

The value of μ_0 is exactly $4\pi \times 10^{-7}$ H/m

Electromagnetic forces between more generally shaped current carrying wires and magnets are governed by a complex set of equations. A full discussion of these physical laws is beyond the scope of this course, and will be covered in EN51.

Applications of electromagnetic forces

Electromagnetic forces are widely exploited in the design of electric motors; force actuators; solenoids; and electromagnets. All these applications are based upon the principle that a current-carrying wire in a magnetic field is subject to a force. The magnetic field can either be induced by a permanent magnet (as in a DC motor); or can be induced by passing a current through a second wire (used in some DC motors, and all AC motors). The general trends of forces in electric motors follow Ampere's law: the force exerted by the motor increases linearly with electric current in the armature; increases roughly in proportion to the length of wire used to wind the armature, and depends on the geometry of the motor.



Two examples of DC motor - the picture on the right is cut open to show the windings. You can find more information on motors at <http://my.execpc.com/~rheadley/magmotor.htm>

Hydrostatic and buoyancy forces

When an object is immersed in a stationary fluid, its surface is subjected to a *pressure*. The pressure is actually induced in the fluid by gravity: the pressure at any depth is effectively supporting the weight of fluid above that depth.

A pressure is a *distributed force*. If a pressure p acts on a surface, a small piece of the surface with area dA is subjected to a force

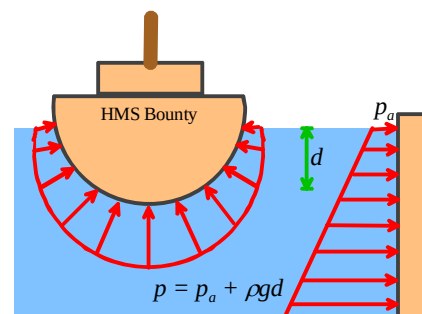
$$d\mathbf{F} = -p \, dA \, \mathbf{n}$$

where $\mathbf{n}$ is a unit vector perpendicular to the surface. The total force on a surface must be calculated by integration. We will show how this is done shortly.

The pressure in a stationary fluid varies linearly with depth below the fluid surface

$$p = p_a + \rho g d$$

where p_a is atmospheric pressure (often neglected as it's generally small compared with the second term); ρ is the fluid density; g is the acceleration due to gravity; and d is depth below the fluid surface.



The wing area $A_w = cL$ where c is the chord of the wing (see the picture) and L is its length, is used in defining both the lift and drag coefficient.

The variation of C_L and C_D with *angle of attack* α are crucial in the design of aircraft. For reasonable values of α (below stall - say less than 10 degrees) the behavior can be approximated by

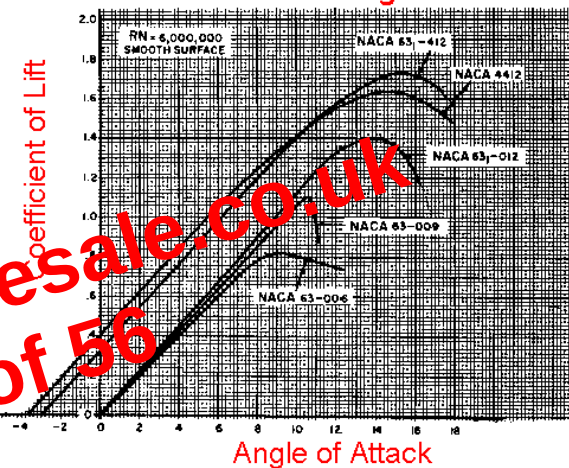
$$C_L = k_L \alpha$$

$$C_D = k_{Dp} + k_{DI} \alpha^2$$

where k_L , k_{Dp} and k_{DI} are more or less constant for any given airfoil shape, for practical ranges of Reynolds number. The first term in the drag coefficient, k_{Dp} , represents *parasite drag* – due to viscous drag and some pressure drag. The second term $k_{DI} \alpha^2$ is called *induced drag*, and is an undesirable by-product of lift.

