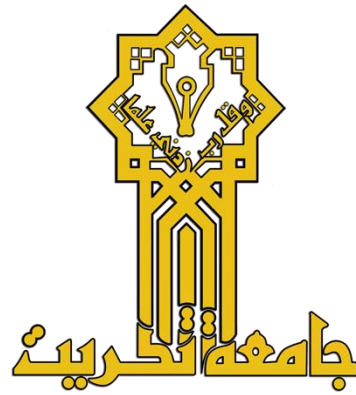




DIGITAL DEVICES AND LOGIC FAMILIES

ELECTRICAL DEPARTMENT

FOURTH STAGE

INSTRUCTOR: DR. ANAS MQDAD

LECTURE FIVE

DATA STORAGE

OUTLINES

The Random-Access Memory (RAM)

A RAM is a read/write memory in which data can be written into or read from any selected address in any sequence. When a data unit is written into a given address in the RAM, the data unit previously stored at that address is replaced by the new data unit. When a data unit is read from a given address in the RAM, the data unit remains stored and is not erased by the read operation. This nondestructive read operation can be viewed as copying the content of an address while leaving the content intact. A RAM is typically used for short-term data storage because it cannot retain stored data when power is turned off.

The RAM Family

- ❑ The two major categories of RAM are the static RAM (SRAM) and the dynamic RAM (DRAM). **SRAMs** generally use latches as storage elements and can therefore store data indefinitely as long as dc power is applied.
- ❑ **DRAMs** use capacitors as storage elements and cannot retain data very long without the capacitors being recharged by a process called **refreshing**. Both SRAMs and DRAMs will lose stored data when dc power is removed and, therefore, are classified as volatile memories.

- ❑ Data can be read much faster from SRAMs than from DRAMs. However, DRAMs can store much more data than SRAMs for a given physical size and cost because the DRAM cell is much simpler and more cells can be crammed into a given chip area than in the SRAM.

- ❑ The basic types of SRAM are the asynchronous SRAM and the synchronous SRAM with a burst feature. The basic types of DRAM are the Fast Page Mode DRAM (FPM DRAM), the Extended Data Out DRAM (EDO DRAM), the Burst EDO DRAM (BEDO DRAM), and the synchronous DRAM (SDRAM). These are shown in Figure 1

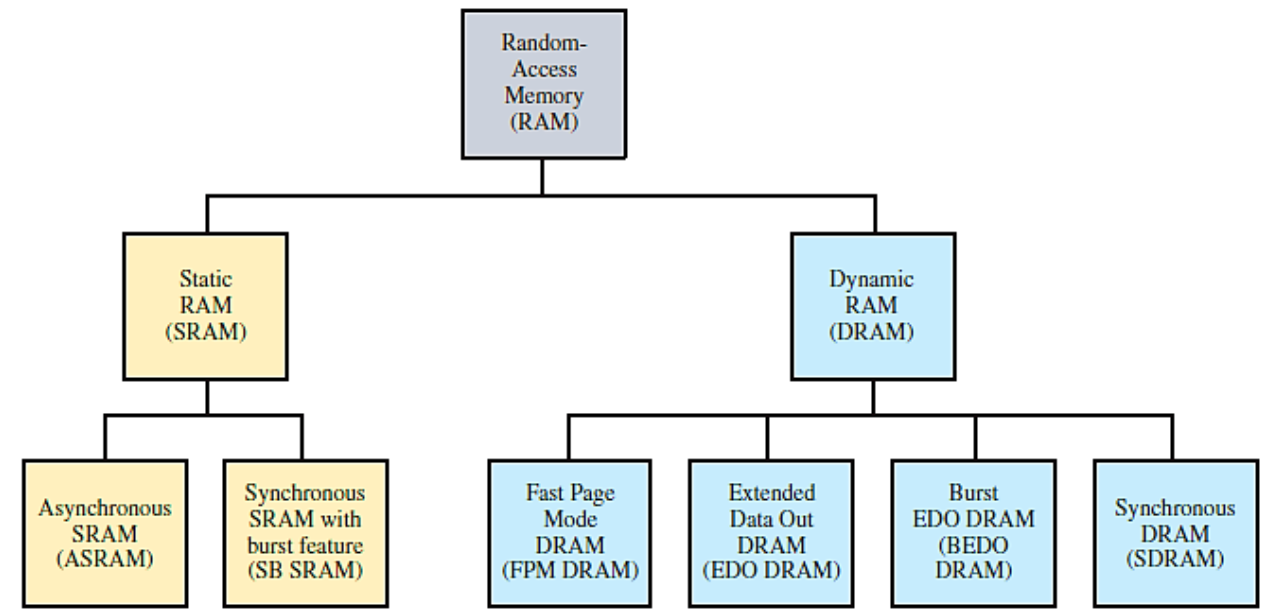


FIGURE 1 The RAM family.

Static RAMs (SRAMs)

Memory Cell

- ❑ All SRAMs are characterized by latch memory cells. As long as dc power is applied to a **static memory** cell, it can retain a 1 or 0 state indefinitely. If power is removed, the stored data bit is lost.
- ❑ Figure 2 shows a basic SRAM latch memory cell. The cell is selected by an active level on the Select line and a data bit (1 or 0) is written into the cell by placing it on the Data in line. A data bit is read by taking it off the Data out line.

Static Memory Cell Array

- ❑ The memory cells in a SRAM are organized in rows and columns, as illustrated in Figure 3 for the case of an $n \times 4$ array. All the cells in a row share the same Row Select line. Each set of Data in and Data out lines go to each cell in a given column and are connected to a single data line that serves as both an input and output (Data I/O) through the data input and data output buffers.

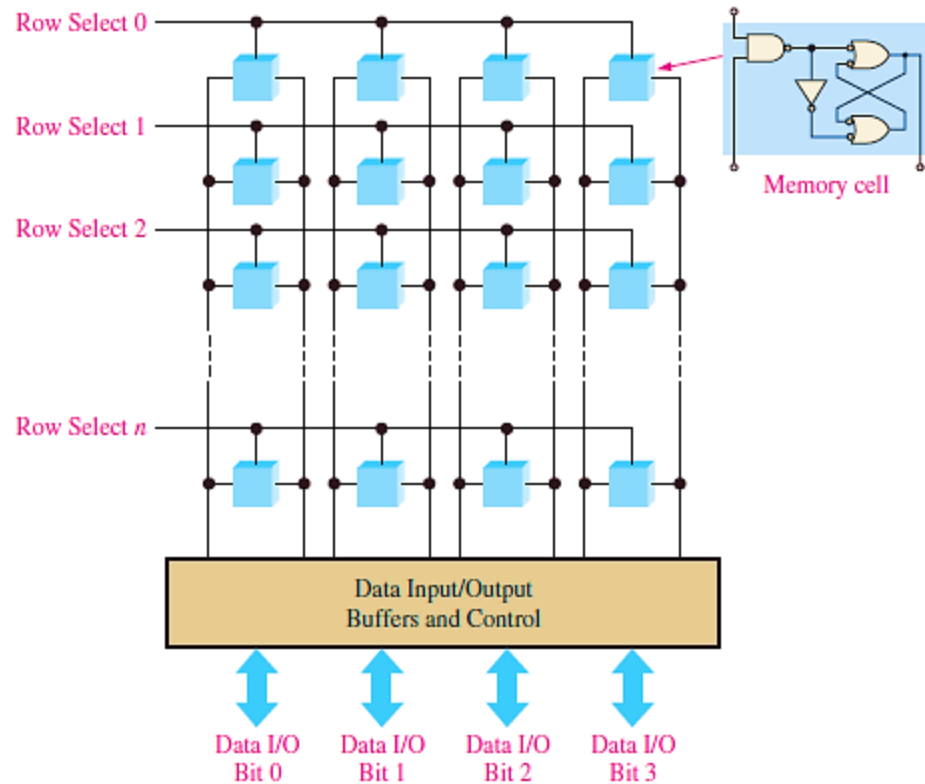


FIGURE 3 Basic SRAM array.

