
4

Switch Realization

We have seen in previous chapters that the switching elements of the buck, boost, and several other dc-dc converters can be implemented using a transistor and diode. One might wonder why this is so, and how to realize semiconductor switches in general. These are worthwhile questions to ask, and switch implementation can depend on the power processing function being performed. The switches of inverters and cycloconverters require more complicated implementations than those of dc-dc converters. Also, the way in which a semiconductor switch is implemented can alter the behavior of a converter in ways not predicted by the ideal-switch analysis of the previous chapters—an example is the discontinuous conduction mode treated in the next chapter. The realization of switches using transistors and diodes is the subject of this chapter.

Semiconductor power devices behave as single-pole single-throw (SPST) switches, represented ideally in Fig. 4.1. So, although we often draw converter schematics using ideal single-pole double-throw (SPDT) switches as in Fig. 4.2(a), the schematic of Fig. 4.2(b) containing SPST switches is more realistic. The realization of a SPDT switch using two SPST switches is not as trivial as it might at first seem, because Fig. 4.2(a) and 4.2(b) are not exactly equivalent. It is possible for both SPST switches to be simultaneously in the on state or in the off state, leading to behavior not predicted by the SPDT switch of Fig. 4.2(a). In addition, it is possible for the switch state to depend on the applied voltage or current waveforms—a familiar example is the diode. Indeed, it is common for these phenomena to occur in converters operating at light load, or occasionally at heavy load, leading to the discontinuous conduction mode previously mentioned. The converter properties are then significantly modified.

How an ideal switch can be realized using semiconductor devices depends on the polarity of the voltage that the devices must block in the off state, and on the polarity of the current that the devices

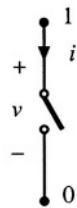


Fig. 4.1 SPST switch, with defined voltage and current polarities.

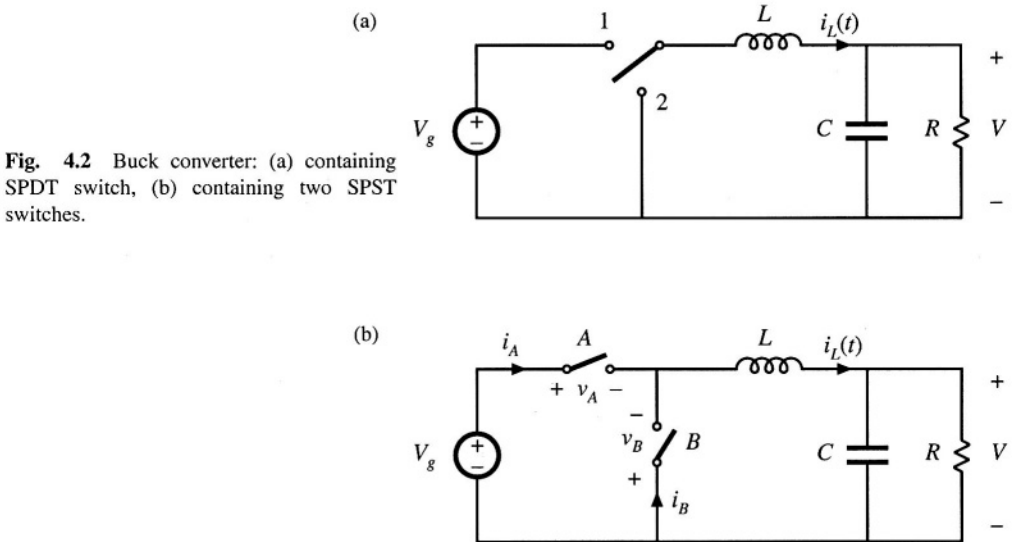


Fig. 4.2 Buck converter: (a) containing SPDT switch, (b) containing two SPST switches.

must conduct in the on state. For example, in the dc–dc buck converter of Fig. 4.2(b), switch A must block positive voltage V_g when in the off state, and must conduct positive current i_L when in the on state. If, for all intended converter operating points, the current and blocking voltage lie in a single quadrant of the plane as illustrated in Fig. 4.3, then the switch can be implemented in a simple manner using a transistor or a diode. Use of single-quadrant switches is common in dc–dc converters. Their operation is discussed briefly here.

In inverter circuits, two-quadrant switches are required. The output current is ac, and hence is sometimes positive and sometimes negative. If this current flows through the switch, then its current is ac, and the semiconductor switch realization is more complicated. A two-quadrant SPST switch can be realized using a transistor and diode. The dual case also sometimes occurs, in which the switch current is always positive, but the blocking voltage is ac. This type of two-quadrant switch can be constructed using a different arrangement of a transistor and diode. Cycloconverters generally require four-quadrant switches, which are capable of blocking ac voltages and conducting ac currents. Realizations of these elements are also discussed in this chapter.

Next, the synchronous rectifier is examined. The reverse-conducting capability of the metal oxide semiconductor field-effect transistor (MOSFET) allows it to be used where a diode would normally be required. If the MOSFET on-resistance is sufficiently small, then its conduction loss is less than that obtained using a diode. Synchronous rectifiers are sometimes used in low-voltage high-current applications to obtain improved efficiency. Several basic references treating single-, two-, and four-quadrant

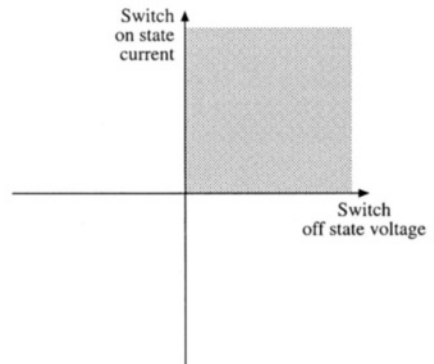


Fig. 4.3 A single-quadrant switch is capable of conducting currents of a single polarity, and of blocking voltages of a single polarity.

switches are listed at the end of this chapter [1–8].

Several power semiconductor devices are briefly discussed in Section 4.2. Majority-carrier devices, including the MOSFET and Schottky diode, exhibit very fast switching times, and hence are preferred when the off state voltage levels are not too high. Minority-carrier devices, including the bipolar junction transistor (BJT), insulated-gate bipolar transistor (IGBT), and thyristors [gate turn-off (GTO) and MOS-controlled thyristor (MCT)] exhibit high breakdown voltages with low forward voltage drops, at the expense of reduced switching speed.

Having realized the switches using semiconductor devices, switching loss can next be discussed. There are a number of mechanisms that cause energy to be lost during the switching transitions [11]. When a transistor drives a clamped inductive load, it experiences high instantaneous power loss during the switching transitions. Diode stored charge further increases this loss, during the transistor turn-on transition. Energy stored in certain parasitic capacitances and inductances is lost during switching. Parasitic ringing, which decays before the end of the switching period, also indicates the presence of switching loss. Switching loss increases directly with switching frequency, and imposes a maximum limit on the operating frequencies of practical converters.

4.1 SWITCH APPLICATIONS

4.1.1 Single-Quadrant Switches

The ideal SPST switch is illustrated in Fig. 4.1. The switch contains power terminals 1 and 0, with current and voltage polarities defined as shown. In the on state, the voltage v is zero, while the current i is zero in the off state. There is sometimes a third terminal C , where a control signal is applied. Distinguishing features of the SPST switch include the control method (active vs. passive) and the region of the i – v plane in which they can operate.

A passive switch does not contain a control terminal C . The state of the switch is determined by the waveforms $i(t)$ and $v(t)$ applied to terminals 0 and 1. The most common example is the diode, illustrated in Fig. 4.4. The ideal diode requires that $v(t) \leq 0$ and $i(t) \geq 0$. The diode is off ($i = 0$) when $v < 0$, and is on ($v = 0$) when $i > 0$. It can block negative voltage but not positive voltage. A passive SPST switch can be realized using a diode provided that the intended operating points [i.e., the values of $v(t)$ and $i(t)$ when the switch is in the on and off states] lie on the diode characteristic of Fig. 4.4(b).

The conducting state of an active switch is determined by the signal applied to the control terminal C . The state does not directly depend on the waveforms $v(t)$ and $i(t)$ applied to terminals 0 and 1. The BJT, MOSFET, IGBT, GTO, and MCT are examples of active switches. Idealized characteristics $i(t)$ vs. $v(t)$ for the BJT and IGBT are sketched in Fig. 4.5. When the control terminal causes the transistor to be in the off state, $i = 0$ and the device is capable of blocking positive voltage: $v \geq 0$. When the control terminal causes the transistor to be in the on state, $v = 0$ and the device is capable of conducting positive current: $i \geq 0$. The reverse-conducting and

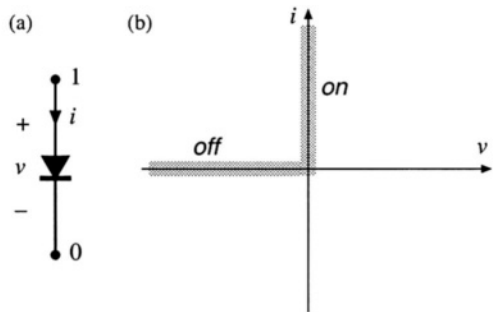


Fig. 4.4 Diode symbol (a), and its ideal characteristic (b).

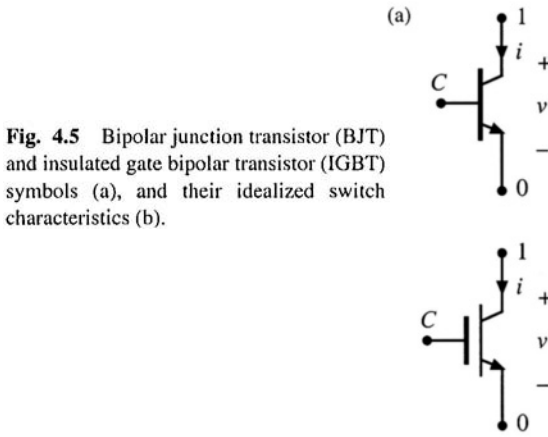
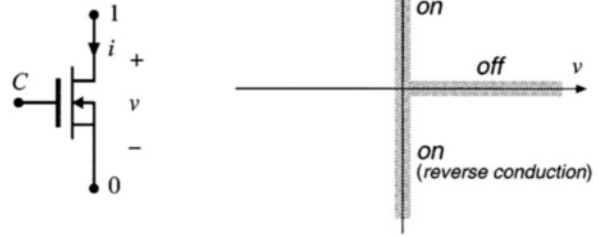


Fig. 4.5 Bipolar junction transistor (BJT) and insulated gate bipolar transistor (IGBT) symbols (a), and their idealized switch characteristics (b).

Fig. 4.6 Power MOSFET symbol (a), and its idealized switch characteristics (b).



reverse-blocking characteristics of the BJT and IGBT are poor or nonexistent, and have essentially no application in the power converter area. The power MOSFET (Fig. 4.6) has similar characteristics, except that it is able to conduct current in the reverse direction. With one notable exception (the synchronous rectifier discussed later), the MOSFET is normally operated with $i \geq 0$, in the same manner as the BJT and IGBT. So an active SPST switch can be realized using a BJT, IGBT, or MOSFET, provided that the intended operating points lie on the transistor characteristic of Fig. 4.5(b).

To determine how to implement an SPST switch using a transistor or diode, one compares the switch operating points with the i - v characteristics of Figs. 4.4(b), 4.5(b), and 4.6(b). For example, when it is intended that the SPOT switch of Fig. 4.2(a) be in position 1, SPST switch A of Fig. 4.2(b) is closed, and SPST switch B is opened. Switch A then conducts the positive inductor current, $i_A = i_L$, and switch B must block negative voltage, $v_B = -V_R$. These switch operating points are illustrated in Fig. 4.7. Likewise, when it is intended that the SPDT switch of Fig. 4.2(a) be in position 2, then SPST switch A is opened and switch B is closed. Switch B then conducts the positive inductor current, $i_B = i_L$, while switch A blocks positive voltage, $v_A = V_R$.

By comparison of the switch A operating points of Fig. 4.7(a) with Figs. 4.5(b) and 4.6(b), it can be seen that a transistor (BJT, IGBT, or MOSFET) could be used, since switch A must block positive voltage and conduct positive current. Likewise, comparison of Fig. 4.7(b) with Fig. 4.4(b) reveals that switch B can be implemented using a diode, since switch B must block negative voltage and conduct positive current. Hence a valid switch realization is given in Fig. 4.8.

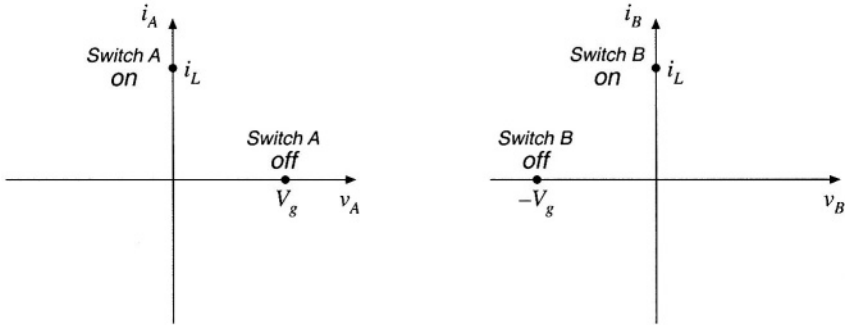


Fig. 4.7 Operating points of switch A, (a), and switch B, (b), in the buck converter of Fig. 4.2(b).

Figure 4.8 is an example of a single-quadrant switch realization: the devices are capable of conducting current of only one polarity, and blocking voltage of only one polarity. When the controller turns the transistor on, the diode becomes reverse-biased since $v_B = -V_g$. It is required that V_g be positive; otherwise, the diode will be forward-biased. The transistor conducts current i_L . This current should also be positive, so that the transistor conducts in the forward direction.

When the controller turns the transistor off, the diode must turn on so that the inductor current can continue to flow. Turning the transistor off causes the inductor current $i_L(t)$ to decrease. Since $v_L(t) = L di_L(t)/dt$, the inductor voltage becomes sufficiently negative to forward-bias the diode, and the diode turns on. Diodes that operate in this manner are sometimes called *freewheeling diodes*. It is required that i_L be positive; otherwise, the diode cannot be forward-biased since $i_B = i_L$. The transistor blocks voltage V_g ; this voltage should be positive to avoid operating the transistor in the reverse blocking mode.

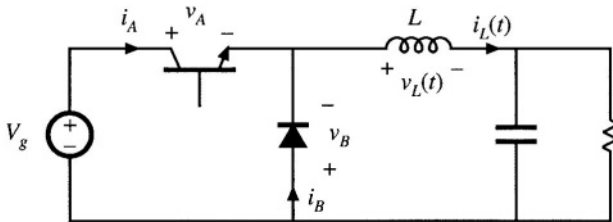


Fig. 4.8 Implementation of the SPST switches of Fig. 4.2(b) using a transistor and diode.

4.1.2 Current-Bidirectional Two-Quadrant Switches

In any number of applications such as dc-ac inverters and servo amplifiers, it is required that the switching elements conduct currents of both polarities, but block only positive voltages. A current-bidirectional two-quadrant SPST switch of this type can be realized using a transistor and diode, connected in an anti-parallel manner as in Fig. 4.9.

The MOSFET of Fig. 4.6 is also a two-quadrant switch. However, it should be noted here that

Fig. 4.9 A current-bidirectional two-quadrant SPST switch: (a) implementation using a transistor and antiparallel diode, (b) idealized switch characteristics.

