field implies that its curl is zero. See [Appendix F](#) for a discussion of the various vector operators in different coordinate systems.

2.1 Line integral of the electric field

Suppose that $\mathbf{E}$ is the field of some stationary distribution of electric charges. Let P_1 and P_2 denote two points anywhere in the field. The line integral of E between the two points is $\int_{P_1}^{P_2} \mathbf{E} \cdot d\mathbf{s}$, taken along some path that runs from P_1 to P_2 , as shown in [Fig. 2.1](#). This means: divide the chosen path into short segments, each segment being represented by a vector connecting its ends; take the scalar product of the path-segment vector with the field $\mathbf{E}$ at that place; add these products up for the whole path. The integral as usual is to be regarded as the limit of this sum as the segments are made shorter and more numerous without limit.

Let's consider the field of a point charge q and some paths running from point P_1 to point P_2 in that field. Two different paths are shown in [Fig. 2.2](#). It is easy to compute the line integral of $\mathbf{E}$ along path A, which is made up of a radial segment running outward from P_1 and an arc of

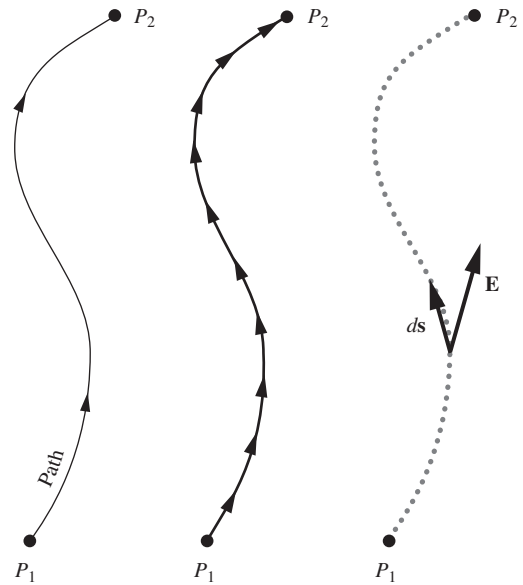


Figure 2.1.
Showing the division of the path into path elements ds .

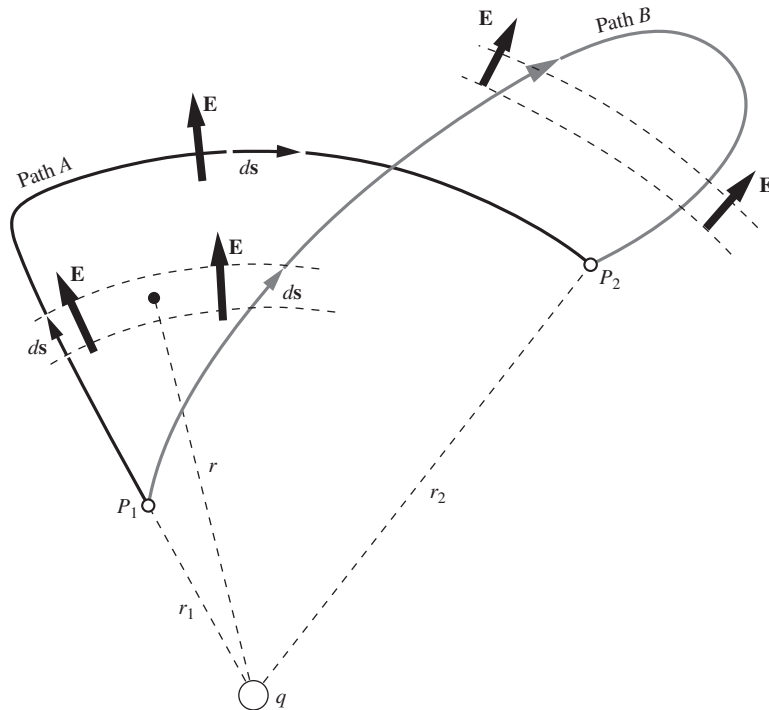


Figure 2.2.
The electric field $\mathbf{E}$ is that of a positive point charge q . The line integral of $\mathbf{E}$ from P_1 to P_2 along path A has the value $(q/4\pi\epsilon_0)(1/r_1 - 1/r_2)$. It will have exactly the same value if calculated for path B, or for any other path from P_1 to P_2 .

radius r_2 . Along the radial segment of path A, $\mathbf{E}$ and $d\mathbf{s}$ are parallel, the magnitude of $\mathbf{E}$ is $q/4\pi\epsilon_0 r^2$, and $\mathbf{E} \cdot d\mathbf{s}$ is simply $(q/4\pi\epsilon_0 r^2) ds$. Thus the line integral on that segment is

$$\int_{r_1}^{r_2} \frac{q dr}{4\pi\epsilon_0 r^2} = \frac{q}{4\pi\epsilon_0} \left(\frac{1}{r_1} - \frac{1}{r_2} \right). \quad (2.1)$$

The second leg of path A, the circular segment, gives zero because $\mathbf{E}$ is perpendicular to $d\mathbf{s}$ everywhere on that arc. The entire line integral is therefore

$$\int_{P_1}^{P_2} \mathbf{E} \cdot d\mathbf{s} = \frac{q}{4\pi\epsilon_0} \left(\frac{1}{r_1} - \frac{1}{r_2} \right). \quad (2.2)$$

Now look at path B. Because $\mathbf{E}$ is radial with magnitude $q/4\pi\epsilon_0 r^2$, $\mathbf{E} \cdot d\mathbf{s} = (q/4\pi\epsilon_0 r^2) dr$ even when $d\mathbf{s}$ is not radially oriented. The corresponding pieces of path A and path B indicated in the diagram make identical contributions to the integral. The part of path B that loops beyond r_2 makes a net contribution of zero; contributions from corresponding outgoing and incoming parts cancel. For the entire line integral, path B will give the same result as path A. As there is nothing special about path B, Eq. (2.1) must hold for *any* path running from P_1 to P_2 .

Here we have essentially repeated, in different language, the argument in Section 1.5, illustrated in Fig. 1.5, concerning the work done in moving one point charge near another. But now we are interested in the total electric field produced by any distribution of charges. One more step will bring us to an important conclusion. The line integral of the sum of fields equals the sum of the line integrals of the fields calculated separately. Or, stated more carefully, if $\mathbf{E} = \mathbf{E}_1 + \mathbf{E}_2 + \dots$, then

$$\int_{P_1}^{P_2} \mathbf{E} \cdot d\mathbf{s} = \int_{P_1}^{P_2} \mathbf{E}_1 \cdot d\mathbf{s} + \int_{P_1}^{P_2} \mathbf{E}_2 \cdot d\mathbf{s} + \dots, \quad (2.3)$$

where the same path is used for all the integrations. Now any electrostatic field can be regarded as the sum of a number (possibly enormous) of point-charge fields, as expressed in Eq. (1.20) or Eq. (1.22). Therefore if the line integral from P_1 to P_2 is independent of path for each of the point-charge fields $\mathbf{E}_1, \mathbf{E}_2, \dots$, the total field $\mathbf{E}$ must have this property:

The line integral $\int_{P_1}^{P_2} \mathbf{E} \cdot d\mathbf{s}$ for any given electrostatic field $\mathbf{E}$ has the same value for all paths from P_1 to P_2 .