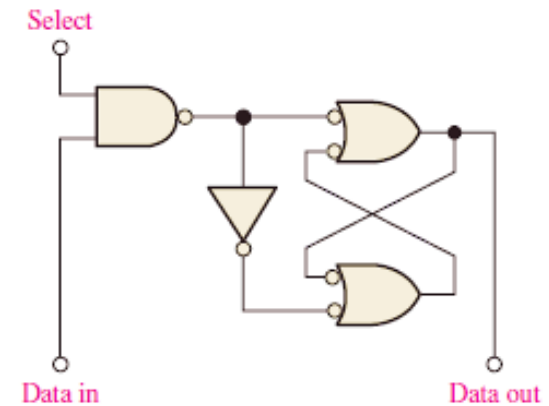


FIGURE 2 A typical SRAM latch memory cell.

- ❑ To write a data unit, in this case a nibble (4 bits), into a given row of cells in the memory array, the Row Select line is taken to its active state and four data bits are placed on the Data I/O lines. The Write line is then taken to its active state, which causes each data bit to be stored in a selected cell in the associated column. To read a data unit, the Read line is taken to its active state, which causes the four data bits stored in the selected row to appear on the Data I/O lines.

Basic Asynchronous SRAM Organization

- ❑ An asynchronous SRAM is one in which the operation is not synchronized with a system clock. To illustrate the general organization of a SRAM, a 32k X 8 bit memory is used. A logic symbol for this memory is shown in Figure 4.
- ❑ In the READ mode, the eight data bits that are stored in a selected address appear on the data output lines. In the WRITE mode, the eight data bits that are applied to the data input lines are stored at a selected address. The data input and data output lines (I/O0 through I/O7) share the same lines. During READ, they act as output lines (O0 through O7) and during WRITE they act as input lines (I0 through I7).

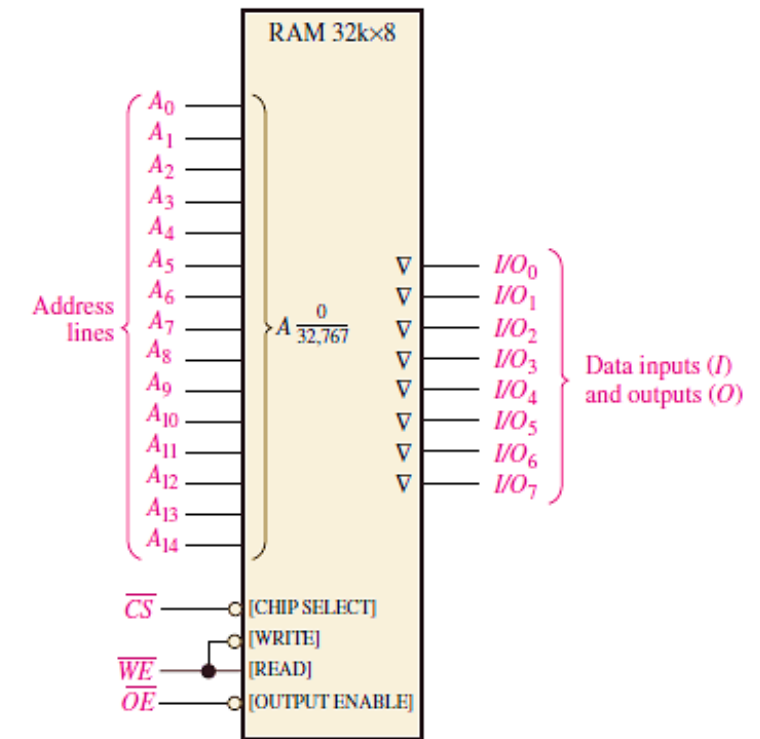


FIGURE 4 Logic diagram for an asynchronous 32k * 8 SRAM.

Tri-state Outputs and Buses

- ❑ Tri-state buffers in a memory allow the data lines to act as either input or output lines and connect the memory to the data bus in a computer. These buffers have three output states:
- ❑ HIGH (1), LOW (0), and HIGH-Z (open). Tri-state outputs are indicated on logic symbols by a small inverted triangle (◡), as shown in Figure 4, and are used for compatibility with bus structures such as those found in microprocessor-based systems.
- ❑ Physically, a **bus** is one or more conductive paths that serve to interconnect two or more functional components of a system or several diverse systems. Electrically, a bus is a collection of specified voltage levels and/or current levels and signals that allow various devices to communicate and work properly together.
- ❑ A microprocessor is connected to memories and input/output devices by certain bus structures. An address bus allows the microprocessor to address the memories, and a data bus provides for transfer of data between the microprocessor, the memories, and the input/ output devices such as monitors, printers, keyboards, and modems. A control bus allows the microprocessor to control data transfers and timing for the various components.

Memory Array

- ❑ SRAM chips can be organized in single bits, nibbles (4 bits), bytes (8 bits), or multiple bytes (words with 16, 24, 32 bits, etc.). Figure 5 shows the organization of a small 32k * 8 SRAM. The memory cell array is arranged in 256 rows and 128 columns, each with 8 bits, as shown in part (a).
- ❑ There are actually $2^{15} = 32,768$ addresses and each address contains 8 bits. The capacity of this example memory is 32,768 bytes (typically expressed as 32 kB). Although small by today's standards, this memory serves to introduce the basic concepts.
- ❑ The SRAM in Figure 5(b) works as follows. First, the chip select, CS, must be LOW for the memory to operate. (Other terms for chip select are enable or chip enable.) Eight of the fifteen address lines are decoded by the row decoder to select one of the 256 rows. Seven of the fifteen address lines are decoded by the column decoder to select one of the 128 8-bit columns.

