



**INSTITUTO FEDERAL DE
EDUCAÇÃO CIÊNCIA E
TECNOLOGIA FARROUPILHA**



APOSTILA DE HARDWARE I

Professor: Heleno Cabral

- Campus Alegrete -

*Mesmo que tu já tenhas feito uma
longa caminhada, há sempre um
caminho a fazer.*
(Sto. Agostinho).

NOTA DO PROFESSOR

Com a evolução da humanidade e o avanço tecnológico, os computadores ganharam muito espaço dentro do cotidiano do homem, isto se deve ao fato de que as pessoas estão convergindo-se tecnologicamente, ou seja, buscando as maneiras mais fáceis de realizar uma tarefa ou uma maneira mais fácil de comunicarem, tudo hoje em dia está em torno da informação e a informação está totalmente ligada ao computador.

Esta apostila irá abordar algumas técnicas de montagem, configuração e reconhecimento de possíveis erros que possam ocorrer com o funcionamento desta máquina tão importante que é o computador, segundo CEFAS (2011), desde a sua **Primeira Geração** onde J.P. Eckert e John Mauchly, da Universidade da Pensilvânia, inauguraram o novo computador em 14 de fevereiro de 1946. O ENIAC era mil vezes mais rápido do que qualquer máquina anterior, resolvendo 5 mil adições e subtrações, 350 multiplicações ou 50 divisões por segundo. E tinha o dobro do tamanho do Mark I: encheu 40 gabinetes com 100 mil componentes, incluindo cerca de 17 mil válvulas eletrônicas. Pesava 27 toneladas e media 5,50 x 24,40 m e consumia 150 kW.

Em 1945 Von Neumann sugeriu que o sistema binário fosse adotado em todos os computadores, e que as instruções e dados fossem compilados e armazenados internamente no computador, na sequência correta de utilização. Estas sugestões tornaram-se a base filosófica para projetos de computadores. (Atualmente pesquisam-se computadores "não Von Neumann", que funcionam com *fuzzy logic*, lógica confusa) A partir dessas idéias, e da lógica matemática ou álgebra de Boole, introduzida por Boole no início do século XIX, é que Mauchly e Eckert projetaram e construíram o EDVAC, Electronic Discrete Variable Automatic Computer, completado em 1952, que foi a primeira máquina comercial eletrônica de processamento de dados do mundo.

Nesta caminhada surgem os computadores a transistores nos anos 50, pesando 150 kg, com consumo inferior a 1.500 W e maior capacidade que seus antecessores valvulados, eis a **Segunda Geração**.

Em 1957 o matemático Von Neumann colaborou para a construção de um computador avançado, o qual, por brincadeira, recebeu o nome de MANIAC, Mathematical Analyser Numerator Integrator and Computer. Em janeiro de 1959 a Texas Instruments anuncia ao mundo uma criação de Jack Kilby: o circuito integrado. Enquanto uma pessoa de nível médio levaria cerca de cinco minutos para multiplicar dois números de dez dígitos, o MARK I o fazia em cinco segundos, o ENIAC em dois milésimos de segundo, um computador transistorizado em cerca de quatro bilionésimos de segundo, e, uma máquina de terceira geração em menos tempo ainda. A **Terceira Geração** de computadores é da década de 60, com a introdução dos circuitos integrados. O Burroughs B-2500 foi um dos primeiros. Enquanto o ENIAC podia armazenar vinte números de dez dígitos, estes armazenavam milhões de números. Surgem então, conceitos como memória virtual, multiprogramação e sistemas operacionais complexos.

O primeiro minicomputador comercial surgiu em 1965, o PDP-5, lançado pela americana DEC, *Digital Equipment Corporation*. Dependendo de sua configuração e acessórios ele podia ser adquirido pelo acessível preço de US\$ 18.000,00. Em 1970 a INTEL Corporation introduziu no mercado um tipo novo de circuito integrado: o microprocessador. O primeiro foi o 4004, de quatro bits. Foi seguido pelo 8008, em 1972, o difundidíssimo 8080, o 8085, etc. A partir daí surgem os microcomputadores.

Para muitos, a **Quarta Geração** surge com os chips VLSI, de integração em muito larga escala. As coisas começam a acontecer com maior rapidez e frequência. Em 1972 Bushnell lança o vídeo game Atari. Kildall lança o CP/M em 1974. O primeiro kit de microcomputador, o ALTAIR 8800 em 1974/5. Em 1975 Paul Allen e Bill Gates criam a Microsoft e o primeiro software para microcomputador: uma adaptação BASIC para o ALTAIR.

Nesse período surgiu também o processamento distribuído, o disco ótico e o a grande difusão do microcomputador, que passou a ser utilizado para processamento de texto, cálculos auxiliados. Em 1982, surgiu o 286 usando memória de 30 pinos e slots ISA de 16 bits, já vinha equipado com memória cache, para auxiliar o processador em suas funções. Utilizava ainda monitores CGA em alguns raros modelos estes monitores eram coloridos mas a grande maioria era verde, laranja ou cinza. O 386 já contava com placas VGA que podiam atingir até 256 cores desde que o monitor também suportasse essa configuração. O 486 DX em 1989 já utilizava o co-processador matemático embutido no próprio processador, houve também uma melhora sensível na velocidade devido o advento da memória de 72 pinos, muito mais rápida que sua antepassada de 30 pinos e das placas PCI de 32 bits duas vezes mais velozes que as placas ISA. Os equipamentos já tinham capacidade para as placas SVGA que poderiam atingir até 16 milhões de cores, porém este artifício seria usado comercialmente mais para frente com o advento do Windows 95. Neste momento iniciava uma grande debandada para as pequenas redes como, a Novel e a Lantastic que rodariam perfeitamente nestes equipamentos, substituindo os "micrões" que rodavam em sua grande maioria o sistema UNIX. Esta substituição era extremamente viável devido à diferença brutal de preço entre estas máquinas. A **Quinta Geração** surge com a exigência crescente de uma maior capacidade de processamento e armazenamento de dados. Sistemas especialistas, sistemas multimídia (combinação de textos, gráficos, imagens e sons), banco de dados distribuídos e redes neurais, são apenas alguns exemplos dessas necessidades. Uma das principais características dessa geração são a simplificação e miniaturização do computador, além de melhor desempenho e maior capacidade de armazenamento. Tudo isso, com os preços cada vez mais acessíveis. A tecnologia VLSI foi sendo substituída pela ULSI (ULTRA LARGE SCALE INTEGRATION). O conceito de processamento começou a partir para os processadores paralelos, ou seja, a execução de muitas operações simultaneamente pelas máquinas. A partir de 1997 não houve grandes novidades, sendo que as mudanças ficaram por conta dos cada vez mais velozes processadores e seus derivados como Dual - Core, Core2Duo, X2, X3, i3, i5, i7 e por aí vai...