The points P_2 and P_1 may coincide. In that case the paths are all closed curves, among them paths of vanishing length. This leads to the following corollary:

The line integral $\int \mathbf{E} \cdot d\mathbf{s}$ around any closed path in an electrostatic field is zero.

By *electrostatic field* we mean, strictly speaking, the electric field of stationary charges. Later on, we shall encounter electric fields in which the line integral is *not* path-independent. Those fields will usually be associated with rapidly moving charges. For our present purposes we can say that, if the source charges are moving slowly enough, the field $\mathbf{E}$ will be such that $\int \mathbf{E} \cdot d\mathbf{s}$ is practically path-independent. Of course, if $\mathbf{E}$ itself is varying in time, the $\mathbf{E}$ in $\int \mathbf{E} \cdot d\mathbf{s}$ must be understood as the field that exists over the whole path at a given instant of time. With that understanding we can talk meaningfully about the line integral in a changing electrostatic field.

2.2 Potential difference and the potential function

Because the line integral in the electrostatic field is path-independent, we can use it to define a scalar quantity ϕ_{21} , without specifying any particular path:

$$\phi_{21} = - \int_{P_1}^{P_2} \mathbf{E} \cdot d\mathbf{s}. \quad (2.4)$$

With the minus sign included here, ϕ_{21} is the work per unit charge done by an external agency in moving a positive charge from P_1 to P_2 in the field $\mathbf{E}$. (The external agency must supply a force $\mathbf{F}_{\text{ext}} = -q\mathbf{E}$ to balance the electrical force $\mathbf{F}_{\text{elec}} = q\mathbf{E}$; hence the minus sign.) Thus ϕ_{21} is a single-valued scalar function of the two positions P_1 and P_2 . We call it the *electric potential difference* between the two points.

In our SI system of units, potential difference is measured in joule/coulomb. This unit has a name of its own, the *volt*:

$$1 \text{ volt} = 1 \frac{\text{joule}}{\text{coulomb}}. \quad (2.5)$$

One joule of work is required to move a charge of one coulomb through a potential difference of one volt. In the Gaussian system of units, potential difference is measured in erg/esu. This unit also has a name of its own, the statvolt (“stat” comes from “electrostatic”). As an exercise, you can use the $1 \text{ C} \approx 3 \cdot 10^9 \text{ esu}$ relation from [Section 1.4](#) to show that one volt is equivalent to approximately 1/300 statvolt. These two relations are accurate to better than 0.1 percent, thanks to the accident that c is that

close to $3 \cdot 10^8$ m/s. [Appendix C](#) derives the conversion factors between all of the corresponding units in the SI and Gaussian systems. Further discussion of the exact relations between SI and Gaussian electrical units is given in [Appendix E](#), which takes into account the definition of the meter in terms of the speed of light.

Suppose we hold P_1 fixed at some reference position. Then ϕ_{21} becomes a function of P_2 only, that is, a function of the spatial coordinates x, y, z . We can write it simply $\phi(x, y, z)$, without the subscript, if we remember that its definition still involves agreement on a reference point P_1 . We can say that ϕ is the potential associated with the vector field $\mathbf{E}$. It is a scalar function of position, or a scalar field (they mean the same thing). Its value at a point is simply a number (in units of work per unit charge) and has no direction associated with it. Once the vector field $\mathbf{E}$ is given, the potential function ϕ is determined, except for an arbitrary additive constant allowed by the arbitrariness in our choice of P_1 .

Example Find the potential associated with the electric field described in [Fig. 2.3](#), the components of which are $E_x = Ky$, $E_y = Kx$, $E_z = 0$, with K a constant. This is a possible electrostatic field; we will see why in [Section 2.17](#). Some field lines are shown.

Solution Since $E_z = 0$, the potential will be independent of z and we need consider only the xy plane. Let x_1, y_1 be the coordinates of P_1 , and x_2, y_2 the coordinates of P_2 . It is convenient to locate P_1 at the origin: $x_1 = 0, y_1 = 0$. To evaluate $-\int \mathbf{E} \cdot d\mathbf{s}$ from this reference point to a general point (x_2, y_2) it is easiest to use a path like the dashed path ABC in [Fig. 2.3](#):

$$\phi(x_2, y_2) = - \int_{(0,0)}^{(x_2, y_2)} \mathbf{E} \cdot d\mathbf{s} = - \int_{(0,0)}^{(x_2, 0)} E_x dx - \int_{(x_2, 0)}^{(x_2, y_2)} E_y dy. \quad (2.6)$$

The first of the two integrals on the right is zero because E_x is zero along the x axis. The second integration is carried out at constant x , with $E_y = Kx_2$:

$$- \int_{(x_2, 0)}^{(x_2, y_2)} E_y dy = - \int_0^{y_2} Kx_2 dy = -Kx_2 y_2. \quad (2.7)$$

There was nothing special about the point (x_2, y_2) so we can drop the subscripts:

$$\phi(x, y) = -Kxy \quad (2.8)$$

for any point (x, y) in this field, with zero potential at the origin. Any constant could be added to this. That would only mean that the reference point to which zero potential is assigned had been located somewhere else.

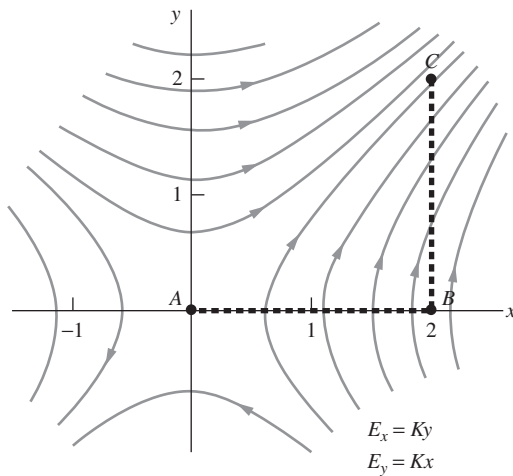


Figure 2.3. A particular path, ABC , in the electric field $E_x = Ky$, $E_y = Kx$. Some field lines are shown.

