



Introduction to Power Electronic Converters for Motor Drives

1. INTRODUCTION

In this chapter we look at examples of the power converter circuits which are widely used with motor drives, providing either d.c. or a.c. outputs, and working from either a d.c. (battery) supply, or from the conventional (50 or 60 Hz) utility supply. The coverage is not intended to be exhaustive, but rather to highlight the most important features and aspects of behavior which recur in many types of drive converter.

Although there are many different types of converter, all except very low-power ones are based on electronic switching. The need to adopt a switching strategy is emphasized in the first example, where the consequences are explored in some depth. We will see that switching is essential in order to achieve high-efficiency power conversion, but that the resulting waveforms are inevitably less than ideal from the point of view of the motor and the power supply.

The examples have been chosen to illustrate typical practice, so for each the most commonly used switching devices (e.g. thyristor, transistor) are shown. In many cases, several different switching devices may be suitable (see later), so we should not identify a particular circuit as being the exclusive preserve of a particular device.

Before discussing particular circuits it will be useful to take an overall look at a typical drive system, so that the role of the converter can be seen in its proper context.

1.1 General arrangement of drive

A complete drive system is shown in block diagram form in [Figure 2.1](#).

The job of the converter is to draw electrical energy from the utility supply (at constant voltage and frequency) and supply electrical energy to the motor at whatever voltage and frequency are necessary to achieve the desired mechanical output. In [Figure 2.1](#), the ‘demanded’ output is the speed of the motor, but equally it could be the torque, the position of the motor shaft, or some other system variable.

Except in the very simplest converter (such as a basic diode rectifier), there are two distinct parts to the converter. The first is the power stage, through which the

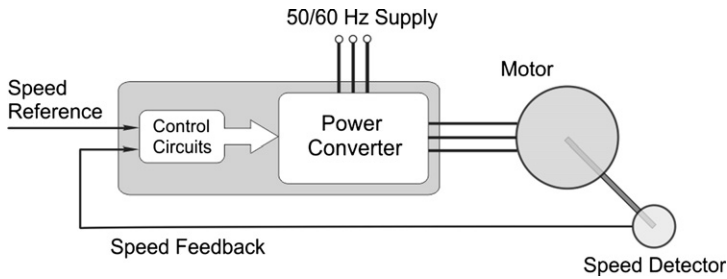


Figure 2.1 General arrangement of speed-controlled drive.

energy flows to the motor, and the second is the control section, which regulates the power flow. Low-power control signals tell the converter what it is supposed to be doing, while other low-power feedback signals are used to measure what is actually happening. By comparing the demand and feedback signals, and adjusting the output accordingly, the target output is maintained.

The basic arrangement shown in [Figure 2.1](#) is clearly a speed control system, because the signal representing the demand or reference quantity is speed, and we note that it is a ‘closed-loop’ system because the quantity that is to be controlled is measured and fed back to the controller so that action can be taken if the two signals do not correspond. All drives employ some form of closed-loop (feedback) control, so readers who are unfamiliar with the basic principles might find it helpful to read the introduction in [Appendix 1](#).

In later chapters we will explore the internal workings and control arrangements at greater length, but it is worth mentioning that all drives employ current feedback in order to control the motor torque, and that in all except high-performance drives it is unusual to find external transducers, which can account for a significant fraction of the total cost of the drive system. Instead of measuring the actual speed using, for example, a shaft-mounted tachogenerator as shown in [Figure 2.1](#), speed is more likely to be derived from sampled measurements of motor voltages, currents, and frequency, used in conjunction with a stored mathematical model of the motor.

A characteristic of power electronic converters which is shared with most electrical systems is that they have very little capacity for storing energy. This means that any sudden change in the power supplied by the converter to the motor must be reflected in a sudden increase in the power drawn from the supply. In most cases this is not a serious problem, but it does have two drawbacks. First, sudden increases in the current drawn from the supply will cause momentary drops in the supply voltage, because of the effect of the supply impedance. These voltage ‘spikes’ will appear as unwelcome distortion to other users on the same supply. And secondly, there may be an enforced delay before the supply can furnish extra power. For example, with a single-phase utility

supply, there can be no sudden increase in the power supply at the instant where the utility voltage is zero, because instantaneous power is necessarily zero at this point in the cycle since the voltage is itself zero.

It would be ideal if the converter could store at least enough energy to supply the motor for several cycles of the 50/60 Hz supply, so that short-term energy demands could be met instantly, thereby reducing rapid fluctuations in the power drawn from the mains. But unfortunately this is just not economic: most converters do have a small store of energy in their smoothing inductors and capacitors, but the amount is not sufficient to buffer the supply sufficiently to shield it from anything more than very-short-term fluctuations.

2. VOLTAGE CONTROL – D.C. OUTPUT FROM D.C. SUPPLY

In Chapter 1 we saw that control of the basic d.c. machine is achieved by controlling the current in the conductor, which is readily done by variation of the voltage. A controllable voltage source is therefore a key element of a motor drive, as we will see in later chapters.

However, for the sake of simplicity we will begin by exploring the problem of controlling the voltage across a $2\ \Omega$ resistive load, fed from a 12 V constant-voltage source such as a battery. Three different methods are shown in Figure 2.2, in which the circle on the left represents an ideal 12 V d.c. source, the tip of the arrow indicating the positive terminal. Although this set-up is not quite the same as if the load was a d.c. motor, the conclusions which we draw are more or less the same.

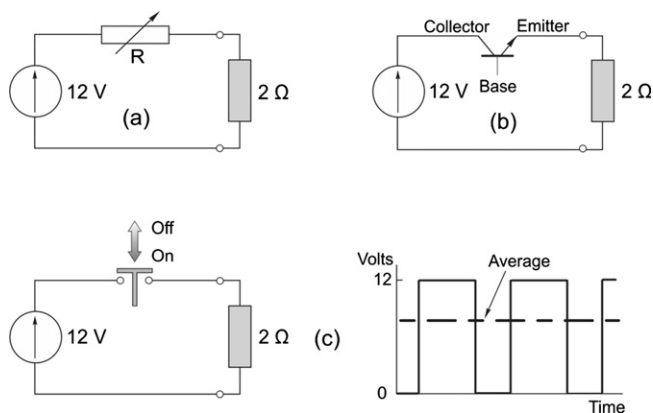


Figure 2.2 Methods of obtaining a variable-voltage output from a constant-voltage source.

Method (a) uses a variable resistor (R) to absorb whatever fraction of the battery voltage is not required at the load. It provides smooth (albeit manual) control over the full range from 0 to 12 V, but the snag is that power is wasted in the control resistor. For example, if the load voltage is to be reduced to 6 V, the resistor R must be set to $2\ \Omega$, so that half of the battery voltage is dropped across R . The current will be 3 A, the load power will be 18 W, and the power dissipated in R will also be 18 W. In terms of overall power conversion efficiency (i.e. useful power delivered to the load divided by total power from the source), the efficiency is a very poor 50%. If R is increased further, the efficiency falls still lower, approaching zero as the load voltage tends to zero. This method of control is therefore unacceptable for motor control, except perhaps in applications such as toy racing cars.

Method (b) is much the same as (a) except that a transistor is used instead of a manually operated variable resistor. The transistor in [Figure 2.2\(b\)](#) is connected with its collector and emitter terminals in series with the voltage source and the load resistor. The transistor is a variable resistor, of course, but a rather special one in which the effective collector–emitter resistance can be controlled over a wide range by means of the base–emitter current. The base–emitter current is usually very small, so it can be varied by means of a low-power electronic circuit (not shown in [Figure 2.2](#)) whose losses are negligible in comparison with the power in the main (collector–emitter) circuit.

Method (b) shares the drawback of method (a) above, i.e. the efficiency is very low. But even more seriously, the ‘wasted’ power (up to a maximum of 18 W in this case) is burned off inside the transistor, which therefore has to be large, well-cooled, and hence expensive. Transistors are never operated in this ‘linear’ way when used in power electronics, but are widely used as switches, as discussed below.

2.1 Switching control

The basic ideas underlying a switching power regulator are shown by the arrangement in [Figure 2.2\(c\)](#), which uses a mechanical switch. By operating the switch repetitively and varying the ratio of ‘on’ to ‘off’ time, the average load voltage can be varied continuously between 0 V (switch off all the time) through 6 V (switch on and off for half of each cycle) to 12 V (switch on all the time).

The circuit shown in [Figure 2.2\(c\)](#) is often referred to as a ‘chopper’, because the battery supply is ‘chopped’ on and off. A constant repetition frequency is normally used, and the width of the on pulse is varied to control the mean output voltage (see the waveform in [Figure 2.2](#)): this is known as ‘pulse-width modulation’ (PWM).

The main advantage of the chopper circuit is that no power is wasted, and the efficiency is thus 100%. When the switch is on, current flows through it, but the voltage across it is zero because its resistance is negligible. The power dissipated in the switch is therefore zero. Likewise, when ‘off’, the current through it is zero, so

although the voltage across the switch is 12 V, the power dissipated in it is again zero.

The obvious disadvantage is that by no stretch of the imagination could the load voltage be seen as ‘good’ d.c.: instead it consists of a mean or ‘d.c.’ level, with a superimposed ‘a.c.’ component. Bearing in mind that we really want the load to be a d.c. motor, rather than a resistor, we are bound to ask whether the pulsating voltage will be acceptable. Fortunately, the answer is yes, provided that the chopping frequency is high enough. We will see later that the inductance of the motor causes the current to be much smoother than the voltage, which means that the motor torque fluctuates much less than we might suppose; and the mechanical inertia of the motor filters the torque ripples so that the speed remains almost constant, at a value governed by the mean (or d.c.) level of the chopped waveform.