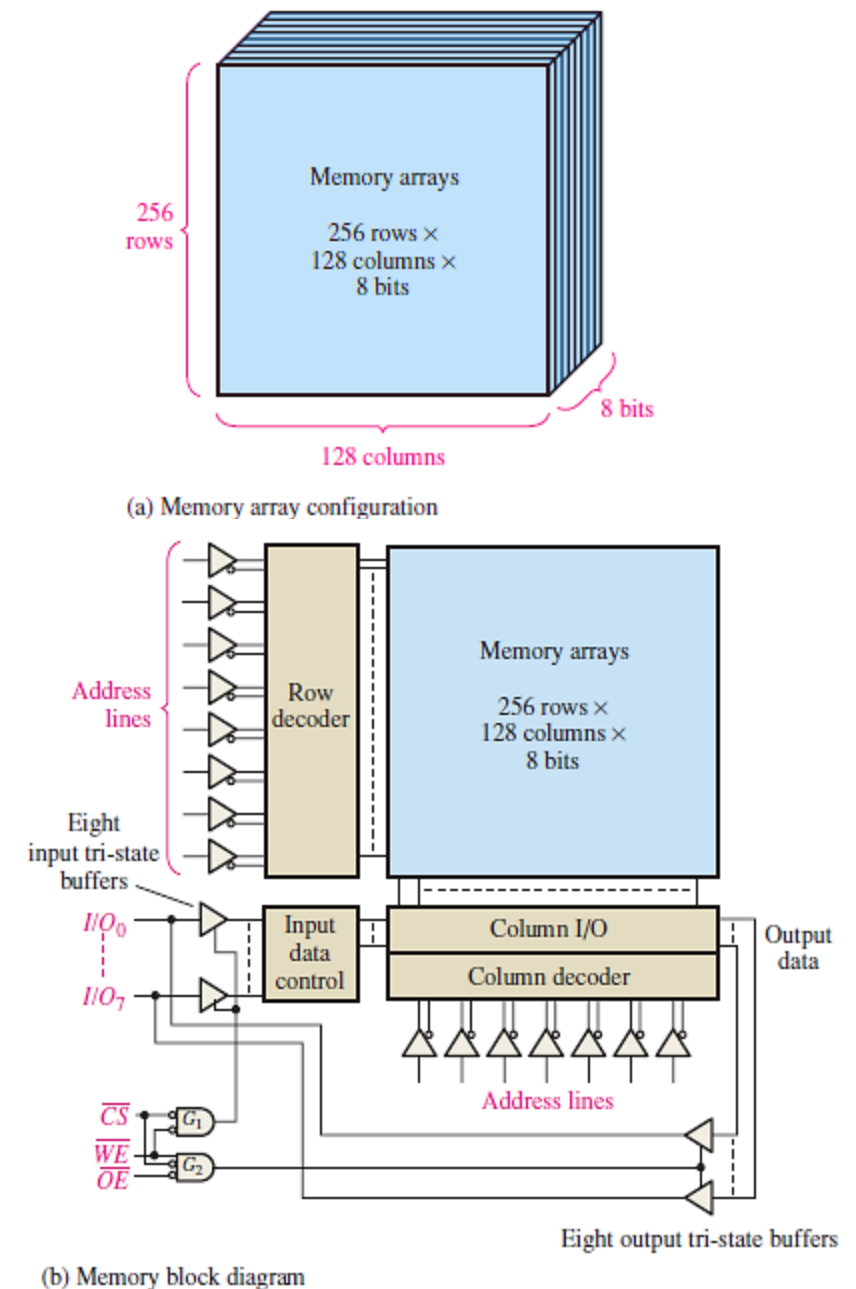


FIGURE 5 Basic organization of an asynchronous 32k * 8 SRAM.

Read

- ❑ In the READ mode, the write enable input, $\overline{WE}$, is HIGH and the output enable, $\overline{OE}$, is LOW. The input tri-state buffers are disabled by gate G_1 , and the column output tri-state buffers are enabled by gate G_2 .
- ❑ Therefore, the eight data bits from the selected address are routed through the column I/O to the data lines (I/O_0 through I/O_7), which are acting as data output lines.

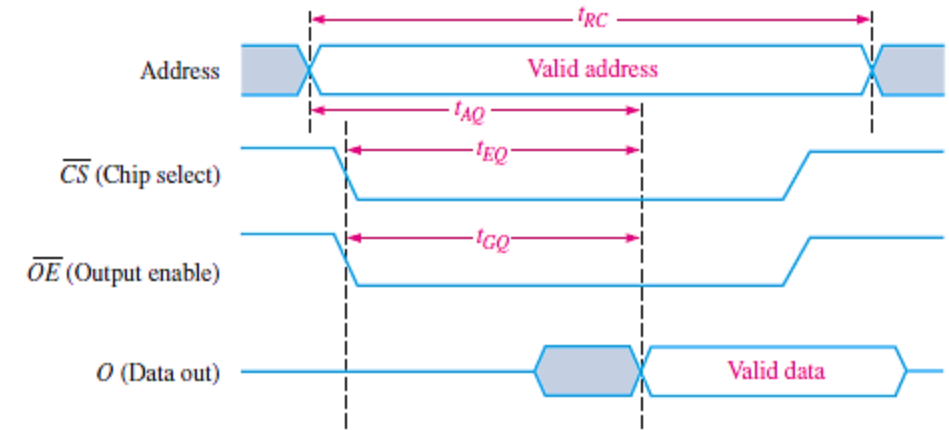
Write

- ❑ In the WRITE mode, $\overline{WE}$ is LOW and $\overline{OE}$ is HIGH. The input tri-state buffers are enabled by gate G_1 , and the output tri-state buffers are disabled by gate G_2 . Therefore, the eight input data bits on the data lines are routed through the input data control and the column I/O to the selected address and stored.

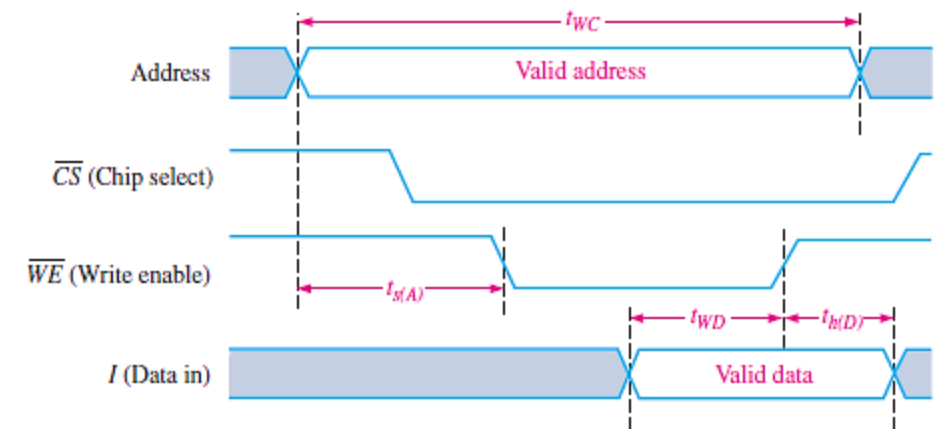
Read and Write Cycles

□ Figure 6 shows typical timing diagrams for a memory read cycle and a write cycle. For the read cycle shown in part (a), a valid address code is applied to the address lines for a specified time interval called the *read cycle time*, t_{RC} . Next, the chip select ($\overline{CS}$) and the output enable ($\overline{OE}$) inputs go LOW.

□ One time interval after the $\overline{OE}$ input goes LOW, a valid data byte from the selected address appears on the data lines. This time interval is called the output enable access time, t_{GQ} . Two other access times for the read cycle are the address access time, t_{AQ} , measured from the beginning of a valid address to the appearance of valid data on the data lines and the chip enable access time, t_{EQ} , measured from the HIGH-to-LOW transition of $\overline{CS}$ (to the appearance of valid data on the data lines).



(a) Read cycle ($\overline{WE}$ HIGH)



(b) Write cycle ($\overline{WE}$ LOW)

FIGURE 6 Timing diagrams for typical read and write cycles for the SRAM in Figure 5.

- ❑ During each read cycle, one unit of data, a byte in this case, is read from the memory. For the write cycle shown in Figure 6 (b), a valid address code is applied to the address lines for a specified time interval called the *write cycle time*, t_{wc} . Next, the chip select ($\overline{CS}$) and the write enable ($\overline{WE}$) inputs go LOW.
- ❑ The required time interval from the beginning of a valid address until the $\overline{WE}$ input goes LOW is called the *address setup time*, $t_{s(A)}$. The time that the $\overline{WE}$ input must be LOW is the write pulse width. The time that the input $\overline{WE}$ must remain LOW after valid data are applied to the data inputs is designated t_{wd} ; the time that the valid input data must remain on the data lines after the $\overline{WE}$ input goes HIGH is the *data hold time*, $t_{h(D)}$. During each write cycle, one unit of data is written into the memory.

Synchronous SRAM with Burst Feature

- ❑ Unlike the asynchronous SRAM, a synchronous SRAM is synchronized with the system clock. For example, in a computer system, the synchronous SRAM operates with the same clock signal that operates the microprocessor so that the microprocessor and memory are synchronized for faster operation.
- ❑ The fundamental concept of the synchronous feature of a SRAM can be shown with Figure 7, which is a simplified block diagram of a 32k * 8 memory for purposes of illustration.

- ❑ The synchronous SRAM is similar to the asynchronous SRAM in terms of the memory array, address decoder, and read/write and enable inputs. The basic difference is that the synchronous SRAM uses clocked registers to synchronize all inputs with the system clock. The address, the read/write input, the chip enable, and the input data are all latched into their respective registers on an active clock pulse edge. Once this information is latched, the memory operation is in sync with the clock.
- ❑ For the purpose of simplification, a notation for multiple parallel lines or bus lines is introduced in Figure 7, as an alternative to drawing each line separately. A set of parallel lines can be indicated by a single heavy line with a slash and the number of separate lines in the set. For example, the following notation represents a set of 8 parallel lines:



- ❑ The address bits A_0 through A_{14} are latched into the Address register on the positive edge of a clock pulse. On the same clock pulse, the state of the write enable ($\overline{WE}$) line and chip select ($\overline{CS}$) are latched into the Write register and the Enable register respectively.