LISTA DE FIGURAS

Figura 1. Gabinete Desktop. Fonte: http://2.bp.blogspot.com/_6BbTEzXQA98/TAP6eFpwNSI/AAAAAAAAAJE/os3ndzHJTus/s1600/a4000-topless.jpg	10
Figura 2. Gabinete minitorre. Fonte: http://www.netstorageti.com.br/images/52268353.jpg	10
Figura 3. Processadores AMD e Intel. Fonte: http://www.precos.com.pt/processadoresc3139/amd/phenom-ii-x4-630-2-8ghz-am3-p38567377/	11
Figura 4. Memórias DDR2 e DDR3. Fonte: http://www.cheatsbrasil.org/local/hardware/968-o-que-memoria-ram.html	12
Figura 5. Imagem de um HD aberto. Fonte: http://www.infowester.com/hds1	13
Figura 6. Placa de vídeo Trident 9440. Fonte: [CEFAS, 2011]	14
Figura 7. Polígonos e imagem finalizada. Fonte: [CEFAS, 2011]	14
Figura 8. Placa da ATi FireGL, altíssimo poder gráfico. Fonte: http://www.moledeaprender.com.br/o-que-e-uma-placa-de-video-de-computador/	15
Figura 9. Placa-mãe. Fonte: http://www.gigabyte.com.tw	15
Figura 10. Fonte de alimentação. Fonte: http://helpcenterpc.com.br/loja/view_product.php?product=ATX	17
Figura 11. Cooler. Fonte: http://guia.mercadolivre.com.br/cooler-6489-VGP	18
Figura 12. Placa de Fax/MODEM. Fonte: http://www.dausacker.com.br	18
Figura 13. Placa de Som. Fonte: http://todaoferta.uol.com.br/guia/placa-de-som-escolhendo-a-certa-para-seu-computador.html	19
Figura 14. Placa de Rede RJ-45. Fonte: http://imei12joana.blogspot.com/	20
Figura 15. Cabos <i>Flats</i> . Fonte: http://toniinfo.com/flat-cable/	20
Figura 16. Teclado. Fonte: http://juniorberhends.blogspot.com/2011/02/verdadeira-funcao-das-teclas.html	21
Figura 17. Mouse. Fonte: http://avitricedazicadopantano.blogspot.com/2009/09/historia-do-mouse-de-computador.html	22
Figura 18. Monitor LCD com tecnologia Led. Fonte: http://www.mundotecno.info/noticias/aoc-v17-monitor-lcd-com-12-milimetros-de-espessura	22
Figura 19. Impressora a Jato de Tinta com visor LCD. Fonte: http://www.forumpcs.com.br/comunidade/viewtopic.php?t=106211	23
Figura 20. Scanner de Mesa. Fonte: http://ccinfo.hdfree.com.br/loja/listdeprodutos/scanner.htm	24

Figura 21. Caixas de Som. Fonte: http://www.netserv19.com/ecommerce_site/produto_94623_1598_Kit-Multimidia-Teclado-Mult--Mouse-Optico-e-Caixa-de-Som-300w	24
Figura 22. WebCam. Fonte: http://baixandobemaki.blogspot.com/2010/09/fake-webcam.html	25
Figura 23. Porta Serial DB-9. Fonte: http://ulisses16.blogspot.com/2007_11_01_archive.html	28
Figura 24. Porta Paralela DB-25. Fonte: http://pt.wikipedia.org/wiki/Interface_paralela	28
Figura 25. Placa IR. Fonte: http://www.comofaco.com/2006/03/23/construindo-uma-interface-infra-vermelho/	29
Figura 26. Porta Joystick DB-15.....	29
Figura 27. Porta USB. Fonte: http://milos-djacic.com/m4uhow2/6.php	29
Figura 28. Conector SCSI. Fonte: http://www.datapro.net/techinfo/scsi_doc.html	30
Figura 29. Conector IDE. Fonte: http://www.refrigeracao.net/Cursos/eletronica/conector.htm	30
Figura 30. Slot ISA. Fonte: [CEFAS, 2011]	30
Figura 31. Slot AGP e PCIs. [CEFAS, 2011]	31
Figura 32. Placa AGP de 3.3V e placa AGP universal. Fonte: [CEFAS, 2011]	30
Figura 33. Slots PCI-e. Fonte: http://en.wikipedia.org/wiki/File:PCIExpress.jpg	35
Figura 34. Ponte norte do chipset (com o dissipador removido). Fonte: [CEFAS, 2011]	36
Figura 35. Ponte sul. Fonte: [CEFAS, 2011]	36
Figura 36. Ciclos DDR2	39
Figura 37. Memória DDR3 da Corsair. Fonte: http://buscaparaguai.com.br/memoria-corsair-ddr3-1gb-1333mhz-26793-produto-no-paraguai	40

LISTA DE TABELAS

Tabela 1. Lista de Siglas	10
Tabela 2. Quadro de ordens	47

SUMÁRIO

NOTA DO PROFESSOR	3
LISTA DE FIGURAS	5
LISTA DE TABELAS	7
LISTA DE SIGLAS	10
1- COMO UM PC FUNCIONA	11
2- HARDWARE	13
2.1- O que é Hardware	13
2.1.1- Gabinete	13
2.1.2- Processador	14
2.1.3- Memória RAM	14
2.1.4- Disco Rígido	15
2.1.5- Placa de Vídeo	16
2.1.6- Placa-Mãe	18
2.1.7- Fonte de Alimentação	19
2.1.8- <i>Cooler</i>	20
2.1.9- Placa Fax/MODEM	21
2.1.10- Placa de Som	21
2.1.11- Placa de Rede	22
2.1.12- Cabo <i>Flat</i>	22
2.1.13- Teclado	23
2.1.14- Mouse	24
2.1.15- Monitor	24
2.1.16- Impressora	25
2.1.17- Scanner	26
2.1.18- Kit Multimídia	27
2.1.19- <i>WebCam</i>	27
3- COMPONENTES DA PLACA-MÃE	28
3.1- BIOS, CMOS e POST	28
3.1.1- BIOS	28
3.1.2- CMOS	28
3.1.3- POST	29

3.2- Interfaces	30
3.2.1- Interface Serial	30
3.2.2- Interface Paralela	30
3.2.3- Interface IR	31
3.2.4- Interface <i>Joystick</i> ou <i>Game</i>	31
3.2.5- Interface USB	31
3.2.6- Interface SCSI	31
3.2.7- Interface IDE	32
3.3- Slots de Barramento	33
3.3.1- ISA	33
3.3.2- PCI	34
3.3.3- AGP	35
3.3.4- PCI Express	37
3.3.5- <i>Chipsets</i>	38
3.4- Tipos de Memórias	40
3.4.1- Memórias SDRAM	40
3.4.2- Memórias DDR2	41
3.4.3- Memórias DDR3	42
4- DISCOS RAID	45
4.1- O que é RAID?	45
4.2- Os Níveis do RAID	45
5- SISTEMAS NUMÉRICOS	48
5.1- Sistema Numérico Decimal	48
5.2- Sistema Numérico Binário	48
5.3- Sistema Numérico Hexadecimal	49
REFERÊNCIAS BIBLIOGRÁFICAS	50

LISTA DE SIGLAS

BIOS	Basic Input/Output System – Sistema básico de entrada e saída.
CMOS	Complementary Metal-Oxide Semiconductor – Semicondutor Óxido De Metal Complementar. Um chip fabricado para reproduzir as funções de outros chips, como chips de memória ou microprocessadores. Esta tecnologia propicia um menos consumo de energia.
EEPROM	Electrically Erasable Programmable Read-Only Memory – Memória de Apenas Leitura Apagável Eletricamente e Programável. Chip que pode ser reprogramado após ter seu conteúdo apagado eletricamente.
EPROM	Erasable Programmable Read-Only Memory – Memória de Apenas Leitura Apagável e Programável. Chip que pode ser reprogramado após ter seu conteúdo apagado por luz ultravioleta.
FLASH	Memória que mantém suas informações mesmo após desligado o microcomputador. Exemplos: BIOS em FLASH EEPROM que possibilita a atualização do BIOS via software, cartões PCMCIA utilizados como expansão de discos nos notebooks, cartões de memória e/ou pendrives.
PROM	Programmable Read-Only Memory – Memória de Apenas Leitura Programável. Chip que já sei programado de fábrica para funcionar com determinado microcomputador.
RAM	Random Access Memory – Memória de Acesso Randômico. Memória do tipo volátil.
ROM	Read-Only Memory – Memória de Apenas Leitura. Memória que não perde seu conteúdo quando o microcomputador é desligado. Utilizada para armazenar funções e rotinas necessárias para controles de funcionamento básico.
SRAM	Static Random Access Memory – Memória Estática de Acesso Randômico. Tem as características da memória RAM, acrescidas do avanço tecnológico de não necessitar da renovação constante pelo processador o que faz dela uma memória mais rápida que as RAM.

