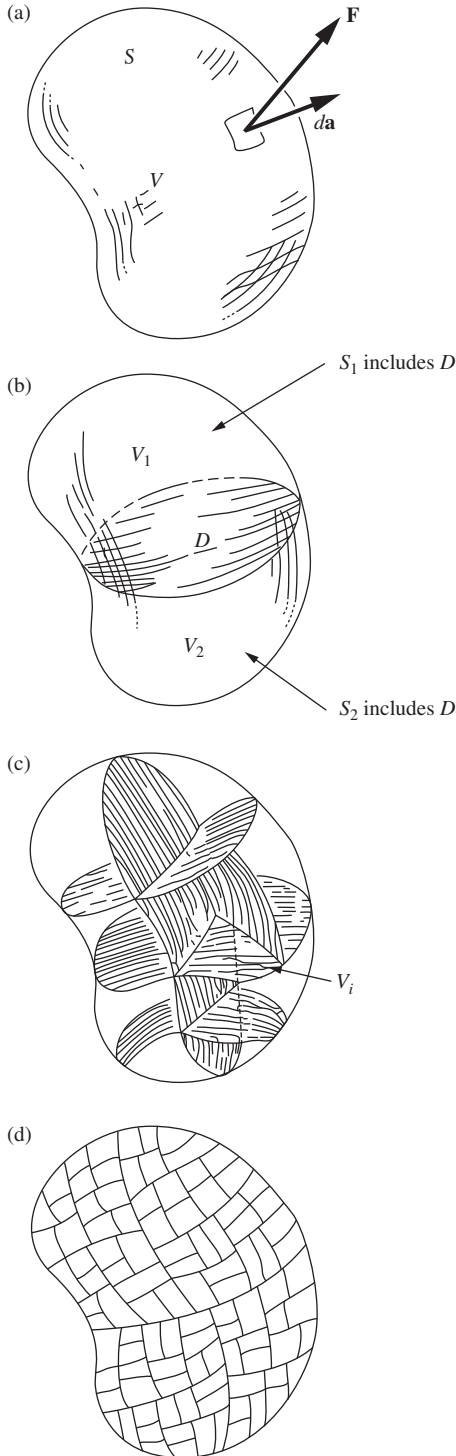




2.8 Divergence of a vector function

The electric field has a definite direction and magnitude at every point. It is a vector function of the coordinates, which we have often indicated by writing $\mathbf{E}(x, y, z)$. What we are about to say can apply to any vector function, not just to the electric field; we shall use another symbol, $\mathbf{F}(x, y, z)$, as a reminder of that. In other words, we shall talk mathematics rather than physics for a while and call $\mathbf{F}$ simply a general vector function. We shall keep to three dimensions, however.

Consider a finite volume V of some shape, the surface of which we shall denote by S . We are already familiar with the notion of the total flux Φ emerging from S . It is the value of the surface integral of $\mathbf{F}$ extended over the whole of S :

$$\Phi = \int_S \mathbf{F} \cdot d\mathbf{a}. \quad (2.42)$$

In the integrand $d\mathbf{a}$ is the infinitesimal vector whose magnitude is the area of a small element of S and whose direction is the outward-pointing normal to that little patch of surface, indicated in Fig. 2.21(a).

Now imagine dividing V into two parts by a surface, or a diaphragm, D that cuts through the “balloon” S , as in Fig. 2.21(b). Denote the two parts of V by V_1 and V_2 and, treating them as distinct volumes, compute the surface integral over each separately. The boundary surface S_1 of V_1 includes D , and so does S_2 . It is pretty obvious that the sum of the two surface integrals

$$\int_{S_1} \mathbf{F} \cdot d\mathbf{a}_1 + \int_{S_2} \mathbf{F} \cdot d\mathbf{a}_2 \quad (2.43)$$

will equal the original integral over the whole surface expressed in Eq. (2.42). The reason is that any given patch on D contributes with one sign to the first integral and the same amount with opposite sign to the second, the “outward” direction in one case being the “inward” direction in the other. In other words, any flux *out* of V_1 , through this surface D , is flux *into* V_2 . The rest of the surface involved is identical to that of the original entire volume.