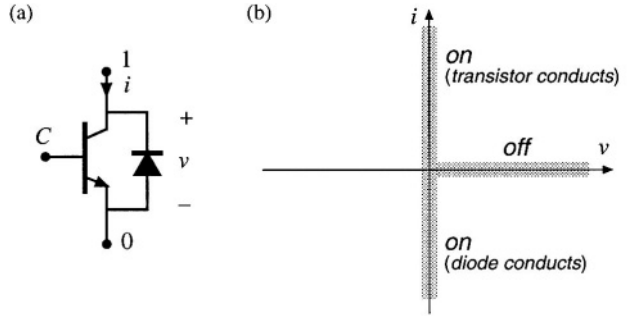
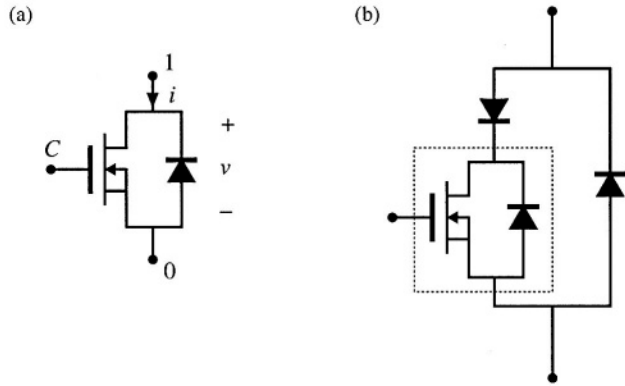


Fig. 4.10 The power MOSFET inherently contains a built-in body diode: (a) equivalent circuit, (b) addition of external diodes to prevent conduction of body diode.



practical power MOSFETs inherently contain a built-in diode, often called the *body diode*, as illustrated in Fig. 4.10. The switching speed of the body diode is much slower than that of the MOSFET. If the body diode is allowed to conduct, then high peak currents can occur during the diode turn-off transition. Most MOSFETs are not rated to handle these currents, and device failure can occur. To avoid this situation, external series and antiparallel diodes can be added as in Fig. 4.10(b). Power MOSFETs can be specifically designed to have a fast-recovery body diode, and to operate reliably when the body diode is allowed to conduct the rated MOSFET current. However, the switching speed of such body diodes is still somewhat slow, and significant switching loss due to diode stored charge (discussed later in this chapter) can occur.

A SPDT current-bidirectional two-quadrant switch can again be derived using two SPST switches as in Fig. 4.2(b). An example is given in Fig. 4.11. This converter operates from positive and negative dc supplies, and can produce an ac output voltage $v(t)$ having either polarity. Transistor Q_2 is driven with the complement of the Q_1 drive signal, so that Q_1 conducts during the first subinterval $0 < t < DT_s$, and Q_2 conducts during the second subinterval $DT_s < t < T_s$.

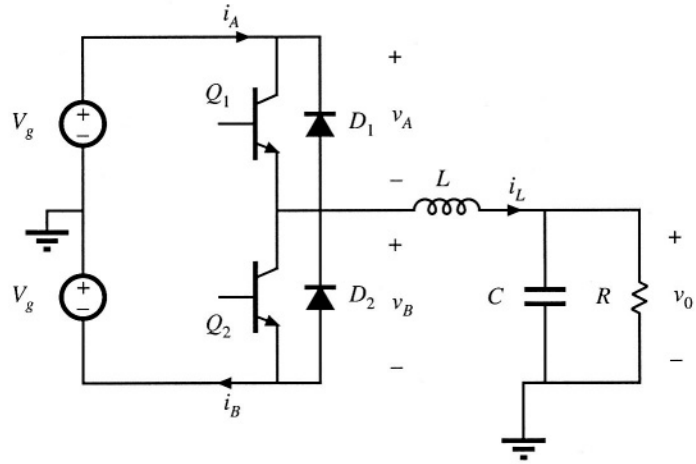
It can be seen from Fig. 4.11 that the switches must block voltage $2V_g$. It is required that V_g be positive; otherwise, diodes D_1 and D_2 will conduct simultaneously, shorting out the source.

It can be shown via inductor volt-second balance that

$$v_0 = (2D - 1)V_g \quad (4.1)$$

This equation is plotted in Fig. 4.12. The converter output voltage v_0 is positive for $D > 0.5$, and negative for $D < 0.5$. By sinusoidal variation of the duty cycle,

Fig. 4.11 Inverter circuit using two-quadrant switches.



$$D(t) = 0.5 + D_m \sin(\omega t) \quad (4.2)$$

with D_m being a constant less than 0.5, the output voltage becomes sinusoidal. Hence this converter could be used as a dc-ac inverter.

The load current is given by v_0/R ; in equilibrium, this current coincides with the inductor current i_L .

$$i_L = \frac{v_0}{R} = (2D - 1) \frac{V_g}{R} \quad (4.3)$$

The switches must conduct this current. So the switch current is also positive when $D > 0.5$, and negative when $D < 0.5$. With high-frequency duty cycle variations, the L - C filter may introduce a phase lag into the inductor current waveform, but it is nonetheless true that switch currents of both polarities occur. So the switch must operate in two quadrants of the plane, as illustrated in Fig. 4.13. When i_L is positive, Q_1 and D_2 alternately conduct. When i_L is negative, Q_2 and D_1 alternately conduct.

A well-known dc-3 ϕ ac inverter circuit, the *voltage-source inverter* (VSI), operates in a similar manner. As illustrated in Fig. 4.14, the VSI contains three two-quadrant SPDT switches, one per phase. These switches block the dc input voltage V_g , and must conduct the output ac phase currents i_a , i_b , and i_c .

Fig. 4.12 Output voltage vs. duty cycle, for the inverter of Fig. 4.11. This converter can produce both positive and negative output voltages.

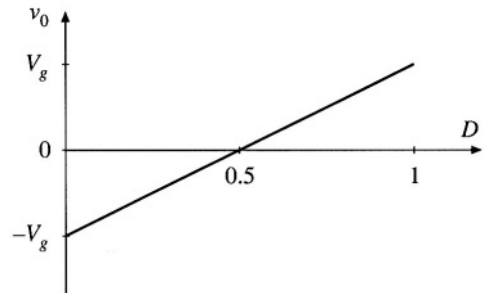


Fig. 4.13 The switches in the inverter of Fig. 4.11 must be capable of conducting both positive and negative current, but need block only positive voltage.

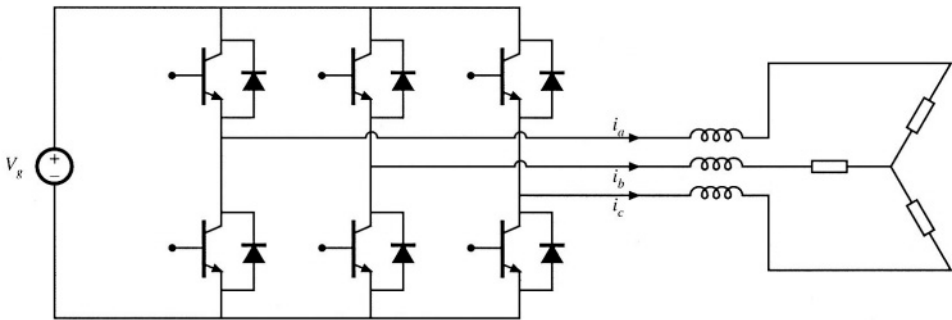
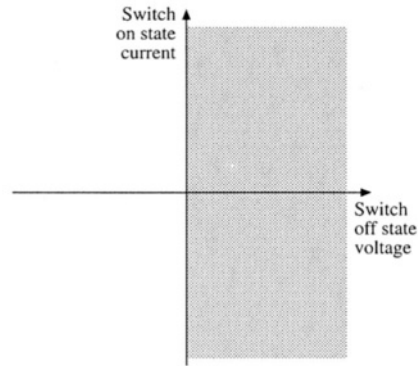


Fig. 4.14 The dc-3 ϕ ac voltage-source inverter requires two-quadrant switches.

respectively.

Another current-bidirectional two-quadrant switch example is the bidirectional battery charger/discharger illustrated in Fig. 4.15. This converter can be used, for example, to interface a battery to the main power bus of a spacecraft. Both the dc bus voltage v_{bus} and the battery voltage v_{batt} are always positive. The semiconductor switch elements block positive voltage v_{bus} . When the battery is being charged, i_L is positive, and Q_1 and D_2 alternately conduct current. When the battery is being discharged, i_L is negative, and Q_2 and D_1 alternately conduct. Although this is a dc-dc converter, it requires two-quadrant switches because the power can flow in either direction.

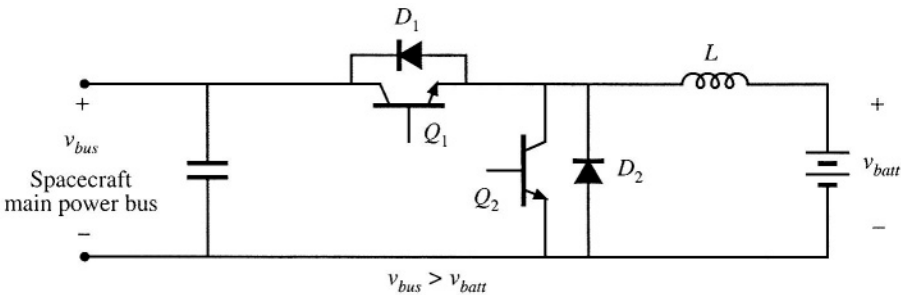


Fig. 4.15 Bidirectional battery charger/discharger, based on the dc-dc buck converter.

Fig. 4.16 Voltage-bidirectional two-quadrant switch properties.

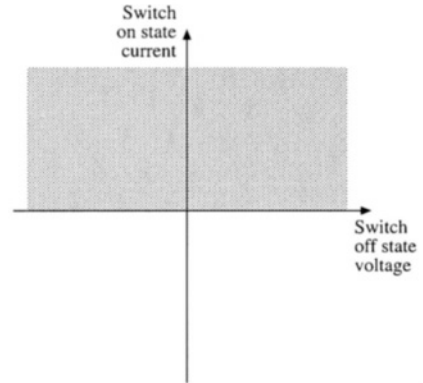
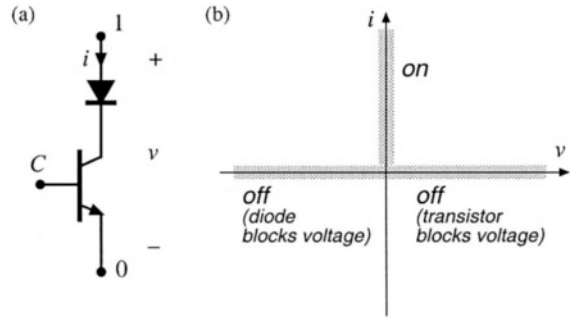


Fig. 4.17 A voltage-bidirectional two-quadrant SPST switch: (a) implementation using a transistor and series diode, (b) idealized switch characteristics.



4.1.3 Voltage-Bidirectional Two-Quadrant Switches

Another type of two-quadrant switch, having the voltage-bidirectional properties illustrated in Fig. 4.16, is sometimes required. In applications where the switches must block both positive and negative voltages, but conduct only positive current, an SPST switch can be constructed using a series-connected transistor and diode as in Fig. 4.17. When it is intended that the switch be in the off state, the controller turns the transistor off. The diode then blocks negative voltage, and the transistor blocks positive voltage. The series connection can block negative voltages up to the diode voltage rating, and positive voltages up to the transistor voltage rating. The silicon-controlled rectifier is another example of a voltage-bidirectional two-quadrant switch.

A converter that requires this type of two-quadrant switch is the dc-3 ϕ ac buck-boost inverter shown in Fig. 4.18 [4]. If the converter functions in inverter mode, so that the inductor current $i_L(t)$ is always positive, then all switches conduct only positive current. But the switches must block the output ac line-to-line voltages, which are sometimes positive and sometimes negative. Hence voltage-bidirectional two-quadrant switches are required.

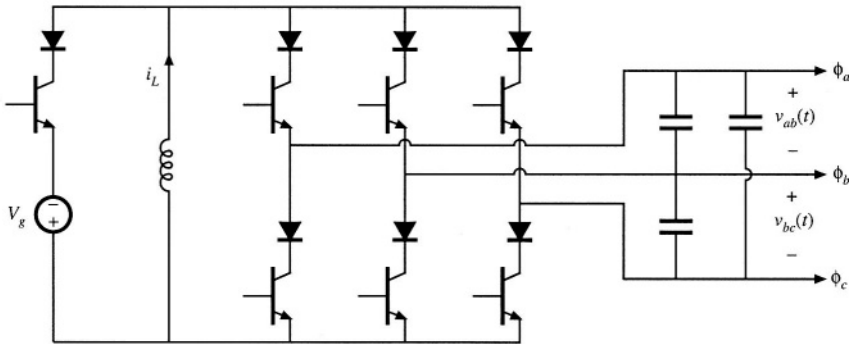


Fig. 4.18 DC-3φac buck-boost inverter.

4.1.4 Four-Quadrant Switches

The most general type of switch is the four-quadrant switch, capable of conducting currents of either polarity and blocking voltages of either polarity, as in Fig. 4.19. There are several ways of constructing a four-quadrant switch. As illustrated in Fig. 4.20(b), two current-bidirectional two-quadrant switches described in Section 4.1.2 can be connected back-to-back. The transistors are driven on and off simultaneously. Another approach is the antiparallel connection of two voltage-bidirectional two-quadrant switches described in Section 4.1.3, as in Fig. 4.20(a). A third approach, using only one transistor but additional diodes, is given in Fig. 4.20(c).

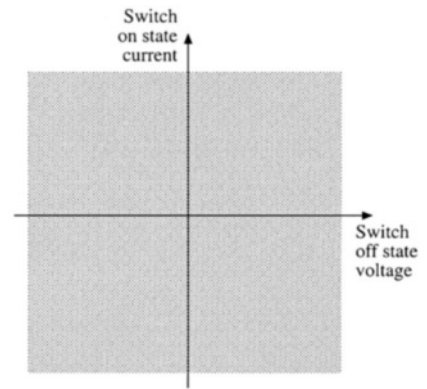


Fig. 4.19 A four-quadrant switch can conduct either polarity of current, and can block either polarity of voltage.

Cycloconverters are a class of converters requiring four-quadrant switches. For example, a 3φac-to-3φac matrix converter is illustrated in Fig. 4.21. Each of the nine SPST switches is realized using one of the semiconductor networks of Fig. 4.20. With proper control of the switches, this converter can produce a three-phase output of variable fre-

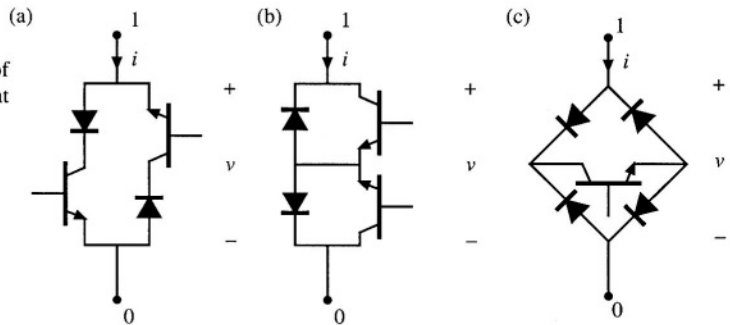


Fig. 4.20 Three ways of implementing a four-quadrant SPST switch.

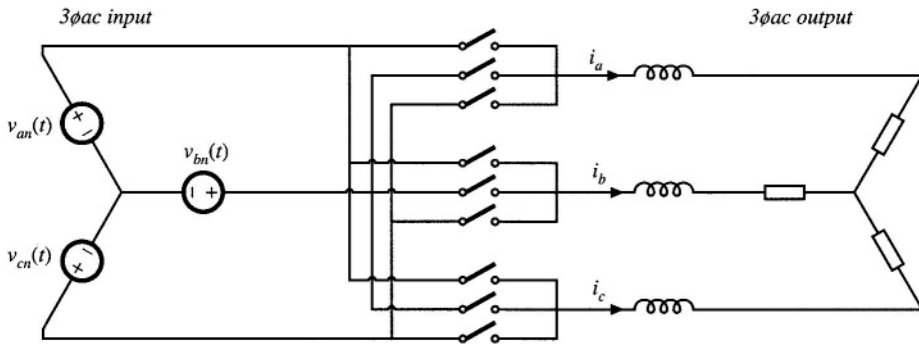


Fig. 4.21 A 3 ϕ ac–3 ϕ ac matrix converter, which requires nine SPST four-quadrant switches.

quency and voltage, from a given three-phase ac input. Note that there are no dc signals in this converter: all of the input and output voltages and currents are ac, and hence four-quadrant switches are necessary.