The graphs on the right, (taken from 'Aerodynamics for Naval Aviators, H.H. Hurt, U.S. Naval Air Systems Command reprint') shows some experimental data for lift coefficient C_L as a function of AOA (that's angle of attack, but you're engineers now so you have to talk in code to maximize your nerd factor. That's NF). The data suggest that $k_L \approx 0.1 \text{deg}^{-1}$, and in fact a simple model known as 'thin airfoil theory' predicts that lift coefficient should vary by 2π per radian (that works out as $0.1096/\text{degree}$)

Coefficient of Lift vs Angle of Attack



The induced drag coefficient k_{DI} can be estimated from the formula

$$k_{DI} = \frac{c^2 k_L^2}{\pi e A_w}$$

where $k_L \approx 0.1 \text{deg}^{-1}$, L is the length of the wing and c is its width; while e is a constant known as the 'Oswald efficiency factor.' The constant e is always less than 1 and is of order 0.9 for a high performance wing (eg a jet aircraft or glider) and of order 0.7 for el cheapo wings.

The parasite drag coefficient k_{Dp} is of order 0.05 for the wing of a small general aviation aircraft, and of order 0.005 or lower for a commercial airliner.

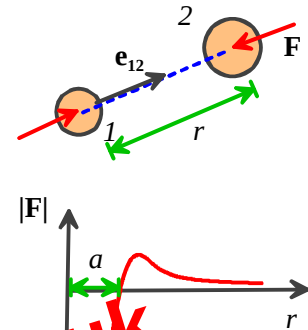
Interatomic forces

Engineers working in the fields of nanotechnology, materials design, and bio/chemical engineering are often interested in calculating the motion of molecules or atoms in a system.

They do this using 'Molecular Dynamics,' which is a computer method for integrating the equations of motion for every atom in the solid. The equations of motion are just Newton's law – $\mathbf{F} = m\mathbf{a}$ for each atom – but for the method to work, it is necessary to calculate the forces acting on the atoms. Specifying these forces is usually the most difficult part of the calculation.

The forces are computed using empirical force laws, which are either determined experimentally, or (more often) by means of quantum-mechanical calculations. In the simplest models, the atoms are assumed to interact through **pair forces**. In this case

- The forces exerted by two interacting atoms depends only on their relative positions, and is independent on the position of other atoms in the solid
- The forces act along the line connecting the atoms.
- The magnitude of the force is a function of the distance between them. The function is chosen so that (i) the force is repulsive when the atoms are close together; (ii) the force is zero at the equilibrium interatomic spacing; (iii) there is some critical distance where the attractive force has its maximum value (see the figure); and (iv) the force drops to zero when the atoms are far apart.



Various functions are used to specify the detailed shape of the force-separation law. A common one is the so-called 'Lennard Jones' function, which gives the force acting on atom (1) as

$$\mathbf{F}^{(1)} = -12E \left[\left(\frac{a}{r} \right)^{13} - \left(\frac{a}{r} \right)^7 \right] \mathbf{e}_{12}$$

Here a is the equilibrium separation between the atoms, and E is the total bond energy – the amount of work required to separate the bond by stretching it from initial length a to infinity.

This function was originally intended to model the atoms in a Noble gas – like He or Ar, etc. It is sometimes used in simple models of liquids and glasses. It would not be a good model of a metal, or covalently bonded solids. In fact, for these materials pair potentials don't work well, because the force exerted between two atoms depends not just on the relative positions of the two atoms themselves, but also on the positions of other nearby atoms. More complicated functions exist that can account for this kind of behavior, but there is still a great deal of uncertainty in the choice of function for a particular material.

- (1) Does the connection allow the two connected solids move relative to each other? If so, what is the direction of motion? There can be no component of reaction force along the direction of relative motion.
- (2) Does the connection allow the two connected solids rotate relative to each other? If so, what is the axis of relative rotation? There can be no component of reaction moment parallel to the axis of relative rotation.
- (3) For certain types of joint, a more appropriate question may be 'Is it really healthy/legal for me to smoke this?'

2.4.3 Drawing free body diagrams with constraint forces

When we solve problems with constraints, we are nearly always interested in analyzing forces in a structure containing many parts, or the motion of a machine with a number of separate moving components. Solving this kind of problem is not difficult – but it is very complicated because of the large number of forces involved and the large number of equations that must be solved to determine them. To avoid making mistakes, it is critical to use a systematic, and logical, procedure for drawing free body diagrams and labeling forces.