Tabela 1 – Lista de Siglas

1- COMO UM PC FUNCIONA?

O primeiro PC foi lançado em 1981, pela IBM. A plataforma PC não é a primeira nem será a última plataforma de computadores pessoais, mas ela é de longe a mais usada e provavelmente continuará assim por mais algumas décadas. Para a maioria das pessoas, "PC" é sinônimo de computador. Começando do básico, existem duas maneiras de representar uma informação: analogicamente ou digitalmente. Uma música gravada em uma antiga fita K7 é armazenada de forma analógica, codificada na forma de uma grande onda de sinais magnéticos, que podem assumir um número virtualmente ilimitado de frequências. Quando a fita é tocada, o sinal magnético é amplificado e novamente convertido em som, gerando uma espécie de "eco" do áudio originalmente gravado. O grande problema é que o sinal armazenado na fita se degrada com o tempo, e existe sempre certa perda de qualidade ao fazer cópias [CEFAS, 2011].

Ao tirar várias cópias sucessivas, cópia da cópia, você acabava com uma versão muito degradada da música original. Ao digitalizar a mesma música, transformando-a em um arquivo MP3, você pode copiá-la do PC para o MP3 player, e dele para outro PC, sucessivamente, sem causar qualquer degradação. Você pode perder alguma qualidade ao digitalizar o áudio, ou ao comprimir a faixa original, gerando o arquivo MP3, mas a partir daí pode reproduzir o arquivo indefinidamente e fazer cópias exatas. Isso é possível devido à própria natureza do sistema digital, que permite armazenar qualquer informação na forma de uma sequência de valores positivos e negativos, ou seja, na forma de uns e zeros.

O número 181, por exemplo, pode ser representado digitalmente como 10110101; uma foto digitalizada é transformada em uma grande grade de pixels e um valor de 8, 16 ou 24 bits é usado para representar cada um; um vídeo é transformado em uma sequência de imagens, também armazenadas na forma de pixels e assim por diante. A grande vantagem do uso do sistema binário é que ele permite armazenar informações com uma grande confiabilidade, em praticamente qualquer tipo de mídia; já que qualquer informação é reduzida a combinações de apenas dois valores diferentes. A informação pode ser armazenada de forma magnética, como no caso dos HDs; de forma óptica, como no caso dos CDs e DVDs ou até mesmo na forma de impulsos elétricos, como no caso dos chips de memória flash.

Cada UM ou ZERO processado ou armazenado é chamado de "bit", contração de "*binary digit*" ou "dígito binário". Um conjunto de 8 bits forma um byte, e um conjunto de 1024 bytes forma um kilobyte (ou KByte). O número 1024 foi escolhido por ser a potência de 2 mais próxima de 1000. É mais fácil para os computadores trabalharem com múltiplos de dois do que usar o sistema decimal como nós. Um conjunto de 1024 kbytes forma um megabyte e um conjunto de 1024 megabytes forma um gigabyte. Os próximos múltiplos são o terabyte (1024 gigabytes) e o petabyte (1024 terabytes), exabyte, zettabyte e o yottabyte, que equivale a 1.208.925.819.614.629.174.706.176 bytes. É provável que, com a evolução da informática, daqui a algumas décadas surja algum tipo de unidade de armazenamento capaz de armazenar um yottabyte inteiro, mas atualmente ele é um número quase inatingível. Para armazenar um yottabyte inteiro, usando tecnologia atual, seria necessário construir uma estrutura colossal de servidores. Imagine que, para manter os custos baixos, fosse adotada uma estratégia estilo Google, usando PCs comuns, com HDs IDE/SATA. Cada PC seria equipado com dois HDs de 2 TB. Estes PCs seriam então organizados em enormes racks, onde cada rack teria espaço para 256 PCs. Os PCs de cada rack seriam ligados a um conjunto de switches e cada grupo de switches seria ligado a um grande roteador. Uma vez ligados em

rede, os 256 PCs seriam configurados para atuar como um enorme cluster, trabalhando como se fossem um único sistema. Construiríamos então um enorme galpão, capaz de comportar 1024 desses racks, construindo uma malha de switches e roteadores capaz de ligá-los em rede com um desempenho minimamente aceitável. Esse galpão precisa de um sistema de refrigeração colossal, sem falar da energia consumida por quase meio milhão de PCs dentro dele, por isso construímos uma usina hidrelétrica para alimentá-lo, represando um rio próximo. Com tudo isso, conseguiríamos montar uma estrutura computacional capaz de armazenar 1 exabyte. Ainda precisaríamos construir mais 1.048.576 mega datacenters como esse para chegar a 1 yottabyte. Se toda a humanidade se dividisse em grupos de 6.000 pessoas e cada grupo fosse capaz de construir um ao longo de sua vida, deixando de lado outras necessidades existenciais, poderíamos chegar lá. Mas... voltando à realidade, usamos também os termos kbit, megabit e gigabit, para representar conjuntos de 1024 bits. Como um byte corresponde a 8 bits, um megabyte corresponde a 8 megabits e assim por diante. Quando você compra uma placa de rede de "100 megabits" está na verdade levando para a casa uma placa que transmite 12.5 megabytes por segundo, pois cada byte tem 8 bits. Quando vamos abreviar, também existe diferença. Quando estamos falando de kbytes ou megabytes, abreviamos respectivamente como KB e MB, sempre com o B maiúsculo.

Por outro lado, quando estamos falando de kbits ou megabits abreviamos da mesma forma, porém usando o B minúsculo: Kb, Mb e assim por diante. Parece só um daqueles detalhes sem importância, mas esta é uma fonte de muitas confusões. Se alguém anuncia no jornal que está vendendo uma "placa de rede de 1000 MB", está dando a entender que a placa trabalha a 8000 megabits e não a 1000.

2- HARDWARE

2.1- O que é Hardware?

O hardware, circuitaria, material ou ferramental é a parte física do computador, ou seja, é o conjunto de componentes eletrônicos, circuitos integrados e placas, que se comunicam através de barramentos [WIKIPÉDIA, 2011]. Em complemento ao hardware, o software é a parte lógica, ou seja, o conjunto de instruções e dados processado pelos circuitos eletrônicos do hardware.

Toda interação dos usuários de computadores modernos é realizada através do software, que é a camada, colocada sobre o hardware, que transforma o computador em algo útil para o ser humano. Além de todos os componentes de hardware, o computador também precisa de um software chamado Sistema Operacional. O Sistema Operacional torna o computador utilizável. Ele é o responsável por gerenciar os dispositivos de hardware do computador (como memória, unidade de disco rígido, unidade de CD) e oferecer o suporte para os outros programas funcionarem (como processador de texto, planilhas, etc.).