4.1.5 Synchronous Rectifiers

The ability of the MOSFET channel to conduct current in the reverse direction makes it possible to employ a MOSFET where a diode would otherwise be required. When the MOSFET is connected as in Fig. 4.22(a) [note that the source and drain connections are reversed from the connections of Fig. 4.6(a)], the characteristics of Fig. 4.22(b) are obtained. The device can now block negative voltage and conduct positive current, with properties similar to those of the diode in Fig. 4.4. The MOSFET must be controlled such that it operates in the on state when the diode would normally conduct, and in the off state when the diode would be reverse-biased.

Thus, we could replace the diode in the buck converter of Fig. 4.8 with a MOSFET, as in Fig. 4.23. The BJT has also been replaced with a MOSFET in the figure. MOSFET Q_2 is driven with the complement of the Q_1 control signal.

The trend in computer power supplies is reduction of output voltage levels, from 5 V to 3.3 V and lower. As the output voltage is reduced, the diode conduction loss increases; in consequence, the diode conduction loss is easily the largest source of power loss in a 3.3 V power supply. Unfortunately, the diode junction contact potential limits what can be done to reduce the forward voltage drop of diodes. Schottky diodes having reduced junction potential can be employed; nonetheless, low-voltage power

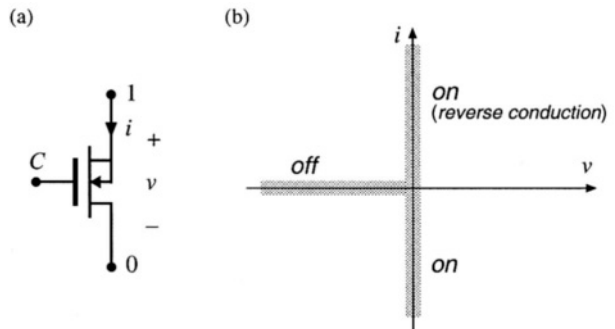


Fig. 4.22 Power MOSFET connected as a synchronous rectifier, (a), and its idealized switch characteristics, (b).

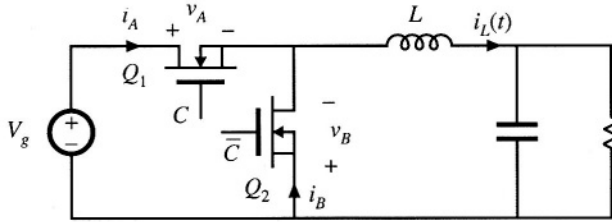


Fig. 4.23 Buck converter, implemented using a synchronous rectifier.

supplies containing diodes that conduct the output current must have low efficiency.

A solution is to replace the diodes with MOSFETs operated as synchronous rectifiers. The conduction loss of a MOSFET having on-resistance R_{on} and operated with rms current is I_{rms} , is $I_{rms}^2 R_{on}$. The on-resistance can be decreased by use of a larger MOSFET. So the conduction loss can be reduced as low as desired, if one is willing to pay for a sufficiently large device. Synchronous rectifiers find widespread use in low-voltage power supplies.

4.2 A BRIEF SURVEY OF POWER SEMICONDUCTOR DEVICES

The most fundamental challenge in power semiconductor design is obtaining a high breakdown voltage, while maintaining low forward voltage drop and on-resistance. A closely related issue is the longer switching times of high-voltage low-on-resistance devices. The tradeoff between breakdown voltage, on-resistance, and switching times is a key distinguishing feature of the various power devices.

The breakdown voltage of a reverse-biased p - n junction and its associated depletion region is a function of doping level: obtaining a high breakdown voltage requires low doping concentration, and hence high resistivity, in the material on at least one side of the junction. This high-resistivity region is usually the dominant contributor to the on-resistance of the device, and hence high-voltage devices must have higher on-resistance than low-voltage devices. In *majority carrier* devices, including the MOSFET and Schottky diode, this accounts for the first-order dependence of on-resistance on rated voltage. However, *minority carrier* devices, including the diffused-junction p - n diode, the bipolar junction transistor (BJT), the insulated-gate bipolar transistor (IGBT), and the thyristor family (SCR, GTO, MCT), exhibit another phenomenon known as *conductivity modulation*. When a minority-carrier device operates in the on state, minority carriers are injected into the lightly doped high-resistivity region by the forward-biased p - n junction. The resulting high concentration of minority carriers effectively reduces the apparent resistivity of the region, reducing the on-resistance of the device. Hence, minority-carrier devices exhibit lower on-resistances than comparable majority-carrier devices.

However, the advantage of decreased on-resistance in minority-carrier devices comes with the disadvantage of decreased switching speed. The conducting state of any semiconductor device is controlled by the presence or absence of key charge quantities within the device, and the turn-on and turn-off switching times are equal to the times required to insert or remove this controlling charge. Devices operating with conductivity modulation are controlled by their injected minority carriers. The total amount of controlling minority charge in minority-carrier devices is much greater than the charge required to control an equivalent majority-carrier device. Although the mechanisms for inserting and removing the controlling charge of the various devices can differ, it is nonetheless true that, because of their large amounts of minority charge, minority-carrier devices exhibit switching times that are significantly longer than those of majority-carrier devices. In consequence, majority-carrier devices find application at lower volt-

age levels and higher switching frequencies, while the reverse is true of minority-carrier devices.

Modern power devices are fabricated using up-to-date processing techniques. The resulting small feature size allows construction of highly interdigitated devices, whose unwanted parasitic elements are less significant. The resulting devices are more rugged and well-behaved than their predecessors.

A detailed description of power semiconductor device physics and switching mechanisms is beyond the scope of this book. Selected references on power semiconductor devices are listed in the reference section [9-19].

4.2.1 Power Diodes

As discussed above, the diffused-junction p - n diode contains a lightly doped or intrinsic high-resistivity region, which allows a high breakdown voltage to be obtained. As illustrated in Fig. 4.24(a), this region comprises one side of the p - n junction (denoted n^-); under reverse-biased conditions, essentially all of the applied voltage appears across the depletion region inside the n^- region. On-state conditions are illustrated in Fig. 4.24(b). Holes are injected across the forward-biased junction, and become minority carriers in the n^- region. These minority carriers effectively reduce the apparent resistivity of the n^- region via conductivity modulation. Essentially all of the forward current $i(t)$ is comprised of holes that diffuse across the p - n region, and then recombine with electrons from the n region.

Typical switching waveforms are illustrated in Fig. 4.25. The familiar exponential i - v character-

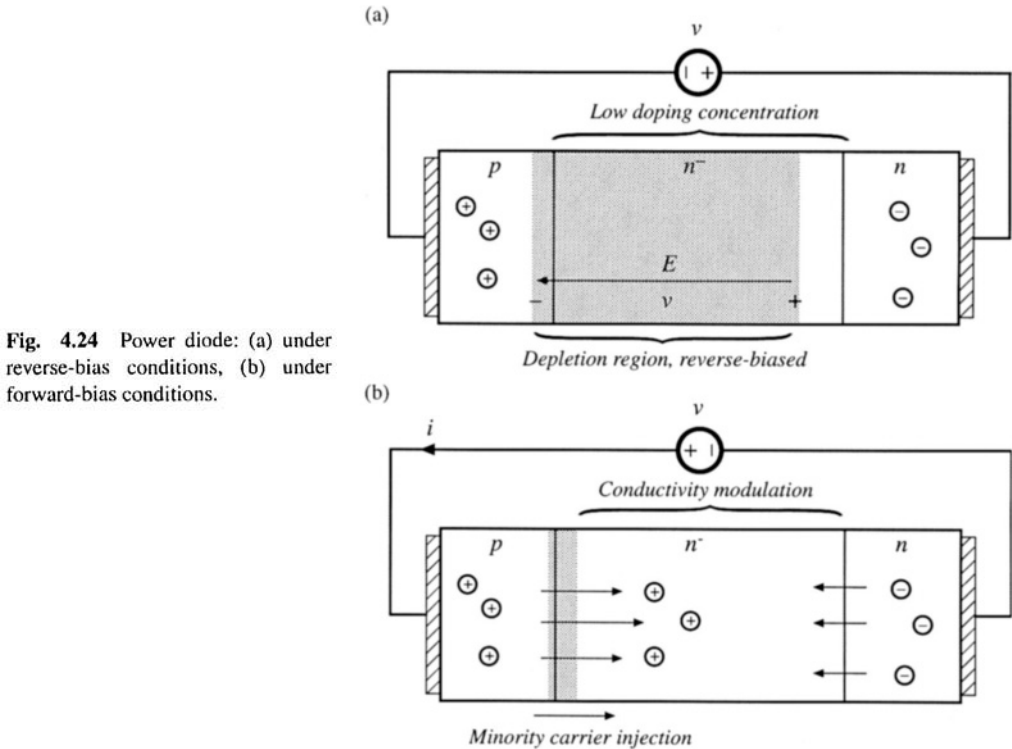
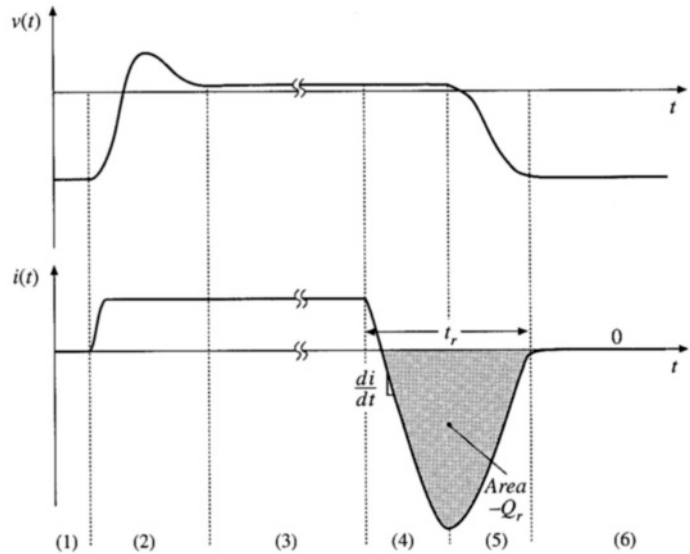


Fig. 4.24 Power diode: (a) under reverse-bias conditions, (b) under forward-bias conditions.

Fig. 4.25 Diode voltage and current waveforms. Interval (1): off state. Interval (2): turn-on transition. Interval (3): on state. Intervals (4) and (5): turn-off transition. Interval (6): off state.



istic of the p - n diode is an equilibrium relation. During transients, significant deviations from the exponential characteristic are observed; these deviations are associated with changes in the stored minority charge. As illustrated in Fig. 4.25, the diode operates in the off state during interval (1), with zero current and negative voltage. At the beginning of interval (2), the current increases to some positive value. This current charges the effective capacitance of the reverse-biased diode, supplying charge to the depletion region and increasing the voltage $v(t)$. Eventually, the voltage becomes positive, and the diode junction becomes forward-biased. The voltage may rise to a peak value of several volts, or even several tens of volts, reflecting the somewhat large resistance of the lightly doped n^- region. The forward-biased p - n^- junction continues to inject minority charge into the n^- region. As the total minority charge in the n^- region increases, conductivity modulation of the n^- region causes its effective resistance to decrease, and hence the forward voltage drop $v(t)$ also decreases. Eventually, the diode reaches equilibrium, in which the minority carrier injection rate and recombination rate are equal. During interval (3), the diode operates in the on state, with forward voltage drop given by the diode static i - v characteristic.

The turn-off transient is initiated at the beginning of interval (4). The diode remains forward-biased while minority charge is present in the vicinity of the diode p - n^- junction. Reduction of the stored minority charge can be accomplished either by active means, via negative terminal current, or by passive means, via recombination. Normally, both mechanisms occur simultaneously. The charge Q_r contained in the negative portion of the diode turn-off current waveform is called the *recovered charge*. The portion of Q_r occurring during interval (4) is actively-removed minority charge. At the end of interval (4), the stored minority charge in the vicinity of the p - n^- junction has been removed, such that the diode junction becomes reverse-biased and is able to block negative voltage. The depletion region effective capacitance is then charged during interval (5) to the negative off-state voltage. The portion of Q_r occurring during interval (5) is charge supplied to the depletion region, as well as minority charge that is actively removed from remote areas of the diode. At the end of interval (5), the diode is able to block the entire applied reverse voltage. The length of intervals (4) and (5) is called the *reverse recovery time* t_r . During interval (6), the diode operates in the off state. The diode turn-off transition, and its influence on switching loss in a PWM converter, is discussed further in Section 4.3.2.

Table 4.1 Characteristics of several commercial power rectifier diodes

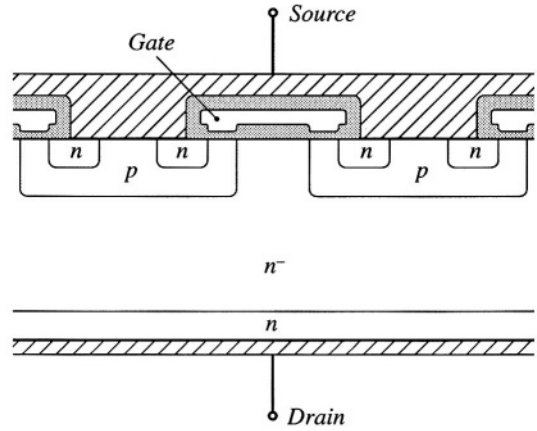
Part number	Rated maximum voltage	Rated average current	V_F (typical)	t_r (max)
Fast recovery rectifiers				
1N3913	400 V	30 A	1.1 V	400 ns
SD453N25S20PC	2500 V	400 A	2.2 V	3 μ s
Ultra-fast recovery rectifiers				
MUR815	150 V	8 A	0.975 V	35 ns
MUR1560	600 V	15 A	1.2 V	60 ns
RHRU100120	1200 V	100 A	2.6 V	60 ns
Schottky rectifiers				
MBR6030L	30 V	60 A	0.48 V	
444CNQ045	45 V	440 A	0.69 V	
30CPQ150	150 V	30 A	1.19 V	

Diodes are rated according to the length of their reverse recovery time t_{rr} . *Standard recovery* rectifiers are intended for 50 Hz or 60 Hz operation; reverse recovery times of these devices are usually not specified. *Fast recovery* rectifiers and *ultrafast recovery* rectifiers are intended for use in converter applications. The reverse recovery time t_{rr} , and sometimes also the recovered charge Q_{rr} , are specified by manufacturers of these devices. Ratings of several commercial devices are listed in Table 4.1.

Schottky diodes are essentially majority-carrier devices whose operation is based on the rectifying characteristic of a metal-semiconductor junction. These devices exhibit negligible minority stored charge, and their switching behavior can be adequately modeled simply by their depletion-region capacitance and equilibrium exponential i - v characteristic. Hence, an advantage of the Schottky diode is its fast switching speed. An even more important advantage of Schottky diodes is their low forward voltage drops, especially in devices rated 45 V or less. Schottky diodes are restricted to low breakdown voltages; very few commercial devices are rated to block 100 V or more. Their off-state reverse currents are considerably higher than those of p - n junction diodes. Characteristics of several commercial Schottky rectifiers are also listed in Table 4.1.

