

Introduction to Power MOSFETs

What is a Power MOSFET?

We all know how to use a diode to implement a switch. But we can only switch with it, not gradually control the signal flow. Furthermore, a diode acts as a switch depending on the direction of signal flow; we can't program it to pass or block a signal. For such applications involving either "flow control" or programmable on/off switching we need a 3-terminal device...and Bardeen & Brattain heard us and "invented" (almost by accident, like many other great discoveries!) the bipolar transistor.

Structurally it is implemented with only two junctions back-to-back (no big deal; we were probably making common cathodes - same structure - long before Bardeen). But functionally it is a totally different device which acts like a "faucet" controlling the flow of emitter current - and the "hand" manipulating the faucet is the base current. A bipolar transistor is therefore a current-controlled device.

The Field Effect Transistor (FET), although structurally different, provides the same "faucet" function. The difference: the FET is voltage-controlled; one doesn't need base current but voltage to exercise flow control. The bipolar transistor was born in 1947; the FET (at least the concept) came soon after, in 1948 from another pair of illustrious parents: Shockley and Pearson. The terminals are called DRAIN instead of COLLECTOR, GATE instead of BASE and SOURCE instead of EMITTER to differentiate it from his older bipolar "cousin". The FET

comes in two major variants, optimized for different types of applications: the JFET (junction FET) used in small-signal processing and the MOSFET (metal-oxide-semiconductor FET) mainly used in linear or switching power applications.

Why Did They Need to Invent the Power MOSFET?

When scaled-up for power applications the bipolar transistor starts showing some annoying limitations. Sure, you can still find it in your washing machine, in your air conditioner and refrigerator but these are "low" power applications for us, the average consumer, who can tolerate a certain degree of inefficiency in his appliances. Transistors are still used in some UPSs, motor controls or welding robots but their usage is practically limited to less than 10kHz and they are rapidly disappearing from the "technology edge" applications where overall efficiency is the "key" parameter (SMPSS, sophisticated motor controls, converters, to name a few).

Being a bipolar device, the transistor relies on the minority carriers injected in the base to "defeat" recombination and be re-injected in the collector. In order to sustain a large collector current we want to inject many of them in the base from the emitter side and, if possible, recuperate all of them at the base/collector boundary (meaning that recombination in the base should be kept at a minimum). But this means that when we want the transistor switched off, there will be a considerable amount of minority carriers in the base with a low recombination factor to be taken care of before the switch can close - in other words the stored charge problem associated with all minority carrier devices limiting the maximum

operating speed. The major advantage of the FET now comes to light: being a majority carrier device there is no stored minority charge therefore it can work at much higher frequencies. The switching delays characteristic to mosfets are rather a consequence of the charging and discharging of the parasitic capacitors. One may say: I see the need for a fast switching mosfet in high frequency applications but why should I use such device in my relatively slow switching circuitry? The answer is straightforward: improved efficiency. The device sees both high current and high voltage during the interval in which switching occurs; a faster device will therefore experience proportionally less energy loss. In many applications this advantage alone more than compensates for the slightly higher conduction losses associated with higher voltage mosfets: switch-mode power supplies (smps) operating beyond 150 kHz would not be possible without them.

The bipolar transistor is current driven; in fact the more current we want to drive, the more current we need to supply to the base because the gain (ratio of the collector and base currents) drops significantly as the collector current (IC) increases. One consequence is that the bipolar transistor starts dissipating significant control power, reducing the overall efficiency of the circuitry. To make things worse this drawback is accentuated at higher operating temperatures. Another consequence is the need for rather complicated base drive circuitry capable of fast current sourcing and sinking. Not the (MOS)FET; this device has practically zero current consumption in the gate; even at 125°C the typical gate current stays below 100 nA. Once the parasitic capacitances are charged, only the very low leakage currents have to be provided by the drivers.

Add to this the circuit simplicity resulting from driving a device with voltage rather than current and you'll spot another reason why the (MOS)FET is so appealing to the design engineer.

Another major advantage is the nonexistence of a secondary breakdown mechanism. Try to block a lot of power with a bipolar transistor; local defects unavoidable in any semiconductor structure will act to concentrate the current, the result will be localized heating of the silicon. Since the temperature coefficient of the resistivity is negative the local defect will act as low resistance current path, directing even more current into it, self heating even more until non-reversible destruction occurs. The MOSFET has a positive resistivity thermal coefficient. On one hand this can be perceived as the disadvantage of an increased $R_{DS(on)}$ at elevated temperatures - this important parameter roughly doubles between 25 C and 125 C due to carrier mobility reduction. On the other hand this same phenomenon brings a significant advantage: any defect trying to act as described above would actually divert current from it - one would have "cooling-spots" instead of the "hot-spots" characteristic to bipolar devices!