O termo hardware não se refere apenas aos computadores pessoais, mas também aos equipamentos embarcados em produtos que necessitam de processamento computacional, como o dispositivos encontrados em equipamentos hospitalares, automóveis, aparelhos celulares, entre outros.

2.1.1- Gabinete: O gabinete é a caixa metálica ou plástica que abriga o PC. Essa caixa metálica pode ter vários formatos, sendo os mais usuais o desktop e o minitorre.



Figura 1 – Gabinete desktop

No desktop a placa-mãe é instalada horizontalmente à mesa onde o gabinete será apoiado, enquanto que no minitorre a placa-mãe é instalada perpendicularmente à mesa onde o gabinete será apoiado. Como o modelo mais vendido de gabinete é o minitorre, iremos utilizar este modelo de gabinete em nossas explicações.



Figura 2 – Gabinete minitorre

Muitos leigos chamam o gabinete do micro de CPU. CPU é a sigla de *Central Processing Unit* (Unidade Central de Processamento), sinônimo para o processador do micro, e não para o gabinete. Chamar um micro de CPU é tão estranho quanto chamar um carro somente de motor. Independentemente se o gabinete é desktop ou minitorre, internamente ele deverá ter o espaço para alocar corretamente uma placa-mãe.

Exist(em) (iam) basicamente dois formatos de placa-mãe no mercado: AT e ATX; Dessa forma, deveria escolher um gabinete de acordo com a placa-mãe que iria utilizar. Atualmente todas as placas-mãe encontradas no mercado utilizam o layout ATX com variações para ATX12V.

2.1.2- Processador: O processador é o cérebro do micro [MORIMOTO, 2010] encarregado de processar a maior parte das informações. Ele é também o componente onde são usadas as tecnologias de fabricação mais recentes. Existem no mundo apenas quatro grandes empresas com tecnologia para fabricar processadores competitivos para micros PC: a Intel (que domina mais de 60% do mercado), a AMD (que disputa diretamente com a Intel), a VIA (que fabrica os chips VIA C3 e C7, embora em pequenas quantidades) e a IBM, que esporadicamente fabrica processadores para outras empresas, como a Transmeta.

O Athlon II X4 ou o Core i7 são os componentes mais complexos e frequentemente os mais caros, mas ele não pode fazer nada sozinho. Como todo cérebro, ele precisa de um corpo, que é formado pelos outros componentes do micro, incluindo memória, HD, placa de vídeo e de rede, monitor, teclado e mouse.

Dentre as empresas citadas a AMD começou produzindo processadores 386 e 486, muito similares aos da Intel, porém mais baratos. Quando a Intel lançou o Pentium, que exigia o uso de novas placas-mãe, a AMD lançou o "5x86", um 486 de 133 MHz, que foi bastante popular, servindo como uma opção barata de upgrade. Embora o "5x86" e o clock de 133 MHz dessem a entender que se tratava de um processador com um desempenho similar a um Pentium 133, o desempenho era muito inferior, mal concorrendo com um Pentium 66. Este foi o primeiro de uma série de exemplos, tanto do lado da AMD, quanto do lado da Intel, em que existiu uma diferença gritante entre o desempenho de dois processadores do mesmo clock. Embora seja um item importante, a frequência de operação não é um indicador direto do desempenho do processador.

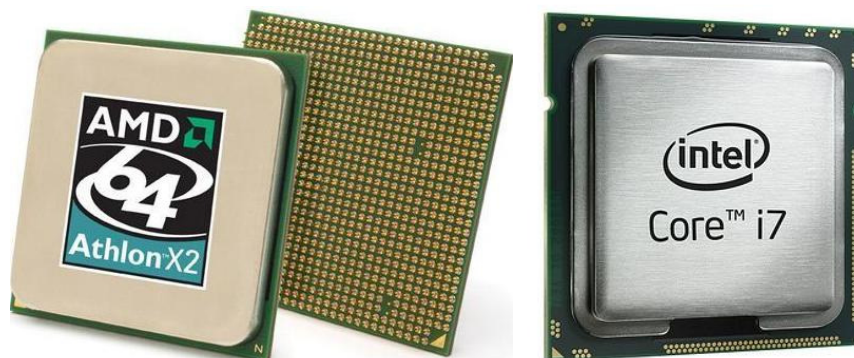


Figura 3 – Processadores AMD e Intel

2.1.3- Memória RAM: Logo após o processador, temos a memória RAM, usada por ele para armazenar os arquivos e programas que estão sendo executados, como uma espécie de mesa de trabalho. A quantidade de memória RAM disponível tem um grande efeito sobre o desempenho, já que sem memória RAM suficiente o sistema passa a usar memória

swap (virtual), que é muito mais lenta. A principal característica da memória RAM é que ela é volátil, ou seja, os dados se perdem ao reiniciar o micro. É por isso que ao ligar é necessário sempre refazer todo o processo de carregamento, em que o sistema operacional e aplicativos usados são transferidos do HD para a memória, onde podem ser executados pelo processador.

Os chips de memória são vendidos na forma de pentes, variando capacidade, barramento e velocidade. Normalmente as placas possuem de dois a quatro encaixes disponíveis. Tratando de memórias pode-se instalar um pente de 2 GB junto com o de 1 GB que veio no micro para ter um total de 3 GB, por exemplo. Ao contrário do processador, que é extremamente complexo, os chips de memória são formados pela repetição de uma estrutura bem simples, formada por um par de um transistor e um capacitor. Um transistor solitário é capaz de processar um único bit de cada vez, e o capacitor permite armazenar a informação por certo tempo. Essa simplicidade faz com que os pentes de memória sejam muito mais baratos que os processadores, principalmente se levarmos em conta o número de transistores. Um pente de 1 GB é geralmente composto por 8 chips, cada um deles com um total de 1024 megabits, o que equivale a 1024 milhões de transistores. Um Athlon 64 X2 tem "apenas" 233 milhões e custa bem mais caro que um pente de memória. Mais recentemente, temos assistido a uma nova migração, com a introdução dos pentes de memória DDR2 e DDR3. Neles, o barramento de acesso à memória trabalha ao dobro da frequência dos chips de memória propriamente ditos. Isso permite que sejam realizadas duas operações de leitura por ciclo, acessando dois endereços diferentes. Como a capacidade de realizar duas transferências por ciclo introduzida nas memórias DDR foi preservada, as memórias DDR2/DDR3 são capazes de realizar um total de 4 operações de leitura por ciclo, uma marca impressionante.

Apesar de toda a evolução a memória RAM continua sendo muito mais lenta que o processador. Para atenuar a diferença, são usados dois níveis de cache, incluídos no próprio processador: o cache L1 e o cache L2. O cache L1 é extremamente rápido, trabalhando próximo à frequência nativa do processador. Na verdade, os dois trabalham na mesma frequência, mas são necessários alguns ciclos de clock para que a informação armazenada no L1 chegue até as unidades de processamento.



Figura 4 – À esquerda temos a DDR3 com dissipador de calor e logo acima a DDR2

2.1.4- Disco Rígido: Também chamado de *hard disk*, HD ou até mesmo de *winchester*. Ele serve como unidade de armazenamento permanente, guardando dados e programas. O HD armazena os dados em discos magnéticos que mantêm a gravação por vários anos. Os discos giram a uma grande velocidade e um conjunto de cabeças de leitura, instaladas em um braço móvel faz o trabalho de gravar ou acessar os dados em qualquer posição nos discos. Junto com o CD-ROM, o HD é um dos poucos componentes mecânicos ainda usados nos micros atuais e, justamente por isso, é o que normalmente dura menos tempo (em média de três a cinco anos de uso contínuo) e que inspira mais cuidados.