We can keep on subdividing until our internal partitions have divided V into a large number of parts, $V_1, \dots, V_i, \dots, V_N$, with surfaces $S_1, \dots, S_i, \dots, S_N$. No matter how far this is carried, we can still be sure that

$$\sum_{i=1}^N \int_{S_i} \mathbf{F} \cdot d\mathbf{a}_i = \int_S \mathbf{F} \cdot d\mathbf{a} = \Phi. \quad (2.44)$$

Figure 2.21.

(a) A volume V enclosed by a surface S is divided (b) into two pieces enclosed by S_1 and S_2 . No matter how far this is carried, as in (c) and (d), the sum of the surface integrals over all the pieces equals the original surface integral over S , for any vector function $\mathbf{F}$.

What we are after is this: in the limit as N becomes enormous we want to identify something which is characteristic of a particular small region – and, ultimately, of the neighborhood of a point. Now the surface integral

$$\int_{S_i} \mathbf{F} \cdot d\mathbf{a}_i \quad (2.45)$$

over one of the small regions is *not* such a quantity, for if we divide everything again, so that N becomes $2N$, this integral divides into two terms, each smaller than before since their sum is constant. In other words, as we consider smaller and smaller volumes in the same locality, the surface integral over one such volume gets steadily smaller. But we notice that, when we divide, the volume is also divided into two parts that sum to the original volume. This suggests that we look at the ratio of surface integral to volume for an element in the subdivided space:

$$\frac{\int_{S_i} \mathbf{F} \cdot d\mathbf{a}_i}{V_i}. \quad (2.46)$$

It seems plausible that for N large enough, that is, for sufficiently fine-grained subdivision, we can halve the volume every time we halve the surface integral, so we find that, with continuing subdivision of any particular region, this ratio approaches a limit. If so, this limit is a property characteristic of the vector function $\mathbf{F}$ in that neighborhood. We call it the *divergence* of $\mathbf{F}$, written $\text{div } \mathbf{F}$. That is, the value of $\text{div } \mathbf{F}$ at any point is defined as

$$\text{div } \mathbf{F} \equiv \lim_{V_i \rightarrow 0} \frac{1}{V_i} \int_{S_i} \mathbf{F} \cdot d\mathbf{a}_i \quad (2.47)$$

where V_i is a volume including the point in question, and S_i , over which the surface integral is taken, is the surface of V_i . We must include the proviso that the limit exists and is independent of our method of subdivision. For the present we shall assume that this is true.

The meaning of $\text{div } \mathbf{F}$ can be expressed in this way: $\text{div } \mathbf{F}$ is the flux out of V_i , per unit of volume, in the limit of infinitesimal V_i . It is a scalar quantity, obviously. It may vary from place to place, its value at any particular location (x, y, z) being the limit of the ratio in Eq. (2.47) as V_i is chopped smaller and smaller while always enclosing the point (x, y, z) . So $\text{div } \mathbf{F}$ is simply a scalar function of the coordinates.

2.9 Gauss's theorem and the differential form of Gauss's law

If we know this scalar function of position, $\text{div } \mathbf{F}$, we can work our way right back to the surface integral over a large volume. We first write

Eq. (2.44) in this way:

$$\int_S \mathbf{F} \cdot d\mathbf{a} = \sum_{i=1}^N \int_{S_i} \mathbf{F} \cdot d\mathbf{a}_i = \sum_{i=1}^N V_i \left[\frac{\int_{S_i} \mathbf{F} \cdot d\mathbf{a}_i}{V_i} \right]. \quad (2.48)$$

In the limit $N \rightarrow \infty$, $V_i \rightarrow 0$, the term in brackets becomes the divergence of $\mathbf{F}$, and the sum goes into a volume integral:

$$\boxed{\int_S \mathbf{F} \cdot d\mathbf{a} = \int_V \operatorname{div} \mathbf{F} dv} \quad (\text{Gauss's theorem}). \quad (2.49)$$

This result is called *Gauss's theorem*, or the *divergence theorem*. It holds for any vector field for which the limit involved in Eq. (2.47) exists. Note that the entire content of the theorem is contained in Eq. (2.44), which itself is simply the statement that the fluxes cancel in pairs over the interior boundaries of all the little regions. The other steps in the proof were the multiplication by 1 in the form of V_i/V_i , the use of the definition in Eq. (2.47), and the conversion of an infinite sum to an integral. None of these steps contains much content.