Obviously a mechanical switch would be unsuitable, and could not be expected to last long when pulsed at high frequency. So an electronic power switch is used instead. The first of many devices to be used for switching was the bipolar junction transistor (BJT), so we will begin by examining how such devices are employed in chopper circuits. If we choose a different device, such as a metal oxide semiconductor field effect transistor (MOSFET) or an insulated gate bipolar transistor (IGBT), the detailed arrangements for turning the device on and off will be different, but the main conclusions we draw will be much the same.

2.2 Transistor chopper

As noted earlier, a transistor is effectively a controllable resistor, i.e. the resistance between collector and emitter depends on the current in the base–emitter junction. In order to mimic the operation of a mechanical switch, the transistor would have to be able to provide infinite resistance (corresponding to an open switch) or zero resistance (corresponding to a closed switch). Neither of these ideal states can be reached with a real transistor, but both can be approximated closely.

The typical relationship between the collector–emitter voltage and the collector current for a range of base currents rising from zero is shown in Figure 2.3. The bulk of the diagram represents the so-called ‘linear’ region, where the transistor exhibits the remarkable property of the collector current remaining more or less constant for a wide range of collector–emitter voltages: when the transistor operates in this region there will be significant power loss. For power electronic applications, we want the device to behave like a switch, so we operate on the margins of the diagram, where either the voltage or the current is close to zero, and the heat released inside the device is therefore very low.

The transistor will be ‘off’ when the base–emitter current (I_b) is zero. Viewed from the main (collector–emitter) circuit, its resistance will be very high, as shown by the region *Oa* in Figure 2.3.

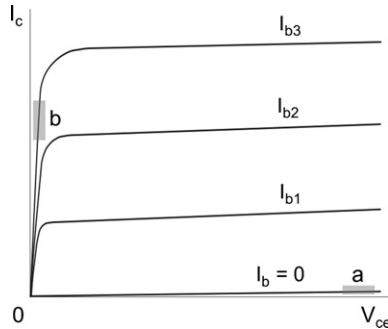


Figure 2.3 Transistor characteristics showing high-resistance (cut-off) region *Oa* and low-resistance (saturation) region *Ob*. Typical ‘off’ and ‘on’ operating states are shown by the shaded areas *a* and *b*, respectively.

Under this ‘cut-off’ condition, only a tiny current (I_c) can flow from the collector to the emitter, regardless of the voltage (V_{ce}) between the collector and emitter. The power dissipated in the device will therefore be negligible, giving an excellent approximation to an open switch.

To turn the transistor fully ‘on’, a base–emitter current must be provided. The base current required will depend on the prospective collector–emitter current, i.e. the current in the load. The aim is to keep the transistor ‘saturated’ so that it has a very low resistance, corresponding to the region *Ob* in Figure 2.3. In the example shown in Figure 2.2, if the resistance of the transistor is very low, the current in the circuit will be almost 6 A, so we must make sure that the base–emitter current is sufficiently large to ensure that the transistor remains in the saturated condition when $I_c = 6$ A.

Typically in a bipolar transistor (BJT) the base current will need to be around 5–10% of the collector current to keep the transistor in the saturation region: in the example (Figure 2.2), with the full load current of 6 A flowing, the base current might be 400 mA, and the collector–emitter voltage might be say 0.33 V, giving an on-state dissipation of 2 W in the transistor when the load power is 72 W. The power conversion efficiency is not 100%, as it would be with an ideal switch, but it is acceptable.

We should note that the on-state base–emitter voltage is very low, which, coupled with the small base current, means that the power required to drive the transistor is very much less than the power being switched in the collector–emitter circuit. Nevertheless, to switch the transistor in the regular pattern shown in Figure 2.2, we obviously need a base current waveform which goes on and off periodically, and we might wonder how we obtain this ‘control’ signal. In most modern drives the signal originates from a microprocessor, many of which are designed with PWM auxiliary functions which can be used for this purpose. Depending on the base circuit power requirements of the main switching transistor,

it may be possible to feed it directly from the microprocessor, but it is more usual to see additional transistors interposed between the signal source and the main device to provide the required power amplification.

Just as we have to select mechanical switches with regard to their duty, we must be careful to use the right power transistor for the job in hand. In particular, we need to ensure that when the transistor is 'on', we don't exceed the safe current, or else the active semiconductor region of the device will be destroyed by overheating. And we must make sure that the transistor is able to withstand whatever voltage appears across the collector-emitter junction when it is in the 'off' condition. If the safe voltage is exceeded, the transistor will break down, and be permanently 'on'.

A suitable heatsink will be a necessity. We have already seen that some heat is generated when the transistor is on, and at low switching rates this is the main source of unwanted heat. But at high switching rates, 'switching loss' can also be very important.

Switching loss is the heat generated in the finite time it takes for the transistor to go from on to off or vice versa. The base-drive circuitry will be arranged so that the switching takes place as fast as possible, but in practice it will seldom take less than a few microseconds. During the switch-on period, for example, the current will be building up, while the collector-emitter voltage will be falling towards zero. The peak power reached can therefore be large, before falling to the relatively low on-state value. Of course the total energy released as heat each time the device switches is modest because the whole process happens so quickly. Hence if the switching rate is low (say once every second) the switching power loss will be insignificant in comparison with the on-state power. But at high switching rates, when the time taken to complete the switching becomes comparable with the on time, the switching power loss can easily become dominant. In drives, switching rates from hundreds of hertz to tens of kilohertz are used: higher frequencies would be desirable from the point of view of smoothness of supply, but cannot be used because the resultant high switching loss becomes unacceptable.

2.3 Chopper with inductive load – overvoltage protection

So far we have looked at chopper control of a resistive load, but in a drives context the load will usually mean the winding of a machine, which will invariably be inductive.

Chopper control of inductive loads is much the same as for resistive loads, but we have to be careful to prevent the appearance of dangerously high voltages each time the inductive load is switched 'off'. The root of the problem lies with the energy stored in the magnetic field of the inductor. When an inductance L carries a current I , the energy stored in the magnetic field (W) is given by

$$W = \frac{1}{2}LI^2 \quad (2.1)$$

If the inductor is supplied via a mechanical switch, and we open the switch with the intention of reducing the current to zero instantaneously, we are in effect seeking to destroy the stored energy. This is not possible, and what happens is that the energy is dissipated in the form of a spark across the contacts of the switch.

The appearance of a spark indicates that there is a very high voltage which is sufficient to break down the surrounding air. We can anticipate this by remembering that the voltage and current in an inductance are related by the equation

$$V_L = L \frac{di}{dt} \quad (2.2)$$

which shows that the self-induced voltage is proportional to the rate of change of current, so when we open the switch in order to force the current to zero quickly, a very large voltage is created in the inductance. This voltage appears across the terminals of the switch and, if sufficient to break down the air, the resulting arc allows the current to continue to flow until the stored magnetic energy is dissipated as heat in the arc.

Sparking across a mechanical switch is unlikely to cause immediate destruction, but when a transistor is used sudden death is certain unless steps are taken to tame the stored energy. The usual remedy lies in the use of a ‘freewheel diode’ (sometimes called a flywheel diode), as shown in [Figure 2.4](#).

A diode is a one-way valve as far as current is concerned: it offers very little resistance to current flowing from anode to cathode (i.e. in the direction of the broad arrow in the symbol for a diode), but blocks current flow from cathode to

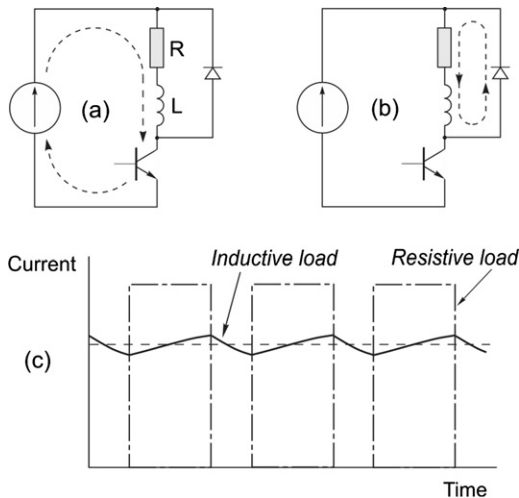


Figure 2.4 Operation of chopper-type voltage regulator.

anode. Actually, when a power diode conducts in the forward direction, the voltage drop across it is usually not all that dependent on the current flowing through it, so the reference above to the diode ‘offering little resistance’ is not strictly accurate because it does not obey Ohm’s law. In practice the volt-drop of power diodes (most of which are made from silicon) is around 0.7 V, regardless of the current rating.

In the circuit of Figure 2.4(a), when the transistor is on, current (I) flows through the load, but not through the diode, which is said to be reverse-biased (i.e. the applied voltage is trying – unsuccessfully – to push current down through the diode). During this period the voltage across the inductance is positive, so the current increases, thereby increasing the stored energy.

When the transistor is turned off, the current through it and the battery drops very quickly to zero. But the stored energy in the inductance means that its current cannot suddenly disappear. So since there is no longer a path through the transistor, the current diverts into the only other route available, and flows upwards through the low-resistance path offered by the diode, as shown in Figure 2.4(b).