Another important characteristic of a power semiconductor device is whether its on-resistance and forward voltage drop exhibit a positive temperature coefficient. Such devices, including the MOSFET and IGBT, are advantageous because multiple chips can be easily paralleled, to obtain high-current modules. These devices also tend to be more rugged and less susceptible to hot-spot formation and second-breakdown problems. Diodes cannot be easily connected in parallel, because of their negative temperature coefficients: an imbalance in device characteristics may cause one diode to conduct more current than the others. This diode becomes hotter, which causes it to conduct even more of the total current. In consequence, the current does not divide evenly between the paralleled devices, and the current rating of one of the devices may be exceeded. Since BJTs and thyristors are controlled by a diode junction, these devices also exhibit negative temperature coefficients and have similar problems when operated in parallel. Of course, it is possible to parallel any type of semiconductor device; however, use of matched devices, a common thermal substrate, and/or external circuitry may be required to cause the on-state currents of the devices to be equal.

Fig. 4.26 Cross-section of DMOS n -channel power MOSFET structure. Crosshatched regions are metallized contacts. Shaded regions are insulating silicon dioxide layers.



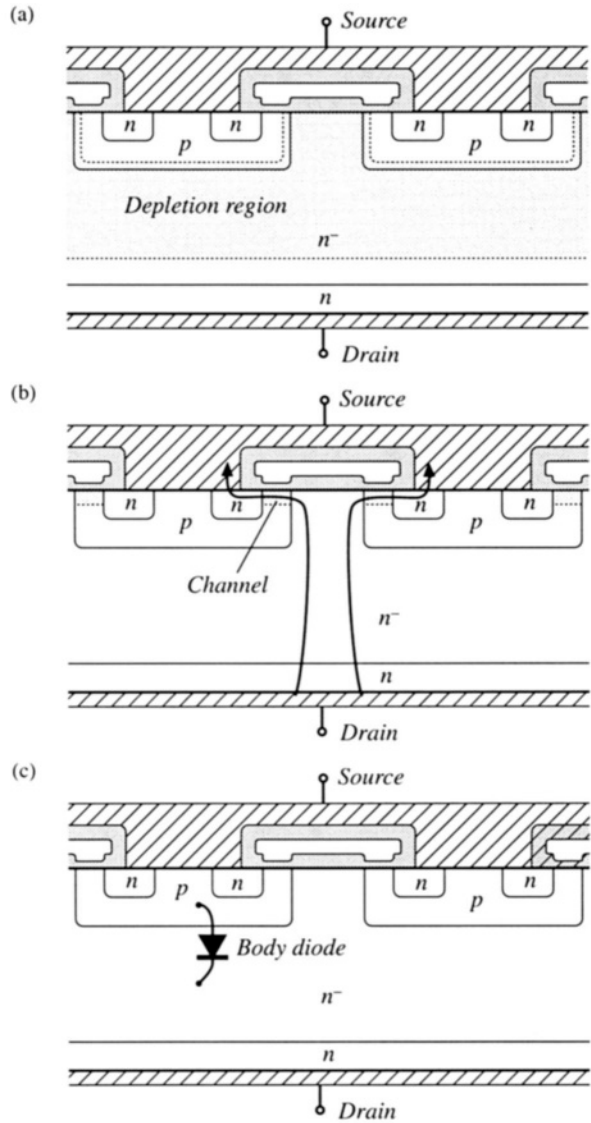
4.2.2 Metal Oxide Semiconductor Field-Effect Transistor (MOSFET)

The power MOSFET is a modern power semiconductor device having gate lengths close to one micron. The power device is comprised of many small parallel-connected enhancement-mode MOSFET cells, which cover the surface of the silicon die. A cross-section of one cell is illustrated in Fig. 4.26. Current flows vertically through the silicon wafer: the metallized drain connection is made on the bottom of the chip, while the metallized source connection and polysilicon gate are on the top surface. Under normal operating conditions, in which $v_{ds} \geq 0$, both the p - n and p - n^- junctions are reverse-biased. Figure 4.27(a) illustrates operation of the device in the off state. The applied drain-to-source voltage then appears across the depletion region of the p - n^- junction. The n^- region is lightly doped, such that the desired breakdown voltage rating is attained. Figure 4.27(b) illustrates operation in the on state, with a sufficiently large positive gate-to-source voltage. A channel then forms at the surface of the p region, underneath the gate. The drain current flows through the n^- region, channel, n region, and out through the source contact. The on-resistance of the device is the sum of the resistances of the n^- region, the channel, the source and drain contacts, etc. As the breakdown voltage is increased, the on-resistance becomes dominated by the resistance of the n^- region. Since there are no minority carriers to cause conductivity modulation, the on-resistance increases rapidly as the breakdown voltage is increased to several hundred volts and beyond.

The p - n^- junction is called the *body diode*; as illustrated in Fig. 4.27(c), this junction forms an effective diode in parallel with the MOSFET channel. The body diode can become forward-biased when the drain-to-source voltage $v_{ds}(t)$ is negative. This diode is capable of conducting the full rated current of the MOSFET. However, most MOSFETs are not optimized with respect to the speed of their body diodes, and the large peak currents that flow during the reverse recovery transition of the body diode can cause device failure. Several manufacturers produce MOSFETs that contain fast recovery body diodes; these devices are rated to withstand the peak currents during the body diode reverse recovery transition.

Typical n -channel MOSFET static switch characteristics are illustrated in Fig. 4.28. The drain current is plotted as a function of the gate-to-source voltage, for various values of drain-to-source voltage. When the gate-to-source voltage is less than the threshold voltage V_{th} , the device operates in the off state. A typical value of V_{th} is 3 V. When the gate-to-source voltage is greater than 6 or 7 V, the device operates in the on state; typically, the gate is driven to 12 or 15 V to ensure minimization of the forward voltage drop. In the on state, the drain-to-source voltage V_{DS} is roughly proportional to the drain current

Fig. 4.27 Operation of the power MOSFET: (a) in the off state, v_{ds} appears across the depletion region in the n^- region; (b) current flow through the conducting channel in the on state; (c) body diode due to the p - n^- junction.

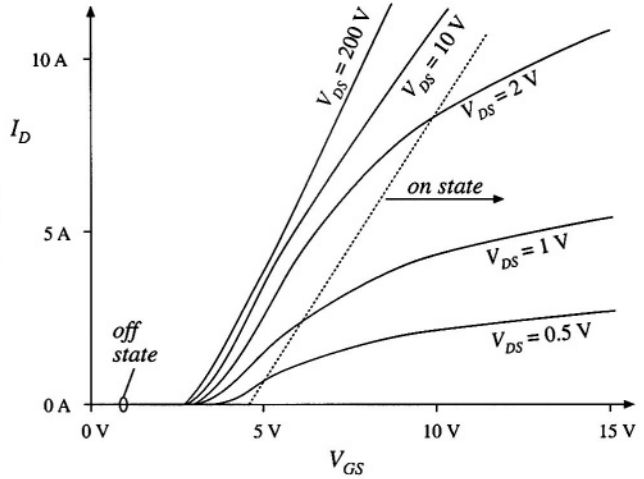


I_D . The MOSFET is able to conduct peak currents well in excess of its average current rating, and the nature of the static characteristics is unchanged at high current levels. Logic-level power MOSFETs are also available, which operate in the on state with a gate-to-source voltage of 5 V. A few p -channel devices can be obtained, but their properties are inferior to those of equivalent n -channel devices.

The on-resistance and forward voltage drop of the MOSFET have a positive temperature coefficient. This property makes it relatively easy to parallel devices. High current MOSFET modules are available, containing several parallel-connect chips.

The major capacitances of the MOSFET are illustrated in Fig. 4.29. This model is sufficient for

Fig. 4.28 Typical static characteristics of a power MOSFET. Drain current I_D is plotted vs. gate-to-source voltage V_{GS} , for various values of drain-to-source voltage V_{DS} .



qualitative understanding of the MOSFET switching behavior; more accurate models account for the parasitic junction field-effect transistor inherent in the DMOS geometry. Switching times of the MOSFET are determined essentially by the times required for the gate driver to charge these capacitances. Since the drain current is a function of the gate-to-source voltage, the rate at which the drain current changes is dependent on the rate at which the gate-to-source capacitance is charged by the gate drive circuit. Likewise, the rate at which the drain voltage changes is a function of the rate at which the gate-to-drain capacitance is charged. The drain-to-source capacitance leads directly to switching loss in PWM converters, since the energy stored in this capacitance is lost during the transistor turn-on transition. Switching loss is discussed in Section 4.3.

The gate-to-source capacitance is essentially linear. However, the drain-to-source and gate-to-drain capacitances are strongly nonlinear: these incremental capacitances vary as the inverse square root of the applied capacitor voltage. For example, the dependence of the incremental drain-to-source capacitance can be written in the form

$$C_{ds}(v_{ds}) = \frac{C_0}{\sqrt{1 + \frac{v_{ds}}{V_0}}} \quad (4.4)$$

where C_0 and V_0 are constants that depend on the construction of the device. These capacitances can easily vary by several orders of magnitude as v_{ds} varies over its normal operating range. For $v_{ds} \gg V_0$, Eq.

Fig. 4.29 MOSFET equivalent circuit which accounts for the body diode and effective terminal capacitances.

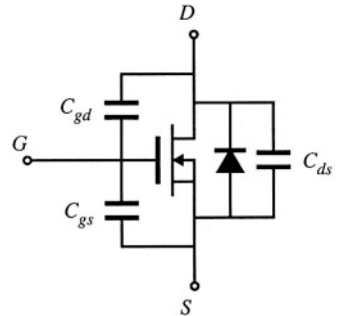


Table 4.2 Characteristics of several commercial n -channel power MOSFETs

Part number	Rated maximum voltage	Rated average current	R_{on}	Q_g (typical)
IRFZ48	60 V	50 A	0.018 Ω	110 nC
IRF510	100 V	5.6 A	0.54 Ω	8.3 nC
IRF540	100 V	28 A	0.077 Ω	72 nC
APT10M25BNR	100 V	75 A	0.025 Ω	171 nC
IRF740	400 V	10 A	0.55 Ω	63 nC
MTM15N40E	400 V	15 A	0.3 Ω	110 nC
APT5025BN	500 V	23 A	0.25 Ω	83 nC
APT1001RBNR	1000 V	11 A	1.0 Ω	150 nC

(4.4) can be approximated as

$$C_{ds}(v_{ds}) \approx C_0 \sqrt{\frac{V_0}{v_{ds}}} = \frac{C_0'}{\sqrt{v_{ds}}} \quad (4.5)$$

These expressions are used in Section 4.3.3 to determine the switching loss due to energy stored in C_{ds} .

Characteristics of several commercially available power MOSFETs are listed in Table 4.2. The gate charge Q_g is the charge that the gate drive circuit must supply to the MOSFET to raise the gate voltage from zero to some specified value (typically 10 V), with a specified value of off state drain-to-source voltage (typically 80% of the rated V_{DS}). The total gate charge is the sum of the charges on the gate-to-drain and the gate-to-source capacitance. The total gate charge is to some extent a measure of the size and switching speed of the MOSFET.

Unlike other power devices, MOSFETs are usually not selected on the basis of their rated average current. Rather, on-resistance and its influence on conduction loss are the limiting factors, and MOSFETs typically operate at average currents somewhat less than the rated value.

MOSFETs are usually the device of choice at voltages less than or equal to approximately 400 to 500 V. At these voltages, the forward voltage drop is competitive or superior to the forward voltage drops of minority-carrier devices, and the switching speed is significantly faster. Typical switching times are in the range 50 ns to 200 ns. At voltages greater than 400 to 500 V, minority-carrier devices having lower forward voltage drops, such as the IGBT, are usually preferred. The only exception is in applications where the high switching speed overrides the increased cost of silicon required to obtain acceptably low conduction loss.

4.2.3 Bipolar Junction Transistor (BJT)

A cross-section of an NPN power BJT is illustrated in Fig. 4.30. As with other power devices, current flows vertically through the silicon wafer. A lightly doped n^- region is inserted in the collector, to obtain the desired voltage breakdown rating. The transistor operates in the off state (cutoff) when the p - n base-emitter junction and the p - n^- base-collector junction are reverse-biased; the applied collector-to-emitter voltage then appears essentially across the depletion region of the p - n^- junction. The transistor operates in the on state (saturation) when both junctions are forward-biased; substantial minority charge is then present in the p and n^- regions. This minority charge causes the n^- region to exhibit a low on-

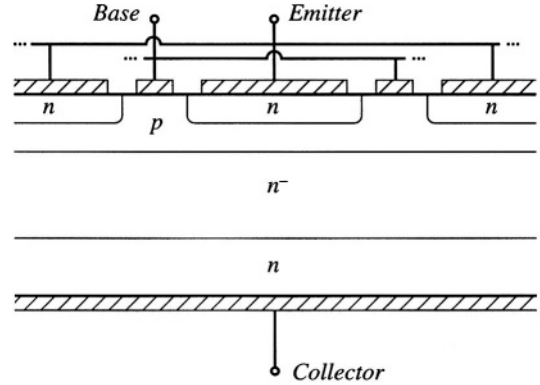


Fig. 4.30 Power BJT structure. Crosshatched regions are metallized contacts.

resistance via the conductivity modulation effect. Between the off state and the on state is the familiar active region, in which the p - n base-emitter junction is forward-biased and the p - n^- base-collector junction is reverse-biased. When the BJT operates in the active region, the collector current is proportional to the base region minority charge, which in turn is proportional (in equilibrium) to the base current. There is in addition a fourth region of operation known as *quasi-saturation*, occurring between the active and saturation regions. Quasi-saturation occurs when the base current is insufficient to fully saturate the device; hence, the minority charge present in the n^- region is insufficient to fully reduce the n^- region resistance, and high transistor on-resistance is observed.

Consider the simple switching circuit of Fig. 4.31. Figure 4.32 contains waveforms illustrating the BJT turn-on and turn-off transitions. The transistor operates in the off state during interval (1), with the base-emitter junction reverse-biased by the source voltage $v_s(t) = -V_{s1}$. The turn-on transition is initiated at the beginning of interval (2), when the source voltage changes to $v_s(t) = +V_{s2}$. Positive current is then supplied by source v_s to the base of the BJT. This current first charges the capacitances of the depletion regions of the reverse-biased base-emitter and base-collector junctions. At the end of interval (2), the base-emitter voltage exceeds zero sufficiently for the base-emitter junction to become forward-biased. The length of interval (2) is called the *turn-on delay time*. During interval (3), minority charge is injected across the base-emitter junction from the emitter into the base region; the collector current is proportional to this minority base charge. Hence during interval (3), the collector current increases. Since the transistor drives a resistive load R_L , the collector voltage also decreases during interval (3). This causes the voltage to reduce across the reverse-biased base-collector depletion region (Miller) capacitance. Increasing the base current I_{B1} (by reducing R_B or increasing V_{s2}) increases the rate of change of both the base region minority charge and the charge in the Miller capacitance. Hence, increased I_{B1} leads to a decreased turn-on switching time.

Near or at the end of interval (3), the base-collector p - n^- junction becomes forward-biased. Minority carriers are then injected into the n^- region, reducing its effective resistivity. Depending on the device geometry and the magnitude of the base current, a *voltage tail* [interval (4)] may be observed as

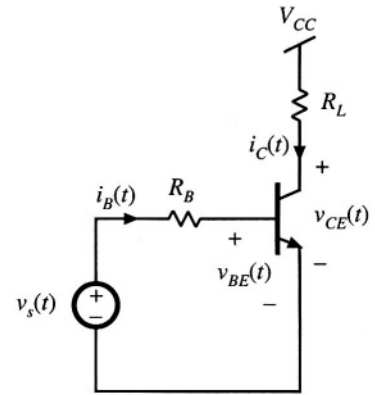


Fig. 4.31 Circuit for BJT switching time example.