Example (Potential due to a uniform sphere) A sphere has radius R and uniform volume charge density ρ . Use the results from the example in [Section 1.11](#) to find the potential for all values of r , both inside and outside the sphere. Take the reference point P_1 to be infinitely far away.

Solution From the example in Section 1.11, the magnitude of the (radial) electric field inside the sphere is $E(r) = \rho r/3\epsilon_0$, and the magnitude outside is $E(r) = \rho R^3/3\epsilon_0 r^2$. Equation (2.4) tells us that the potential equals the negative of the line integral of the field, from P_1 (which we are taking to be at infinity) down to a given radius r . The potential outside the sphere is therefore

$$\phi_{\text{out}}(r) = - \int_{\infty}^r E(r') dr' = - \int_{\infty}^r \frac{\rho R^3}{3\epsilon_0 r'^2} dr' = \frac{\rho R^3}{3\epsilon_0 r}. \quad (2.9)$$

In terms of the total charge in the sphere, $Q = (4\pi R^3/3)\rho$, this potential is simply $\phi_{\text{out}}(r) = Q/4\pi\epsilon_0 r$. This is as expected, because we already knew that the potential energy of a charge q due to the sphere is $qQ/4\pi\epsilon_0 r$. And the potential ϕ equals the potential energy per unit charge.

To find the potential inside the sphere, we must break the integral into two pieces:

$$\begin{aligned} \phi_{\text{in}}(r) &= - \int_{\infty}^R E(r') dr' - \int_R^r E(r') dr' = - \int_{\infty}^R \frac{\rho R^3}{3\epsilon_0 r'^2} dr' - \int_R^r \frac{\rho r'}{3\epsilon_0} dr' \\ &= \frac{\rho R^3}{3\epsilon_0 R} - \frac{\rho}{6\epsilon_0} (r^2 - R^2) = \frac{\rho R^2}{2\epsilon_0} - \frac{\rho r^2}{6\epsilon_0}. \end{aligned} \quad (2.10)$$

Note that Eqs. (2.9) and (2.10) yield the same value of ϕ at the surface of the sphere, namely $\phi(R) = \rho R^2/3\epsilon_0$. So ϕ is continuous across the surface, as it should be. (The field is everywhere finite, so the line integral over an infinitesimal interval must yield an infinitesimal result.) The slope of ϕ is also continuous, because $E(r)$ (which is the negative derivative of ϕ , because ϕ is the negative integral of E) is continuous. A plot of $\phi(r)$ is shown in Fig. 2.4.

The potential at the center of the sphere is $\phi(0) = \rho R^2/2\epsilon_0$, which is 3/2 times the value at the surface. So if you bring a charge in from infinity, it takes 2/3 of your work to reach the surface, and then 1/3 to go the extra distance of R to the center.

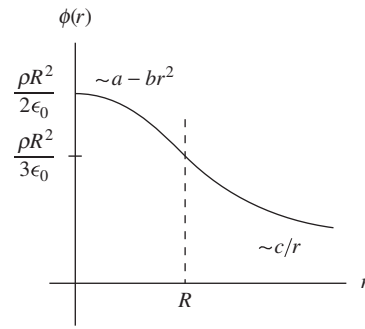


Figure 2.4. The potential due to a uniform sphere of charge.

We must be careful not to confuse the potential ϕ associated with a given field $\mathbf{E}$ with the potential energy of a system of charges. The potential energy of a system of charges is the total work required to assemble it, starting with all the charges far apart. In Eq. (1.14), for example, we expressed U , the potential energy of the charge system in Fig. 1.6. The electric potential $\phi(x, y, z)$ associated with the field in Fig. 1.6 would be the work per unit charge required to move a unit positive test charge from some chosen reference point to the point (x, y, z) in the field of that structure of nine charges.

2.3 Gradient of a scalar function

Given the electric field, we can find the electric potential function. But we can also proceed in the other direction; from the potential we can derive the field. It appears from Eq. (2.4) that the field is in some sense the *derivative* of the potential function. To make this idea precise we introduce the *gradient* of a scalar function of position. Let $f(x, y, z)$ be

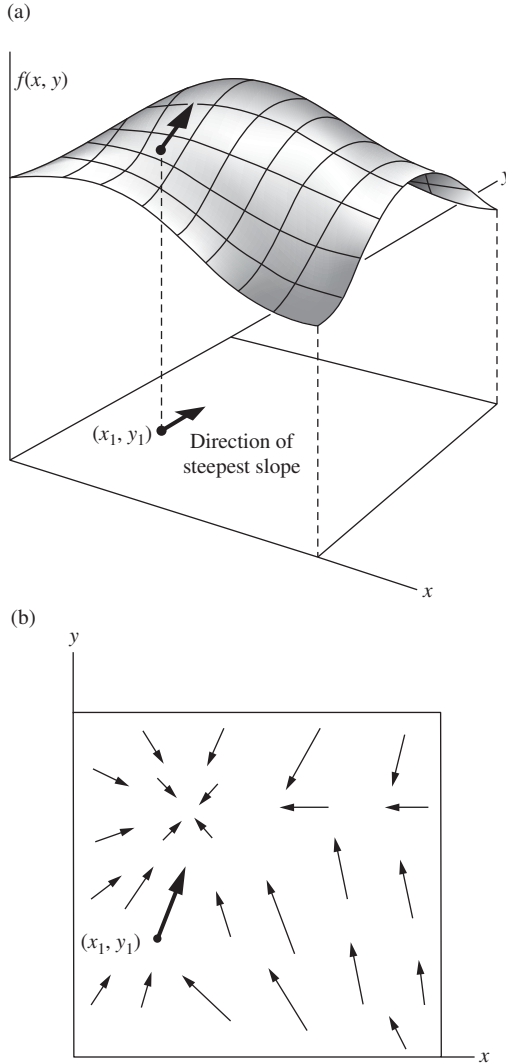


Figure 2.5. The scalar function $f(x, y)$ is represented by the surface in (a). The arrows in (b) represent the vector function, $\text{grad } f$.

some continuous, differentiable function of the coordinates. With its partial derivatives $\partial f/\partial x$, $\partial f/\partial y$, and $\partial f/\partial z$ we can construct at every point in space a vector, the vector whose x , y , z components are equal to the respective partial derivatives.¹ This vector we call the *gradient* of f , written “ $\text{grad } f$,” or ∇f :

$$\nabla f \equiv \hat{\mathbf{x}} \frac{\partial f}{\partial x} + \hat{\mathbf{y}} \frac{\partial f}{\partial y} + \hat{\mathbf{z}} \frac{\partial f}{\partial z}. \quad (2.13)$$

∇f is a vector that tells how the function f varies in the neighborhood of a point. Its x component is the partial derivative of f with respect to x , a measure of the rate of change of f as we move in the x direction. The direction of the vector ∇f at any point is the direction in which one must move from that point to find the most rapid increase in the function f . Suppose we were dealing with a function of two variables only, x and y , so that the function could be represented by a surface in three dimensions. Standing on that surface at some point, we see the surface rising in some direction, sloping downward in the opposite direction. There is a direction in which a short step will take us higher than a step of the same length in any other direction. The gradient of the function is a vector in that direction of steepest ascent, and its magnitude is the slope measured in that direction.