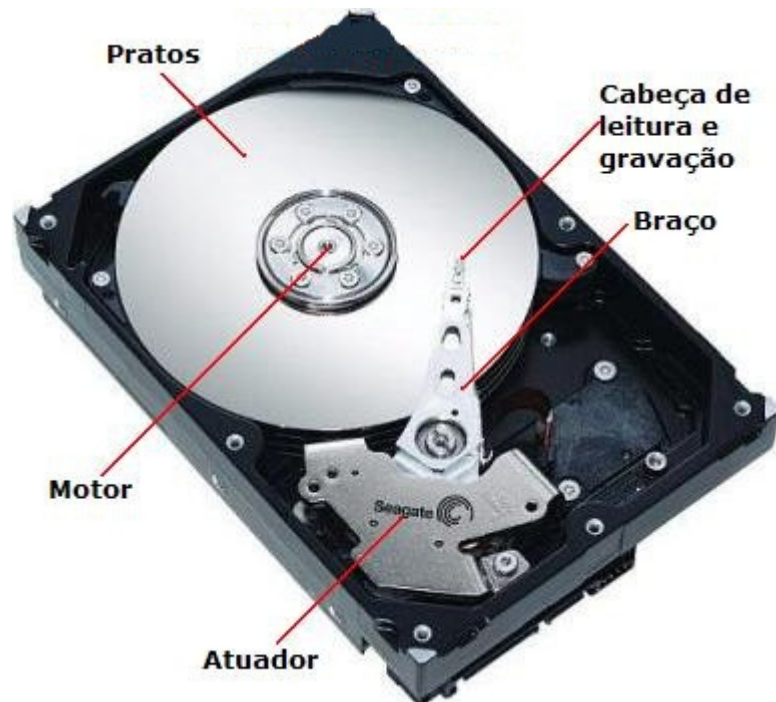


Figura 5 – Imagem de um HD aberto

Os discos magnéticos dos HDs são selados, pois a superfície magnética onde são armazenados os dados é extremamente fina e sensível. Qualquer grão de poeira que chegasse aos discos poderia causar danos à superfície, devido à enorme velocidade de rotação dos discos. Fotos em que o HD aparece aberto são apenas ilustrativas, no mundo real ele é apenas uma caixa fechada sem tanta graça. Apesar disso, é importante notar que os HDs não são fechados hermeticamente, muito menos a vácuo, como muitos pensam. Um pequeno filtro permite que o ar entra e saia, fazendo com que a pressão interna seja sempre igual à do ambiente. O ar é essencial para o funcionamento do HD, já que ele é necessário para criar o "colchão de ar" que evita que as cabeças de leitura toquem os discos. Existem alguns padrões como SATA, SCSI, PATA, e também variações na velocidade geralmente entre 5400 a 7200 rpm (rotação por minuto).

2.1.5- Placa de vídeo: Depois do processador, memória e HD, a placa de vídeo é provavelmente o componente mais importante do PC. Originalmente, as placas de vídeo eram dispositivos simples, que se limitavam a mostrar o conteúdo da memória de vídeo no monitor. A memória de vídeo continha um simples bitmap da imagem atual, atualizada pelo processador, e o RAMDAC (um conversor digital-analógico que faz parte da placa de vídeo) lia a imagem periodicamente e a enviava ao monitor. A resolução máxima suportada pela placa de vídeo era limitada pela quantidade de memória de vídeo. Na época, memória era um artigo caro, de forma que as placas vinham com apenas 1 ou 2 MB. As placas de 1 MB permitiam usar no máximo 800x600 com 16 bits de cor, ou 1024x768 com 256 cores. Estavam limitadas ao que cabia na memória de vídeo. Como mostra a Figura 6, a placa da Trident modelo 9440, era uma placa de vídeo muito comum no início dos anos 90; uma curiosidade é que ela foi uma das poucas placas de vídeo "atualizáveis" da história. Ela vinha com apenas dois chips de memória, totalizando 1 MB, mas era possível instalar mais dois, totalizando 2 MB. Hoje em dia, atualizar a memória da placa de vídeo é impossível, já que as placas utilizam módulos BGA, que podem ser instalados apenas em fábrica.

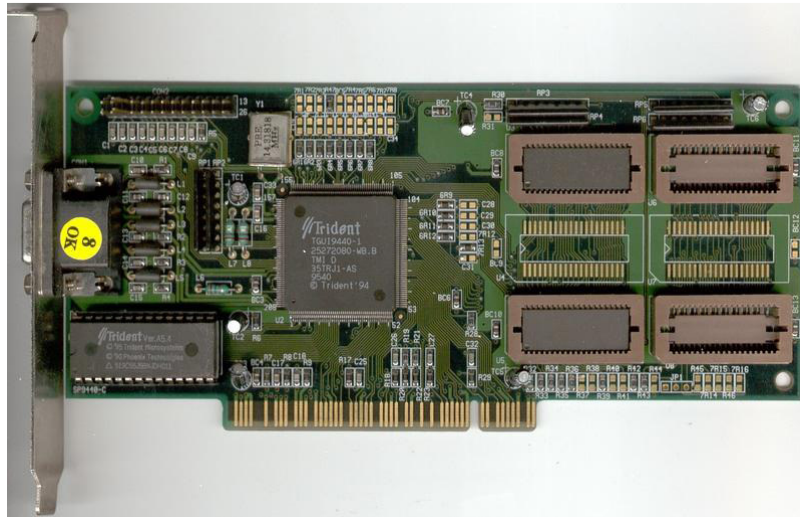


Figura 6 – Placa de vídeo Trident 9440

Em seguida, as placas passaram a suportar recursos de aceleração, que permitem fazer coisas como mover janelas ou processar arquivos de vídeo de forma a aliviar o processador principal. Esses recursos melhoraram bastante a velocidade de atualização da tela (em 2D), tornando o sistema bem mais responsivo. Finalmente, as placas deram o passo final, passando a suportar recursos 3D. Imagens em três dimensões são formadas por polígonos, formas geométricas como triângulos e retângulos em diversos formatos. Qualquer objeto em um game 3D é formado por um grande número destes polígonos. Cada polígono tem sua posição na imagem, um tamanho e cor específicos. O "processador" incluído na placa, responsável por todas estas funções é chamado de GPU (*Graphics Processing Unit*, ou unidade de processamento gráfico).



Figura 7 – Polígonos e imagem finalizada

As placas *onboard* [CEFAS, 2011] são soluções bem mais simples, onde o GPU é integrado ao próprio chipset da placa-mãe e, em vez de utilizar memória dedicada, como nas placas *offboard*, utiliza parte da memória RAM principal, que é "roubada" do sistema. Mesmo uma placa antiga, como a GeForce 4 Ti4600, tem 10.4 GB/s de barramento com a memória de vídeo, enquanto ao usar um pente de memória DDR PC 3200, temos apenas 3.2 GB/s de barramento na memória principal, que ainda por cima precisa ser compartilhado entre o vídeo e o processador principal. O processador lida bem com isso, graças aos caches L1 e L2, mas a

placa de vídeo realmente não tem para onde correr. É por isso que os chipsets de vídeo onboard são normalmente bem mais simples: mesmo um chip caro e complexo não ofereceria um desempenho muito melhor, pois o grande limitante é o acesso à memória.