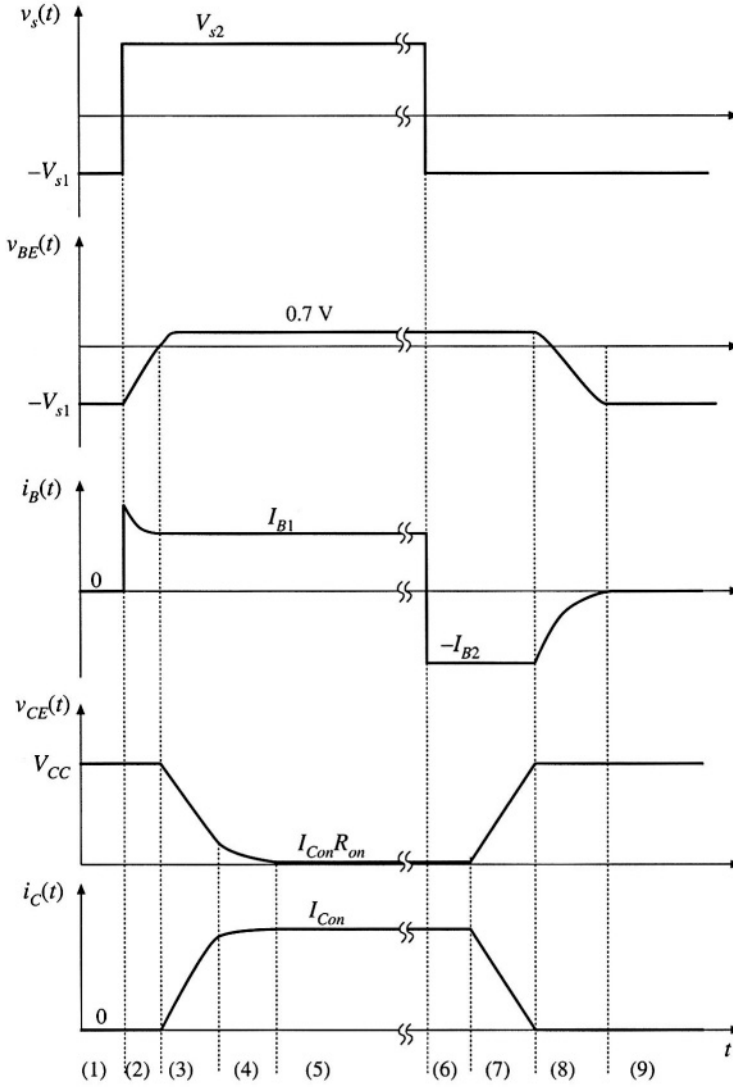
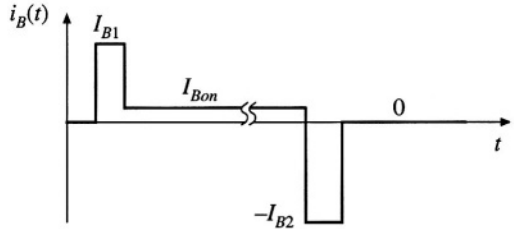


Fig. 4.32 BJT turn-on and turn-off transition waveforms.

the apparent resistance of the n^- region decreases via conductivity modulation. The BJT reaches on-state equilibrium at the beginning of interval (5), with low on-resistance and with substantial minority charge present in both the n^- and p regions. This minority charge significantly exceeds the amount necessary to support the active region conduction of the collector current I_{Con} ; its magnitude is a function of $I_{B1} - I_{Con}/\beta$, where β is the active-region current gain.

The turn-off process is initiated at the beginning of interval (6), when the source voltage changes to $v_s(t) = -V_{s1}$. The base-emitter junction remains forward-biased as long as minority carriers are present in its vicinity. Also, the collector current continues to be $i_C(t) = I_{Con}$ as long as the minority

Fig. 4.33 Ideal base current waveform for minimization of switching times.



charge exceeds the amount necessary to support the active region conduction of I_{Con} , that is, as long as *excess charge* is present. So during interval (6), a negative base current flows equal to $-I_{B2} = (-V_{s1} - v_{BE}(t))/R_B$. This negative base current actively removes the total stored minority charge. Recombination further reduces the stored minority charge. Interval (6) ends when all of the excess minority charge has been removed. The length of interval (6) is called the *storage time*. During interval (7), the transistor operates in the active region. The collector current $i_C(t)$ is now proportional to the stored minority charge. Recombination and the negative base current continue to reduce the minority base charge, and hence the collector decreases. In addition, the collector voltage increases, and hence the base current must charge the Miller capacitance. At the end of interval (7), the minority stored charge is equal to zero, and the base-emitter junction can become reverse-biased. The length of interval (7) is called the *turn-off time* or *fall time*. During interval (8), the reverse-biased base-emitter junction capacitance is discharged to voltage $-V_{s1}$. During interval (9), the transistor operates in equilibrium, in the off state.

It is possible to turn the transistor off using $I_{B2} = 0$; for example, we could let V_{s1} be approximately zero. However, this leads to very long storage and turn-off switching times. If $I_{B2} = 0$, then all of the stored minority charge must be removed passively, via recombination. From the standpoint of minimizing switching times, the base current waveform of Fig. 4.33 is ideal. The initial base current I_{B1} is large in magnitude, such that charge is inserted quickly into the base, and the turn-on switching times are short. A compromise value of equilibrium on state current I_{Bon} is chosen, to yield a reasonably low collector-to-emitter forward voltage drop, while maintaining moderate amounts of excess stored minority charge and hence keeping the storage time reasonably short. The current $-I_{B2}$ is large in magnitude, such that charge is removed quickly from the base and hence the storage and turn-off switching times are minimized.

Unfortunately, in most BJTs, the magnitudes of I_{B1} and I_{B2} must be limited because excessive values lead to device failure. As illustrated in Fig. 4.34, the base current flows laterally through the *p*

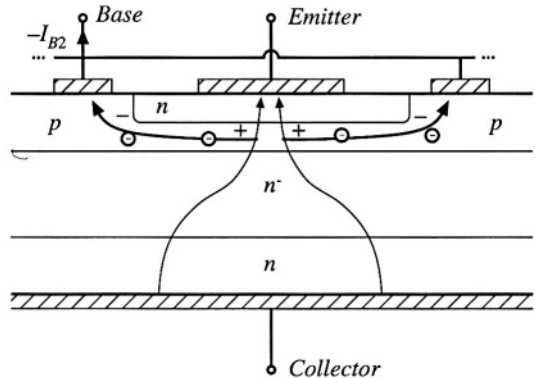


Fig. 4.34 A large I_{B2} leads to focusing of the emitter current away from the base contacts, due to the voltage induced by the lateral base region current.

region. This current leads to a voltage drop in the resistance of the p material, which influences the voltage across the base-emitter junction. During the turn-off transition, the base current $-I_{B2}$ causes the base-emitter junction voltage to be greater in the center of the base region, and smaller at the edges near the base contacts. This causes the collector current to focus near the center of the base region. In a similar fashion, a large I_{B1} causes the collector current to crowd near the edges of the base region during the turn-on transition. Since the collector-to-emitter voltage and collector current are simultaneously large during the switching transitions, substantial power loss can be associated with current focusing. Hence hot spots are induced at the center or edge of the base region. The positive temperature coefficient of the base-emitter junction current (corresponding to a negative temperature coefficient of the junction voltage) can then lead to thermal runaway and device failure. Thus, to obtain reliable operation, it may be necessary to limit the magnitudes of I_{B1} and I_{B2} . It may also be necessary to add external *snubber* networks which reduce the instantaneous transistor power dissipation during the switching transitions.

Steady-state characteristics of the BJT are illustrated in Fig. 4.35. In Fig. 4.35(a), the collector current I_C is plotted as a function of the base current I_B for various values of collector-to-emitter voltage V_{CE} . The cutoff, active, quasi-saturation, and saturation regions are identified. At a given collector cur-

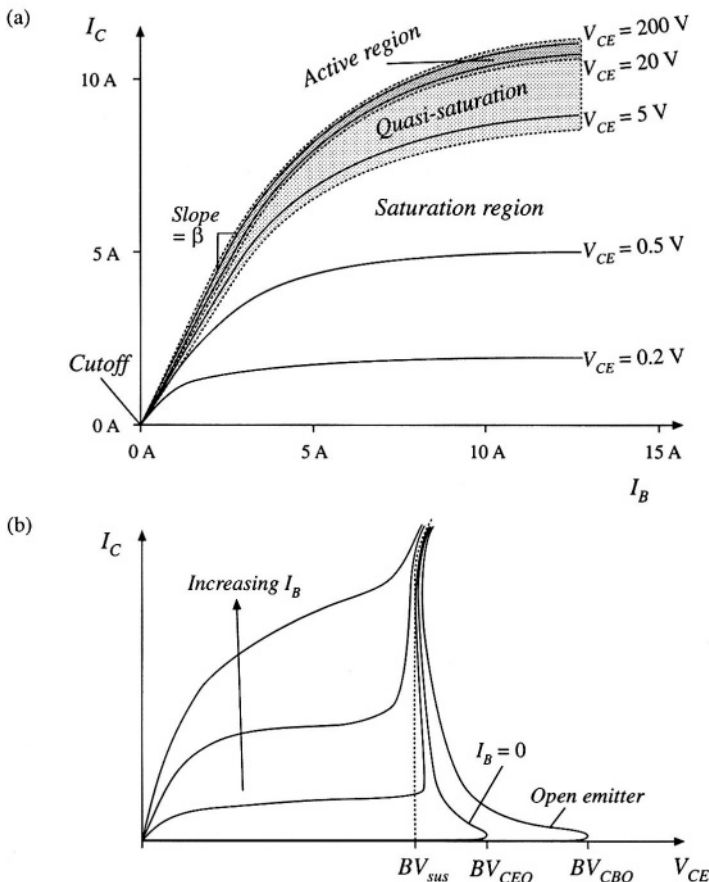


Fig. 4.35 BJT static characteristics: (a) I_C vs. I_B , illustrating the regions of operation; (b) I_C vs. V_{CE} , illustrating voltage breakdown characteristics.

rent I_C to operate in the saturation region with minimum forward voltage drop, the base current I_B must be sufficiently large. The slope dI_C/dI_B in the active region is the current gain β . It can be seen that β decreases at high current—near the rated current of the BJT, the current gain decreases rapidly and hence it is difficult to fully saturate the device. Collector current I_C is plotted as a function of collector-to-emitter voltage V_{CE} in Fig. 4.35(b), for various values of I_B . The breakdown voltages BV_{SUS} , BV_{CEO} , and BV_{CBO} are illustrated. BV_{CBO} is the avalanche breakdown voltage of the base-collector junction, with the emitter open-circuited or with sufficiently negative base current. BV_{CEO} is the somewhat smaller collector-emitter breakdown voltage observed when the base current is zero; as avalanche breakdown is approached, free carriers are created that have the same effect as a positive base current and that cause the breakdown voltage to be reduced. BV_{SUS} is the breakdown voltage observed with positive base current. Because of the high instantaneous power dissipation, breakdown usually results in destruction of the BJT. In most applications, the off state transistor voltage must not exceed BV_{CEO} .

High-voltage BJTs typically have low current gain, and hence Darlington-connected devices (Fig. 4.36) are common. If transistors Q_1 and Q_2 have current gains β_1 and β_2 , respectively, then the Darlington-connected device has the substantially increased current gain $\beta_1 + \beta_2 + \beta_1\beta_2$. In a monolithic Darlington device, transistors Q_1 and Q_2 are integrated on the same silicon wafer. Diode D_1 speeds up the turn-off process, by allowing the base driver to actively remove the stored charge of both Q_1 and Q_2 during the turn-off transition.

At voltage levels below 500 V, the BJT has been almost entirely replaced by the MOSFET in power applications. It is also being displaced in higher voltage applications, where new designs utilize faster IGBTs or other devices.

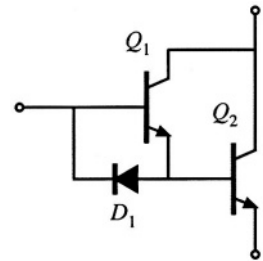


Fig. 4.36 Darlington-connected BJTs, including diode for improvement of turn-off times.

4.2.4 Insulated Gate Bipolar Transistor (IGBT)

A cross-section of the IGBT is illustrated in Fig. 4.37. Comparison with Fig. 4.26 reveals that the IGBT and power MOSFET are very similar in construction. The key difference is the p region connected to the collector of the IGBT. So the IGBT is a modern four-layer power semiconductor device having a MOS gate.

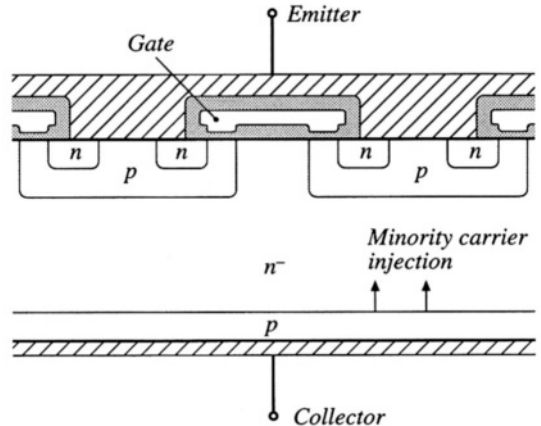


Fig. 4.37 IGBT structure. Crosshatched regions are metallized contacts. Shaded regions are insulating silicon dioxide layers.

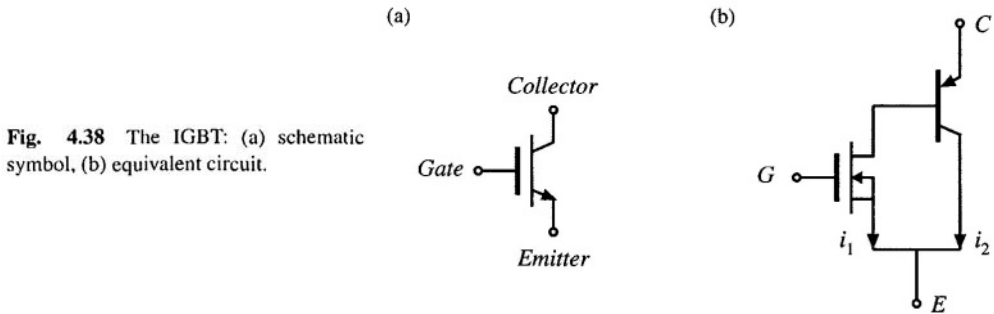


Fig. 4.38 The IGBT: (a) schematic symbol, (b) equivalent circuit.

The function of the added p region is to inject minority charges into the n^- region while the device operates in the on state, as illustrated in Fig. 4.37. When the IGBT conducts, the p - n^- junction is forward-biased, and the minority charges injected into the n^- region cause conductivity modulation. This reduces the on-resistance of the n^- region, and allows high-voltage IGBTs to be constructed which have low forward voltage drops. As of 1999, IGBTs rated as low as 600 V and as high as 3300 V are readily available. The forward voltage drops of these devices are typically 2 to 4 V, much lower than would be obtained in equivalent MOSFETs of the same silicon area.

Several schematic symbols for the IGBT are in current use; the symbol illustrated in Fig. 4.38(a) is the most popular. A two-transistor equivalent circuit for the IGBT is illustrated in Fig. 4.38(b). The IGBT functions effectively as an n -channel power MOSFET, cascaded by a PNP emitter-follower BJT. The physical locations of the two effective devices are illustrated in Fig. 4.39. It can be seen that there are two effective currents: the effective MOSFET channel current i_1 , and the effective PNP collector current i_2 .

The price paid for the reduced voltage drop of the IGBT is its increased switching times, especially during the turn-off transition. In particular, the IGBT turn-off transition exhibits a phenomenon known as *current tailing*. The effective MOSFET can be turned off quickly, by removing the gate charge such that the gate-to-emitter voltage is negative. This causes the channel current i_1 to quickly become zero. However, the PNP collector current i_2 continues to flow as long as minority charge is present in the n^- region. Since there is no way to actively remove the stored minority charge, it slowly decays via recombination. So i_2 slowly decays in proportion to the minority charge, and a current tail is observed. The length of the current tail can be reduced by introduction of recombination centers in the n^- region, at

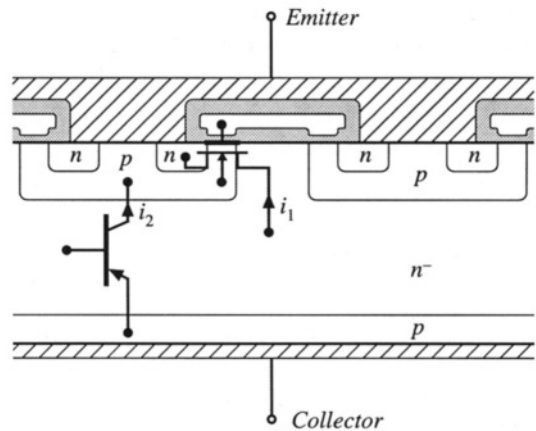


Fig. 4.39 Physical locations of the effective MOSFET and PNP components of the IGBT.