Figure 2.5 may help you to visualize this. Suppose some particular function of two coordinates x and y is represented by the surface $f(x, y)$ sketched in Fig. 2.5(a). At the location (x_1, y_1) the surface rises most steeply in a direction that makes an angle of about 80° with the positive x direction. The gradient of $f(x, y)$, ∇f , is a vector function of x and y . Its character is suggested in Fig. 2.5(b) by a number of vectors at various points in the two-dimensional space, including the point (x_1, y_1) . The vector function ∇f defined in Eq. (2.13) is simply an extension of this idea to three-dimensional space. (Be careful not to confuse Fig. 2.5(a) with real three-dimensional xyz space; the third coordinate there is the value of the function $f(x, y)$.)

As one example of a function in three-dimensional space, suppose f is a function of r only, where r is the distance from some fixed point O . On a sphere of radius r_0 centered about O , $f = f(r_0)$ is constant. On a slightly larger sphere of radius $r_0 + dr$ it is also constant, with the value $f = f(r_0 + dr)$. If we want to make the change from $f(r_0)$ to $f(r_0 + dr)$,

¹ We remind the reader that a partial derivative with respect to x , of a function of x , y , z , written simply $\partial f/\partial x$, means the rate of change of the function with respect to x with the other variables y and z held constant. More precisely,

$$\frac{\partial f}{\partial x} = \lim_{\Delta x \rightarrow 0} \frac{f(x + \Delta x, y, z) - f(x, y, z)}{\Delta x}. \quad (2.11)$$

As an example, if $f = x^2 y z^3$,

$$\frac{\partial f}{\partial x} = 2xy z^3, \quad \frac{\partial f}{\partial y} = x^2 z^3, \quad \frac{\partial f}{\partial z} = 3x^2 y z^2. \quad (2.12)$$

the shortest step we can make is to go radially (as from A to B) rather than from A to C , in Fig. 2.6. The “slope” of f is thus greatest in the radial direction, so ∇f at any point is a radially pointing vector. In fact $\nabla f = \hat{\mathbf{r}}(df/dr)$ in this case, $\hat{\mathbf{r}}$ denoting, for any point, a unit vector in the radial direction. See Section F.2 in Appendix F for further discussion of the gradient.

2.4 Derivation of the field from the potential

It is now easy to see that the relation of the scalar function f to the vector function ∇f is the same, except for a minus sign, as the relation of the potential ϕ to the field $\mathbf{E}$. Consider the value of ϕ at two nearby points, (x, y, z) and $(x + dx, y + dy, z + dz)$. The change in ϕ , going from the first point to the second, is, in first-order approximation,

$$d\phi = \frac{\partial \phi}{\partial x} dx + \frac{\partial \phi}{\partial y} dy + \frac{\partial \phi}{\partial z} dz. \quad (2.14)$$

On the other hand, from the definition of ϕ in Eq. (2.4), the change can also be expressed as

$$d\phi = -\mathbf{E} \cdot d\mathbf{s}. \quad (2.15)$$

The infinitesimal vector displacement $d\mathbf{s}$ is just $\hat{\mathbf{x}} dx + \hat{\mathbf{y}} dy + \hat{\mathbf{z}} dz$. Thus if we identify $\mathbf{E}$ with $-\nabla\phi$, where $\nabla\phi$ is defined via Eq. (2.13), then Eqs. (2.14) and (2.15) become identical. So the electric field is the negative of the gradient of the potential:

$$\mathbf{E} = -\nabla\phi \quad (2.16)$$

The minus sign came in because the electric field points from a region of greater potential toward a region of lesser potential, whereas the vector $\nabla\phi$ is defined so that it points in the direction of increasing ϕ .

To show how this works, we go back to the example of the field in Fig. 2.3. From the potential given by Eq. (2.8), $\phi = -Kxy$, we can recover the electric field we started with:

$$\mathbf{E} = -\nabla(-Kxy) = -\left(\hat{\mathbf{x}} \frac{\partial}{\partial x} + \hat{\mathbf{y}} \frac{\partial}{\partial y}\right)(-Kxy) = K(\hat{\mathbf{x}}\mathbf{y} + \hat{\mathbf{y}}\mathbf{x}). \quad (2.17)$$

2.5 Potential of a charge distribution

We already know the potential that goes with a single point charge, because we calculated the work required to bring one charge into the neighborhood of another in Eq. (1.9). The potential at any point, in the field of an isolated point charge q , is just $q/4\pi\epsilon_0 r$, where r is the distance

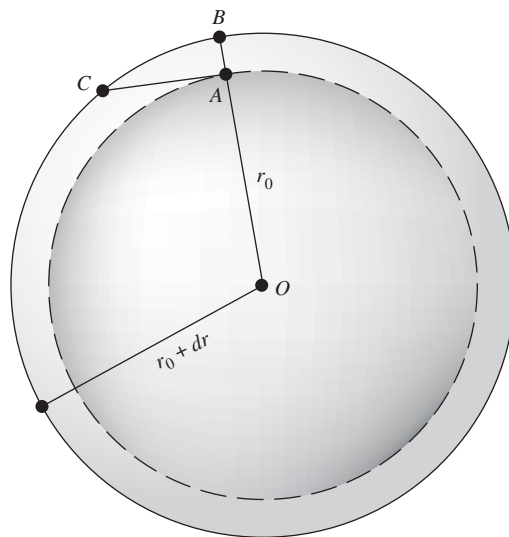


Figure 2.6. The shortest step for a given change in f is the radial step AB , if f is a function of r only.