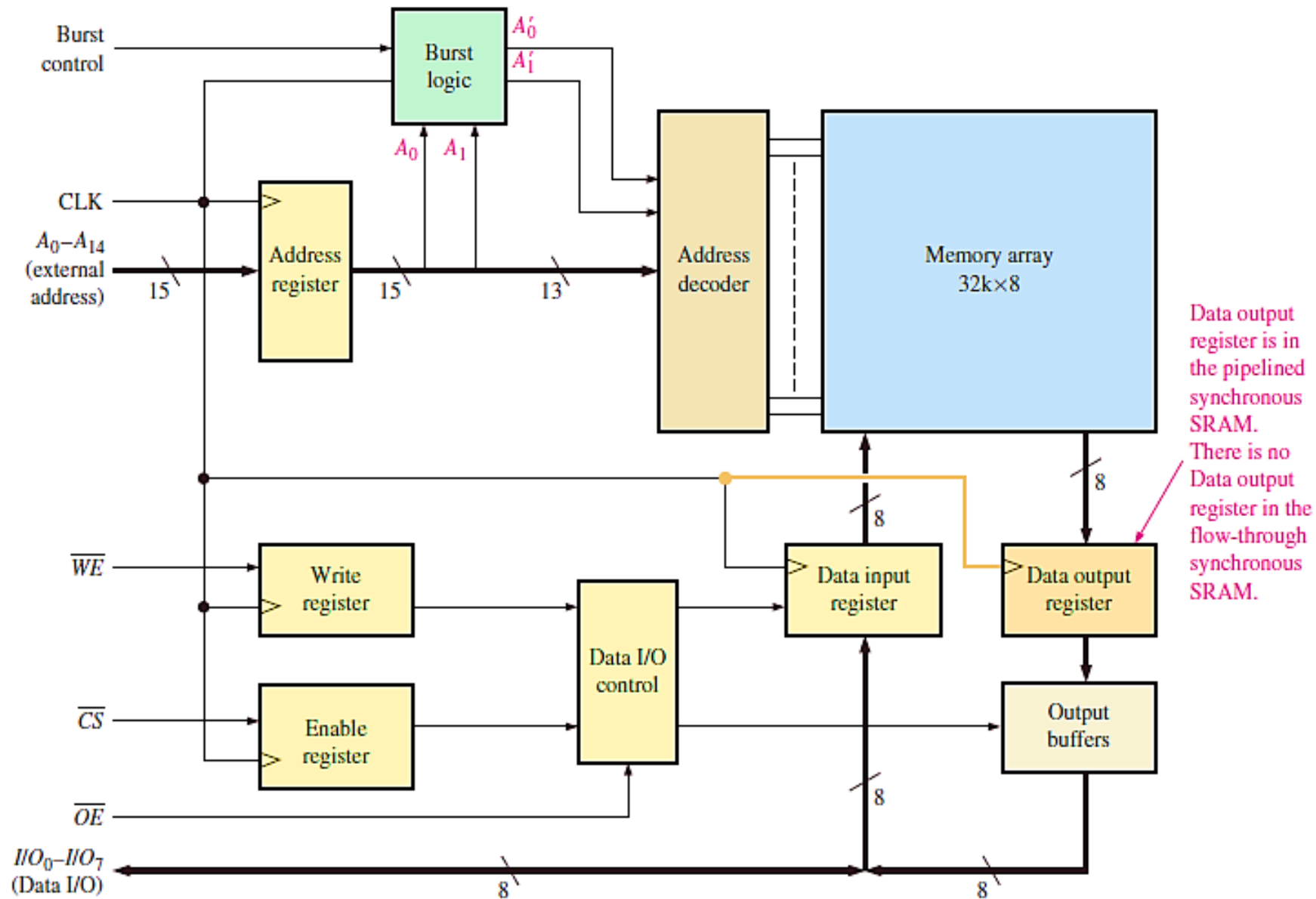


FIGURE 7 A basic block diagram of a synchronous SRAM with burst feature.

- ❑ These are one-bit registers or simply flip-flops. Also, on the same clock pulse the input data are latched into the Data input register for a Write operation, and data in a selected memory address are latched into the Data output register for a Read operation, as determined by the Data I/O control based on inputs from the Write register, Enable register, and the Output enable ($\overline{OE}$).

- ❑ Two basic types of synchronous SRAM are the *flow through* and the *pipelined*. The flow-through synchronous SRAM does not have a Data output register, so the output data flow asynchronously to the data I/O lines through the output buffers. The **pipelined** synchronous SRAM has a Data output register, as shown in Figure 7, so the output data are synchronously placed on the data I/O lines.

The Burst Feature

❑ As shown in Figure 7, synchronous SRAMs normally have an address burst feature, which allows the memory to read or write up to four sequential locations using a single address. When an external address is latched in the address register, the two lowest-order address bits, A_0 and A_1 , are applied to the burst logic. This produces a sequence of four internal addresses by adding 00, 01, 10, and 11 to the two lowest-order address bits on successive clock pulses. The sequence always begins with the base address, which is the external address held in the address register.

❑ The address burst logic in a typical synchronous SRAM consists of a binary counter and exclusive-OR gates, as shown in Figure 8. For 2-bit burst logic, the internal burst address sequence is formed by the base address bits A_2 – A_{14} plus the two burst address bits A'_1 and A'_0 .

❑ To begin the burst sequence, the counter is in its 00 state and the two lowest-order address bits are applied to the inputs of the XOR gates. Assuming that A_0 and A_1 are both 0, the internal address sequence in terms of its two lowest-order bits is 00, 01, 10, and 11.

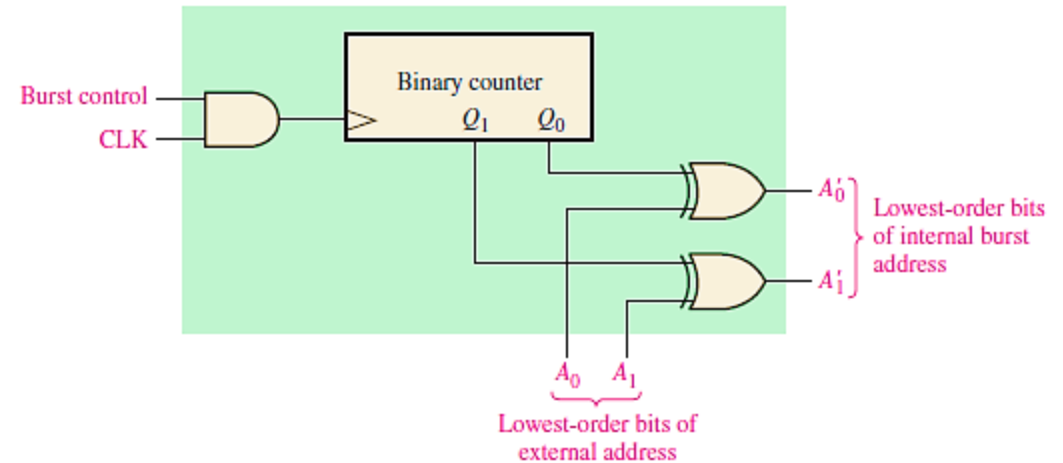


FIGURE 8 Address burst logic.

Cache Memory

- ❑ One of the major applications of SRAMs is in cache memories in computers. **Cache memory** is a relatively small, high-speed memory that stores the most recently used instructions or data from the larger but slower main memory. Cache memory can also use dynamic RAM (DRAM), which is discussed next.
- ❑ Typically, SRAM is several times faster than DRAM. Overall, a cache memory gets stored information to the microprocessor much faster than if only high-capacity DRAM is used. Cache memory is basically a cost-effective method of improving system performance without having to resort to the expense of making all of the memory faster.
- ❑ The concept of cache memory is based on the idea that computer programs tend to get instructions or data from one area of main memory before moving to another area. Basically, the cache controller “guesses” which area of the slow dynamic memory the CPU will need next and moves it to the cache memory so that it is ready when needed. If the cache controller guesses right, the data are immediately available to the microprocessor. If the cache controller guesses wrong, the CPU must go to the main memory and wait much longer for the correct instructions or data. Fortunately, the cache controller is right most of the time.