Obviously the current no longer has a battery to drive it, so it cannot continue to flow indefinitely. During this period the voltage across the inductance is negative, and the current reduces. If the transistor were left ‘off’ for a long period, the current would continue to ‘freewheel’ only until the energy originally stored in the inductance is dissipated as heat, mainly in the load resistance but also in the diode’s own (low) resistance. In normal chopping, however, the cycle restarts long before the current has fallen to zero, giving a current waveform as shown in Figure 2.4(c). Note that the current rises and falls exponentially with a time-constant of L/R , though it never reaches anywhere near its steady-state value in Figure 2.4. The sketch corresponds to the case where the time-constant is much longer than one switching period, in which case the current becomes almost smooth, with only a small ripple. In a d.c. motor drive this is just what we want, since any fluctuation in the current gives rise to torque pulsations and consequent mechanical vibrations. (The current waveform that would be obtained with no inductance is also shown in Figure 2.4: the mean current is the same but the rectangular current waveform is clearly much less desirable, given that ideally we would like constant d.c.)

The freewheel (or flywheel) diode was introduced to prevent dangerously high voltages from appearing across the transistor when it switches off an inductive load, so we should check that this has been achieved. When the diode conducts, the volt-drop across it is small – typically 0.7 volts. Hence while the current is freewheeling, the voltage at the collector of the transistor is only 0.7 volts above the battery voltage. This ‘clamping’ action therefore limits the voltage across the transistor to a safe value, and allows inductive loads to be switched without damage to the switching element.

We should acknowledge that in this example the discussion has focused on steady-state operation, when the current at the end of every cycle is the same, and it

never falls to zero. We have therefore sidestepped the more complex matter of how we get from start-up to the steady state, and we have also ignored the so-called ‘discontinuous current’ mode. We will touch on the significant consequences of discontinuous operation in drives in later chapters.

We can draw some important conclusions which are valid for all power electronic converters from this simple example. First, efficient control of voltage (and hence power) is only feasible if a switching strategy is adopted. The load is alternately connected and disconnected from the supply by means of an electronic switch, and any average voltage up to the supply voltage can be obtained by varying the mark/space ratio. Secondly, the output voltage is not smooth d.c., but contains unwanted a.c. components which, though undesirable, are tolerable in motor drives. And finally, the load current waveform will be smoother than the voltage waveform if – as is the case with motor windings – the load is inductive.

2.4 Boost converter

The previous section dealt with the so-called step-down or buck converter, which provides an output voltage less than the input. However, if the motor voltage is higher than the supply (for example, in an electric vehicle driven by a 240 V motor from a battery of 48 V), a step-up or boost converter is required. Intuitively this seems a tougher challenge altogether, and we might expect first to have to convert the d.c. to a.c. so that a transformer could be used, but in fact the basic principle of transferring ‘packets’ of energy to a higher voltage is very simple and elegant, using the circuit shown in Figure 2.5. Operation of this circuit is worth discussing because it again illustrates features common to power-electronic converters.

As is usual, the converter operates repetitively at a rate determined by the frequency of switching on and off the transistor (T). During the ‘on’ period

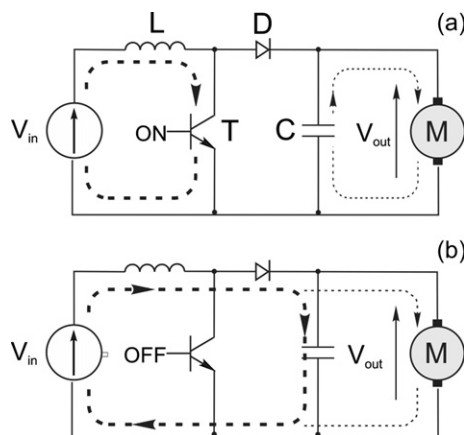


Figure 2.5 Boost converter. The transistor is switched on in (a) and off in (b).

(Figure 2.5(a)), the input voltage (V_{in}) is applied across the inductor (L), causing the current in the inductor to rise linearly, thereby increasing the energy stored in its magnetic field. Meanwhile, the motor current is supplied by the storage capacitor (C), the voltage of which falls only a little during this discharge period. In Figure 2.5(a), the input current is drawn to appear larger than the output (motor) current, for reasons that will soon become apparent. Recalling that the aim is to produce an output voltage greater than the input, it should be clear that the voltage across the diode (D) is negative (i.e. the potential is higher on the right than on the left in Figure 2.5), so the diode does not conduct and the input and output circuits are effectively isolated from each other.

When the transistor turns off, the current through it falls rapidly to zero, and the situation is much the same as in the step-down converter where we saw that, because of the stored energy in the inductor, any attempt to reduce its current results in a self-induced voltage trying to keep the current going. So during the ‘off’ period, the inductor voltage rises extremely rapidly until the potential at the left side of the diode is slightly greater than V_{out} , the diode then conducts and the inductor current flows into the parallel circuit consisting of the capacitor (C) and the motor, the latter continuing to draw a steady current, while the major share of the inductor current goes to recharge the capacitor. The resulting voltage across the inductor is negative, and of magnitude $V_{out} - V_{in}$, so the inductor current begins to reduce and the extra energy that was stored in the inductor during the ‘on’ time is transferred into the capacitor. Assuming that we are in the steady state (i.e. power is being supplied to the motor at a constant voltage and current), the capacitor voltage will return to its starting value at the start of the next ‘on’ time.

If the losses in the transistor and the storage elements are neglected, it is easy to show that the converter functions like an ideal transformer, with output power equal to input power, i.e.

$$V_{in} \times I_{in} = V_{out} \times I_{out}, \text{ or } \frac{V_{out}}{V_{in}} = \frac{I_{in}}{I_{out}}$$

We know that the motor voltage V_{out} is higher than V_{in} , so not surprisingly the quid pro quo is that the motor current is less than the input current, as indicated by line thickness in Figure 2.5.

Also if the capacitor is large enough to hold the output voltage very nearly constant throughout, it is easy to show that the step-up ratio is given by

$$\frac{V_{out}}{V_{in}} = \frac{1}{1 - D}$$

where D is the duty ratio, i.e. the proportion of each cycle for which the transistor is on. Thus for the example at the beginning of this section, where $V_{out}/V_{in} = 240/48 = 5$, the duty ratio is 0.8, i.e. the transistor has to be on for 80% of each cycle.

The advantage of switching at a high frequency becomes clear when we recall that the capacitor has to store enough energy to supply the output during the ‘on’ period, so if, as is often the case, we want the output voltage to remain almost constant, it should be clear that the capacitor must store a good deal more energy than it gives out each cycle. Given that the size and cost of capacitors depends on the energy they have to store, it is clearly better to supply small packets of energy at a high rate, rather than use a lower frequency that requires more energy to be stored. Conversely, the switching and other losses increase with frequency, so a compromise is inevitable.

3. D.C. FROM A.C. – CONTROLLED RECTIFICATION

The vast majority of drives of all types draw their power from a constant-voltage 50 or 60 Hz utility supply, and in nearly all converters the first stage consists of a rectifier which converts the a.c. to a crude form of d.c. Where a constant-voltage (i.e. unvarying average) ‘d.c.’ output is required, a simple (uncontrolled) diode rectifier is sufficient. But where the mean d.c. voltage has to be controllable (as in a d.c. motor drive to obtain varying speeds), a controlled rectifier is used.

Many different converter configurations based on combinations of diodes and thyristors are possible, but we will focus on ‘fully controlled’ converters in which all the rectifying devices are thyristors, because they are predominant in modern motor drives. Half-controlled converters are used less frequently, so will not be covered here.

From the user’s viewpoint, interest centers on the following questions:

- How is the output voltage controlled?
- What does the converter output voltage look like? Will there be any problems if the voltage is not pure d.c.?
- How does the range of the output voltage relate to the utility supply voltage?
- How is the converter and motor drive ‘seen’ by the supply system? What is the power-factor, and is there waveform distortion and interference to other users?

We can answer these questions without going too thoroughly into the detailed workings of the converter. This is just as well, because understanding all the ins and outs of converter operation is beyond our scope. On the other hand, it is well worth trying to understand the essence of the controlled rectification process, because it assists in understanding the limitations which the converter puts on drive performance (see Chapter 4, etc.). Before tackling the questions posed above, however, it is obviously necessary to introduce the thyristor.

3.1 The thyristor

The thyristor is an electronic switch, with two main terminals (anode and cathode) and a ‘switch-on’ terminal (gate), as shown in [Figure 2.6](#). Like in a diode, current

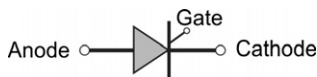


Figure 2.6 Circuit diagram of thyristor.

can only flow in the forward direction, from anode to cathode. But unlike a diode, which will conduct in the forward direction as soon as forward voltage is applied, the thyristor will continue to block forward current until a small current pulse is injected into the gate–cathode circuit, to turn it on or ‘fire’ it. After the gate pulse is applied, the main anode–cathode current builds up rapidly, and as soon as it reaches the ‘latching’ level, the gate pulse can be removed and the device will remain ‘on’.

Once established, the anode–cathode current cannot be interrupted by any gate signal. The non-conducting state can only be restored after the anode–cathode current has reduced to zero, and has remained at zero for the turn-off time (typically 100–200 μs).