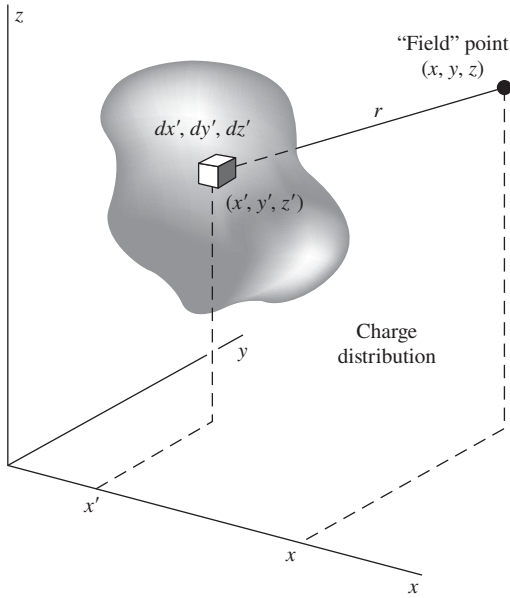


Figure 2.7. Each element of the charge distribution $\rho(x', y', z')$ contributes to the potential ϕ at the point (x, y, z) . The potential at this point is the sum of all such contributions; see Eq. (2.18).

from the point in question to the source q , and where we have assigned zero potential to points infinitely far from the source.

Superposition must work for potentials as well as fields. If we have several sources, the potential function is simply the sum of the potential functions that we would have for each of the sources present alone – *providing* we make a consistent assignment of the zero of potential in each case. If all the sources are contained in some finite region, it is always possible, and usually the simplest choice, to put zero potential at infinite distance. If we adopt this rule, the potential of any charge distribution can be specified by the integral

$$\phi(x, y, z) = \int_{\text{all sources}} \frac{\rho(x', y', z') dx' dy' dz'}{4\pi\epsilon_0 r}, \quad (2.18)$$

where r is the distance from the volume element $dx' dy' dz'$ to the point (x, y, z) at which the potential is being evaluated (Fig. 2.7). That is, $r = [(x - x')^2 + (y - y')^2 + (z - z')^2]^{1/2}$. Notice the difference between this and the integral giving the electric field of a charge distribution; see Eq. (1.22). Here we have r in the denominator, not r^2 , and the integral is a scalar not a vector. From the scalar potential function $\phi(x, y, z)$ we can always find the electric field by taking the negative gradient of ϕ , according to Eq. (2.16).

In the case of a discrete distribution of source charges, the above integral is replaced by a sum over all the charges, indexed by i :

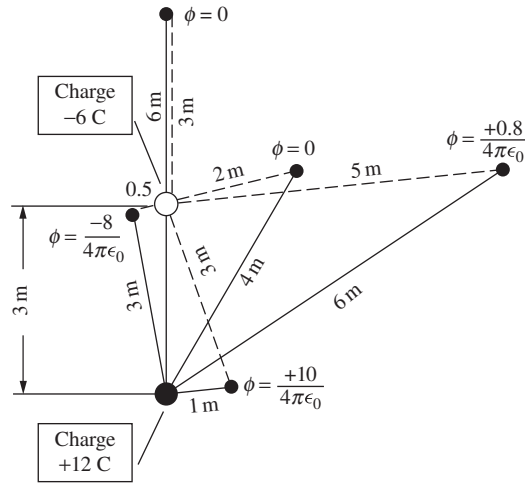
$$\phi(x, y, z) = \sum_{\text{all sources}} \frac{q_i}{4\pi\epsilon_0 r}, \quad (2.19)$$

where r is the distance from the charge q_i to the point (x, y, z) .

Example (Potential of two point charges) Consider a very simple example, the potential of the two point charges shown in Fig. 2.8. A positive charge of $12 \mu\text{C}$ is located 3 m away from a negative charge, $-6 \mu\text{C}$. (The “ μ ” prefix stands for “micro,” or 10^{-6} .) The potential at any point in space is the sum of the potentials due to each charge alone. The potentials for some selected points in space are given in the diagram. No vector addition is involved here, only the algebraic addition of scalar quantities. For instance, at the point on the far right, which is 6 m from the positive charge and 5 m from the negative charge, the potential has the value

$$\begin{aligned} \frac{1}{4\pi\epsilon_0} \left(\frac{12 \cdot 10^{-6} \text{ C}}{6 \text{ m}} + \frac{-6 \cdot 10^{-6} \text{ C}}{5 \text{ m}} \right) &= \frac{0.8 \cdot 10^{-6} \text{ C/m}}{4\pi\epsilon_0} \\ &= 7.2 \cdot 10^3 \text{ J/C} = 7.2 \cdot 10^3 \text{ V}, \end{aligned} \quad (2.20)$$

where we have used $1/4\pi\epsilon_0 \approx 9 \cdot 10^9 \text{ N m}^2/\text{C}^2$ (and also $1 \text{ N m} = 1 \text{ J}$). The potential approaches zero at infinite distance. It would take $7.2 \cdot 10^3 \text{ J}$ of work

**Figure 2.8.**

The electric potential ϕ at various points in a system of two point charges. ϕ goes to zero at infinite distance and is given in units of volts, or joules per coulomb.

to bring a unit positive charge in from infinity to a point where $\phi = 7.2 \cdot 10^3\text{ V}$. Note that two of the points shown on the diagram have $\phi = 0$. The net work done in bringing in any charge to one of these points would be zero. You can see that there must be an infinite number of such points, forming a surface in space surrounding the negative charge. In fact, the locus of points with any particular value of ϕ is a surface – an *equipotential surface* – which would show on our two-dimensional diagram as a curve.

There is one restriction on the use of Eq. (2.18): it may not work unless all sources are confined to some finite region of space. A simple example of the difficulty that arises with charges distributed out to infinite distance is found in the long charged wire whose field $\mathbf{E}$ we studied in Section 1.12. If we attempt to carry out the integration over the charge distribution indicated in Eq. (2.18), we find that the integral diverges – we get an infinite result. No such difficulty arose in finding the electric field of the infinitely long wire, because the contributions of elements of the line charge to the field decrease so rapidly with distance. Evidently we had better locate the zero of potential somewhere close to home, in a system that has charges distributed out to infinity. Then it is simply a matter of calculating the difference in potential ϕ_{21} , between the general point (x, y, z) and the selected reference point, using the fundamental relation, Eq. (2.4).