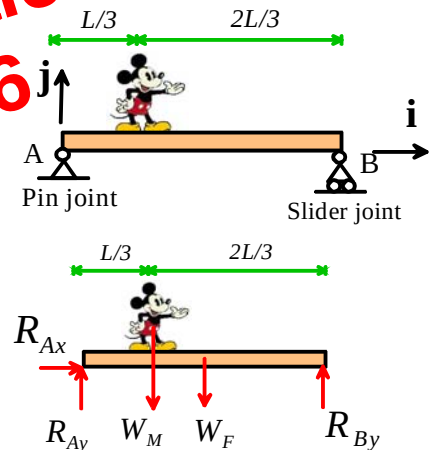
The procedure is best illustrated by means of some simple Mickey Mouse examples. When drawing free body diagrams yourself, you will find it helpful to consult Section 4.3.4 for the nature of reaction forces associated with various constraints.

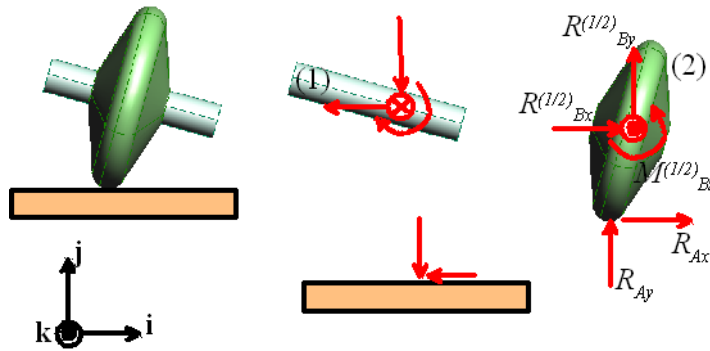
2D Mickey-mouse problem 1. The figure shows Mickey Mouse standing on a beam supported by a pin joint at one end and a slider joint at the other.

We consider Mickey and the floor together as the system of interest. We draw a picture of the system, isolated from its surroundings (disconnect all the joints, remove contacts, etc). In the picture, all the joints and connections are replaced by forces, following the rules outlined in the preceding section.

Notice how we've introduced variables to denote the unknown force components. It is sensible to use a convention that allows you to quickly identify both the position and direction associated with each variable. It is a good idea to use double subscripts – the first subscript shows *where* the force acts, the second shows its direction. Forces are *always* taken to be positive if they act along the positive x , y and z directions.

We've used the fact that A is a pin joint, and therefore exerts both vertical and horizontal forces; while B is a roller joint, and exerts only a vertical force. *Note that we always, always draw all admissible forces on the FBD, even if we suspect that some components may turn out later to be zero.* For example, it's fairly clear that $R_{Ax} = 0$ in this example, but it would be incorrect to leave off this force. This is especially important in dynamics problems where your intuition regarding forces is very often incorrect.





The forces and moments shown are the *only nonzero components of reaction force*.

The missing force and moment components can be shown to be zero by considering force and moment balance for the wheel. The details are left as an exercise.

Finally, a word of caution.

You can only use these shortcuts if:

1. The wheel's weight is negligible;
2. The wheel rotates freely (no bearing friction, and nothing driving the wheel);
3. There is only one contact point on the wheel.

If any of these conditions are violated you must solve the problem by applying all the proper reaction forces at contacts and bearings, and drawing a separate free body diagram for the wheel.

Preview from Notesale.co.uk
Page 47 of 56

2.5 Friction Forces

Friction forces act wherever two solids touch. It is a type of contact force – but rather more complicated than the contact forces we’ve dealt with so far.

It’s worth reviewing our earlier discussion of contact forces. When we first introduced contact forces, we said that the nature of the forces acting at a contact depends on three things:

- (4) Whether the contact is lubricated, i.e. whether friction acts at the contact
- (5) Whether there is significant rolling resistance at the contact
- (6) Whether the contact is *conformal*, or *nonconformal*.

We have so far only discussed two types of contact (a) fully lubricated (frictionless) contacts; and (b) ideally rough (infinite friction) contacts.

Remember that for a frictionless contact, only one component of force acts on the two contacting solids, as shown in the picture on the left below. In contrast, for an ideally rough (infinite friction) contact, three components of force are present as indicated on the figure on the right.