Figura 8 – Placa da ATi FireGL, altíssimo poder gráfico

2.1.6- Placa-mãe: A placa-mãe é o componente mais importante do micro [MORIMOTO, 2010], pois é ela a responsável pela comunicação entre todos os componentes. Pela enorme quantidade de chips, trilhas, capacitores e encaixes, a placa-mãe também é o componente que, de uma forma geral, mais dá defeitos. É comum que um slot PCI pare de funcionar (embora os outros continuem normais), que instalar um pente de memória no segundo soquete faça o micro passar a travar, embora o mesmo pente funcione perfeitamente no primeiro e assim por diante. A maior parte dos problemas de instabilidade e travamentos são causados por problemas diversos na placa-mãe, por isso ela é o componente que deve ser escolhido com mais cuidado.

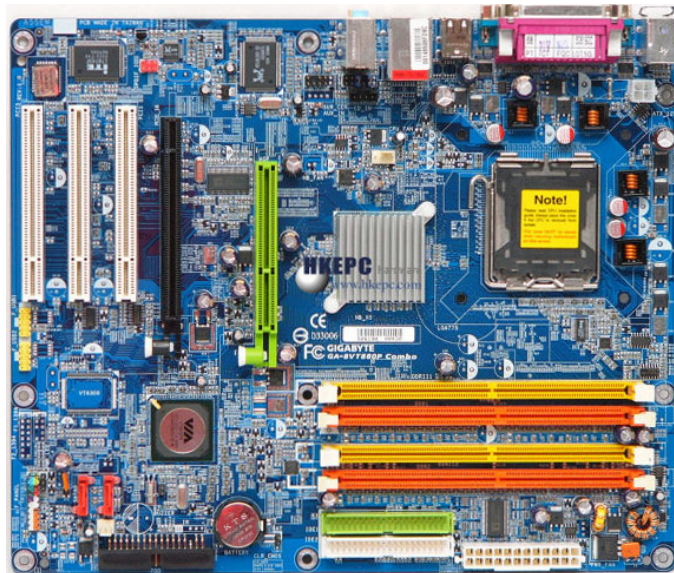


Figura 9 – Placa-mãe

Em geral, vale mais a pena investir numa boa placa-mãe e economizar nos demais componentes, do que o contrário. A qualidade da placa-mãe é de longe mais importante que o desempenho do processador. Você mal vai perceber uma diferença de 20% no *clock* do processador, mas com certeza vai perceber se o seu micro começar a travar ou se a placa de vídeo onboard não tiver um bom suporte no Linux, por exemplo. Ao montar um PC de baixo

custo, economize primeiro no processador, depois na placa de vídeo, som e outros periféricos. Deixe a placa-mãe por último no corte de despesas.

Não se baseie apenas na marca da placa na hora de comprar, mas também no fornecedor. Como muitos componentes entram no país ilegalmente, "via Paraguai", é muito comum que lotes de placas remanufaturadas ou defeituosas acabem chegando ao mercado. Muitas pessoas compram esses lotes, vende por um preço um pouco abaixo do mercado e depois desaparecem. Outras lojas simplesmente vão vendendo placas que sabem ser defeituosas até acharem algum cliente que não reclame. Muitas vezes os travamentos da placa são confundidos com "paus do Windows", de forma que sempre aparece algum desavisado que não percebe o problema. Antigamente existia a polêmica entre as placas com ou sem componentes onboard. Hoje em dia isso não existe mais, pois todas as placas vêm com som e rede onboard. Apenas alguns modelos não trazem vídeo onboard, atendendo ao público que vai usar uma placa 3D offboard e prefere uma placa mais barata ou com mais slots PCI do que com o vídeo que, de qualquer forma, não vai usar. Os conectores disponíveis na placa estão muito relacionados ao nível de atualização do equipamento. Placas atuais incluem conectores PCI Express x16, usados para a instalação de placas de vídeo offboard, slots PCI Express x1 e slots PCI, usados para a conexão de periféricos diversos. Placas antigas não possuem slots PCI Express nem portas SATA, oferecendo no lugar um slot AGP para a conexão da placa de vídeo e duas ou quatro portas IDE para a instalação dos HDs e drives ópticos. Tem-se ainda soquetes para a instalação dos módulos de memória, o soquete do processador, conector para a fonte de alimentação e o painel traseiro, que agrupa os encaixes dos componentes onboard, incluindo o conector VGA ou DVI de vídeo, conectores de som, conector da rede e as portas USB. O soquete (ou *slot*) para o processador é a principal característica da placa-mãe, pois indica com quais processadores ela é compatível.

Por exemplo, não pode ser instalado um Athlon X2 em uma placa soquete A (que é compatível com os antigos Athlons, Durons e Semprons), nem muito menos encaixar um Sempron numa placa soquete 478, destinada aos Pentium 4 e Celerons antigos. O soquete é na verdade apenas um indício de diferenças mais "estruturais" na placa, incluindo o chipset usado, o layout das trilhas de dados, etc. É preciso desenvolver uma placa quase que inteiramente diferente para suportar um novo processador. Existem dois tipos de portas para a conexão do HD: as portas IDE tradicionais, de 40 pinos (chamadas de PATA, de "Parallel ATA") e os conectores SATA (Serial ATA), que são muito menores. Muitas placas recentes incluem um único conector PATA e quatro conectores SATA. Outras incluem as duas portas IDE tradicionais e dois conectores SATA, e algumas já passam a trazer apenas conectores SATA, deixando de lado os conectores antigos. Tudo isso é montado dentro do gabinete, que contém outro componente importante: a fonte de alimentação.

2.1.7- Fonte de alimentação: A função da fonte é transformar a corrente alternada da tomada em corrente contínua (AC) já nas tensões corretas, usadas pelos componentes. Ela serve também como uma última linha de defesa contra picos de tensão e instabilidade na corrente, depois do *nobreak* ou estabilizador. Embora quase sempre relegada a último plano, a fonte é outro componente essencial num PC atual. Com a evolução das placas de vídeo e dos processadores, os PCs consomem cada vez mais energia. Na época dos 486, as fontes mais vendidas tinham 200 watts ou menos, enquanto as atuais têm a partir de 450 watts. Existem ainda fontes de maior capacidade, especiais para quem quer usar duas placas 3D de ponta em SLI, que chegam a oferecer 1000 watts! Evite comprar fontes muito baratas e, ao montar um micro mais potente invista numa fonte de maior capacidade. outro fator importante é o aterramento, porém frequentemente esquecido. O fio terra funciona como uma rota de fuga para picos de tensão provenientes da rede elétrica. A eletricidade flui de uma forma similar à água: vai sempre pelo caminho mais fácil. Sem ter para onde ir, um raio vai

torrar o estabilizador, a fonte de alimentação e, com um pouco mais de azar, a placa-mãe e o resto do micro. O fio terra evita isso, permitindo que a eletricidade escoe por um caminho mais fácil, deixando todo o equipamento intacto.



Figura 10 – Fonte de alimentação

Nas grandes cidades, é relativamente raro que os micros realmente queimem por causa de raios, pois os transformadores e disjuntores oferecem uma proteção razoável. Mas, pequenos picos de tensão são responsáveis por pequenos danos nos pentes de memória e outros componentes sensíveis, danos que se acumulam, comprometendo a estabilidade e abreviando a vida útil do equipamento.