Example (Potential of a long charged wire) To see how this goes in the case of the infinitely long charged wire, let us arbitrarily locate the reference point P_1 at a distance r_1 from the wire. Then to carry a charge from P_1 to

any other point P_2 at distance r_2 requires the work per unit charge, using Eq. (1.39):

$$\begin{aligned}\phi_{21} &= - \int_{P_1}^{P_2} \mathbf{E} \cdot d\mathbf{s} = - \int_{r_1}^{r_2} \left(\frac{\lambda}{2\pi\epsilon_0 r} \right) dr \\ &= - \frac{\lambda}{2\pi\epsilon_0} \ln r_2 + \frac{\lambda}{2\pi\epsilon_0} \ln r_1.\end{aligned}\quad (2.21)$$

This shows that the electric potential for the charged wire can be taken as

$$\phi = - \frac{\lambda}{2\pi\epsilon_0} \ln r + \text{constant}.\quad (2.22)$$

The constant, $(\lambda/2\pi\epsilon_0) \ln r_1$ in this case, has no effect when we take $-\text{grad } \phi$ to get back to the field $\mathbf{E}$. In this case,

$$\mathbf{E} = -\nabla\phi = -\hat{\mathbf{r}} \frac{d\phi}{dr} = \frac{\lambda\hat{\mathbf{r}}}{2\pi\epsilon_0 r}.\quad (2.23)$$

2.6 Uniformly charged disk

Let us now study the electric potential and field around a uniformly charged disk. This is a charge distribution like that discussed in Section 1.13, except that it has a limited extent. The flat disk of radius a in Fig. 2.9 carries a positive charge spread over its surface with the constant density σ , in C/m^2 . (This is a single sheet of charge of infinitesimal thickness, not two layers of charge, one on each side. That is, the total charge in the system is $\pi a^2 \sigma$.) We shall often meet surface charge distributions in the future, especially on metallic conductors. However, the object just described is *not* a conductor; if it were, as we shall soon see, the charge could not remain uniformly distributed but would redistribute itself, crowding more toward the rim of the disk. What we have is an insulating disk, like a sheet of plastic, upon which charge has been “sprayed” so that every square meter of the disk has received, and holds fixed, the same amount of charge.

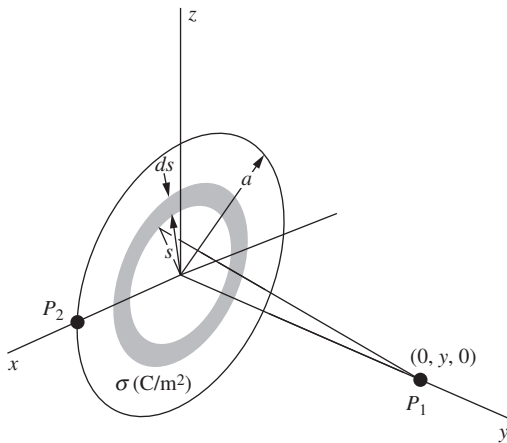


Figure 2.9. Finding the potential at a point P_1 on the axis of a uniformly charged disk.

Example (Potential on the axis) Let us find the potential due to our uniformly charged disk, at some point P_1 on the axis of symmetry, which we have made the y axis. All charge elements in a thin, ring-shaped segment of the disk lie at the same distance from P_1 . If s denotes the radius of such an annular segment and ds is its width, its area is $2\pi s ds$. The amount of charge it contains, dq , is therefore $dq = \sigma 2\pi s ds$. Since all parts of this ring are the same distance away from P_1 , namely, $r = \sqrt{y^2 + s^2}$, the contribution of the ring to the potential at P_1 is $dq/4\pi\epsilon_0 r = \sigma s ds / (2\epsilon_0 \sqrt{y^2 + s^2})$. To get the potential due to the whole disk, we have to integrate over all such rings:

$$\phi(0, y, 0) = \int \frac{dq}{4\pi\epsilon_0 r} = \int_0^a \frac{\sigma s ds}{2\epsilon_0 \sqrt{y^2 + s^2}} = \frac{\sigma}{2\epsilon_0} \sqrt{y^2 + s^2} \Big|_0^a.\quad (2.24)$$

Putting in the limits, we obtain

$$\phi(0, y, 0) = \frac{\sigma}{2\epsilon_0} \left(\sqrt{y^2 + a^2} - y \right) \quad \text{for } y > 0. \quad (2.25)$$

A minor point deserves a comment. The result we have written down in Eq. (2.25) holds for all points on the *positive* y axis. It is obvious from the physical symmetry of the system (there is no difference between one face of the disk and the other) that the potential must have the same value for negative and positive y , and this is reflected in Eq. (2.24), where only y^2 appears. But in writing Eq. (2.25) we made a choice of sign in taking the square root of y^2 , with the consequence that it holds only for positive y . The correct expression for $y < 0$ is obtained by the other choice of root and is given by

$$\phi(0, y, 0) = \frac{\sigma}{2\epsilon_0} \left(\sqrt{y^2 + a^2} + y \right) \quad \text{for } y < 0. \quad (2.26)$$

In view of this, we should not be surprised to find a kink in the plot of $\phi(0, y, 0)$ at $y = 0$. Indeed, the function has an abrupt change of slope there, as we see in Fig. 2.10, where we have plotted as a function of y the potential on the axis. The potential at the center of the disk is

$$\phi(0, 0, 0) = \frac{\sigma a}{2\epsilon_0}. \quad (2.27)$$

This much work would be required to bring a unit positive charge in from infinity, by any route, and leave it sitting at the center of the disk.

The behavior of $\phi(0, y, 0)$ for very large y is interesting. For $y \gg a$ we can approximate Eq. (2.25) as follows:

$$\sqrt{y^2 + a^2} - y = y \left[\left(1 + \frac{a^2}{y^2} \right)^{1/2} - 1 \right] = y \left[1 + \frac{1}{2} \left(\frac{a^2}{y^2} \right) + \cdots - 1 \right] \approx \frac{a^2}{2y}. \quad (2.28)$$

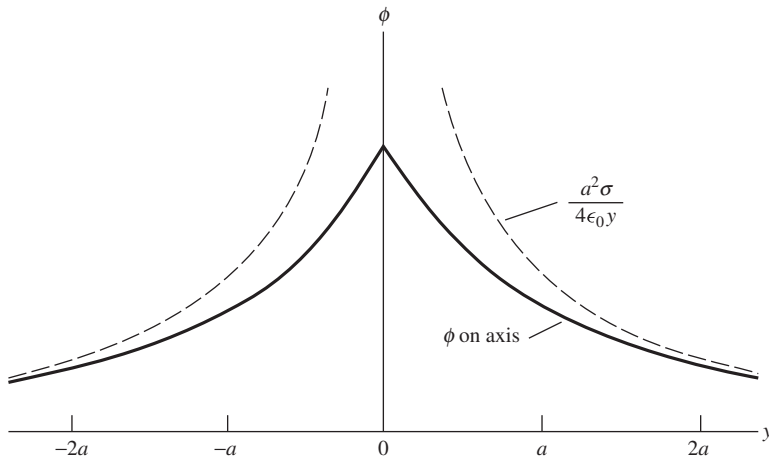


Figure 2.10.

A graph of the potential on the axis. The dashed curve is the potential of a point charge $q = \pi a^2 \sigma$.