Let us see what Eq. (2.49) implies for the electric field $\mathbf{E}$. We have Gauss's law, Eq. (1.31), which assures us that

$$\int_S \mathbf{E} \cdot d\mathbf{a} = \frac{1}{\epsilon_0} \int_V \rho dv. \quad (2.50)$$

If the divergence theorem holds for any vector field, it certainly holds for $\mathbf{E}$:

$$\int_S \mathbf{E} \cdot d\mathbf{a} = \int_V \operatorname{div} \mathbf{E} dv. \quad (2.51)$$

Equations (2.50) and (2.51) hold for *any* volume we care to choose – of any shape, size, or location. Comparing them, we see that this can only be true if, at every point,

$$\boxed{\operatorname{div} \mathbf{E} = \frac{\rho}{\epsilon_0}} \quad (2.52)$$

If we adopt the divergence theorem as part of our regular mathematical equipment from now on, we can regard Eq. (2.52) simply as an alternative statement of Gauss's law. It is Gauss's law in differential form, that is, stated in terms of a local relation between charge density and electric field.

Example (Field and density in a sphere) Let's use the result from the example in Section 1.11 to verify that Eq. (2.52) holds both inside and outside a sphere with radius R and uniform density ρ . Spherical coordinates are of course the most convenient ones to use here, given that we are dealing with a sphere. For the purposes of this example we will simply accept the expression given in

Eq. (F.3) in Appendix F for the divergence (also written as $\nabla \cdot \mathbf{E}$) in spherical coordinates. This appendix explains how to derive the various vector operators, including the divergence, in the common systems of coordinates (Cartesian, cylindrical, spherical). You are encouraged to read it in parallel with this chapter. In Section 2.10 we give a detailed derivation of the form of the divergence in Cartesian coordinates.

Since the electric field due to the sphere has only an r component, Eq. (F.3) tells us that the divergence of $\mathbf{E}$ is $\text{div } \mathbf{E} = (1/r^2)\partial(r^2 E_r)/\partial r$. Inside the sphere, we have $E_r = \rho r/3\epsilon_0$ from Eq. (1.35), so

$$\text{div } \mathbf{E}_{\text{in}} = \frac{1}{r^2} \frac{\partial}{\partial r} \left(r^2 \frac{\rho r}{3\epsilon_0} \right) = \frac{1}{r^2} \frac{\rho r^2}{\epsilon_0} = \frac{\rho}{\epsilon_0}, \quad (2.53)$$

as desired. Outside the sphere, the field is $E_r = \rho R^3/3\epsilon_0 r^2$ from Eq. (1.34), which equals the standard $Q/4\pi\epsilon_0 r^2$ result when written in terms of the total charge Q . However, the exact form doesn't matter here. All that matters is that E_r is proportional to $1/r^2$, because then

$$\text{div } \mathbf{E}_{\text{out}} \propto \frac{1}{r^2} \frac{\partial}{\partial r} \left(r^2 \frac{1}{r^2} \right) = 0. \quad (2.54)$$

This agrees with Eq. (2.52) because $\rho = 0$ outside the sphere. Of course, it is no surprise that these relations worked out – we originally derived E_r from Gauss's law, and Eq. (2.52) is simply the differential form of Gauss's law.

Although we used spherical coordinates in this example, Eq. (2.52) must still be true for any choice of coordinates. The task of Exercise 2.68 is to redo this example in Cartesian coordinates. If you are uneasy about invoking the above form of the divergence in spherical coordinates, you should solve Exercise 2.68 after reading the following section.

2.10 The divergence in Cartesian coordinates

While Eq. (2.47) is the fundamental definition of *divergence*, independent of any system of coordinates, it is useful to know how to calculate the divergence of a vector function when we are given its explicit form. Suppose a vector function $\mathbf{F}$ is expressed as a function of Cartesian coordinates x , y , and z . That means that we have three scalar functions, $F_x(x, y, z)$, $F_y(x, y, z)$, and $F_z(x, y, z)$. We'll take the region V_i in the shape of a little rectangular box, with one corner at the point (x, y, z) and sides Δx , Δy , and Δz , as in Fig. 2.22(a). Whether some other shape will yield the same limit is a question we must face later.