2.1.8- Cooler: Do inglês que significa refrigerador, na informática é o conjunto de dissipação térmica instalado sobre o processador que é responsável pela diminuição do calor. Consiste basicamente em 2 componentes:

1- Microventilador (ventilador de pequena dimensão responsável pelo fluxo de ar). Também conhecido como *FAN* (ventilador em português).

2- Dissipador (peça em cobre ou alumínio responsável pela transferência de calor).

O excesso de calor gerado pelo processador é transferido para o dissipador, este recebe diretamente o ar ambiente impulsionado pela ventoinha que mantém num processo contínuo a baixa temperatura, essencial para o funcionamento adequado do processador. Ligar o computador sem a presença do cooler causa danos irreparáveis ao processador (queima instantânea).

Com o rápido avanço dos processadores e grande variedade de modelos, existem muitos tipos de cooler no mercado, veja como identificar o cooler correto para sua necessidade. Primeiramente identifique o tipo de *socket* utilizado em seu computador e a velocidade do processador. Os *sockets* mais atuais são: 462, 754, 423, 478, 775, 939, AM2. Pode ser encontrado este número no próprio *socket* (suporte) do processador.



Figura 11 – Cooler

2.1.9- Placa Fax/MODEM: Tratando-se de comunicação, nada mais lógico do que usar as linhas telefônicas, largamente disponíveis para realizar a comunicação entre computadores. Porém, usando linhas telefônicas comuns enfrentamos um pequeno problema: os computadores trabalham com sinais digitais, neles qualquer informação será armazenada e processada na forma de 0s ou 1s. As linhas telefônicas por sua vez são analógicas, sendo adequadas para a transmissão de voz, mas não para a transmissão de dados.

Justamente para permitir a comunicação entre computadores utilizando linhas telefônicas comuns, foram criados os modems. Modem é a contração de **modulador e demodulador** e se refere a um aparelho capaz de transformar sinais digitais em sinais analógicos que são transmitidos pela linha telefônica e, em seguida, novamente transformados em sinais digitais pelo modem receptor.

Os modems apresentaram uma notável evolução na última década. Os primeiros modems eram capazes de transmitir apenas 300 bits de dados por segundo, enquanto que os atuais são capazes de manter conexões com velocidades em torno de 56 Kbits por segundo.

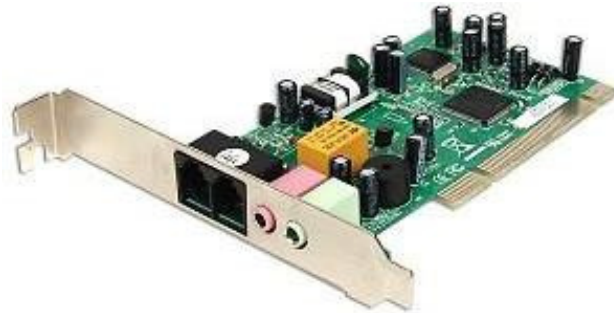


Figura 12 – Placa de Fax/MODEM

2.1.10- Placa de Som: Segundo a Wikipédia, a placa de som é um dispositivo de hardware que envia e recebe sinais sonoros entre equipamentos de som e um computador executando um processo de conversão com um mínimo de qualidade e também para gravação e edição.

Antes que se pensasse em utilizar placas, com processadores dedicados, os primeiros IBM PC/AT já vinham equipados com um dispositivo para gerar som, que se mantém até hoje nos seus sucessores, os *speakers*, pequenos alto-falantes, apesar dos PCs atuais contarem com complexos sistemas de som tridimensional de altíssima resolução.



Figura 13 – Placa de Som

O funcionamento destes dispositivos era, e ainda é, bem primitivo. Um oscilador programável recebe um valor pelo qual dividirá a frequência base, e um flip-flop, liga e

desliga o alto-falante. Não há como controlar o volume, mas isso não impede que ao utilizar-se de recursos de algoritmos bastante complexos, um programador possa conseguir um razoável controle. Tanto o beep inicial que afirma que as rotinas de inicialização do computador foram concluídas com sucesso, quando os beeps informando falhas neste processo, e as músicas dos jogos são gerados do mesmo modo.

2.1.11- Placa de Rede: Uma placa de rede (também chamada adaptador de rede ou NIC) [TORRES, 2010] é um dispositivo de hardware responsável pela comunicação entre os computadores em uma rede. Segundo a Wikipédia, a placa de rede é o hardware que permite aos computadores conversarem entre si através da rede. Sua função é controlar todo o envio e recebimento de dados através da rede.

Cada arquitetura de rede exige um tipo específico de placa de rede; sendo as arquiteturas mais comuns a rede em anel Token Ring e a tipo Ethernet. Além da arquitetura usada, as placas de rede à venda no mercado diferenciam-se também pela taxa de transmissão, cabos de rede suportados e barramento utilizado (onboard, PCI ou USB).

Cabos diferentes exigem encaixes diferentes na placa de rede. O mais comum em placas Ethernet, é a existência de dois encaixes, uma para cabos de par trançado e outro para cabos coaxiais (atualmente muito pouco utilizado, vindo na placa apenas para o conector RJ-45 par trançado). Placas que trazem encaixes para mais de um tipo de cabo são chamadas placas combo. A existência de 2 ou 3 conectores serve apenas para assegurar a compatibilidade da placa com vários cabos de rede diferentes. Naturalmente, você só poderá utilizar um conector de cada vez, sem esquecer também as atuais placas *wireless* (rede sem fio).



Figura 14 – Placa de Rede RJ-45

2.1.12- Flat Cable: O *Flat Cable* (cabo liso em português) [BATTISTI, 2009], é um tipo de cabo responsável pela conexão de dispositivos como HDs, unidades de disquete, unidades de CDs/DVDs entre outros com a placa-mãe. Os cabos *flats* são formados por fios finos que são unidos paralelamente. Esses fios são conhecidos como vias, e é por estas vias que ocorre a transmissão dos dados.

Outra informação importante é que existe mais um tipo de “cabo *flat*”. Que muda de tamanho, quantidade de conectores e vias de acordo com a necessidade. Por exemplo, cabos flat de HDs podem ter 40 ou 80 vias, enquanto o cabo flat do disquete tem apenas 34 vias, sem esquecer dos cabos para drives SATA de 7 vias dos quais 3 são para o aterramento. As outras 4 vias restantes separadas em 2 pares e blindadas são utilizados, sendo um para cada canal de comunicação, que são diferentes e independentes, explicando o modo de operação

full-duplex. Em um dos fios do canal de comunicação, transmite-se o dado propriamente dito, no outro é transmitido uma cópia do dado, mas com o “sinal invertido”, garantindo a integridade dos dados.