Table 4.3 Characteristics of several commercial IGBTs

Part number	Rated maximum voltage	Rated average current	V_F (typical)	t_f (typical)
Single-chip devices				
HGTP12N60A4	600 V	23 A	2.0 V	70 ns
HGTG32N60E2	600 V	32 A	2.4 V	0.62 μ s
HGTG30N120D2	1200 V	30 A	3.2 V	0.58 μ s
Multiple-chip modules				
CM400HA-12E	600 V	400 A	2.7 V	0.3 μ s
CM300HA-24E	1200 V	300 A	2.7 V	0.3 μ s
CM800HA-34H	1700 V	800 A	3.3 V	0.6 μ s
High voltage modules				
CM 800HB-50H	2500 V	800 A	3.15 V	1.0 μ s
CM 600HB-90H	4500 V	900 A	3.3 V	1.2 μ s

the expense of a somewhat increased on-resistance. The current gain of the effective PNP transistor can also be minimized, causing i_1 to be greater than i_2 . Nonetheless, the turn-off switching time of the IGBT is significantly longer than that of the MOSFET, with typical turn-off times in the range 0.5 μ s to 5 μ s. Switching loss induced by IGBT current tailing is discussed in Section 4.3.1. The switching frequencies of PWM converters containing IGBTs are typically in the range 1 to 30 kHz.

The added $p-n$ diode junction of the IGBT is not normally designed to block significant voltage. Hence, the IGBT has negligible reverse voltage-blocking capability.

Since the IGBT is a four-layer device, there is the possibility of SCR-type latchup, in which the IGBT cannot be turned off by gate voltage control. Recent devices are not susceptible to this problem. These devices are quite robust, hot-spot and current crowding problems are nonexistent, and the need for external snubber circuits is minimal.

The on-state forward voltage drop of the IGBT can be modeled by a forward-biased diode junction, in series with an effective on-resistance. The temperature coefficient of the IGBT forward voltage drop is complicated by the fact that the diode junction voltage has a negative temperature coefficient, while the on-resistance has a positive temperature coefficient. Fortunately, near rated current the on-resistance dominates, leading to an overall positive temperature coefficient. In consequence, IGBTs can be easily connected in parallel, with a modest current derating. Large modules are commercially available, containing multiple parallel-connected chips.

Characteristics of several commercially available single-chip IGBTs and multiple-chip IGBT modules are listed in Table 4.3.

4.2.5 Thyristors (SCR, GTO, MCT)

Of all conventional semiconductor power devices, the silicon-controlled rectifier (SCR) is the oldest, has the lowest cost per rated kVA, and is capable of controlling the greatest amount of power. Devices having voltage ratings of 5000 to 7000 V and current ratings of several thousand amperes are available. In utility dc transmission line applications, series-connected light-triggered SCRs are employed in inverters and rectifiers that interface the ac utility system to dc transmission lines which carry roughly 1 kA and 1 MV. A single large SCR fills a silicon wafer that is several inches in diameter, and is mounted in a hockey-puck-style case.

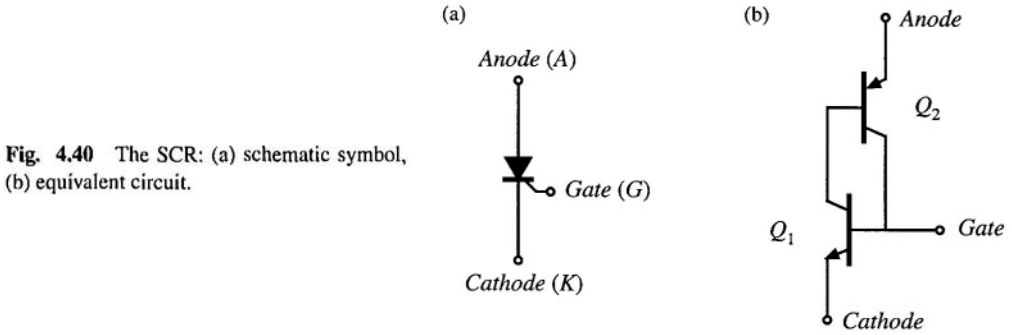


Fig. 4.40 The SCR: (a) schematic symbol, (b) equivalent circuit.

The schematic symbol of the SCR is illustrated in Fig. 4.40(a), and an equivalent circuit containing NPN and PNP BJT devices is illustrated in Fig. 4.40(b). A cross-section of the silicon chip is illustrated in Fig. 4.41. Effective transistor Q_1 is composed of the n , p , and n^- regions, while effective transistor Q_2 is composed of the p , n^- , and p regions as illustrated.

The device is capable of blocking both positive and negative anode-to-cathode voltages. Depending on the polarity of the applied voltage, one of the p - n^- junctions is reverse-biased. In either case, the depletion region extends into the lightly doped n^- region. As with other devices, the desired voltage breakdown rating is obtained by proper design of the n^- region thickness and doping concentration.

The SCR can enter the on state when the applied anode-to-cathode voltage v_{AK} is positive. Positive gate current i_G then causes effective transistor Q_1 to turn on; this in turn supplies base current to effective transistor Q_2 , and causes it to turn on as well. The effective connections of the base and collector regions of transistors Q_1 and Q_2 constitute a positive feedback loop. Provided that the product of the current gains of the two transistors is greater than one, then the currents of the transistors will increase regeneratively. In the on state, the anode current is limited by the external circuit, and both effective transistors operate fully saturated. Minority carriers are injected into all four regions, and the resulting conductivity modulation leads to very low forward voltage drop. In the on state, the SCR can be modeled as a forward-biased diode junction in series with a low-value on-resistance. Regardless of the gate current, the SCR is latched in the on state: it cannot be turned off except by application of negative anode current or negative anode-to-cathode voltage. In phase controlled converters, the SCR turns off at the zero crossing of the converter ac input or output waveform. In forced commutation converters, external commuta-

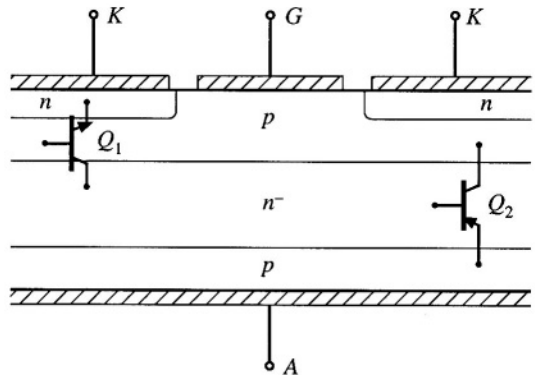
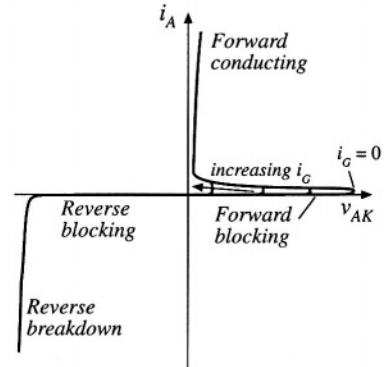


Fig. 4.41 Physical locations of the effective NPN and PNP components of the SCR.

Fig. 4.42 Static i_A - v_{AK} characteristics of the SCR.

tion circuits force the controlled turn-off of the SCR, by reversing either the anode current or the anode-to-cathode voltage.

Static i_A - v_{AK} characteristics of the conventional SCR are illustrated in Fig. 4.42. It can be seen that the SCR is a voltage-bidirectional two-quadrant switch. The turn-on transition is controlled actively via the gate current. The turn-off transition is passive.

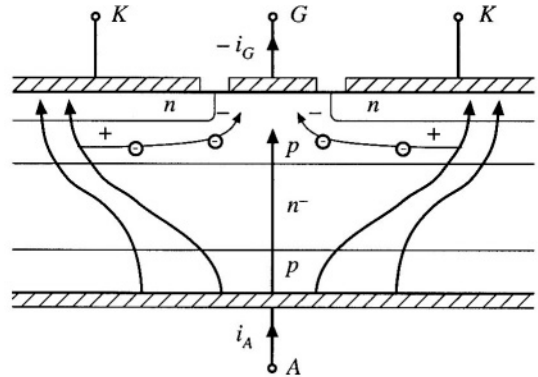
During the turn-off transition, the rate at which forward anode-to-cathode voltage is reapplied must be limited, to avoid retriggering the SCR. The turn-off time t_q is the time required for minority stored charge to be actively removed via negative anode current, and for recombination of any remaining minority charge. During the turn-off transition, negative anode current actively removes stored minority charge, with waveforms similar to diode turn-off transition waveforms of Fig. 4.25. Thus, after the first zero crossing of the anode current, it is necessary to wait for time t_q before reapplying positive anode-to-cathode voltage. It is then necessary to limit the rate at which the anode-to-cathode voltage increases, to avoid retriggering the device. Inverter-grade SCRs are optimized for faster switching times, and exhibit smaller values of t_q .

Conventional SCR wafers have large feature size, with coarse or nonexistent interdigitation of the gate and cathode contacts. The parasitic elements arising from this large feature size lead to several limitations. During the turn-on transition, the rate of increase of the anode current must be limited to a safe value. Otherwise, cathode current focusing can occur, which leads to formation of hot spots and device failure.

The coarse feature size of the gate and cathode structure is also what prevents the conventional SCR from being turned off by active gate control. One might apply a negative gate current, in an attempt to actively remove all of the minority stored charge and to reverse-bias the p - n gate-cathode junction. The reason that this attempt fails is illustrated in Fig. 4.43. The large negative gate current flows laterally through the adjoining p region, inducing a voltage drop as shown. This causes the gate-cathode junction voltage to be smaller near the gate contact, and relatively larger away from the gate contact. The negative gate current is able to reverse-bias only the portion of the gate-cathode junction in the vicinity of the gate contact; the remainder of the gate-cathode junction continues to be forward-biased, and cathode current continues to flow. In effect, the gate contact is able to influence only the nearby portions of the cathode.

The gate turn off thyristor, or GTO, is a modern power device having small feature size. The gate and cathode contacts highly interdigitated, such that the entire gate-cathode p - n junction can be reverse-biased via negative gate current during the turn-off transition. Like the SCR, a single large GTO can fill an entire silicon wafer. Maximum voltage and current ratings of commercial GTOs are lower than those of SCRs.

Fig. 4.43 Negative gate current is unable to completely reverse-bias the gate-cathode junction. The anode current focuses away from the gate contact.



The turn-off gain of a GTO is the ratio of on-state current to the negative gate current magnitude required to switch the device off. Typical values of this gain are 2 to 5, meaning that several hundred amperes of negative gate current may be required to turn off a GTO conducting 1000 A. Also of interest is the maximum controllable on-state current. The GTO is able to conduct peak currents significantly greater than the rated average current; however, it may not be possible to switch the device off under gate control while these high peak currents are present.

The MOS-controlled thyristor, or MCT, is a recent power device in which MOSFETs are integrated onto a highly interdigitated SCR, to control the turn-on and turn-off processes. Like the MOSFET and IGBT, the MCT is a single-quadrant device whose turn-on and turn-off transitions are controlled by a MOS gate terminal. Commercial MCTs are *p*-type devices. Voltage-bidirectional two-quadrant MCTs, and *n*-type MCTs, are also possible.

A cross-section of an MCT containing MOSFETs for control of the turn-on and turn-off transitions is illustrated in Fig. 4.44. An equivalent circuit which explains the operation of this structure is given in Fig. 4.45. To turn the device on, the gate-to-anode voltage is driven negative. This forward-biases *p*-channel MOSFET Q_3 , forward-biasing the base-emitter junction of BJT Q_1 . Transistors Q_1 and Q_2 then latch in the on-state. To turn the device off, the gate-to-anode voltage is driven positive. This forward-biases *n*-channel MOSFET Q_4 , which in turn reverse-biases the base-emitter junction of BJT Q_2 . The BJTs then turn off. It is important that the on-resistance of the *n*-channel MOSFET be small enough

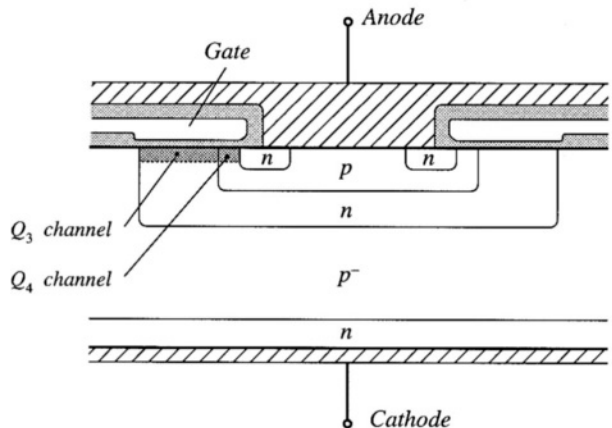
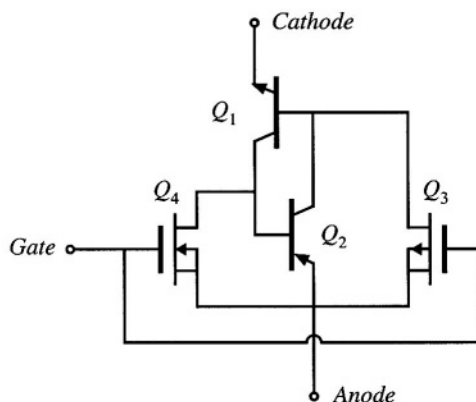


Fig. 4.44 MCT structure. Crosshatched regions are metallized contacts. Lightly shaded regions are insulating silicon dioxide layers.

Fig. 4.45 Equivalent circuit for the MCT.



that sufficient influence on the cathode current is exerted—this limits the maximum controllable on state current (i.e., the maximum current that can be interrupted via gate control).

High-voltage MCTs exhibit lower forward voltage drops and higher current densities than IGBTs of similar voltage ratings and silicon area. However, the switching times are longer. Like the GTO, the MCT can conduct considerable surge currents; but again, the maximum current that can be interrupted via gate control is limited. To obtain a reliable turn-off transition, external snubbers are required to limit the peak anode-to-cathode voltage. A sufficiently fast gate-voltage rise time is also required. To some extent, the MCT is still an emerging device—future generations of MCTs may exhibit considerable improvements in performance and ratings.

4.3 SWITCHING LOSS

Having implemented the switches using semiconductor devices, we can now discuss another major source of loss and inefficiency in converters: switching loss. As discussed in the previous section, the turn-on and turn-off transitions of semiconductor devices require times of tens of nanoseconds to microseconds. During these switching transitions, very large instantaneous power loss can occur in the semiconductor devices. Even though the semiconductor switching times are short, the resulting average power loss can be significant.

Semiconductor devices are charge controlled. For example, the conducting state of a MOSFET is determined by the charge on its gate and in its channel, and the conducting state of a silicon diode or a BJT is determined by the presence or absence of stored minority charge in the vicinity of the semiconductor junctions inside the device. To switch a semiconductor device between the on and off states, the controlling charge must be inserted or removed; hence, the amount of controlling charge influences both the switching times and the switching loss. Charge, and energy, are also stored in the output capacitances of semiconductor devices, and energy is stored in the leakage and stray inductances in the circuit. In most converter circuits, these stored energies are also lost during the switching transitions.