Cache Analogy

There are many analogies that can be used to describe a cache memory, but comparing it to a home refrigerator is perhaps the most effective. A home refrigerator can be thought of as a “cache” for certain food items while the supermarket is the main memory where all foods are kept. Each time you want something to eat or drink, you can go to the refrigerator (cache) first to see if the item you want is there. If it is, you save a lot of time. If it is not there, then you have to spend extra time to get it from the supermarket (main memory).

L1 and L2 Caches

A first-level cache (L1 cache) is usually integrated into the processor chip and has a very limited storage capacity. L1 cache is also known as *primary cache*. A second-level cache (L2 cache) may also be integrated into the processor or as a separate memory chip or set of chips external to the processor; it usually has a larger storage capacity than an L1 cache. L2 cache is also known as *secondary cache*. Some systems may have higher-level caches (L3, L4, etc.), but L1 and L2 are the most common. Also, some systems use a disk cache to enhance the performance of the hard disk because DRAM, although much slower than SRAM, is much faster than the hard disk drive. Figure 9 illustrates L1 and L2 cache memories in a computer system.

Dynamic RAM (DRAM) Memory Cells

- ❑ **Dynamic memory** cells store a data bit in a small capacitor rather than in a latch. The advantage of this type of cell is that it is very simple, thus allowing very large memory arrays to be constructed on a chip at a lower cost per bit. The disadvantage is that the storage capacitor cannot hold its charge over an extended period of time and will lose the stored data bit unless its charge is refreshed periodically
- ❑ To refresh requires additional memory circuitry and complicates the operation of the DRAM. Figure 10 shows a typical DRAM cell consisting of a single MOS transistor (MOSFET) and a capacitor. In this type of cell, the transistor acts as a switch. The basic simplified operation is illustrated in Figure 11 and is as follows.

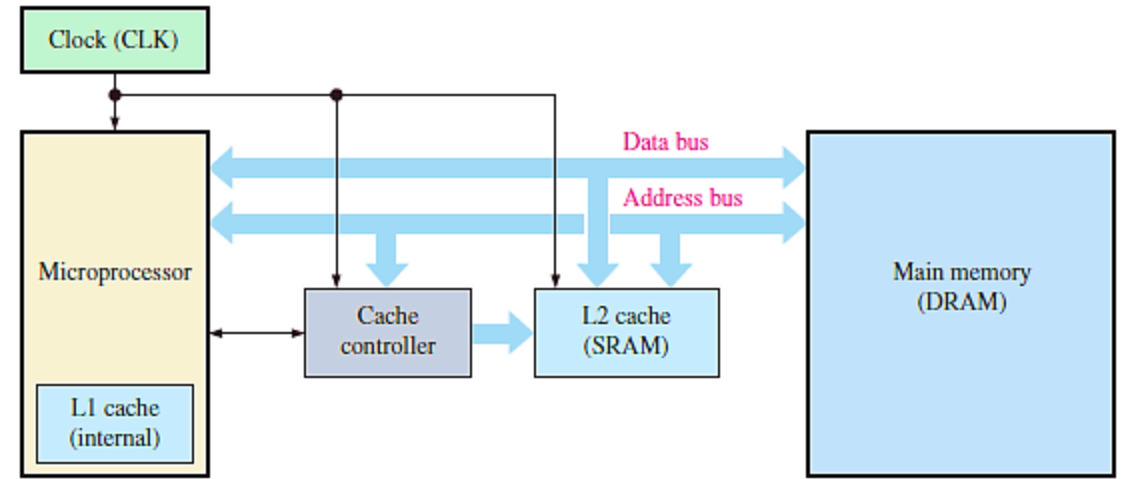


FIGURE 9 Block diagram showing L1 and L2 cache memories in a computer system.

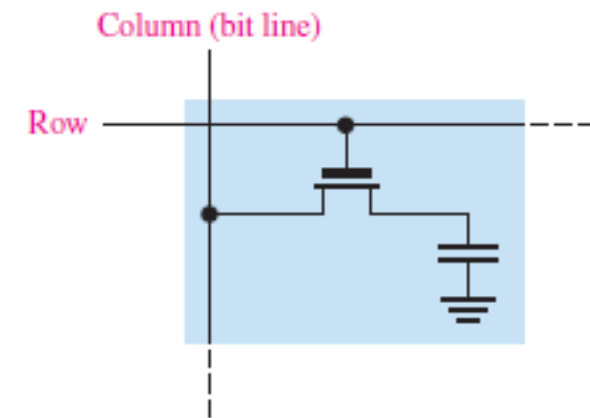
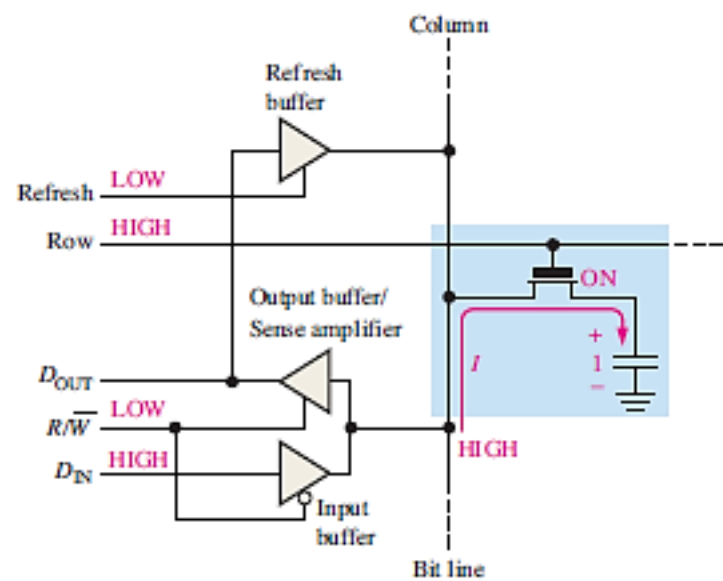
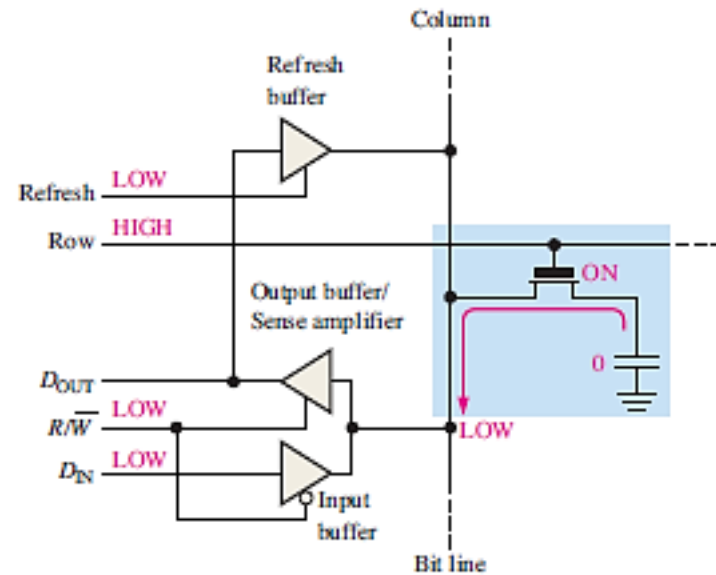


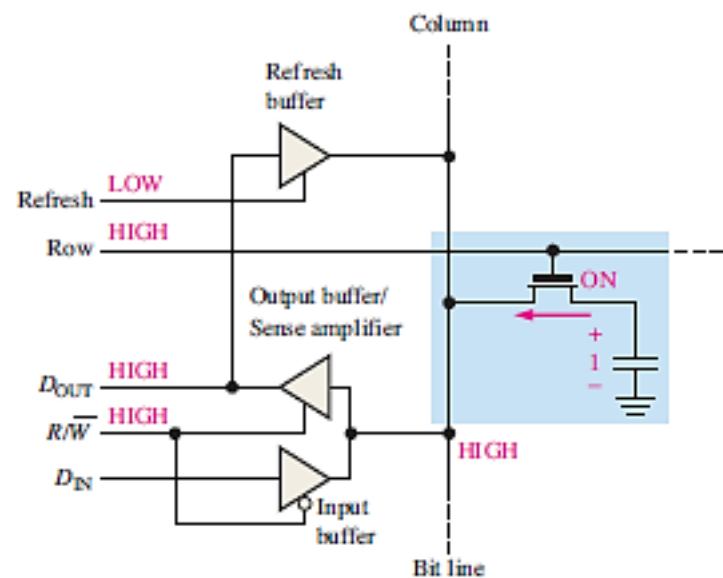
FIGURE 10 A MOS DRAM cell.