Consider two opposite faces of the box, the top and bottom for instance, which would be represented by the $d\mathbf{a}$ vectors $\hat{\mathbf{z}} \Delta x \Delta y$ and $-\hat{\mathbf{z}} \Delta x \Delta y$. The flux through these faces involves only the z component of $\mathbf{F}$, and the net contribution depends on the *difference* between F_z at the top and F_z at the bottom or, more precisely, on the difference between the average of F_z over the top face and the average of F_z over the bottom face of the box. To the first order in small quantities this difference is $(\partial F_z / \partial z) \Delta z$. Figure 2.22(b) will help to explain this. The average value of F_z on the bottom surface of the box, if we consider only first-order

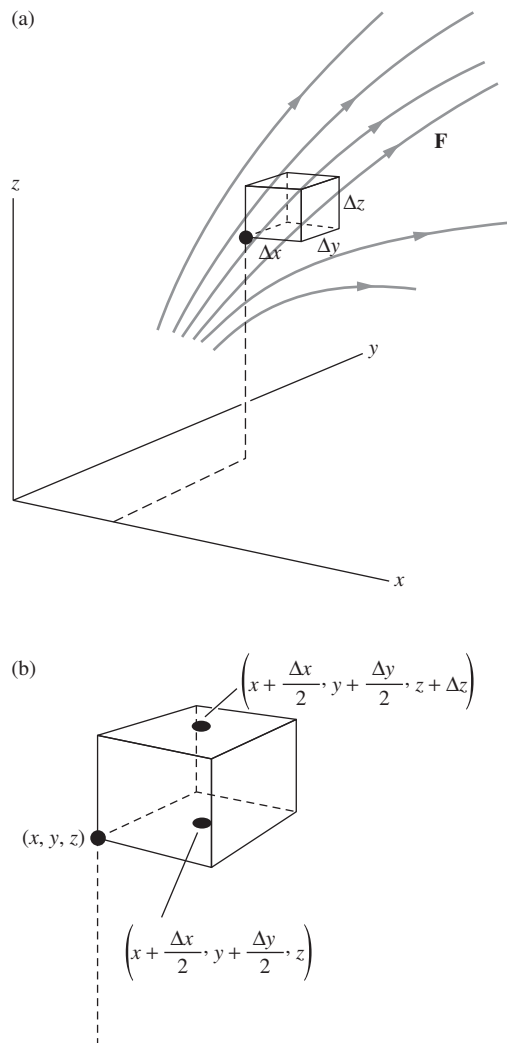


Figure 2.22. Calculation of flux from the box of volume $\Delta x \Delta y \Delta z$.

variations in F_z over this small rectangle, is its value at the center of the rectangle. That value is, to first order³ in Δx and Δy ,

$$F_z(x, y, z) + \frac{\Delta x}{2} \frac{\partial F_z}{\partial x} + \frac{\Delta y}{2} \frac{\partial F_z}{\partial y}. \quad (2.55)$$

For the average of F_z over the top face we take the value at the center of the top face, which to first order in the small displacements is

$$F_z(x, y, z) + \frac{\Delta x}{2} \frac{\partial F_z}{\partial x} + \frac{\Delta y}{2} \frac{\partial F_z}{\partial y} + \Delta z \frac{\partial F_z}{\partial z}. \quad (2.56)$$

The net flux out of the box through these two faces, each of which has the area of $\Delta x \Delta y$, is therefore

$$\begin{aligned} & \underbrace{\Delta x \Delta y \left[F_z(x, y, z) + \frac{\Delta x}{2} \frac{\partial F_z}{\partial x} + \frac{\Delta y}{2} \frac{\partial F_z}{\partial y} + \Delta z \frac{\partial F_z}{\partial z} \right]}_{\text{(flux out of box at top)}} \\ & - \underbrace{\Delta x \Delta y \left[F_z(x, y, z) + \frac{\Delta x}{2} \frac{\partial F_z}{\partial x} + \frac{\Delta y}{2} \frac{\partial F_z}{\partial y} \right]}_{\text{(flux into box at bottom)}}, \end{aligned} \quad (2.57)$$

which reduces to $\Delta x \Delta y \Delta z (\partial F_z / \partial z)$. Obviously, similar statements must apply to the other pairs of sides. That is, the net flux out of the box is $\Delta x \Delta z \Delta y (\partial F_y / \partial y)$ through the sides parallel to the xz plane and $\Delta y \Delta z \Delta x (\partial F_x / \partial x)$ through the sides parallel to the yz plane. Note that the product $\Delta x \Delta y \Delta z$ occurs in all of these expressions. Thus the total flux out of the little box is

$$\Phi = \Delta x \Delta y \Delta z \left(\frac{\partial F_x}{\partial x} + \frac{\partial F_y}{\partial y} + \frac{\partial F_z}{\partial z} \right). \quad (2.58)$$