In this section the major sources of switching loss are described, and a simple method for estimation of their magnitudes is given. For clarity, conduction losses and semiconductor forward voltage drops are neglected throughout this discussion.

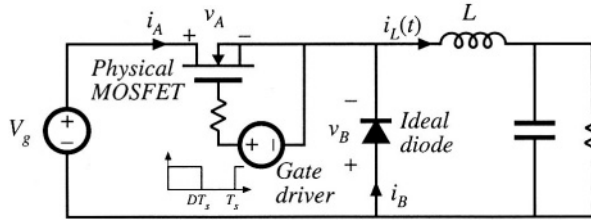


Fig. 4.46 MOSFET driving a clamped inductive load, buck converter example.

4.3.1 Transistor Switching with Clamped Inductive Load

Let's consider first the switching waveforms in the buck converter of Fig. 4.46. Let us treat the diode as ideal, and investigate only the switching loss due to the MOSFET switching times. The MOSFET drain-to-source capacitance is also neglected.

The diode and inductor present a clamped inductive load to the transistor. With such a load, the transistor voltage $v_A(t)$ and current $i_A(t)$ do not change simultaneously. For example, a magnified view of the transistor turn-off-transition waveforms is given in Fig. 4.47. For simplicity, the waveforms are

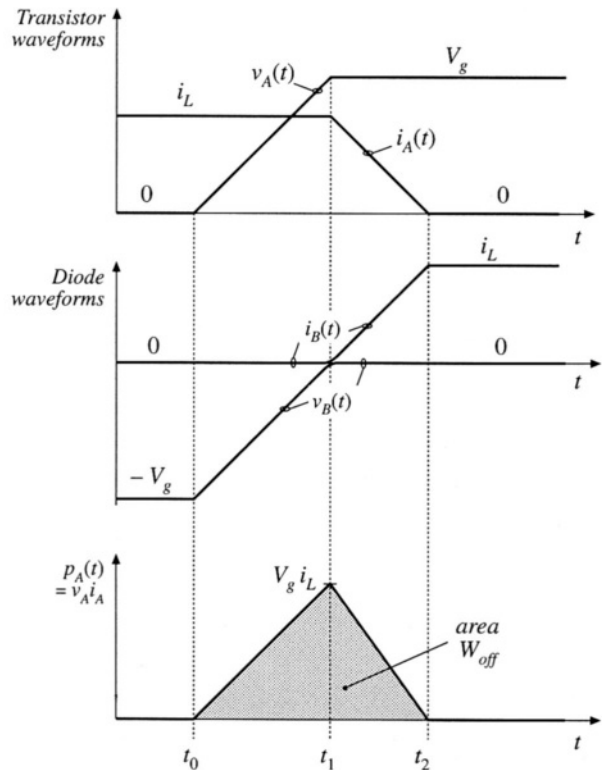


Fig. 4.47 Magnified view of transistor turn-off transition waveforms for the circuit of Fig. 4.46.

approximated as piecewise-linear. The switching times are short, such that the inductor current $i_L(t)$ is essentially constant during the entire switching transition $t_0 < t < t_2$. No current flows through the diode while the diode is reverse-biased, and the diode cannot become forward-biased while its voltage $v_B(t)$ is negative. So first, the voltage $v_A(t)$ across the transistor must rise from zero to V_g . The interval length $(t_1 - t_0)$ is essentially the time required for the gate driver to charge the MOSFET gate-to-drain capacitance. The transistor current $i_A(t)$ is constant and equal to i_L during this interval.

The diode voltage $v_B(t)$ and current $i_B(t)$ are given by

$$\begin{aligned} v_B(t) &= v_A(t) - V_g \\ i_A(t) + i_B(t) &= i_L \end{aligned} \quad (4.6)$$

At time $t = t_1$, when $v_A = V_g$, the diode becomes forward-biased. The current i_L now begins to commute from the transistor to the diode. The interval length $(t_2 - t_1)$ is the time required for the gate driver to discharge the MOSFET gate-to-source capacitance down to the threshold voltage which causes the MOSFET to be in the off state.

The instantaneous power $p_A(t)$ dissipated by the transistor is equal to $v_A(t)i_A(t)$. This quantity is also sketched in Fig. 4.47. The energy W_{off} lost during the transistor turn-off transition is the area under this waveform. With the simplifying assumption that the waveforms are piecewise-linear, then the energy lost is the area of the shaded triangle:

$$W_{off} = \frac{1}{2} V_g i_L (t_2 - t_1) \quad (4.7)$$

This is the energy lost during each transistor turn-off transition in the simplified circuit of Fig. 4.46.

The transistor turn-on waveforms of the simplified circuit of Fig. 4.46 are qualitatively similar to those of Fig. 4.47, with the time axis reversed. The transistor current must first rise from 0 to i_L . The diode then becomes reverse-biased, and the transistor voltage can fall from V_g to zero. The instantaneous transistor power dissipation again has peak value $V_g i_L$, and if the waveforms are piecewise linear, then the energy lost during the turn-on transition W_{on} is given by $0.5 V_g i_L$ multiplied by the transistor turn-on time.

Thus, during one complete switching period, the total energy lost during the turn-on and turn-off transitions is $(W_{on} + W_{off})$. If the switching frequency is f_s , then the average power loss incurred due to switching is

$$P_{sw} = \frac{1}{T_s} \int_{\text{switching transitions}} p_A(t) dt = (W_{on} + W_{off}) f_s \quad (4.8)$$

So the switching loss P_{sw} is directly proportional to the switching frequency.

An example where the loss due to transistor switching times is particularly significant is the current tailing phenomenon observed during the turn-off transition of the IGBT. As discussed in Section 4.2.4, current tailing occurs due to the slow recombination of stored minority charge in the n^- region of the IGBT. This causes the collector current to slowly decay after the gate voltage has been removed.

A buck converter circuit containing an ideal diode and nonideal (physical) IGBT is illustrated in Fig. 4.48. Turn-off transition waveforms are illustrated in Fig. 4.49; these waveforms are similar to the MOSFET waveforms of Fig. 4.47. The diode is initially reverse-biased, and the voltage $v_A(t)$ rises from approximately zero to V_g . The interval length $(t_1 - t_0)$ is the time required for the gate drive circuit to charge the IGBT gate-to-collector capacitance. At time $t = t_1$, the diode becomes forward-biased, and current begins to commute from the IGBT to the diode. The interval $(t_2 - t_1)$ is the time required for the gate drive circuit to discharge the IGBT gate-to-emitter capacitance to the threshold value which causes

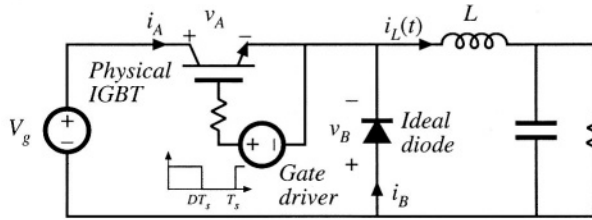


Fig. 4.48 IGBT switching loss example.

the effective MOSFET in Fig. 4.38(b) to be in the off state. This time can be minimized by use of a high-current gate drive circuit which discharges the gate capacitance quickly. However, switching off the effective MOSFET does not completely interrupt the IGBT current $i_A(t)$; current $i_2(t)$ continues to flow through the effective PNP bipolar junction transistor of Fig. 4.38(b) as long as minority carriers continue to exist within its base region. During the interval $t_2 < t < t_3$, the current is proportional to this stored minority charge, and the current tail interval length $(t_3 - t_2)$ is equal to the time required for this remaining stored minority charge to recombine.

The energy W_{off} lost during the turn-off transition of the IGBT is again the area under the instantaneous power waveform, as illustrated in Fig. 4.49. The switching loss can again be evaluated using Eq. (4.8).

The switching times of the IGBT are typically in the vicinity of 0.2 to 2 μs , or several times

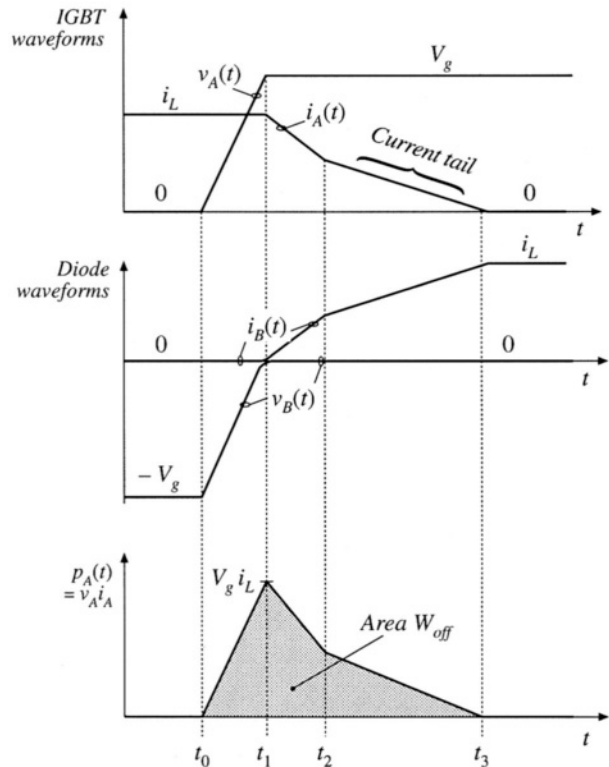


Fig. 4.49 IGBT turn-off transition waveforms for the circuit of Fig. 4.48.

longer than those of the power MOSFET. The resulting switching loss limits the maximum switching frequencies of conventional PWM converters employing IGBTs to roughly 1 to 30 kHz.

4.3.2 Diode Recovered Charge

As discussed previously, the familiar exponential i - v characteristic of the diffused-junction p - n diode is an equilibrium relationship. During switching transients, significant deviations from this characteristic are observed, which can induce transistor switching loss. In particular, during the diode turn-off transient, its stored minority charge must be removed, either actively via negative current $i_B(t)$, or passively via recombination inside the device. The diode remains forward-biased while minority charge is present in the vicinity of the diode semiconductor junction. The initial amount of minority charge is a function of the forward current, and its rate of change, under forward-biased conditions. The turn-off switching time is the time required to remove all of this charge, and to establish a new reverse-biased operating point. This process of switching the diode from the forward-biased to reverse-biased states is called *reverse recovery*.

Again, most diffused-junction power diodes are actually p - n - n^+ or p - i - n devices. The lightly doped or intrinsic region (of the diode and other power semiconductor devices as well) allows large breakdown voltages to be obtained. Under steady-state forward-biased conditions, a substantial amount of stored charge is present in this region, increasing its conductivity and leading to a low diode on-resistance. It takes time to insert and remove this charge, however, so there is a tradeoff between high breakdown voltage, low on-resistance, and fast switching times.

To understand how the diode stored charge induces transistor switching loss, let us consider the buck converter of Fig. 4.50. Assume for this discussion that the transistor switching times are much faster than the switching times of the diode, such that the diode reverse recovery mechanism is the only significant source of switching loss. A magnified view of the transistor-turn-on transition waveforms under these conditions is given in Fig. 4.51.

Initially, the diode conducts the inductor current, and hence some amount of stored minority charge is present in the diode. The transistor is initially in the off state. When the transistor turns on, a negative current flows through the diode; this current actively removes some or most of the diode stored minority charge, while the remainder of the minority charge recombines within the diode. The rate of change of the current is typically limited by the package inductance and other stray inductances present in the external circuit; hence, the peak magnitude of the reverse current depends on the external circuit, and can be many times larger than the forward current i_L . The area within the negative portion of the diode current waveform is the recovered stored charge Q_r , while the interval length $(t_2 - t_0)$ is the reverse recovery time t_r . The magnitude of Q_r is a function of the on state forward current i_L at the initiation of the turn-off process, as well as the circuit-limited rate-of-change of the diode current, $di_B(t)/dt$. During

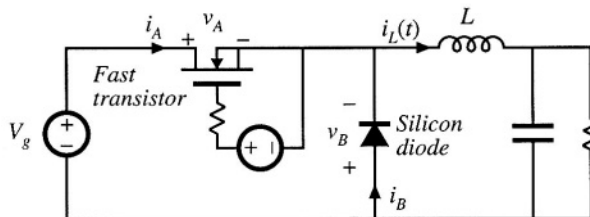


Fig. 4.50 Example, switching loss induced by diode stored charge.

the interval $t_0 < t < t_1$, the diode remains forward-biased, and hence the transistor voltage is V_g . At time $t = t_1$, the stored charge in the vicinity of the p - n or p - i junction is exhausted. This junction becomes reverse-biased, and begins to block voltage. During the interval $t_1 < t < t_2$, the diode voltage decreases to $-V_g$. Some negative diode current continues to flow, removing any remaining stored minority charge as well as charging the depletion layer capacitance. At time $t = t_2$, this current is essentially zero, and the diode operates in steady state under reverse-biased conditions.

Diodes in which the interval length $(t_2 - t_1)$ is short compared to $(t_1 - t_0)$ are called abrupt-recovery or “snappy” diodes. Soft recovery diodes exhibit larger values of $(t_2 - t_1)/(t_1 - t_0)$. When significant package and/or stray inductance is present in series with the diode, ringing of the depletion region capacitance with the package and stray inductances may be observed. If severe, this ringing can cause excess reverse voltage that leads to device failure. External R - C snubber circuits are sometimes necessary for reliable operation. The reverse-recovery characteristics of soft recovery diodes are intended to exhibit less ringing and voltage overshoot. Snubbing of these diodes can be reduced or eliminated.

The instantaneous power $p_A(t)$ dissipated in the transistor is also sketched in Fig. 4.51. The energy lost during the turn-on transition is

$$W_D = \int_{\text{switching transition}} v_A(t) i_A(t) dt \quad (4.9)$$

For an abrupt-recovery diode in which $(t_2 - t_1) \ll (t_1 - t_0)$, this integral can be evaluated in a simple manner. The transistor voltage $v_A(t)$ is then equal to V_g for essentially the entire diode recovery interval. In addition, $i_A = i_L - i_B$. Equation (4.9) then becomes

$$\begin{aligned} W_D &\approx \int_{\text{switching transition}} V_g (i_L - i_B(t)) dt \\ &= V_g i_L t_r + V_g Q_r \end{aligned} \quad (4.10)$$

where the recovered charge Q_r is defined as the integral of the diode current $-i_B(t)$ over the interval $t_0 < t < t_2$. Hence, the diode reverse recovery process leads directly to switching loss W_{Df_s} . This is often the largest single component of switching loss in a conventional switching converter. It can be reduced by use of faster diodes, designed for minimization of stored minority charge.

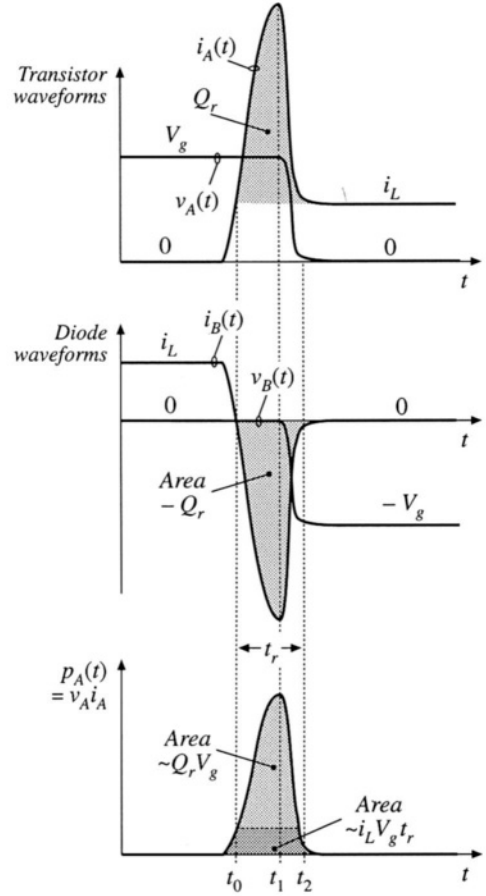
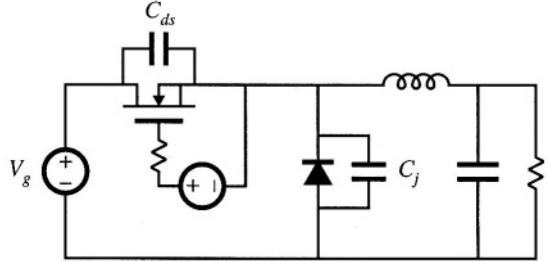


Fig. 4.51 Transistor-turn-on transition waveforms for the circuit of Fig. 4.50.