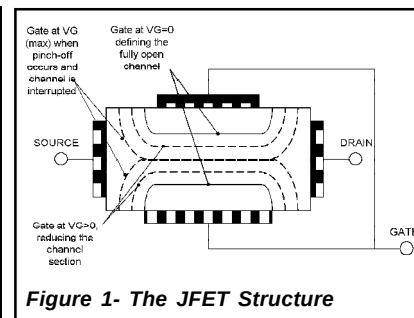
An equally important consequence of this self-cooling mechanism is the ease of paralleling MOSFETs to boost-up the current capability of a device. Bipolar transistors are very

sensitive to paralleling; precautions (emitter ballasting resistors, fast response current-sensing feedback loops) have to be taken for equal sharing of currents, otherwise the device with the lowest saturation voltage would divert most of the current, overheating as described above and ultimately resulting in a short-circuit. Not the MOSFET; they can be paralleled with no other precautions than design insured circuit symmetry and balancing the gates so they open equally allowing the same amount of current in all transistors. The extra bonus is that even if the gates are not balanced and the channels have different degrees of opening, this would still result in a steady state condition with some drain currents being slightly larger than others.

A useful feature appealing to the design engineer is a consequence of the unique structure of the mosfet (see #3 for a more detailed description): the "parasitic" body-diode formed between source and drain. Although it is not optimized for fast switching or low conduction loss, it acts as a clamping diode in inductive load switching applications at no extra cost.

MOSFET Structure

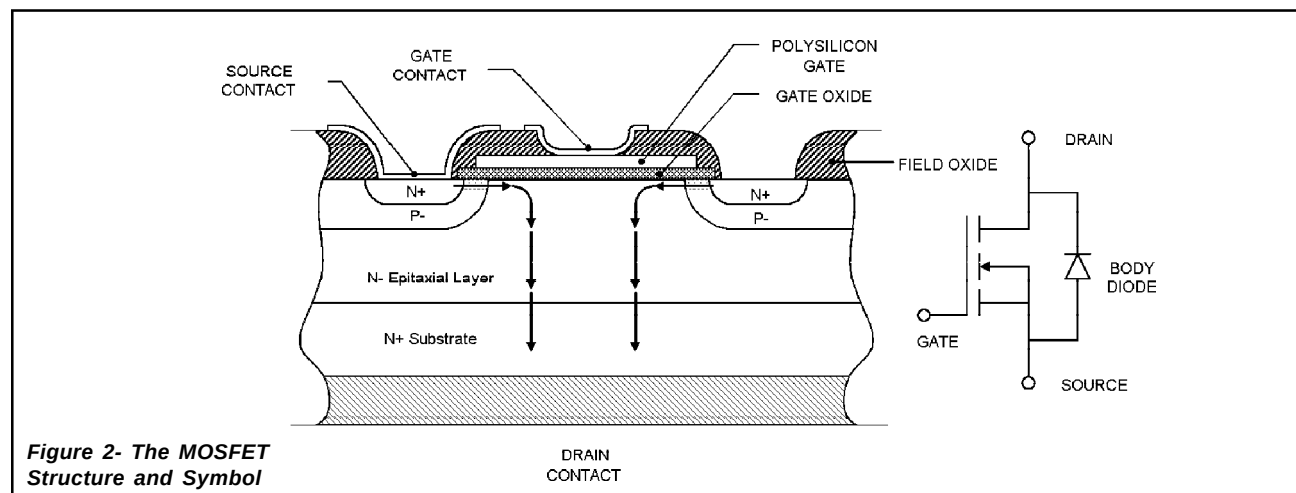
The basic idea of a JFET (fig.1) is to control the current flowing from source to drain by modulating (pinching) the cross-sectional area of the Drain-Source channel. This is achieved using a reverse biased



junction as gate; its (reverse) voltage modulates the depleted region consequently pinching the channel and increasing its resistance by reducing its section. With no voltage applied to the gate the channel resistance is at its lowest value and maximum drain current flows through the device. As gate voltage is increased the two depleted region fronts advance, reducing the drain current by increasing the channel resistance until total pinch-off occurs when the two fronts meet. The channel is now severed and no more current flows.