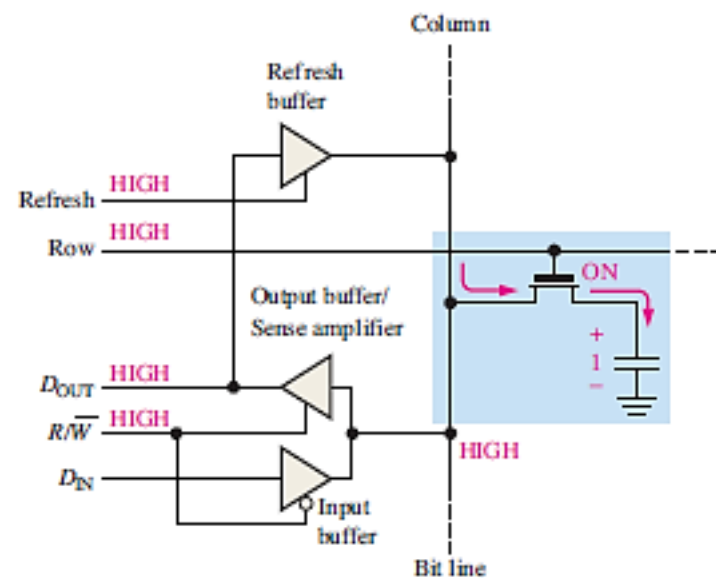
(a) Writing a 1 into the memory cell



(b) Writing a 0 into the memory cell



(c) Reading a 1 from the memory cell



(d) Refreshing a stored 1

FIGURE 11 Basic operation of a DRAM cell.

- ❑ A LOW on the $R/\overline{W}$ line (WRITE mode) enables the tri-state input buffer and disables the output buffer. For a 1 to be written into the cell, the D_{IN} line must be HIGH, and the transistor must be turned on by a HIGH on the row line. The transistor acts as a closed switch connecting the capacitor to the bit line. This connection allows the capacitor to charge to a positive voltage, as shown in Figure 11(a). When a 0 is to be stored, a LOW is applied to the D_{IN} line. If the capacitor is storing a 0, it remains uncharged, or if it is storing a 1, it discharges as indicated in Figure 11 (b). When the row line is taken back LOW, the transistor turns off and disconnects the capacitor from the bit line, thus “trapping” the charge (1 or 0) on the capacitor.
- ❑ To read from the cell, the $R/\overline{W}$ (Read/Write) line is HIGH, enabling the output buffer and disabling the input buffer. When the row line is taken HIGH, the transistor turns on and connects the capacitor to the bit line and thus to the output buffer (sense amplifier), so the data bit appears on the data-output line (D_{OUT}). This process is illustrated in Figure 11(c).
- ❑ For refreshing the memory cell, the $R/\overline{W}$ line is HIGH, the row line is HIGH, and the refresh line is HIGH. The transistor turns on, connecting the capacitor to the bit line. The output buffer is enabled, and the stored data bit is applied to the input of the refresh buffer, which is enabled by the HIGH on the refresh input. This produces a voltage on the bit line corresponding to the stored bit, thus replenishing the capacitor. This is illustrated in Figure 11(d).

DRAM Organization

The major application of DRAMs is in the main memory of computers. The difference between DRAMs and SRAMs is the type of memory cell. As you have seen, the DRAM memory cell consists of one transistor and a capacitor and is much simpler than the SRAM cell. This allows much greater densities in DRAMs and results in greater bit capacities for a given chip area, although much slower access time. Again, because charge stored in a capacitor will leak off, the DRAM cell requires a frequent refresh operation to preserve the stored data bit. This requirement results in more complex circuitry than in a SRAM. Several features common to most DRAMs are now discussed, using a generic 1M x 1 bit DRAM as an example.

Address Multiplexing

DRAMs use a technique called *address multiplexing* to reduce the number of address lines. Figure 12 shows the block diagram of a 1,048,576-bit (1 Mb) DRAM with a 1M x 1 organization. We will focus on the blue blocks to illustrate address multiplexing. The green blocks represent the refresh logic.

The ten address lines are time multiplexed at the beginning of a memory cycle by the row address select ($\overline{RAS}$) and the column address select ($\overline{CAS}$) into two separate 10-bit address fields. First, the 10-bit row address is latched into the row address register. Next, the 10-bit column address is latched into the column address register. The row address and the column address are decoded to select one of the 1,048,576 addresses ($2^{20} = 1,048,576$) in the memory array. The basic timing for the address multiplexing operation is shown in Figure 13.

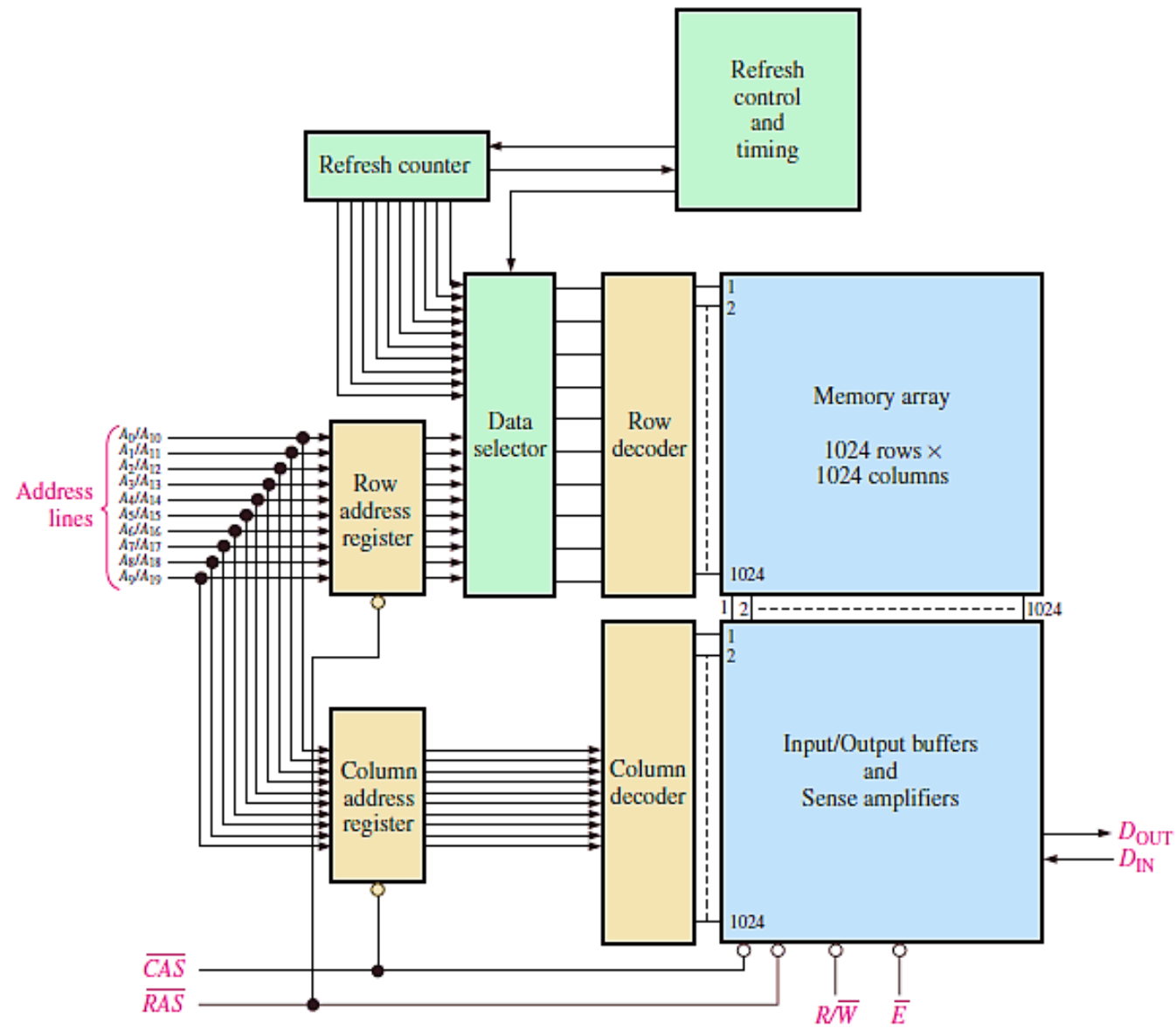


FIGURE 12 Simplified block diagram of a 1M * 1 DRAM.