(a) Reaction forces at a frictionless contact

(b) Reaction forces at an ideally rough contact

All real surfaces lie somewhere between these two extremes. The contacting surfaces will experience both a normal and tangential force. The normal force must be repulsive, but can have an arbitrary magnitude. The tangential forces can act in any direction, but their magnitude is limited. If the tangential forces get too large, the two contacting surfaces will *slip* relative to each other.

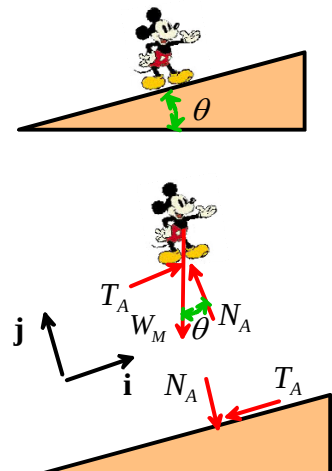
This is why it’s easy to walk up a dry, rough slope, but very difficult to walk up an icy slope. The picture below helps understand how friction forces work. The picture shows the big MM walking up a slope with angle θ , and shows the forces acting on M and the slope. We can relate the normal and tangential forces acting at the contact to Mickey’s weight and the angle θ by doing a force balance

Omitting the tedious details, we find that

$$T_A = W_M \sin \theta$$

$$N_A = W_M \cos \theta$$

Note that a tangential force $T_A = W_M \sin \theta$ must act at the contact. If the tangential force gets too large, then Mickey will start to slip down the slope.



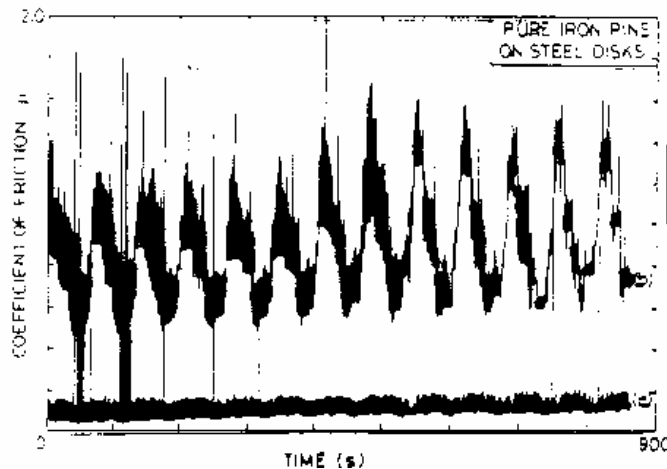
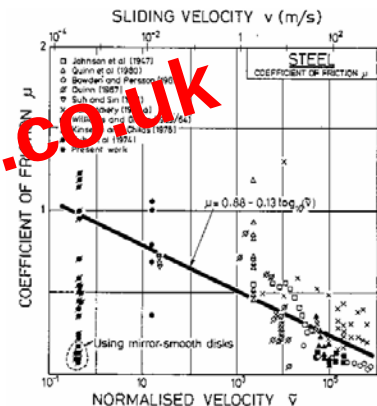
2.5.3 Experimental values for friction coefficient

The table below (taken from 'Engineering Materials' by Ashby and Jones, Pergamon, 1980) lists rough values for friction coefficients for various material pairs.

Material	Approx friction coefficient
Clean metals in air	0.8-2
Clean metals in wet air	0.5-1.5
Steel on soft metal (lead, bronze, etc)	0.1-0.5
Steel on ceramics (sapphire, diamond, ice)	0.1-0.5
Ceramics on ceramics (eg carbides on carbides)	0.05-0.5
Polymers on polymers	0.05-1.0
Metals and ceramics on polymers (PE, PTFE, PVC)	0.04-0.5
Boundary lubricated metals (thin layer of grease)	0.05-0.2
High temperature lubricants (eg graphite)	0.05-0.2
Hydrodynamically lubricated surfaces (full oil film)	0.0001-0.0005

These are rough guides only – friction coefficients for a given material can be highly variable. For example, Lim and Ashby (Cambridge University Internal Report CUED/C-mat./TR.123 January 1986) have catalogued a large number of experimental measurements of friction coefficient for steel on steel, and present the data graphically as shown below. You can see that friction coefficient for steel on steel varies anywhere between 0.0001 to 3.

Friction coefficient can even vary significantly during a measurement. For example, the picture below (from Lim and Ashby, Acta Met 37 3 (1989) p 767) shows the time variation of friction coefficient during a pin-on-disk experiment.