The volume of the box is $\Delta x \Delta y \Delta z$, so the ratio of flux to volume is $\partial F_x / \partial x + \partial F_y / \partial y + \partial F_z / \partial z$, and as this expression does not contain the dimensions of the box at all, it remains as the limit when we let the box shrink. (Had we retained terms proportional to $(\Delta x)^2$, $(\Delta x \Delta y)$, etc., in the calculation of the flux, they would of course vanish on going to the limit.)

Now we can begin to see why this limit is going to be independent of the shape of the box. Obviously it is independent of the proportions of the rectangular box, but that isn't saying much. It is easy to see that it will be the same for any volume that we can make by sticking together little rectangular boxes of any size and shape. Consider the two boxes in Fig. 2.23. The sum of the flux Φ_1 out of box 1 and Φ_2 out of box 2 is not

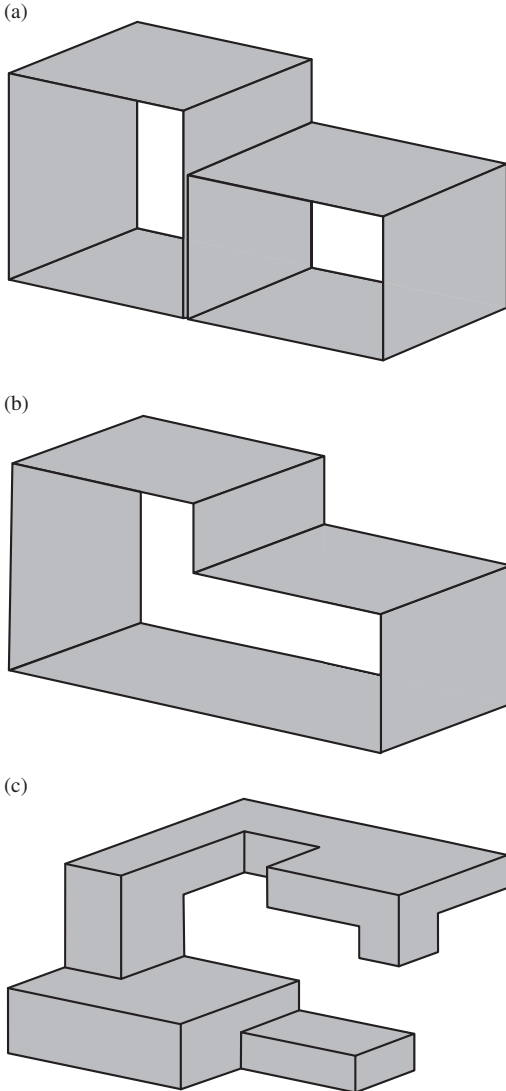


Figure 2.23.
The limit of the flux/volume ratio is independent of the shape of the box.

³ This is simply the beginning of a Taylor expansion of the scalar function F_z , in the neighborhood of (x, y, z) . That is, $F_z(x + a, y + b, z + c) = F_z(x, y, z) + \left(a \frac{\partial}{\partial x} + b \frac{\partial}{\partial y} + c \frac{\partial}{\partial z} \right) F_z + \dots + \frac{1}{n!} \left(a \frac{\partial}{\partial x} + b \frac{\partial}{\partial y} + c \frac{\partial}{\partial z} \right)^n F_z + \dots$. The derivatives are all to be evaluated at (x, y, z) . In our case $a = \Delta x/2$, $b = \Delta y/2$, $c = 0$, and we drop the higher-order terms in the expansion.

changed by removing the adjoining walls to make one box, for whatever flux went through that plane was negative flux for one and positive for the other. So we could have a bizarre shape like Fig. 2.23(c) without affecting the result. We leave it to the reader to generalize further. Tilted surfaces can be taken care of if you first prove that the vector sum of the four surface areas of the tetrahedron in Fig. 2.24 is zero.