When a thyristor is conducting it approximates to a closed switch, with a forward drop of only one or two volts over a wide range of current. Despite the low volt-drop in the ‘on’ state, heat is dissipated, and heatsinks must usually be provided, perhaps with fan cooling. Devices must be selected with regard to the voltages to be blocked and the r.m.s. and peak currents to be carried. Their overcurrent capability is very limited, and it is usual in drives for devices to have to withstand perhaps twice full-load current for a few seconds only. Special fuses must be fitted to protect against heavy fault currents.

The reader may be wondering why we need the thyristor, since in the previous section we discussed how a transistor could be used as an electronic switch. On the face of it the transistor appears even better than the thyristor because it can be switched off while the current is flowing, whereas the thyristor will remain on until the current through it has been reduced to zero by external means. The primary reason for the use of thyristors is that they are cheaper and their voltage and current ratings extend to higher levels than in power transistors. In addition, the circuit configuration in rectifiers is such that there is no need for the thyristor to be able to interrupt the flow of current, so its inability to do so is no disadvantage. Of course there are other circuits (see, for example, the next section dealing with inverters) where the devices need to be able to switch off on demand, in which case the transistor then has the edge over the thyristor.

3.2 Single pulse rectifier

The simplest phase-controlled rectifier circuit is shown in [Figure 2.7](#). When the supply voltage is positive, the thyristor blocks forward current until the gate pulse arrives, and up to this point the voltage across the resistive load is zero. As soon as a firing pulse is delivered to the gate–cathode circuit (not shown in [Figure 2.7](#)) the device turns on, the voltage across it falls to near zero, and the load voltage becomes

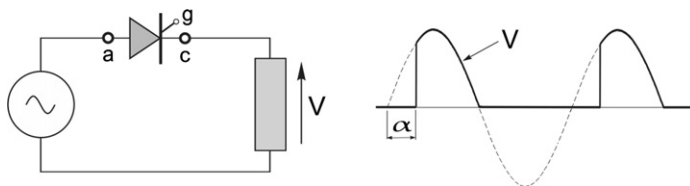


Figure 2.7 Simple single-pulse thyristor-controlled rectifier, with resistive load and firing angle delay α .

equal to the supply voltage. When the supply voltage reaches zero, so does the current. At this point the thyristor regains its blocking ability, and no current flows during the negative half-cycle.

If we neglect the small on-state volt-drop across the thyristor, the load voltage (Figure 2.7) will consist of part of the positive half-cycles of the a.c. supply voltage. It is obviously not smooth, but is ‘d.c.’ in the sense that it has a positive mean value; and by varying the delay angle (α) of the firing pulses the mean voltage can be controlled. With a purely resistive load, the current waveform will simply be a scaled version of the voltage.

This arrangement gives only one peak in the rectified output for each complete cycle of the supply, and is therefore known as a ‘single-pulse’ or half-wave circuit. The output voltage (which ideally we would like to be steady d.c.) is so poor that this circuit is never used in drives. Instead, drive converters use four or six thyristors, and produce much superior output waveforms with two or six pulses per cycle, as will be seen in the following sections.

3.3 Single-phase fully controlled converter – output voltage and control

The main elements of the converter circuit are shown in Figure 2.8. It comprises four thyristors, connected in bridge formation. (The term bridge stems from early four-arm measuring circuits which presumably suggested a bridge-like structure to their inventors.)

The conventional way of drawing the circuit is shown in Figure 2.8(a), while in Figure 2.8(b) it has been redrawn to assist understanding. The top of the load can be connected (via T1) to terminal A of the supply, or (via T2) to terminal B of the supply, and likewise the bottom of the load can be connected either to A or to B via T3 or T4, respectively.

We are naturally interested to find what the output voltage waveform on the d.c. side will look like, and in particular to discover how the mean d.c. level can be controlled by varying the firing delay angle α . The angle α is measured from the point on the waveform when a diode in the same circuit position would start to conduct, i.e. when the anode becomes positive with respect to the cathode.

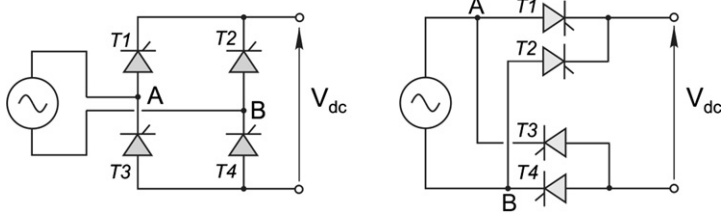


Figure 2.8 Single-phase 2-pulse (full-wave) fully controlled rectifier.

This is not such a simple matter as we might have expected, because it turns out that the mean voltage for a given α depends on the nature of the load. We will therefore look first at the case where the load is resistive, and explore the basic mechanism of phase control. Later, we will see how the converter behaves with a typical motor load.

3.3.1 Resistive load

Thyristors T1 and T4 are fired together when terminal A of the supply is positive, while on the other half-cycle, when B is positive, thyristors T2 and T3 are fired simultaneously. The output voltage and current waveform are shown in Figure 2.9(a) and (b), respectively, the current simply being a replica of the voltage. There are two pulses per mains cycle, hence the description ‘2-pulse’ or full-wave.

At every instant the load is either connected to the mains by the pair of switches T1 and T4, or it is connected the other way up by the pair of switches T2 and T3, or it is disconnected. The load voltage therefore consists of rectified chunks of the incoming supply voltage. It is much smoother than in the single-pulse circuit, though again it is far from pure d.c.

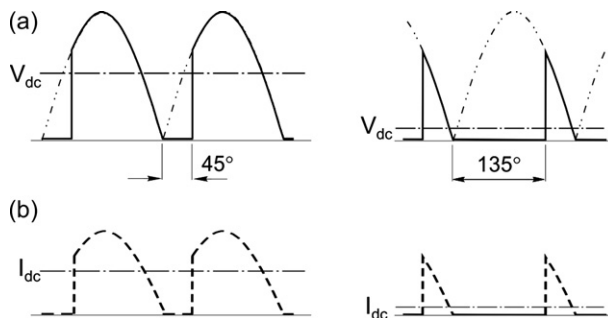


Figure 2.9 Output voltage waveform (a) and current (b) of single-phase fully controlled rectifier with resistive load, for firing angle delays of 45° and 135° .

The waveforms in Figure 2.9 correspond to $\alpha = 45^\circ$ and $\alpha = 135^\circ$, respectively. The mean value, V_{dc} , is shown in each case. It is clear that the larger the delay angle, the lower the output voltage. The maximum output voltage (V_{do}) is obtained with $\alpha = 0^\circ$: this is the same as would be obtained if the thyristors were replaced by diodes, and is given by

$$V_{do} = \frac{2}{\pi} \sqrt{2} V_{rms} \quad (2.3)$$

where V_{rms} is the r.m.s. voltage of the incoming supply. The variation of the mean d.c. voltage with α is given by

$$V_{dc} = \left\{ \frac{1 + \cos \alpha}{2} \right\} V_{do} \quad (2.4)$$

from which we see that with a resistive load the d.c. voltage can be varied from a maximum of V_{do} down to zero by varying α from 0° to 180° .

3.3.2 Inductive (motor) load

As mentioned above, motor loads are inductive, and we have seen earlier that the current cannot change instantaneously in an inductive load. We must therefore expect the behavior of the converter with an inductive load to differ from that with a resistive load, in which the current was seen to change instantaneously.

The realization that the mean voltage for a given firing angle might depend on the nature of the load is a most unwelcome prospect. What we would like is to be able to say that, regardless of the load, we can specify the output voltage waveform once we have fixed the delay angle α . We would then know what value of α to select to achieve any desired mean output voltage. What we find in practice is that once we have fixed α , the mean output voltage with a resistive-inductive load is not the same as with a purely resistive load, and therefore we cannot give a simple general formula for the mean output voltage in terms of α . This is of course very undesirable: if, for example, we had set the speed of our unloaded d.c. motor to the target value by adjusting the firing angle of the converter to produce the correct mean voltage, the last thing we would want is for the voltage to fall when the load current drawn by the motor increased, as this would cause the speed to fall below the target.

Fortunately, however, it turns out that the output voltage waveform for a given α does become independent of the load inductance once there is sufficient inductance to prevent the load current from ever falling to zero. This condition is known as ‘continuous current’, and, happily, many motor circuits do have sufficient self-inductance to ensure that we achieve continuous current. Under continuous current conditions, the output voltage waveform only depends on the firing angle, and not on the actual inductance present. This makes things much more

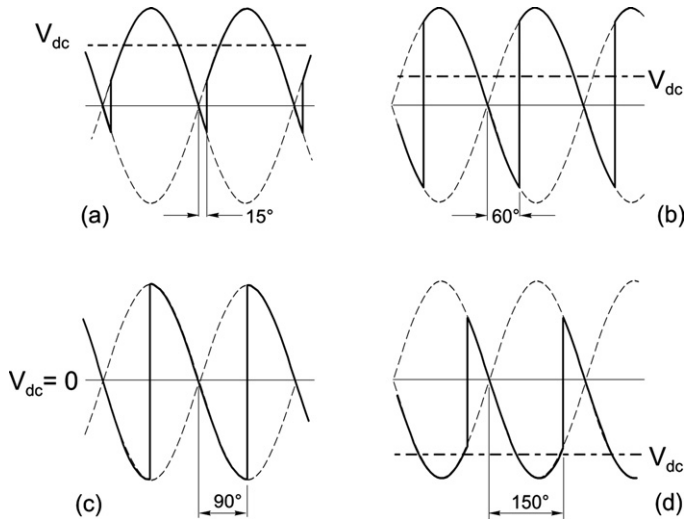


Figure 2.10 Output voltage waveforms of single-phase fully controlled rectifier supplying an inductive (motor) load, for various firing angles.

straightforward, and typical output voltage waveforms for this continuous current condition are shown in [Figure 2.10](#).