Hence

$$\phi(0, y, 0) \approx \frac{a^2 \sigma}{4\epsilon_0 y} \quad \text{for } y \gg a. \quad (2.29)$$

Now $\pi a^2 \sigma$ is the total charge q on the disk, and Eq. (2.29), which can be written as $\pi a^2 \sigma / 4\pi \epsilon_0 y$, is just the expression for the potential due to a point charge of this magnitude. As we should expect, at a considerable distance from the disk (relative to its diameter), it doesn't matter much how the charge is shaped; only the total charge matters, in first approximation. In Fig. 2.10 we have drawn, as a dashed curve, the function $a^2 \sigma / 4\epsilon_0 y$. You can see that the axial potential function approaches its asymptotic form pretty quickly.

It is not quite so easy to derive the potential for general points away from the axis of symmetry, because the definite integral isn't so simple. It proves to be something called an *elliptic integral*. These functions are well known and tabulated, but there is no point in pursuing here mathematical details peculiar to a special problem. However, one further calculation, which is easy enough, may be instructive.

Example (Potential on the rim) We can find the potential at a point on the very edge of the disk, such as P_2 in Fig. 2.11. To calculate the potential at P_2 we can consider first the thin wedge of length R and angular width $d\theta$, as shown. An element of the wedge, the black patch at distance r from P_2 , contains an amount of charge $dq = \sigma r d\theta dr$. Its contribution to the potential at P_2 is therefore $dq / 4\pi \epsilon_0 r = \sigma d\theta dr / 4\pi \epsilon_0$. The contribution of the entire wedge is then $(\sigma d\theta / 4\pi \epsilon_0) \int_0^R dr = (\sigma R / 4\pi \epsilon_0) d\theta$. Now R is $2a \cos \theta$, from the geometry of the right triangle, and the whole disk is swept out as θ ranges from $-\pi/2$ to $\pi/2$. Thus we find the potential at P_2 :

$$\phi = \frac{\sigma a}{2\pi \epsilon_0} \int_{-\pi/2}^{\pi/2} \cos \theta d\theta = \frac{\sigma a}{\pi \epsilon_0}. \quad (2.30)$$

Comparing this with the potential at the center of the disk, $\sigma a / 2\epsilon_0$, we see that, as we should expect, the potential falls off from the center to the edge of the disk. The electric field, therefore, must have an *outward* component in the plane of the disk. That is why we remarked earlier that the charge, if free to move, would redistribute itself toward the rim. To put it another way, our uniformly charged disk is *not* a surface of constant potential, which any conducting surface must be unless charge is moving.²

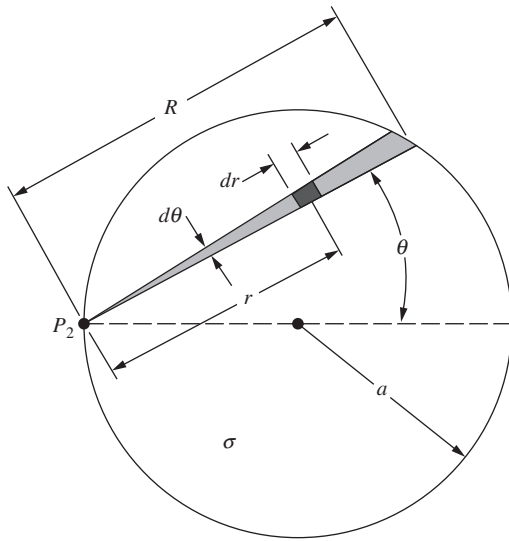


Figure 2.11. Finding the potential at a point P_2 on the rim of a uniformly charged disk.

² The fact that conducting surfaces have to be equipotentials will be discussed thoroughly in Chapter 3.

Let us now examine the electric field due to the disk. For $y > 0$, the field on the symmetry axis can be computed directly from the potential function given in Eq. (2.25):

$$\begin{aligned} E_y &= -\frac{\partial \phi}{\partial y} = -\frac{d}{dy} \frac{\sigma}{2\epsilon_0} \left(\sqrt{y^2 + a^2} - y \right) \\ &= \frac{\sigma}{2\epsilon_0} \left[1 - \frac{y}{\sqrt{y^2 + a^2}} \right] \quad y > 0. \end{aligned} \quad (2.31)$$

To be sure, it is not hard to compute E_y directly from the charge distribution, for points on the axis. We can again slice the disk into concentric rings, as we did prior to Eq. (2.24). But we must remember that $\mathbf{E}$ is a vector and that only the y component survives in the present setup, whereas we did not need to worry about components when calculating the scalar function ϕ above.

As y approaches zero from the positive side, E_y approaches $\sigma/2\epsilon_0$. On the negative y side of the disk, which we shall call the back, $\mathbf{E}$ points in the other direction and its y component E_y is $-\sigma/2\epsilon_0$. This is the same as the field of an infinite sheet of charge of density σ , derived in Section 1.13. It ought to be, for at points close to the center of the disk, the presence or absence of charge out beyond the rim can't make much difference. In other words, any sheet looks infinite if viewed from close up. Indeed, E_y has the value $\sigma/2\epsilon_0$ not only at the center, but also all over the disk.

For large y , we can find an approximate expression for E_y by using a Taylor series approximation as we did in Eq. (2.28). You can show that E_y approaches $a^2\sigma/4\epsilon_0 y^2$, which can be written as $\pi a^2\sigma/4\pi\epsilon_0 y^2$. This is correctly the field due to a point charge with magnitude $\pi a^2\sigma$.

In Fig. 2.12 we show some field lines for this system and also, plotted as dashed curves, the intersections on the yz plane of the surfaces of constant potential. Near the center of the disk these are lens-like surfaces, while at distances much greater than a they approach the spherical form of equipotential surfaces around a point charge.

Figure 2.12 illustrates a general property of field lines and equipotential surfaces. A field line through any point and the equipotential surface through that point *are perpendicular to one another*, just as, on a contour map of hilly terrain, the slope is steepest at right angles to a contour of constant elevation. This must be so, because if the field at any point had a component parallel to the equipotential surface through that point, it would require work to move a test charge along a constant-potential surface.