We conclude that, assuming only that the functions F_x , F_y , and F_z are differentiable, the limit does exist and is given by

$$\operatorname{div} \mathbf{F} = \frac{\partial F_x}{\partial x} + \frac{\partial F_y}{\partial y} + \frac{\partial F_z}{\partial z} \quad (2.59)$$

We can also write the divergence in a very compact form using the “ ∇ ” symbol. From Eq. (2.13) we see that the gradient operator (symbolized by ∇ and often called “del”) can be treated in Cartesian coordinates as a vector consisting of derivatives:

$$\nabla = \hat{\mathbf{x}} \frac{\partial}{\partial x} + \hat{\mathbf{y}} \frac{\partial}{\partial y} + \hat{\mathbf{z}} \frac{\partial}{\partial z}. \quad (2.60)$$

In terms of this vector operator, we can write the divergence in the simple form, as you can quickly verify,

$$\operatorname{div} \mathbf{F} = \nabla \cdot \mathbf{F}. \quad (2.61)$$

If $\operatorname{div} \mathbf{F}$ has a positive value at some point, we find – thinking of $\mathbf{F}$ as a velocity field – a net “outflow” in that neighborhood. For instance, if all three partial derivatives in Eq. (2.59) are positive at a point P , we might have a vector field in that neighborhood something like that suggested in Fig. 2.25. But the field could look quite different and still have positive divergence, for any vector function $\mathbf{G}$ such that $\operatorname{div} \mathbf{G} = 0$ could be superimposed. Thus one or two of the three partial derivatives could be negative, and we might still have $\operatorname{div} \mathbf{F} > 0$. The divergence is a quantity that expresses only one aspect of the spatial variation of a vector field.

Example (Field due to a cylinder) Let’s find the divergence of an electric field that is rather easy to visualize. An infinitely long circular cylinder of radius a is filled with a distribution of positive charge of density ρ . We know from Gauss’s law that outside the cylinder the electric field is the same as that of a line charge on the axis. It is a radial field with magnitude proportional to $1/r$, given by Eq. (1.39) with $\lambda = \rho(\pi a^2)$. The field inside is found by applying Gauss’s law to a cylinder of radius $r < a$. You can do this as an easy problem

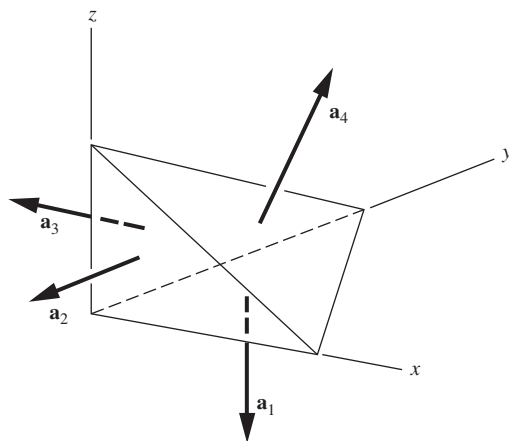


Figure 2.24.

You can prove that $\mathbf{a}_1 + \mathbf{a}_2 + \mathbf{a}_3 + \mathbf{a}_4 = 0$.

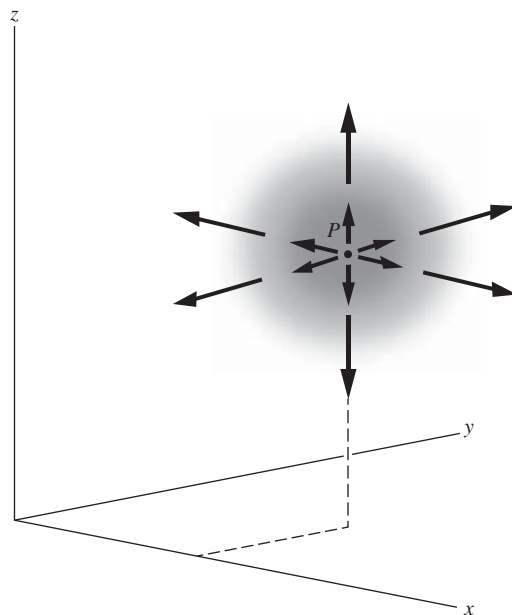


Figure 2.25.

Showing a field that in the neighborhood of point P has a nonzero divergence.