The waveforms in [Figure 2.10](#) show that, as with the resistive load, the larger the delay angle the lower the mean output voltage. However, with the resistive load the output voltage was never negative, whereas we see that, although the mean voltage is positive for values of α below 90° , there are brief periods when the output voltage becomes negative. This is because the inductance smoothes out the current (see [Figure 4.2](#), for example) so that at no time does it fall to zero. As a result, one or other pair of thyristors is always conducting, so at every instant the load is connected directly to the supply, and therefore the load voltage always consists of chunks of the supply voltage.

Rather surprisingly, we see that when α is greater than 90° , the average voltage is negative (though, of course, the current is still positive). The fact that we can obtain a net negative output voltage with an inductive load contrasts sharply with the resistive load case, where the output voltage could never be negative. The combination of negative voltage and positive current means that the power flow is reversed, and energy is fed back to the supply system. We will see later that this facility allows the converter to return energy from the load to the supply, and this is important when we want to use the converter with a d.c. motor in the regenerating mode.

It is not immediately obvious why the current switches over (or ‘commutates’) from the first pair of thyristors to the second pair when the latter are fired,

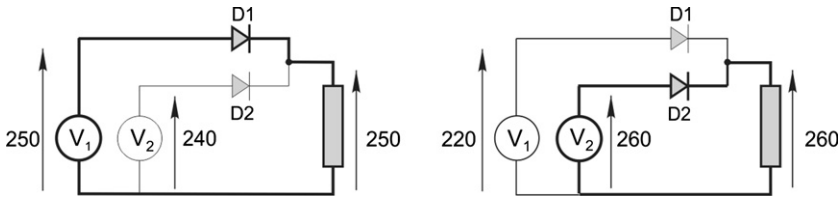


Figure 2.11 Diagram illustrating commutation between diodes: the current flows through the diode with the higher anode potential.

so a brief look at the behavior of diodes in a similar circuit configuration may be helpful at this point. Consider the set-up shown in Figure 2.11, with two voltage sources (each time-varying) supplying a load via two diodes. The question is, what determines which diode conducts, and how does this influence the load voltage?

We can consider two instants as shown in the diagram. On the left, V_1 is 250 V, V_2 is 240 V, and D_1 is conducting, as shown by the heavy line. If we ignore the volt-drop across the diode, the load voltage will be 250 V, and the voltage across diode D_2 will be $240 - 250 = -10$ V, i.e. it is reverse-biased and hence in the non-conducting state. At some other instant (on the right of the diagram), V_1 has fallen to 220 V while V_2 has increased to 260 V: now D_2 is conducting instead of D_1 , again shown by the heavy line, and D_1 is reverse-biased by -40 V. The simple pattern is that the diode with the highest anode potential will conduct, and as soon as it does so it automatically reverse-biases its predecessor. In a single-phase diode bridge, for example, the commutation occurs at the point where the supply voltage passes through zero: at this instant the anode voltage on one pair goes from positive to negative, while on the other pair the anode voltage goes from negative to positive.

The situation in controlled thyristor bridges is very similar, except that before a new device can take over conduction, it must not only have a higher anode potential, but it must also receive a firing pulse. This allows the changeover to be delayed beyond the point of natural (diode) commutation by the angle α , as shown in Figure 2.10. Note that the maximum mean voltage (V_{d0}) is again obtained when α is zero, and is the same as for the resistive load (equation (2.3)). It is easy to show that the mean d.c. voltage is now related to α by

$$V_{dc} = V_{d0} \cos \alpha \quad (2.5)$$

This equation indicates that we can control the mean output voltage by controlling α , though equation (2.5) shows that the variation of mean voltage with α is different from that for a resistive load (equation (2.4)), not least because when α is greater than 90° the mean output voltage is negative.

It is sometimes suggested (particularly by those with a light-current background) that a capacitor could be used to smooth the output voltage, this being common

practice in cheap low-power d.c. supplies. However, the power levels in most drives are such that in order to store enough energy to smooth the voltage waveform over the half-cycle of the utility supply (20 ms at 50 Hz), very bulky and expensive capacitors would be required. Fortunately, as will be seen later, it is not necessary for the voltage to be smooth as it is the current which directly determines the torque, and as already pointed out the current is always much smoother than the voltage because of inductance.

3.4 Three-phase fully controlled converter

The main power elements are shown in Figure 2.12. The 3-phase bridge has only two more thyristors than the single-phase bridge, but the output voltage waveform is vastly better, as shown in Figure 2.13. There are now six pulses of the output voltage per cycle, hence the description ‘6-pulse’. The thyristors are again fired in

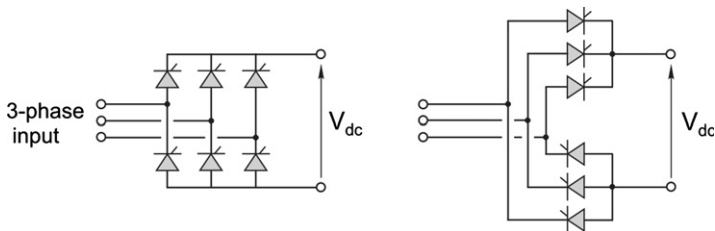


Figure 2.12 Three-phase fully controlled thyristor converter. (The alternative diagram (right) is intended to assist understanding.)

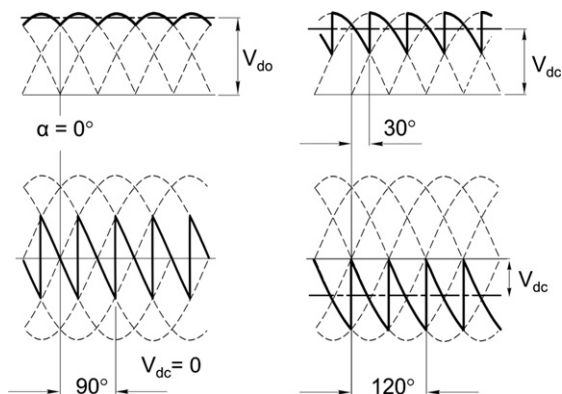


Figure 2.13 Output voltage waveforms for 3-phase fully controlled thyristor converter supplying an inductive (motor) load, for various firing angles from 0° to 120° . The mean d.c. voltage is shown by the horizontal line, except for $\alpha = 90^\circ$ where the mean d.c. voltage is zero.

pairs (one in the top half of the bridge and one – from a different leg – in the bottom half), and each thyristor carries the output current for one-third of the time. As in the single-phase converter, the delay angle controls the output voltage, but now $\alpha = 0$ corresponds to the point at which the phase voltages are equal (see Figure 2.13).

The enormous improvement in the smoothness of the output voltage waveform is clear when we compare Figures 2.13 and 2.10, and it underlines the benefit of choosing a 3-phase converter whenever possible. The very much better voltage waveform also means that the desirable ‘continuous current’ condition is much more likely to be met, and the waveforms in Figure 2.13 have therefore been drawn with the assumption that the load current is in fact continuous. Occasionally, even a 6-pulse waveform is not sufficiently smooth, and some very large drive converters therefore consist of two 6-pulse converters with their outputs in series. A phase-shifting transformer is used to insert a 30° shift between the a.c. supplies to the two 3-phase bridges. The resultant ripple voltage is then 12-pulse.

Returning to the 6-pulse converter, the mean output voltage can be shown to be given by

$$V_{dc} = V_{do} \cos \alpha = \frac{3}{\pi} \sqrt{2} V_{rms} \cos \alpha \quad (2.6)$$

We note that we can obtain the full range of output voltages from $+V_{do}$ to close to $-V_{do}$, so that, as with the single-phase converter, regenerative operation will be possible.

It is probably a good idea at this point to remind the reader that, in the context of this book, our first application of the controlled rectifier will be to supply a d.c. motor. When we examine the d.c. motor drive in Chapter 4, we will see that it is the average or mean value of the output voltage from the controlled rectifier that determines the speed, and it is this mean voltage that we refer to when we talk of ‘the’ voltage from the converter. We must not forget the unwanted a.c. or ripple element, however, as this can be large. For example, we see from Figure 2.13 that to obtain a very low d.c. voltage (to make the motor run very slowly) α will be close to 90° ; but if we were to connect an a.c. voltmeter to the output terminals it could register several hundred volts, depending on the incoming supply voltage!

3.5 Output voltage range

In Chapter 4 we will discuss the use of the fully controlled converter to drive a d.c. motor, so it is appropriate at this stage to look briefly at the typical voltages we can expect. Utility supply voltages vary, but single-phase supplies are usually 220–240 V, and we see from equation (2.3) that the maximum mean d.c. voltage available from a single-phase 240 V supply is 216 V. This is suitable for 180–200 V

motors. If a higher voltage is needed (for say a 300 V motor), a transformer must be used to step up the incoming supply.

Turning now to typical 3-phase supplies, the lowest 3-phase industrial voltages are usually around 380–440 V. (Higher voltages of up to 11 kV are used for large drives, but these will not be discussed here.) So with $V_{\text{rms}} = 400$ V, for example, the maximum d.c. output voltage (equation (2.6)) is 540 volts. After allowances have been made for supply variations and impedance drops, we are not able to rely on obtaining much more than 500–520 V, and it is usual for the motors used with 6-pulse drives fed from 400 V, 3-phase supplies to be rated in the range 430–500 V. (Often the motor's field winding will be supplied from single-phase 230 V, and field voltage ratings are then around 180–200 V, to allow a margin in hand from the theoretical maximum of 207 V referred to earlier.)