Read and Write Cycles

- At the beginning of each read or write memory cycle, $\overline{RAS}$ and $\overline{CAS}$ go active (LOW) to multiplex the row and column addresses into the registers, and decoders. For a read cycle, the $R/\overline{W}$ input is HIGH. For a write cycle, the $R/\overline{W}$ input is LOW. This is illustrated in Figure 14.

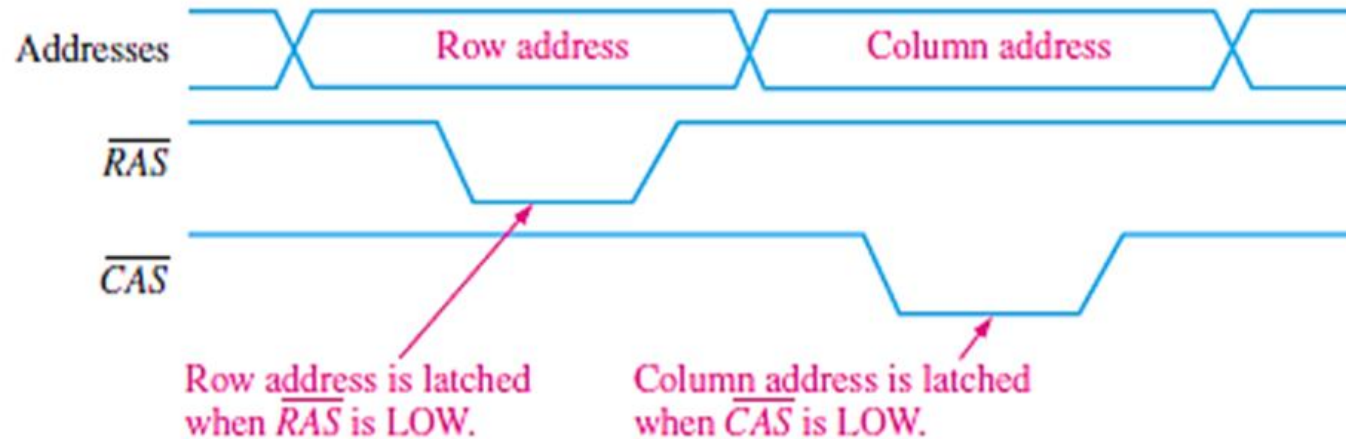
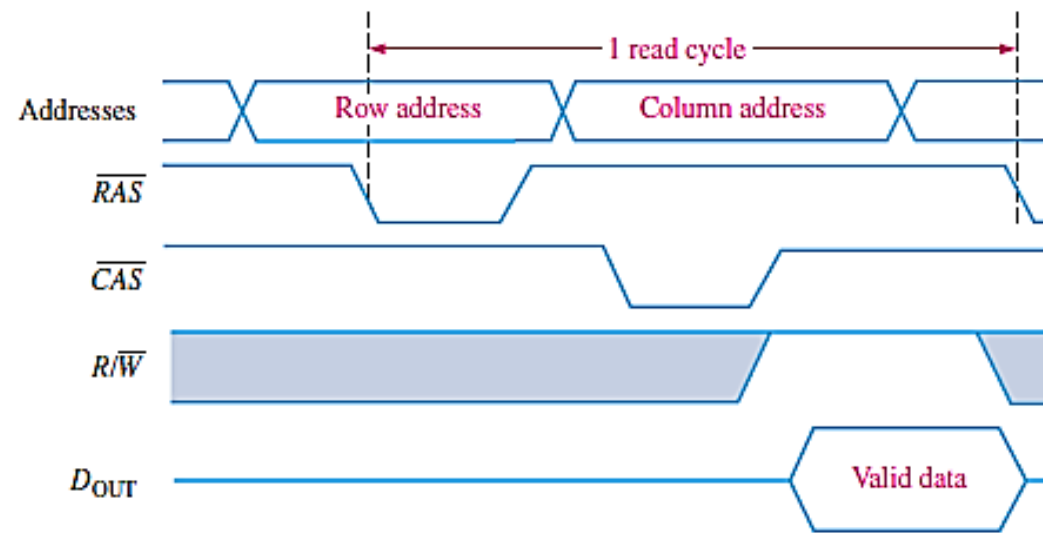
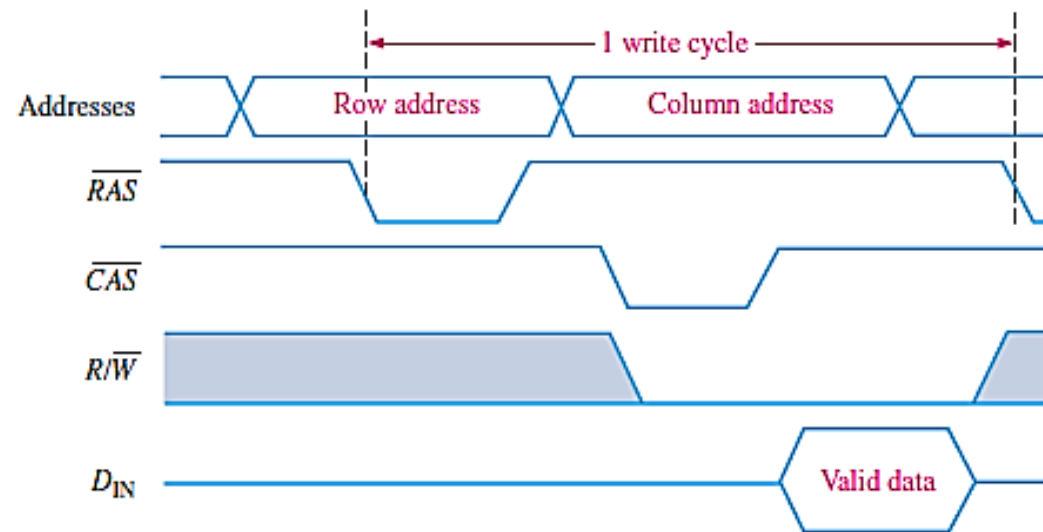


FIGURE 13 Basic timing for address multiplexing.



(a) Read cycle



(b) Write cycle

FIGURE 14 Timing diagrams for normal read and write cycles.

Fast Page Mode

- ❑ In the normal read or write cycle described previously, the row address for a particular memory location is first loaded by an active-LOW $\overline{RAS}$ and then the column address for that location is loaded by an active-LOW $\overline{CAS}$. The next location is selected by another $\overline{RAS}$ followed by a $\overline{CAS}$, and so on.
- ❑ A “page” is a section of memory available at a single row address and consists of all the columns in a row. Fast page mode allows fast successive read or write operations at each column address in a selected row. A row address is first loaded by $\overline{RAS}$ going LOW and remaining LOW while $\overline{CAS}$ is toggled between HIGH and LOW. A single row address is selected and remains selected while $\overline{RAS}$ is active. Each successive $\overline{CAS}$ selects another column in the selected row. So, after a fast page mode cycle, all of the addresses in the selected row have been read from or written into, depending on $R/\overline{W}$.
- ❑ For example, a fast page mode cycle for the DRAM in Figure 12 requires $\overline{CAS}$ to go active 1024 times for each row selected by $\overline{RAS}$.
- ❑ Fast page mode operation for read is illustrated by the timing diagram in Figure 15. When $\overline{CAS}$ goes to its non asserted state (HIGH), it disables the data outputs. Therefore, the transition of $\overline{CAS}$ to HIGH must occur only after valid data are latched by the external system.