Fig. 4.52 The energy stored in the semiconductor output capacitances is lost during the transistor turn-on transition.



4.3.3 Device Capacitances, and Leakage, Package, and Stray Inductances

Reactive elements can also lead to switching loss. Capacitances that are effectively in parallel with switching elements are shorted out when the switch turns on, and any energy stored in the capacitance is lost. The capacitances are charged without energy loss when the switching elements turn off, and the transistor turn-off loss W_{off} computed in Eq. (4.7) may be reduced. Likewise, inductances that are effectively in series with a switching element lose their stored energy when the switch turns off. Hence, series inductances lead to additional switching loss at turn-off, but can reduce the transistor turn-on loss.

The stored energies of the reactive elements can be summed to find the total energy loss per switching period due to these mechanisms. For linear capacitors and inductors, the stored energy is

$$\begin{aligned} W_C &= \sum_{\text{capacitive elements}} \frac{1}{2} C_i V_i^2 \\ W_L &= \sum_{\text{inductive elements}} \frac{1}{2} L_j I_j^2 \end{aligned} \quad (4.11)$$

A common source of this type of switching loss is the output capacitances of the semiconductor switching devices. The depletion layers of reverse-biased semiconductor devices exhibit capacitance which stores energy. When the transistor turns on, this stored energy is dissipated by the transistor. For example, in the buck converter of Fig. 4.52, the MOSFET exhibits drain-to-source capacitance C_{ds} , and the reverse-biased diode exhibits junction capacitance C_j . During the switching transitions these two capacitances are effectively in parallel, since the dc source V_g is effectively a short-circuit at high frequency. To the extent that the capacitances are linear, the energy lost when the MOSFET turns on is

$$W_C = \frac{1}{2} (C_{ds} + C_j) V_g^2 \quad (4.12)$$

Typically, this type of switching loss is significant at voltage levels above 100 V. The MOSFET gate drive circuit, which must charge and discharge the MOSFET gate capacitances, also exhibits this type of loss.

As noted in Section 4.2.2, the incremental drain-to-source capacitance C_{ds} of the power MOSFET is a strong function of the drain-to-source voltage v_{ds} . $C_{ds}(v_{ds})$ follows an approximate inverse-square-root dependence of v_{ds} , as given by Eq. (4.5). The energy stored in C_{ds} at $v_{ds} = V_{DS}$ is

$$W_{Cds} = \int v_{ds} i_C dt = \int_0^{V_{DS}} v_{ds} C_{ds}(v_{ds}) dv_{ds} \quad (4.13)$$

where $i_C = C_{ds}(v_{ds}) dv_{ds}/dt$ is the current in C_{ds} . Substitution of Eq. (4.5) into (4.13) yields

$$W_{C_{ds}} = \int_0^{V_{DS}} C_{ds}'(v_{ds}) \sqrt{v_{ds}} dv_{ds} = \frac{2}{3} C_{ds}(V_{DS}) V_{DS}^2 \quad (4.14)$$

This energy is lost each time the MOSFET switches on. From the standpoint of switching loss, the drain-to-source capacitance is equivalent to a linear capacitance having the value $\frac{4}{3} C_{ds}(V_{DS})$.

The Schottky diode is essentially a majority-carrier device, which does not exhibit a reverse-recovery transient such as in Fig. 4.51. Reverse-biased Schottky diodes do exhibit significant junction capacitance, however, which can be modeled with a parallel capacitor C_j as in Fig. 4.52, and which leads to energy loss at the transistor turn-on transition.

Common sources of series inductance are transformer leakage inductances in isolated converters (discussed in Chapter 6), as well as the inductances of interconnections and of semiconductor device packages. In addition to generating switching loss, these elements can lead to excessive peak voltage stress during the transistor turn-off transition. Interconnection and package inductances can lead to significant switching loss in high-current applications, and leakage inductance is an important source of switching loss in many transformer-isolated converters.

Diode stored minority charge can induce switching loss in the (nonideal) converter reactive elements. As an example, consider the circuit of Fig. 4.53, containing an ideal voltage source $v_i(t)$, an inductor L , a capacitor C (which may represent the diode junction capacitance, or the junction capacitance in parallel with an external capacitor), and a silicon diode. The diode switching processes of many converters can be modeled by a circuit of this form. Many rectifier circuits containing SCRs exhibit similar waveforms. The voltage source produces the rectangular waveform $v_i(t)$ illustrated in Fig. 4.54. This voltage is initially positive, causing the diode to become forward-biased and the inductor current $i_L(t)$ to increase linearly with slope V_1/L . Since the current is increasing, the stored minority charge inside the diode also increases. At time $t = t_1$, the source voltage $v_i(t)$ becomes negative, and the inductor current decreases with slope $di_L/dt = -V_2/L$. The diode stored charge also decreases, but at a slower rate that depends not only on i_L but also on the minority carrier recombination lifetime of the silicon material in the diode. Hence, at time $t = t_2$, when $i_L(t)$ reaches zero, some stored minority charge remains in the diode. So the diode continues to be forward-biased, and the inductor current continues to decrease with the same slope. The negative current for $t > t_2$ constitutes a reverse diode current, which actively removes diode stored charge. At some time later, $t = t_3$, the diode stored charge in the vicinity of the diode junction becomes zero, and the diode junction becomes reverse-biased. The inductor current is now negative, and must flow through the capacitor. The inductor and capacitor then form a series resonant circuit, which rings with decaying sinusoidal waveforms as shown. This ringing is eventually damped out by the parasitic loss elements of the circuit, such as the inductor winding resistance, inductor core loss, and capacitor equivalent series resistance.

The diode recovered charge induces loss in this circuit. During the interval $t_2 < t < t_3$, the minority stored charge Q_r recovered from the diode is

$$Q_r = - \int_{t_2}^{t_3} i_L(t) dt \quad (4.15)$$

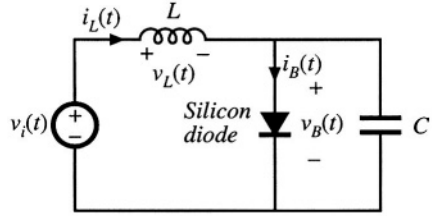


Fig. 4.53 A circuit in which the diode stored charge induces ringing, and ultimately switching loss, in (nonideal) reactive elements.

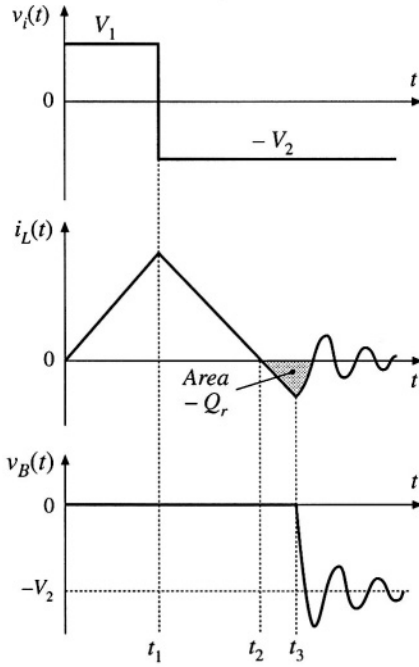


Fig. 4.54 Waveforms of the circuit of Fig. 4.53.

This charge is directly related to the energy stored in the inductor during this interval. The energy W_L stored in the inductor is the integral of the power flowing into the inductor:

$$W_L = \int_{t_2}^{t_3} v_L(t) i_L(t) dt \quad (4.16)$$

During this interval, the applied inductor voltage is

$$v_L(t) = L \frac{di_L(t)}{dt} = -V_2 \quad (4.17)$$

Substitution of Eq. (4.17) into Eq. (4.16) leads to

$$W_L = \int_{t_2}^{t_3} L \frac{di_L(t)}{dt} i_L(t) dt = \int_{t_2}^{t_3} (-V_2) i_L(t) dt \quad (4.18)$$

Evaluation of the integral on the left side yields the stored inductor energy at $t = t_3$, or $Li_L^2(t_3)/2$. The right-side integral is evaluated by noting that V_2 is constant and by substitution of Eq. (4.15), yielding $V_2 Q_r$. Hence, the energy stored in the inductor at $t = t_3$ is

$$W_L = \frac{1}{2} Li_L^2(t_3) = V_2 Q_r \quad (4.19)$$

or, the recovered charge multiplied by the source voltage. For $t > t_3$, the ringing of the resonant circuit formed

by the inductor and capacitor causes this energy to be circulated back and forth between the inductor and capacitor. If parasitic loss elements in the circuit cause the ringing amplitude to eventually decay to zero, then the energy becomes lost as heat in the parasitic elements.

So diode stored minority charge can lead to loss in circuits that do not contain an active switching element. Also, ringing waveforms that decay before the end of the switching period indicate the presence of switching loss.

4.3.4 Efficiency vs. Switching Frequency

Suppose next that we add up all of the energies lost due to switching, as discussed above:

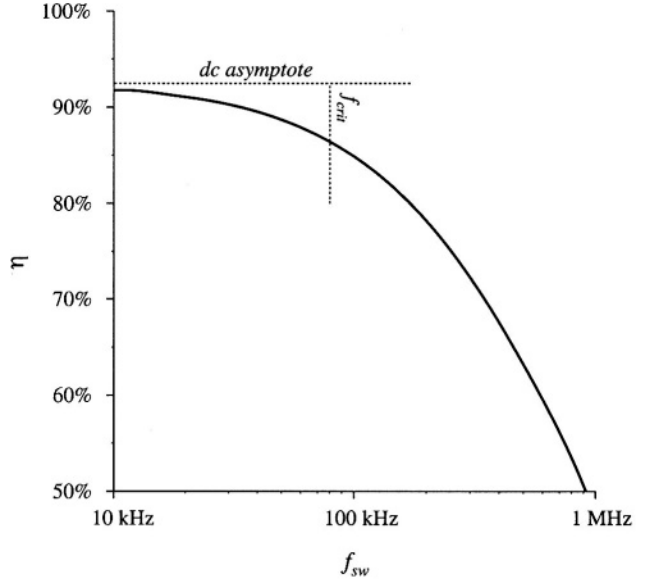
$$W_{tot} = W_{on} + W_{off} + W_D + W_C + W_L + \dots \quad (4.20)$$

This is the energy lost in the switching transitions of one switching period. To obtain the average switching power loss, we must multiply by the switching frequency:

$$P_{sw} = W_{tot} f_{sw} \quad (4.21)$$

Other losses in the converter include the conduction losses P_{cond} , modeled and solved as in Chapter 3,

Fig. 4.55 Efficiency vs. switching frequency, based on Eq. (4.22), using arbitrary choices for the values of loss and load power. Switching loss causes the efficiency to decrease rapidly at high frequency.



and other frequency-independent fixed losses P_{fixed} , such as the power required to operate the control circuit. The total loss is therefore

$$P_{loss} = P_{cond} + P_{fixed} + W_{tot} f_{sw} \quad (4.22)$$

which increases linearly with frequency. At the critical frequency

$$f_{crit} = \frac{P_{cond} + P_{fixed}}{W_{tot}} \quad (4.23)$$

the switching losses are equal to the other converter losses. Below this critical frequency, the total loss is dominated by the conduction and fixed loss, and hence the total loss and converter efficiency are not strong functions of switching frequency. Above the critical frequency, the switching loss dominates the total loss, and the converter efficiency decreases rapidly with increasing switching frequency. Typical dependence of the full-load converter efficiency on switching frequency is plotted in Fig. 4.55, for an arbitrary choice of parameter values. The critical frequency f_{crit} can be taken as a rough upper limit on the switching frequency of a practical converter.

4.4 SUMMARY OF KEY POINTS

1. How an SPST ideal switch can be realized using semiconductor devices depends on the polarity of the voltage that the devices must block in the off state, and on the polarity of the current which the devices must conduct in the on state.
2. Single-quadrant SPST switches can be realized using a single transistor or a single diode, depending on the relative polarities of the off state voltage and on state current.

3. Two-quadrant SPST switches can be realized using a transistor and diode, connected in series (bidirectional-voltage) or in antiparallel (bidirectional-current). Several four-quadrant schemes are also listed here.
4. A “synchronous rectifier” is a MOSFET connected to conduct reverse current, with gate drive control as necessary. This device can be used where a diode would otherwise be required. If a MOSFET with sufficiently low R_{on} is used, reduced conduction loss is obtained.
5. Majority carrier devices, including the MOSFET and Schottky diode, exhibit very fast switching times, controlled essentially by the charging of the device capacitances. However, the forward voltage drops of these devices increases quickly with increasing breakdown voltage.
6. Minority carrier devices, including the BJT, IGBT, and thyristor family, can exhibit high breakdown voltages with relatively low forward voltage drop. However, the switching times of these devices are longer, and are controlled by the times needed to insert or remove stored minority charge.
7. Energy is lost during switching transitions, owing to a variety of mechanisms. The resulting average power loss, or switching loss, is equal to this energy loss multiplied by the switching frequency. Switching loss imposes an upper limit on the switching frequencies of practical converters.
8. The diode and inductor present a “clamped inductive load” to the transistor. When a transistor drives such a load, it experiences high instantaneous power loss during the switching transitions. An example where this leads to significant switching loss is the IGBT and the “current tail” observed during its turn-off transition.
9. Other significant sources of switching loss include diode stored charge and energy stored in certain parasitic capacitances and inductances. Parasitic ringing also indicates the presence of switching loss.

REFERENCES

- [1] R. D. MIDDLEBROOK, S. ČUK, and W. BEHEN, “A New Battery Charger/Discharger Converter,” *IEEE Power Electronics Specialists Conference*, 1978 Record, pp. 251-255, June 1978.
- [2] H. MATSUO and F. KUROKAWA, “New Solar Cell Power Supply System Using a Boost Type Bidirectional Dc-Dc Converter,” *IEEE Power Electronics Specialists Conference*, 1982 Record, pp. 14-19, June 1982.
- [3] M. VENTURINI, “A New Sine-Wave-In Sine-Wave-Out Conversion Technique Eliminates Reactive Elements,” *Proceedings Seventh International Solid-State Power Conversion Conference (Powercon 7)*, pp. E3.1-E3.13, 1980.
- [4] K. D. T. NGO, S. ČUK, and R. D. MIDDLEBROOK, “A New Flyback Dc-to-Three-Phase Converter with Sinusoidal Outputs,” *IEEE Power Electronics Specialists Conference*, 1983 Record, pp. 377-388.
- [5] S. ČUK, “Basics of Switched-Mode Power Conversion: Topologies, Magnetics, and Control,” *Advances in Switched-Mode Power Conversion*, Vol. II, Irvine CA: Teslaco, pp. 279-310, 1983.
- [6] L. GYUGI and B. PELLÉ, *Static Power Frequency Changers: Theory, Performance, and Applications*, New York: Wiley-Interscience, 1976.
- [7] R. S. KAGAN and M. CHI, “Improving Power Supply Efficiency with MOSFET Synchronous Rectifiers,” *Proceedings Ninth International Solid-State Power Conversion Conference (Powercon 9)*, pp. D4.1-D4.9, July 1982.
- [8] R. BLANCHARD and P. E. THIBODEAU, “The Design of a High Efficiency, Low Voltage Power Supply Using MOSFET Synchronous Rectification and Current Mode Control,” *IEEE Power Electronics Special-*