3.6 Firing circuits

Since the gate pulses are only of low power, the gate drive circuitry is simple and cheap. Often a single chip contains all the circuitry for generating the gate pulses, and for synchronizing them with the appropriate delay angle α with respect to the supply voltage. To avoid direct electrical connection between the high voltages in the main power circuit and the low voltages used in the control circuits, the gate pulses are coupled to the thyristor either by small pulse transformers or opto-couplers.

4. A.C. FROM D.C. – INVERSION

The business of getting a.c. from d.c. is known as inversion, and nine times out of ten we would ideally like to be able to produce sinusoidal output voltages of whatever frequency and amplitude we choose, to achieve speed control of a.c. motors. Unfortunately the constraints imposed by the necessity to use a switching strategy mean that we always have to settle for a voltage waveform which is composed of rectangular chunks, and is thus far from ideal. Nevertheless it turns out that a.c. motors are remarkably tolerant, and will operate satisfactorily despite the inferior waveforms produced by the inverter.

4.1 Single-phase inverter

We can illustrate the basis of inverter operation by considering the single-phase example shown in Figure 2.14. This inverter uses IGBTs (see later) as the switching devices, with diodes to provide the freewheel paths needed when the load is inductive.

The input or d.c. side of the inverter (on the left in Figure 2.14) is usually referred to as the 'd.c. link', reflecting the fact that in the majority of cases the d.c. is

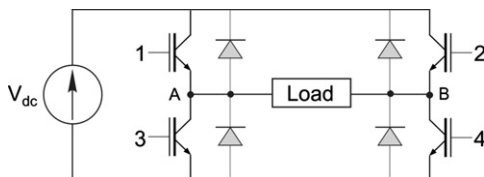


Figure 2.14 Single-phase inverter.

obtained by rectifying the incoming constant-frequency utility supply. The output or a.c. side is taken from terminals A and B in [Figure 2.14](#).

When transistors 1 and 4 are switched on, the load voltage is positive, and equal to the d.c. link voltage, while when 2 and 3 are on it is negative. If no devices are switched on, the output voltage is zero. Typical output voltage waveforms at low and high switching frequencies are shown in [Figures 2.15\(a\)](#) and (b), respectively.

Here each pair of devices is on for one-third of a cycle, and all the devices are off for two periods of one-sixth of a cycle. The output waveform is clearly not a sinewave, but at least it is alternating and symmetrical. The fundamental component is shown dotted in [Figure 2.15](#).

Within each cycle the pattern of switching is regular, and easily generated in software or programmed using appropriate logic circuitry. Frequency variation is obtained by altering the clock frequency controlling the 4-step switching pattern. The effect of varying the switching frequency is shown in [Figure 2.15](#), from which we can see that the amplitude of the fundamental component of voltage remains constant, regardless of frequency. Unfortunately (as explained in Chapter 7) this is not what we want for supplying an induction motor: to prevent the air-gap flux in the motor from falling as the frequency is raised we need to be able to increase the voltage in proportion to the frequency. We will look at voltage control shortly, after a brief digression to discuss the problem of ‘shoot-through’.

Inverters with the configurations shown in [Figures 2.14](#) and [2.17](#) are subject to a potentially damaging condition which can arise if both transistors in one ‘leg’ of the inverter inadvertently turn on simultaneously. This should never happen if the devices are switched correctly, but if something goes wrong and both devices are on together – even for a very short time – they form a short-circuit across the d.c. link. This fault condition is referred to as ‘shoot-through’ because a high current is established very rapidly, destroying the devices. A good inverter therefore includes

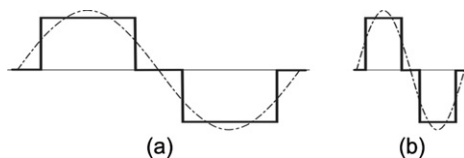


Figure 2.15 Inverter output voltage waveforms – resistive load.

provision for protecting against the possibility of shoot-through, usually by imposing a minimum time-delay, or 'dead time' between one device in the leg going off and the other coming on.

4.2 Output voltage control

There are two ways in which the amplitude of the output voltage can be controlled. First, if the d.c. link is provided from the utility supply via a controlled rectifier or from a battery via a chopper, the d.c. link voltage can be varied. We can then set the amplitude of the output voltage to any value within the range of the link. For a.c. motor drives (see Chapter 7) we can arrange for the link voltage to track the output frequency of the inverter, so that at high output frequency we obtain a high output voltage and vice versa. This method of voltage control results in a simple inverter, but requires a controlled (and thus relatively expensive) rectifier for the d.c. link.

The second method, which predominates in all sizes, achieves voltage control by PWM within the inverter itself. A cheaper uncontrolled rectifier can then be used to provide a constant-voltage d.c. link. The principle of voltage control by PWM is illustrated in Figure 2.16.

At low output frequencies, a low output voltage is usually required, so one of each pair of devices is used to chop the voltage, the mark/space ratio being varied to achieve the desired voltage at the output. The low fundamental voltage component at low frequency is shown as a broken line in Figure 2.16(a). At a higher frequency a higher voltage is needed, so the chopping device is allowed to conduct for a longer fraction of each cycle, giving the higher fundamental output shown in Figure 2.16(b). As the frequency is raised still higher, the separate 'on' periods eventually merge, giving the waveform shown in Figure 2.16(c). Any further increase in frequency takes place without further increase in the output voltage, as shown in Figure 2.16(d).

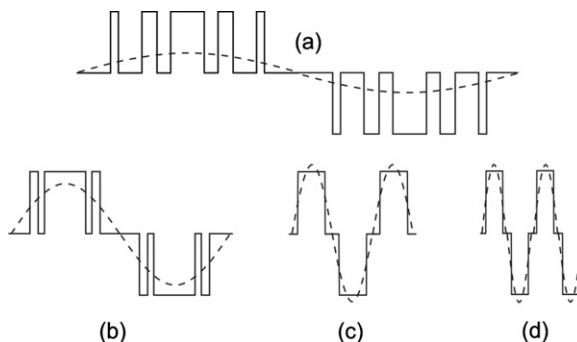


Figure 2.16 Inverter output voltage and frequency control with pulse-width modulation.

When we study a.c. drives later, we will see that the range of frequencies over which the voltage/frequency ratio can be kept constant is sometimes known as the ‘constant-torque’ region, and the upper limit of the range is usually taken to define the ‘base speed’ of the motor. Above this frequency, the inverter can no longer match voltage to frequency, the inverter effectively having run out of steam as far as voltage is concerned. The maximum voltage is thus governed by the link voltage, which must therefore be sufficiently high to provide whatever fundamental voltage the motor needs at its base speed.

Beyond the PWM region the voltage waveform is as shown in Figure 2.16(d): this waveform is usually referred to as ‘quasi-square’, though in the light of the overall object of the exercise (to approximate to a sinewave) a better description might be ‘quasi-sine’.

When supplying an inductive motor load, fast recovery freewheel diodes are needed in parallel with each device. These may be discrete devices, or fitted in a common package with the transistor, or integrated to form a single transistor/diode device.

As mentioned earlier, the switching nature of these converter circuits results in waveforms which contain not only the required fundamental component but also unwanted harmonic voltages. It is particularly important to limit the magnitude of the low-order harmonics because these are most likely to provoke an unwanted torque response from the motor, but the high-order harmonics can lead to acoustic noise if they happen to excite a mechanical resonance.

The number, width and spacing of the pulses are therefore optimized to keep the harmonic content as low as possible. There is an obvious advantage in using a high switching frequency since there are more pulses to play with. Ultrasonic frequencies are now widely used, and, as devices improve, frequencies continue to rise. Most manufacturers claim their particular system is better than the competition, but it is not clear which, if any, is best for motor operation. Some early schemes used comparatively few pulses per cycle, and changed the number of pulses in discrete steps rather than smoothly, which earned them the nickname ‘gear-changers’: their sometimes irritating sound can still be heard in older traction applications.

4.3 Three-phase inverter

Single-phase inverters are seldom used for drives, the vast majority of which use 3-phase motors because of their superior characteristics. A 3-phase output can be obtained by adding only two more switches to the four needed for a single-phase inverter, giving the typical power-circuit configuration shown in Figure 2.17. As usual, a freewheel diode is required in parallel with each transistor to protect against overvoltages caused by an inductive (motor) load. We note that the circuit configuration in Figure 2.17 is the same as for the 3-phase controlled rectifier discussed earlier. We mentioned then that the controlled rectifier could be used to

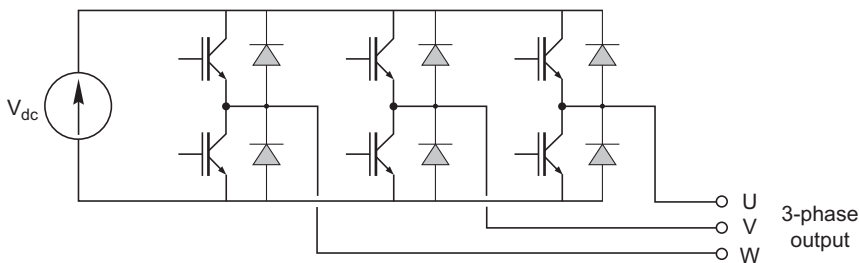


Figure 2.17 Three-phase inverter power circuit.

regenerate, i.e. to convert power from d.c. to a.c., and this is of course ‘inversion’ as we now understand it.