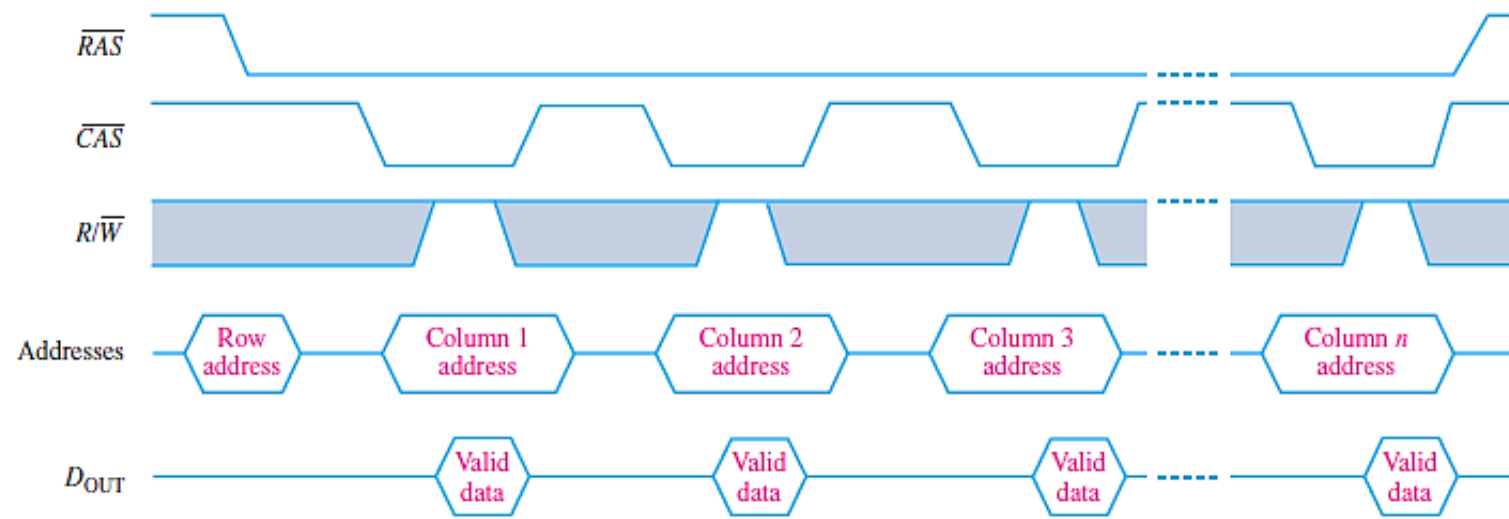


FIGURE 15 Fast page mode timing for a read operation.

Refresh Cycles

- ❑ As you know, DRAMs are based on capacitor charge storage for each bit in the memory array. This charge degrades (leaks off) with time and temperature, so each bit must be periodically refreshed (recharged) to maintain the correct bit state. Typically, a DRAM must be refreshed every several milliseconds, although for some devices the refresh period can be much longer.
- ❑ A read operation automatically refreshes all the addresses in the selected row. However, in typical applications, you cannot always predict how often there will be a read cycle, and so you cannot depend on a read cycle to occur frequently enough to prevent data loss. Therefore, special refresh cycles must be implemented in DRAM systems.
- ❑ Burst refresh and distributed refresh are the two basic refresh modes for refresh operations. In burst refresh, all rows in the memory array are refreshed consecutively each refresh period. For a memory with a refresh period of 8 ms, a burst refresh of all rows occurs once every 8 ms. The normal read and write operations are suspended during a burst refresh cycle.
- ❑ In distributed refresh, each row is refreshed at intervals interspersed between normal read or write cycles. For example, the memory in Figure 12 has 1024 rows. As an example, for an 8 ms refresh period, each row must be refreshed every $8 \text{ ms} / 1024 = 7.8 \text{ ms}$ when distributed refresh is used.

- ❑ The two types of refresh operations are *$\overline{RAS}$ only refresh* and *$\overline{CAS}$ before $\overline{RAS}$ refresh*. $\overline{RAS}$ -only refresh consists of a $\overline{RAS}$ transition to the LOW (active) state, which latches the address of the row to be refreshed while $\overline{CAS}$ remains HIGH (inactive) throughout the cycle. An external counter is used to provide the row addresses for this type of operation. The $\overline{CAS}$ before $\overline{RAS}$ refresh is initiated by $\overline{CAS}$ going LOW before $\overline{RAS}$ goes LOW. This sequence activates an internal refresh counter that generates the row address to be refreshed. This address is switched by the data selector into the row decoder.

Types of DRAMs

Now that you have learned the basic concept of a DRAM, let's briefly look at the major types. These are the *Fast Page Mode (FPM) DRAM*, the *Extended Data Out (EDO) DRAM*, the *Burst Extended Data Out (BEDO) DRAM*, and the *Synchronous (S) DRAM*.

FPM DRAM

- ❑ Fast page mode operation was described earlier. Recall that a page in memory is all of the column addresses contained within one row address. The idea of the **FPM DRAM** is based on the probability that the next several memory addresses to be accessed are in the same row (on the same page). Fortunately, this happens a large percentage of the time.
- ❑ FPM saves time over pure random accessing because in FPM the row address is specified only once for access to several successive column addresses whereas for pure random accessing, a row address is specified for each column address. Recall that in a fast page mode read operation, the $\overline{CAS}$ signal has to wait until the valid data from a given address are accepted (latched) by the external system (CPU) before it can go to its nonasserted state. When $\overline{CAS}$ goes to its nonasserted state, the data outputs are disabled.
- ❑ This means that the next column address cannot occur until after the data from the current column address are transferred to the CPU. This limits the rate at which the columns within a page can be addressed.

EDO DRAM

- ❑ The Extended Data Out DRAM, sometimes called *hyper page mode DRAM*, is similar to the FPM DRAM. The key difference is that the $\overline{CAS}$ signal in the **EDO DRAM** does not disable the output data when it goes to its nonasserted state because the valid data from the current address can be held until $\overline{CAS}$ is asserted again.

- ❑ This means that the next column address can be accessed before the external system accepts the current valid data.
The idea is to speed up the access time.

BEDO DRAM

The Burst Extended Data Out DRAM is an EDO DRAM with address burst capability. Recall from the discussion of the synchronous burst SRAM that the address burst feature allows up to four addresses to be internally generated from a single external address, which saves some access time. This same concept applies to the **BEDO DRAM**.

SDRAM

Faster DRAMs are needed to keep up with the ever-increasing speed of microprocessors. The Synchronous DRAM is one way to accomplish this. Like the synchronous SRAM discussed earlier, the operation of the **SDRAM** is synchronized with the system clock, which also runs the microprocessor in a computer system. The same basic ideas described in relation to the synchronous burst SRAM, also apply to the SDRAM. This synchronized operation makes the SDRAM totally different from the other asynchronous DRAM types. With asynchronous memories, the microprocessor must wait for the DRAM to complete its internal operations. However, with synchronous operation, the DRAM latches addresses, data, and control information from the processor under control of the system clock. This allows the processor to handle other tasks while the memory read or write operations are in progress, rather than having to wait for the memory to do its thing as is the case in asynchronous systems.

DDR SDRAM

DDR stands for double data rate. A DDR SDRAM is clocked on both edges of a clock pulse, whereas a SDRAM is clocked on only one edge. Because of the double clocking, a DDR SDRAM is theoretically twice as fast as an SDRAM. Sometimes the SDRAM is referred to as an SDR SDRAM (single data rate SDRAM) for contrast with the DDR SDRAM.

Thank You