The energy associated with this electric field could be expressed as the integral over all space of $(\epsilon_0/2)E^2 dv$. It is equal to the work done in assembling this distribution, starting with infinitesimal charges far apart. In this particular example, as Exercise 2.56 will demonstrate, that work

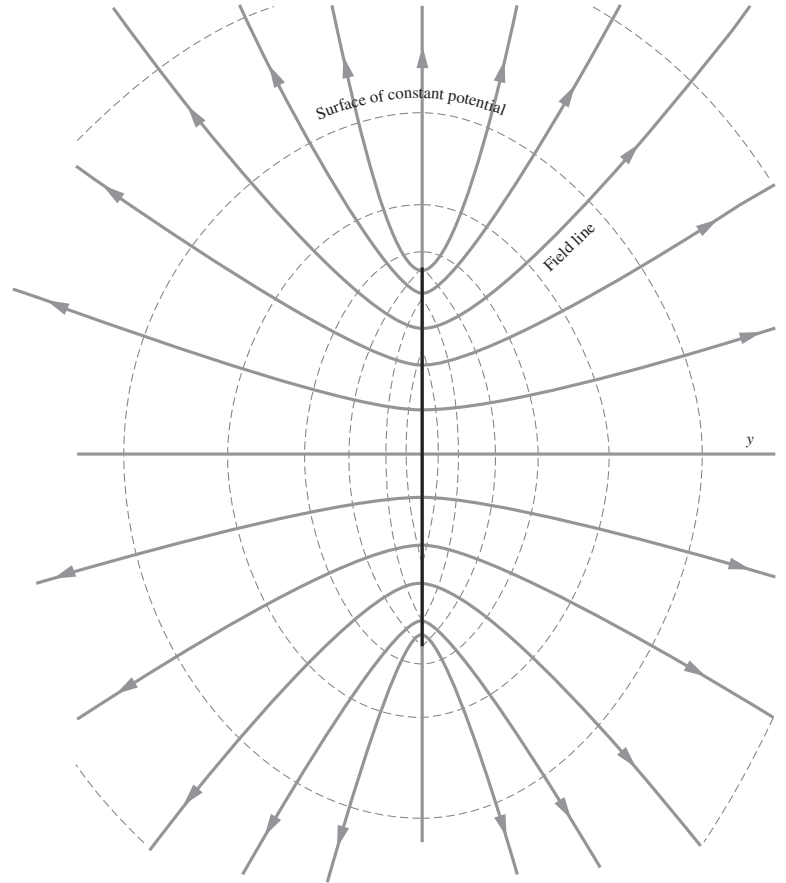


Figure 2.12.

The electric field of the uniformly charged disk. Solid curves are field lines. Dashed curves are intersections, with the plane of the figure, of surfaces of constant potential.

is not hard to calculate directly if we know the potential at the rim of a uniformly charged disk.

There is a general relation between the work U required to assemble a charge distribution $\rho(x, y, z)$ and the potential $\phi(x, y, z)$ of that distribution:

$$U = \frac{1}{2} \int \rho \phi \, dv \quad (2.32)$$

Equation (1.15), which gives the energy of a system of discrete point charges, could have been written in this way:

$$U = \frac{1}{2} \sum_{j=1}^N q_j \sum_{k \neq j} \frac{1}{4\pi\epsilon_0} \frac{q_k}{r_{jk}}. \quad (2.33)$$

The second sum is the potential at the location of the j th charge, due to all the other charges. To adapt this to a continuous distribution we merely

replace q_j with ρdv and the sum over j by an integral, thus obtaining Eq. (2.32).

2.7 Dipoles

Consider a setup with two equal and opposite charges $\pm q$ located at positions $\pm\ell/2$ on the y axis, as shown in Fig. 2.13. This configuration is called a *dipole*. The purpose of this section is to introduce the basics of dipoles. We save further discussion for Chapter 10, where we define the word “dipole” more precisely, derive things in more generality, and discuss examples of dipoles in actual matter. For now we just concentrate on determining the electric field and potential of a dipole. We have all of the necessary machinery at our disposal, so let’s see what we can find.

We will restrict the treatment to points far away from the dipole (that is, points with $r \gg \ell$). Although it is easy enough to write down an exact expression for the potential ϕ (and hence the field $\mathbf{E} = -\nabla\phi$) at any position, the result isn’t very enlightening. But when we work in the approximation of large distances, we obtain a result that, although isn’t exactly correct, is in fact quite enlightening. That’s how approximations work – you trade a little bit of precision for a large amount of clarity.

Our strategy will be to find the potential ϕ in polar (actually spherical) coordinates, and then take the gradient to find the electric field $\mathbf{E}$. We then determine the shape of the field-line and constant-potential curves. To make things look a little cleaner in the calculations below, we write $1/4\pi\epsilon_0$ as k in some intermediate steps.

2.7.1 Calculation of ϕ and $\mathbf{E}$

First note that, since the dipole setup is rotationally symmetric around the line containing the two charges, it suffices to find the potential in an arbitrary plane containing this line. We will use spherical coordinates, which reduce to polar coordinates in a plane because the angle ϕ doesn’t come into play (but note that θ is measured down from the vertical axis). Consider a point P with coordinates (r, θ) , as shown in Fig. 2.14. Let r_1 and r_2 be the distances from P to the two charges. Then the exact expression for the potential at P is (with $k \equiv 1/4\pi\epsilon_0$)

$$\phi_P = \frac{kq}{r_1} - \frac{kq}{r_2}. \quad (2.34)$$

If desired, the law of cosines can be used to write r_1 and r_2 in terms of r , θ , and ℓ .

Let us now derive an approximate form of this result, valid in the $r \gg \ell$ limit. One way to do this is to use the law-of-cosines expressions for r_1 and r_2 ; this is the route we will take in Chapter 10. But for the present purposes a simpler method suffices. In the $r \gg \ell$ limit, a closeup view of the dipole is shown in Fig. 2.15. The two lines from the charges to P are essentially parallel, so we see from the figure that the lengths of

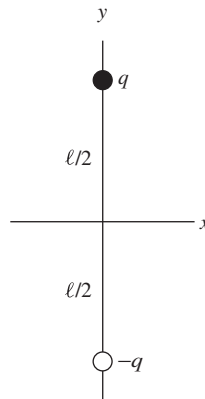


Figure 2.13.
Two equal and opposite charges form a dipole.

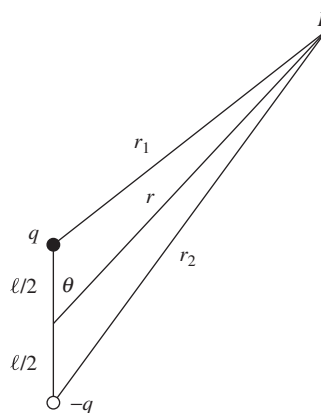


Figure 2.14.
Finding the potential ϕ at point P .

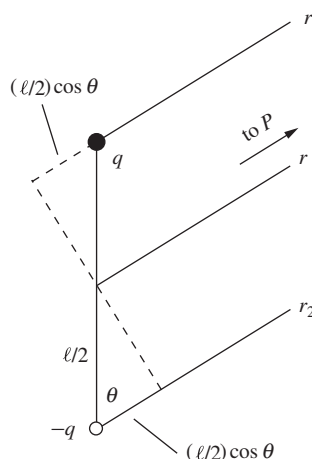


Figure 2.15.
Closeup view of Fig. 2.14.