This circuit forms the basis of the majority of converters for motor drives, and will be explored more fully in Chapter 7. In essence, the output voltage and frequency are controlled in much the same way as for the single-phase inverter discussed in the previous section, but of course the output consists of three identical waveforms displaced by 120° from each other, typically as shown in Figure 2.18.

There are more constraints on switching with this set-up than in the single-phase case. For example, consider the upper and middle waveforms in Figure 2.18: the upper shows the potential of line U with respect to line V (V_{UV}), while the middle shows the potential of line V with respect to line W (V_{VW}). So whenever V_{UV} is positive, the upper right switch and the lower middle switch must be ‘on’. If at this instant V_{VW} is also required to be positive we have a problem because this would require the upper middle switch and the lower left switches to be on, which would lead to a short-circuit through the middle leg. Fortunately, these potential problems are limited to well-defined regions of the cycle and switching patterns include built-in delays to prevent such ‘shoot-through’ faults from occurring.

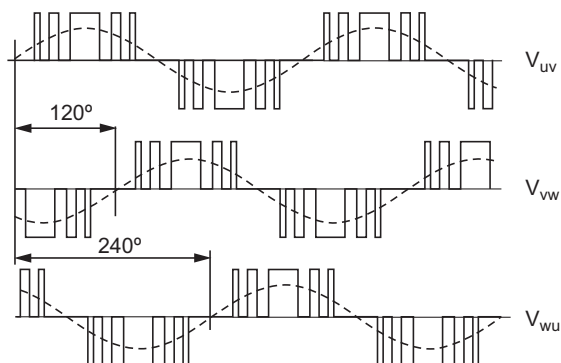


Figure 2.18 Three-phase PWM output voltage waveforms.

5. A.C. FROM A.C.

The converters we have looked at so far have involved a d.c. stage, but the ideal power-electronic converter would allow power conversion in either direction between two systems of any voltage and frequency (including d.c.), and would not involve any intermediate stage, such as a d.c. link. In principle this can be achieved by means of an array of switches that allow any one of a set of input terminals to be connected to any one of a set of output terminals, at any desired instant. The comparatively recent name for such converters is ‘matrix converter’ (see Chapter 8), but here we will take a brief look at a much older embodiment of the principle – the cycloconverter.

5.1 Cycloconverter

The cycloconverter variable-frequency drive has never become very widespread but is still suitable for very large (e.g. 1 MW and above) low-speed induction motor or synchronous motor drives. The cycloconverter is only capable of producing acceptable output waveforms at frequencies well below the utility frequency, but this, coupled with the fact that it is feasible to make large motors with high pole-numbers (e.g. 20), means that a very low-speed direct (gearless) drive becomes practicable. A 20-pole motor, for example, will have a synchronous speed of only 30 rev/min at 5 Hz, making it suitable for mine winders, kilns, crushers, etc.

The principal advantage of the cycloconverter is that naturally commutated devices (thyristors) can be used instead of self-commutating devices, which means that the cost of each device is lower and higher powers can be achieved.

The power conversion circuit for each of the three output phases is the same, so we can consider the simpler question of how to obtain a single variable-frequency supply from a 3-phase supply of fixed frequency and constant voltage. We will see that in essence, the output voltage is synthesized by switching the load directly across whichever phase of the utility supply gives the best approximation to the desired load voltage at each instant of time.

Assuming that the load is an induction motor, we will discover later in the book that the power-factor varies with load but never reaches unity, i.e. that the current is never in phase with the stator voltage. So during the positive half-cycle of the voltage waveform the current will be positive for some of the time, but negative for the remainder, while during the negative voltage half-cycle the current will be negative for some of the time and positive for the rest of the time. This means that the supply to each phase has to be able to handle any combination of both positive and negative voltage and current.

We have already explored how to achieve a variable-voltage d.c. supply using a thyristor converter, which can handle positive currents, but here we need to

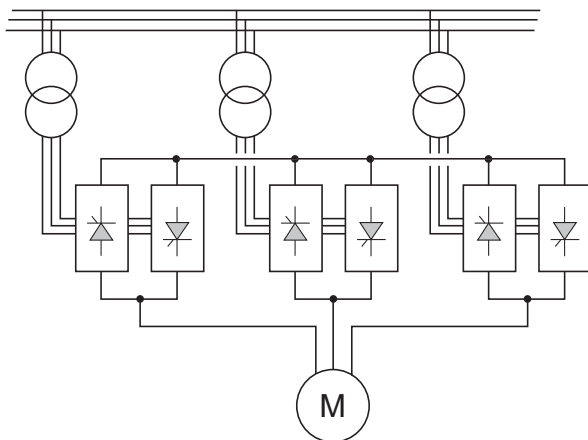


Figure 2.19 Cycloconverter power circuit.

handle negative current as well, so for each of the three motor windings we will need two converters connected in parallel, as shown in Figure 2.19, making a total of 36 thyristors. To avoid short-circuits, isolating transformers are used, again as shown in Figure 2.19.

Previously the discussion focused on the mean or d.c. level of the output voltages, but here we want to provide a low-frequency (preferably sinusoidal) output voltage for an induction motor, and the means for doing this should now be becoming clear. Once we have a double thyristor converter, we can generate a low-frequency sinusoidal output voltage simply by varying the firing angle of the positive-current bridge so that its output voltage increases from zero in a sinusoidal manner with respect to time. Then, when we have completed the positive half-cycle and arrived back at zero voltage, we bring the negative bridge into play and use it to generate the negative half-cycle, and so on.

Hence the output voltage consists of chunks of the incoming supply voltage, and it offers a reasonable approximation to the fundamental-frequency sinewave shown by the dashed line in Figure 2.20.

The output voltage waveform is certainly no worse than the voltage waveform from a d.c. link inverter, and, as we saw in that context, the current waveform in the motor will be much smoother than the voltage waveform, because of the filtering action of the stator leakage inductance. The motor performance will therefore be acceptable, despite the extra losses which arise from the unwanted harmonic components. We should note that, because each phase is supplied from a double converter, the motor can regenerate when required (e.g. to restrain an overhauling load, or to return kinetic energy to the supply when the frequency is lowered to reduce speed).

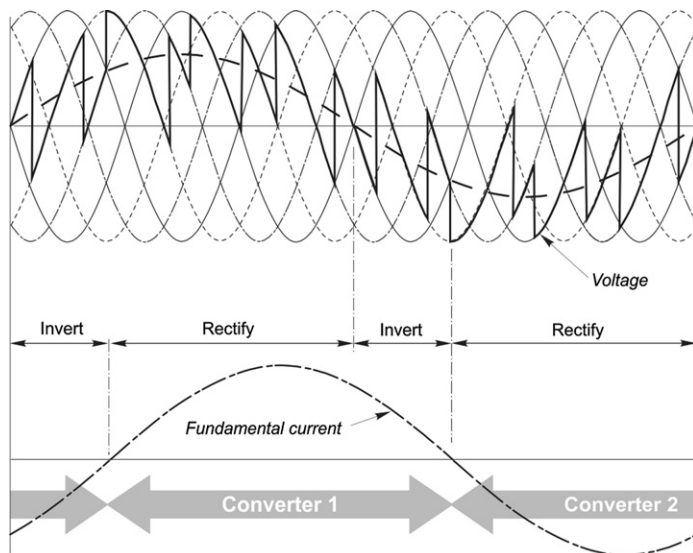


Figure 2.20 Typical output voltage waveform for one phase of 6-pulse cycloconverter supplying an inductive (motor) load. (The output frequency shown in the figure is one-third of the utility frequency, and the amplitude of the fundamental component of the output voltage (shown by the dashed line) is about 75% of the maximum that could be obtained. The fundamental component of the load current is shown in order to define the modes of operation of the converters.)

6. INVERTER SWITCHING DEVICES

As far as the user is concerned, it does not really matter what type of switching device is used inside the inverter, but it is probably helpful to mention the three most important¹ families of devices in current use so that the terminology is familiar and the symbols used for each device can be recognized. The common feature of all three devices is that they can be switched on and off by means of a low-power control signal, i.e. they are self-commutating. We have seen earlier that this ability to be turned on or off on demand is essential in any inverter which feeds a passive load, such as an induction motor.

Each device is discussed briefly below, with a broad indication of its most likely range of application. Because there is considerable overlap between competing devices, it is not possible to be dogmatic and specify which device is best, and the reader should not be surprised to find that one manufacturer may offer a 5 kW inverter which uses MOSFETs while another chooses to use IGBTs.

¹ The gate turn-off thyristor is now seldom used, so is not discussed here.

Power electronics is still developing, and there are other devices (such as those based on silicon carbide) which have yet to emerge onto the drives scene. Trends which continue include the integration of the drive and protection circuitry in the same package as the switching device (or devices); major reductions in device losses; and improved cooling techniques, all of which have already contributed to impressive reductions in product size.

6.1 Bipolar junction transistor (BJT)

Historically the bipolar junction transistor was the first to be used for power switching. Of the two versions (npn and pnp) only the npn has been widely used in inverters for drives, mainly in applications ranging up to a few kW and several hundred volts.