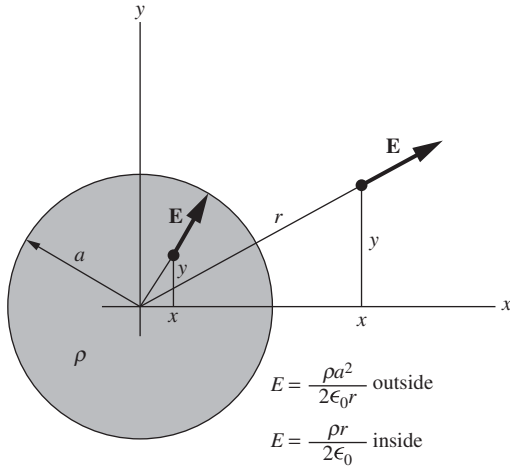


Figure 2.26.
The field inside and outside a uniform cylindrical distribution of charge.

(see Exercise 2.42). You will find that the field inside is directly proportional to r , and of course it is radial also. The exact values are:

$$\begin{aligned} E^{\text{out}} &= \frac{\rho a^2}{2\epsilon_0 r} & \text{for } r > a, \\ E^{\text{in}} &= \frac{\rho r}{2\epsilon_0} & \text{for } r < a. \end{aligned} \quad (2.62)$$

Figure 2.26 is a section perpendicular to the axis of the cylinder. Rectangular coordinates aren't the most natural choice here, but we'll use them anyway to get some practice with Eq. (2.59). With $r = \sqrt{x^2 + y^2}$, the field components are expressed as follows:

$$\begin{aligned} E_x^{\text{out}} &= \left(\frac{x}{r}\right) E^{\text{out}} = \frac{\rho a^2 x}{2\epsilon_0 (x^2 + y^2)} & \text{for } r > a, \\ E_y^{\text{out}} &= \left(\frac{y}{r}\right) E^{\text{out}} = \frac{\rho a^2 y}{2\epsilon_0 (x^2 + y^2)} & \text{for } r > a, \\ E_x^{\text{in}} &= \left(\frac{x}{r}\right) E^{\text{in}} = \frac{\rho x}{2\epsilon_0} & \text{for } r < a, \\ E_y^{\text{in}} &= \left(\frac{y}{r}\right) E^{\text{in}} = \frac{\rho y}{2\epsilon_0} & \text{for } r < a. \end{aligned} \quad (2.63)$$

And E_z is zero everywhere, of course.

Outside the cylinder of charge, $\text{div } \mathbf{E}$ has the value given by

$$\frac{\partial E_x^{\text{out}}}{\partial x} + \frac{\partial E_y^{\text{out}}}{\partial y} = \frac{\rho a^2}{2\epsilon_0} \left[\frac{1}{x^2 + y^2} - \frac{2x^2}{(x^2 + y^2)^2} + \frac{1}{x^2 + y^2} - \frac{2y^2}{(x^2 + y^2)^2} \right] = 0. \quad (2.64)$$

Inside the cylinder, $\text{div } \mathbf{E}$ is

$$\frac{\partial E_x^{\text{in}}}{\partial x} + \frac{\partial E_y^{\text{in}}}{\partial y} = \frac{\rho}{2\epsilon_0} (1 + 1) = \frac{\rho}{\epsilon_0}. \quad (2.65)$$

We expected both results. Outside the cylinder, where there is no charge, the net flux emerging from any volume – large or small – is zero, so the limit of the ratio *flux/volume* is certainly zero. Inside the cylinder we get the result required by the fundamental relation Eq. (2.52).

Having gotten some practice with Cartesian coordinates, let's redo this example in a much quicker manner by using cylindrical coordinates. Since $\mathbf{E}$ has only a radial component, Eq. (F.2) in Appendix F gives the divergence in cylindrical coordinates as $\text{div } \mathbf{E} = (1/r) \partial(rE_r)/\partial r$ (see Section F.3 for the derivation). Inside the cylinder, the field is $E_r = \rho r/2\epsilon_0$, so we quickly find $\text{div } \mathbf{E} = \rho/\epsilon_0$, as above. Outside the cylinder, the field is $E_r = \rho a^2/2\epsilon_0 r$, so we immediately find $\text{div } \mathbf{E} = 0$, which is again correct. All that matters in this latter case is that the field is proportional to $1/r$. Any such field will have $\text{div } \mathbf{E} = 0$.