2.5.4 Static and kinetic friction

Many introductory statics textbooks define two different friction coefficients. One value, known as the *coefficient of static friction* and denoted by μ_s , is used to model static friction in the equation giving the condition necessary to initiate slip at a contact

$$|T| < \mu_s N$$

A second value, known as the *coefficient of kinetic friction*, and denoted by μ_k , is used in the equation for the force required to maintain steady sliding between two surfaces

$$T = \pm \mu_k N$$

I don't like to do this (I'm such a rebel). It is true that for some materials the static friction force can be a bit higher than the kinetic friction force, but this behavior is by no means universal, and in any case the difference between μ_k and μ_s is very small (of the order of 0.05). We've already seen that μ can vary far more than this for a given material pair, so it doesn't make much sense to quibble about such a small difference.

The real reason to distinguish between static and kinetic friction coefficient is to provide a simple explanation for *slip-stick oscillations* between two contacting surfaces. Slip-stick oscillations often occur when we try to do the simple friction experiment shown below.



If the end of the spring is moved *gradually* to the right, the block sticks for a while until the force in the spring gets large enough to overcome friction. At this point, the block jumps to the right and then sticks again, instead of smoothly following the spring. If μ were constant, then this behavior would be impossible. By using $\mu_s > \mu_k$, we can explain it. But if we're not trying to model slip-stick oscillations, it's much more sensible to work with just one value of μ .

In any case, there's a much better way to model slip-stick oscillations, by making μ depend on the velocity of sliding. Most sophisticated models of slip-stick oscillations (e.g. models of earthquakes at faults) do this.

12.6 The microscopic origin of friction forces

Friction is weird. In particular, we need to explain

- (i) why friction forces are independent of the contact area
- (ii) why friction forces are proportional to the normal force.

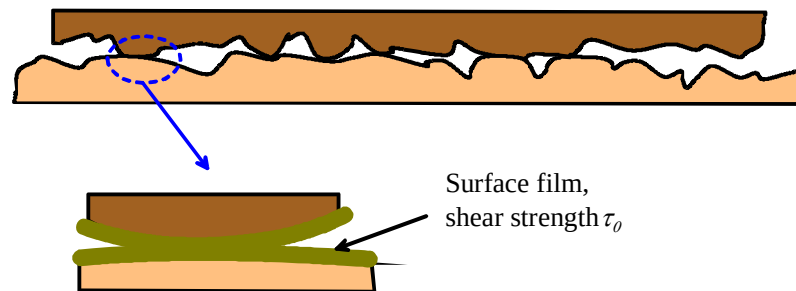
Coulomb grappled with these problems and came up with an incorrect explanation. A truly satisfactory explanation for these observations was only found 20 years or so ago.

To understand friction, we must take a close look at the nature of surfaces. Coulomb/Amontons friction laws are due to two properties of surfaces:

At present, there is no way to measure or calculate the contact C accurately.

This is true for all materials (except for rubbers, which are so compliant that the true contact area is close to the nominal contact area), and is just a consequence of the statistical properties of surface roughness. The reason that the true contact area increases in proportion to the load is that as the surfaces are pushed into contact, the *number* of asperity contacts increases, but the average size of the contacts remains the same, because of the fractal self-similarity of the two surfaces.

Finally, to understand the cause of the Coulomb/Amonton friction law, we need to visualize what happens when two rough surfaces slide against each other.



Each asperity tip is covered with a thin layer of oxide, adsorbed water, or grease. It's possible to remove this film in a lab experiment – in which case friction behavior changes dramatically and no longer follows Coulomb/Amonton law – but for real engineering surfaces it's always present.

The film usually has a low mechanical strength. It will start to deform, and so allow the two asperities to slide past each other when the tangential force per unit area acting on the film reaches the shear strength of the film τ_0 .

The tangential friction force due to shearing the film on the surface of all the contacting asperities is therefore

$$T = \tau_0 A_{true}$$

Combining this with the earlier result for the true contact area gives

$$T = \tau_0 CN$$

$$\Rightarrow \mu = \tau_0 C$$

Thus, the friction force is proportional to the normal force. This simple argument also explains why friction force is independent of contact area; why it is so sensitive to surface films, and why it can be influenced (albeit only slightly) by surface roughness.