The npn version is shown in Figure 2.21: the main (load) current flows into the collector (C) and out of the emitter (E), as shown by the arrow on the device symbol. To switch the device on (i.e. to make the resistance of the collector–emitter circuit low, so that load current can flow) a small current must be caused to flow from the base (B) to the emitter. When the base–emitter current is zero, the resistance of the collector–emitter circuit is very high, and the device is switched off.

The advantage of the bipolar transistor is that when it is turned on, the collector–emitter voltage is low (see point b in Figure 2.3) and hence the power dissipation is small in comparison with the load power, i.e. the device is an efficient power switch. The disadvantage is that although the power required in the base–emitter circuit is tiny in comparison with the load power, it is not insignificant and in the largest power transistors can amount to several tens of watts.

6.2 Metal oxide semiconductor field effect transistor (MOSFET)

In the 1980s the power MOSFET superseded the BJT in inverters for drives. Like the BJT, the MOSFET is a three-terminal device and is available in two versions, the n-channel and the p-channel. The n-channel is the most widely used, and is shown in Figure 2.21. The main (load) current flows into the drain (D) and out of the source (S). (Confusingly, the load current in this case flows in the *opposite*

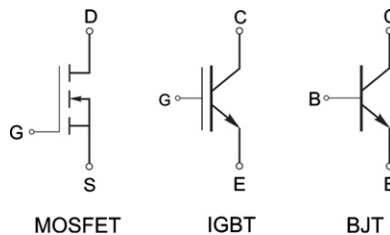


Figure 2.21 Circuit symbols for switching devices.

direction to the arrow on the symbol.) Unlike the BJT, which is controlled by the base *current*, the MOSFET is controlled by the gate-source *voltage*.

To turn the device on, the gate-source voltage must be comfortably above a threshold of a few volts. When the voltage is first applied to the gate, currents flow in the parasitic gate-source and gate-drain capacitances, but once these capacitances have been charged the input current to the gate is negligible, so the steady-state gate drive power is minimal. To turn the device off, the parasitic capacitances must be discharged and the gate-source voltage must be held below the threshold level.

The principal advantage of the MOSFET is that it is a voltage-controlled device which requires negligible power to hold it in the on state. The gate drive circuitry is thus less complex and costly than the base drive circuitry of an equivalent bipolar device. The disadvantage of the MOSFET is that in the 'on' state the effective resistance of the drain-source is higher than for an equivalent bipolar device, so the power dissipation is higher and the device is rather less efficient as a power switch. MOSFETs are used in low and medium power inverters up to a few kW, with voltages generally not exceeding 700 V.

6.3 Insulated gate bipolar transistor (IGBT)

The IGBT (Figure 2.21) is a hybrid device which combines the best features of the MOSFET (i.e. ease of gate turn-on and turn-off from low-power logic circuits) and the BJT (relatively low power dissipation in the main collector-emitter circuit). These obvious advantages give the IGBT the edge over the MOSFET and BJT, and account for its dominance in all but small drives. It is particularly well suited to the medium power, medium voltage range (up to several hundred kW). The path for the main (load) current is from collector to emitter, as in the npn bipolar device.

7. CONVERTER WAVEFORMS, ACOUSTIC NOISE, AND COOLING

In common with most textbooks, the waveforms shown in this chapter (and later in the book) are what we would hope to see under ideal conditions. It makes sense to concentrate on these ideal waveforms from the point of view of gaining a basic understanding, but we ought to be warned that what we see on an oscilloscope may well look rather different, and is often not easy to interpret.

We have seen that the essence of power electronics is the switching process, so it should not come as much of a surprise to learn that in practice the switching is seldom achieved in such a clear-cut fashion as we have assumed. Usually, there will be some sort of high-frequency oscillation or 'ringing' evident, particularly on the voltage waveforms following each transition due to switching. This is due to the effects of stray capacitance and inductance: it will have been anticipated at

the design stage, and steps will have been taken to minimize it by fitting 'snubbing' circuits at the appropriate places in the converter. However, complete suppression of all these transient phenomena is seldom economically worthwhile so the user should not be too alarmed to see remnants of the transient phenomena in the output waveforms.

Acoustic noise is also a matter that can worry newcomers. Most power electronic converters emit whining or humming sounds at frequencies corresponding to the fundamental and harmonics of the switching frequency, though when the converter is used to feed a motor, the sound from the motor is usually a good deal louder than the sound from the converter itself. These sounds are very difficult to describe in words, but typically range from a high-pitched hum through a whine to a piercing whistle. In a.c. drives it is often possible to vary the switching frequency, which may allow specific mechanical resonances to be avoided. If taken above the audible range (which is inversely proportional to age) the sound clearly disappears, but at the expense of increased switching losses, so a compromise has to be sought.

7.1 Cooling of switching devices – thermal resistance

We have seen that by adopting a switching strategy the power loss in the switching devices is small in comparison with the power throughput, so the converter has a high efficiency. Nevertheless almost all the heat which is produced in the switching devices is released in the active region of the semiconductor, which is itself very small and will overheat and break down unless it is adequately cooled. It is therefore essential to ensure that even under the most onerous operating conditions, the temperature of the active junction inside the device does not exceed the safe value.

Consider what happens to the temperature of the junction region of the device when we start from cold (i.e. ambient) temperature and operate the device so that its average power dissipation remains constant. At first, the junction temperature begins to rise, so some of the heat generated is conducted to the metal case, which stores some heat as its temperature rises. Heat then flows into the heatsink (if fitted), which begins to warm up, and heat begins to flow to the surrounding air, at ambient temperature. The temperatures of the junction, case and heatsink continue to rise until eventually an equilibrium is reached when the total rate of loss of heat to ambient temperature is equal to the power dissipation inside the device.

The final steady-state junction temperature thus depends on how difficult it is for the power loss to escape down the temperature gradient to ambient, or in other words on the total 'thermal resistance' between the junction inside the device and the surrounding medium (usually air). Thermal resistance is usually expressed in $^{\circ}\text{C}/\text{watt}$, which directly indicates how much temperature rise will occur in the steady

state for every watt of dissipated power. It follows that for a given power dissipation, the higher the thermal resistance, the higher the temperature rise, so in order to minimize the temperature rise of the device, the total thermal resistance between it and the surrounding air must be made as small as possible.

The device designer aims to minimize the thermal resistance between the semiconductor junction and the case of the device, and provides a large and flat metal mounting surface to minimize the thermal resistance between the case and the heatsink. The converter designer must ensure good thermal contact between the device and the heatsink, usually by a bolted joint smeared with heat-conducting compound to fill any microscopic voids, and must design the heatsink to minimize the thermal resistance to air (or in some cases oil or water). Heatsink design offers the only real scope for appreciably reducing the total resistance, and involves careful selection of the material, size, shape and orientation of the heatsink, and the associated air-moving system (see below).

One drawback of the good thermal path between the junction and case of the device is that the metal mounting surface (or surfaces in the case of the popular hockey-puck package) can be electrically 'live'. This poses a difficulty for the converter designer, because mounting the device directly on the heatsink causes the latter to be dangerous. In addition, several separate isolated heatsinks may be required in order to avoid short-circuits. The alternative is for the devices to be electrically isolated from the heatsink using thin mica spacers, but then the thermal resistance is increased appreciably.

Increasingly devices come in packaged 'modules' with an electrically isolated metal to get round the 'live' problem. The packages contain combinations of transistors, diodes or thyristors, from which various converter circuits can be built up. Several modules can be mounted on a single heatsink, which does not have to be isolated from the enclosure or cabinet. They are available in ratings suitable for converters up to hundreds of kW, and the range is expanding.

7.2 Arrangement of heatsinks and forced-air cooling

The principal factors which govern the thermal resistance of a heatsink are the total surface area, the condition of the surface and the air flow. Many converters use extruded aluminium heatsinks, with multiple fins to increase the effective cooling surface area and lower the resistance, and with a machined face or faces for mounting the devices. Heatsinks are usually mounted vertically to improve natural air convection. Surface finish is important, with black anodized aluminium being typically 30% better than bright. The cooling performance of heatsinks is, however, a complex technical area, with turbulence being very beneficial in forced air-cooled heatsinks.

A typical layout for a medium-power (say 200 kW) converter is shown in [Figure 2.22](#).

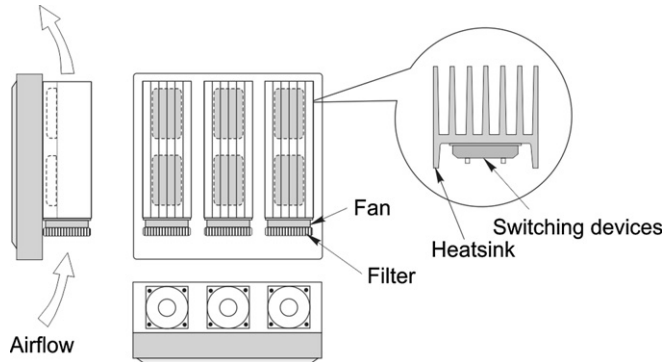


Figure 2.22 Layout of converter showing heatsink and cooling fans.

The fan(s) are positioned at either the top or bottom of the heatsink, and draw external air upwards, assisting natural convection. Even a modest airflow is very beneficial: with an air velocity of only 2 m/s, for example, the thermal resistance is halved as compared with the naturally cooled set-up, which means that for a given temperature rise the heatsink can be half the size of the naturally cooled one. However, large increases in the air velocity bring diminishing returns and also introduce additional noise, which is generally undesirable.