The MOSFET uses a different type of gate mechanism exploiting the properties of the MOS capacitor. By varying the value and the polarity of the bias applied to the top electrode of a MOS structure one can drive the silicon underneath it into enhancement all the way to inversion. Fig.2 shows the simplified structure of an N-channel MOSFET. It is called a vertical, double diffused structure and starts with a heavily concentrated n substrate in order to minimize the

(continued on Page 14)



MicroNotes

Series 901

(continued from Page 13)

bulk portion of the channel resistance. An n- epi layer is grown on it and two successive diffusions are made, a p- zone in which proper bias will generate the channel and an n+ into it defining the source. Next the thin, high quality gate oxide is grown followed by the phosphorous-doped polysilicon thus forming the gate. Contact windows are opened on top defining the source and the gate terminals while the whole bottom of the wafer makes the drain contact. With no bias on the gate the n+ source and the n drain are separated by the p zone and no current flows (transistor is turned-off). If a positive bias is applied to the gate the minority carriers in the p zone (electrons) are attracted to the surface underneath the gate plate. As the bias increases more electrons are being confined to this small space, the local "minority" concentration becomes larger than the hole (p) concentration and "inversion" occurs (meaning that the material immediately under the gate turns from p- to n-type). Now an n "channel" is formed in the p material right under the gate structure connecting the source to the drain; current can now flow. Like in the case of the JFET (although the physical phenomenon is different) the gate (by means of its voltage bias) controls the flow of current between the source and the drain.

There are many MOSFET manufacturers and almost everyone has his own process optimization and his tradename. International Rectifier pioneered the HEXFET, Motorola builds TMOS, Ixys fabricates HiPerFETs and MegaMOS, Siemens has the SiPMOS family of power transistors and Advanced Power Technology the Power MOS IV, to name a few. Whether the process is called VMOS, TMOS or DMOS it has a horizontal gate structure and vertical current flow past the gate.

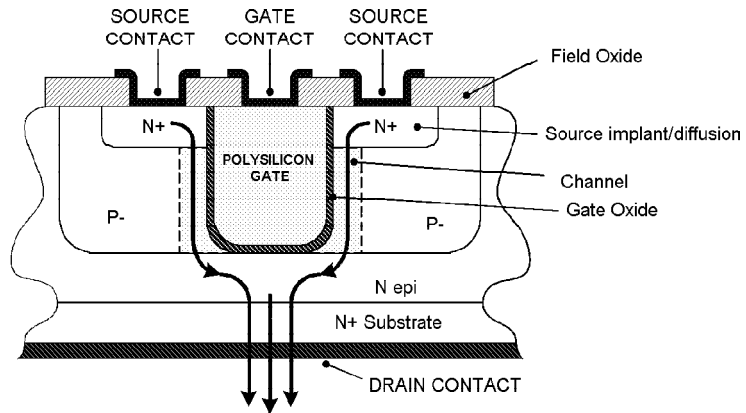


Figure 1- The Trench MOS Structure

The power MOSFET is nothing else but a structure containing a multitude of "cells" like the one described in fig.2 connected in parallel. And like any paralleling of identical resistors the equivalent resistance is 1/n-th of the single cell's RDS(on). The larger the die, the lower its "on" resistance but at the same time the larger its parasitic capacitances and therefore the poorer its switching performance.

If everything is so exactly proportional and predictable, is there any way for improvement? Yes, and the idea is to minimize (scale-down) the size of the basic cell - this way more cells can fit in the same footprint driving the RDS(on) down while maintaining the capacitances. The continuous technological improvement and refinement of the wafer fabrication processes (finer line lithography, better controlled implants, etc.) is responsible for the successively improved MOSFET product generations.

But continuous striving for better processing technology is not the only way to improve the MOSFET; conceptual design changes can result in major performance increase. Such a breakthrough was announced by Philips this November: the development of the TrenchMOS process. The gate structure, instead of being parallel to the die surface, is now built in a trench, perpendicular to the surface, taking much less space and making the current flow truly vertical (see fig.3). The Philips transistors offer 50% size reduction

for the same RDS(on) or a 35% size reduction maintaining the same current handling capability.

Conclusion

This MicroNote helped us refresh our knowledge; we compared the MOSFET to its more known and more used relative, the bipolar transistor. We saw the major advantages the MOSFET has over the BJT. We are also now aware of some trade-offs. The most important conclusion is that overall circuit efficiency is application-specific; one has to closely estimate the balance of conduction and switching losses under all operating conditions and then decide on the device to be used: a regular bipolar or a MOSFET - or maybe an IGBT?

Following MicroNotes will further detail the MOSFET characteristics as well as the IGBT.

For more information on Microsemi's MOSFET capabilities including Power Surface Mount and TO-250 series packaging capabilities, contact Microsemi Santa Ana or Microsemi's Sertech Labs for specific product information.