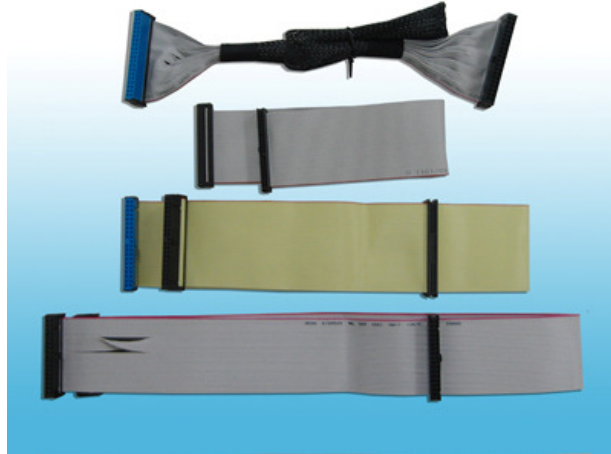


Figura 15 – Cabos *Flats*

2.1.13- Teclado: O teclado de computador é um tipo de periférico utilizado pelo usuário para a entrada manual no sistema de dados e comandos. Possui teclas representando letras, números, símbolos e outras funções, baseado no modelo de teclado das antigas máquinas de escrever. Basicamente, os teclados são projetados para a escrita de textos, onde são usadas para esse meio cerca de 50% delas, embora os teclados sirvam para o controle das funções de um computador e seu sistema operacional. Essas teclas são ligadas a um chip dentro do teclado, responsável por identificar a tecla pressionada e por mandar as informações para o PC. O meio de transporte dessas informações entre o teclado e o computador pode ser sem fio (wireless) ou a cabo (PS/2 e USB).

Os teclados são essencialmente formados por um arranjo de botões retangulares, ou quase retangulares, denominados como teclas. Cada tecla tem um ou mais caracteres impressos ou gravados em baixo relevo em sua face superior, sendo que, aproximadamente, cinquenta por cento das teclas produzem letras, números ou sinais (denominados caracteres). Entretanto, segundo a Wikipédia, em alguns casos, o ato de produzir determinados símbolos requer que duas ou mais teclas sejam pressionadas simultaneamente ou em sequência. Outras teclas não produzem símbolo algum, todavia, afetam o modo como o microcomputador opera ou age sobre o próprio teclado.



Figura 16 – Teclado

Existe uma grande variedade de arranjos diferentes de símbolos nas teclas. Essas características em teclados diferentes surgem porque as diferentes pessoas precisam de um

acesso fácil a símbolos diferentes; tipicamente, isto é, porque elas estão escrevendo em idiomas diferentes, mas existe características de teclado especializadas à matemática, contabilidade, e programas de computação existentes.

O número de teclas em um teclado geralmente varia de 101 a 104 teclas, de certo modo existem até 130 teclas, com muitas teclas programáveis. Também há variantes compactas que têm menos que 90 teclas. Elas normalmente são achadas em *notebooks* ou em computadores de mesa com tamanhos e formas especiais como para canhotos, deficientes físicos, etc.

2.1.14- Mouse: O mouse é um periférico de entrada que, historicamente, se juntou ao teclado como auxiliar no processo de entrada de dados, especialmente em programas com interface gráfica. O rato ou mouse (estrangeirismo, empréstimo do inglês "mouse", que significa "camundongo") tem como função movimentar o cursor (apontador) pela tela do computador. O formato mais comum do cursor é uma seta, contudo, existem opções no sistema operacional e em softwares específicos que permitam a personalização do cursor do mouse. O mouse funciona como um apontador sobre a tela do computador e disponibiliza normalmente quatro tipos de operações: movimento, clique, duplo clique e arrastar e soltar (*drag and drop*).

Existem modelos com um, dois, três ou mais botões cuja funcionalidade depende do ambiente de trabalho e do programa que está a ser utilizado. Claramente, o botão esquerdo é o mais utilizado. O mouse é normalmente ligado ao computador através de uma porta serial, PS2 ou USB (*Universal Serial Bus*). Também existem conexões sem fio, as mais antigas em infravermelho, as atuais em Bluetooth.



Figura 17 – Mouse

Outros dispositivos de entrada competem com o mouse: *touchpads* (usados basicamente em notebooks) e *trackballs*. Também é possível ver o *joystick* como um concorrente, mas eles não são comuns em computadores. É interessante notar que um *trackball* pode ser visto como um mouse de cabeça para baixo.

O mouse por padrão possui pelo menos dois botões, o esquerdo usado para selecionar e clicar (acionar) ícones e o direito realiza funções secundárias como por exemplo exibir as propriedades do objeto apontado. Há ainda na maioria dos mouse um botão *Scroll* (rolagem) em sua parte central, que tem como função principal movimentar a barra de rolagem das janelas.

2.1.15- Monitor: O monitor é um dispositivo de saída do computador, cuja função é transmitir informação ao utilizador através da imagem, estimulando assim a visão.

Os monitores são classificados de acordo com a tecnologia de amostragem de vídeo utilizada na formação da imagem. Atualmente, essas tecnologias são três: CRT , LCD e Plasma e ainda uma mais avançada que altera o cristal líquido por pontos de led, muitas vezes chamados de LCD/Led. À superfície do monitor sobre a qual se projeta a imagem é denominada de tela, ecrã ou écran, dependendo da variação do idioma de seus fabricantes.



Figura 18 – Monitor LCD com tecnologia Led

2.1.16- Impressora: Segundo a Wikipédia, uma impressora ou dispositivo de impressão é um periférico que, quando conectado a um computador ou a uma rede de computadores, tem a função de dispositivo de saída, imprimindo textos, gráficos ou qualquer outro resultado de uma aplicação. Herdando a tecnologia das máquinas de escrever, as impressoras sofreram drásticas mutações ao longo dos tempos:

- a- Impressora de Impacto (matricial):** As impressoras de impacto baseiam-se no princípio da decalcação, ou seja, ao colidir uma agulha ou roda de caracteres contra um fita de tinta dá-se a produção da impressão. As impressoras margarida e impressoras matriciais são exemplos de impressoras de impacto.
- b- Impressora a Jato de Tinta:** Essas impressoras imprimem através de um cartucho de tinta que vai de 2 à 45 ml. Algumas têm uma ótima qualidade de impressão quase se igualando às de Laser. São as impressoras mais utilizadas atualmente.
- c- Impressora Laser:** As impressoras a laser são o topo de gama na área da impressão e seus preços variam enormemente, dependendo do modelo. São o método de impressão preferencial em tipografia e funcionam de modo semelhante ao das fotocopiadoras.
- d- Impressora Térmica:** Embora sejam mais rápidas, mais econômicas e mais silenciosas do que outros modelos de impressoras, as impressoras térmicas praticamente só são utilizadas hoje em dia em aparelhos de fax e máquinas que imprimem cupons fiscais e extratos bancários. O grande problema com este método de impressão é que o papel térmico utilizado desbota com o tempo, obrigando o utilizador a fazer uma fotocópia do mesmo.

- e- **Impressora Térmica a Cera:** uma das melhores resolução, funcionam semelhantes as Lasers, mas em vez de utilizar pó (*toner*), utilizam fitas impregnadas de cera que ao passar no fusor térmico se desprendem e aderem ao papel, devido a utilização das 4 cores primárias, ciano (azul), magenta (rosa), *yellow*(amarela) e preta, uma por vez, torna essa impressora mais lenta e mais cara em relação às lasers, porém a impressão é mais viva e duradoura por ser a prova d'água.



Figura 19 – Impressora a Jato de Tinta com visor LCD

2.1.17- Scanner: Scanner ou Digitalizador é um periférico de entrada responsável por digitalizar imagens, fotos e textos impressos para o computador, um processo inverso ao da impressora. Ele faz varreduras na imagem física gerando impulsos elétricos através de um captador de reflexos. É dividido em duas categorias:

- a- **Digitalizador de Mão:** parecido com um mouse bem grande, no qual deve-se passar por cima do desenho ou texto a ser transferido para o computador. Este tipo não é mais apropriado para trabalhos semi-profissionais devido à facilidade para o aparecimento de ruídos na transferência.
- b- **Digitalizador de Mesa:** parecido com uma fotocopidora, no qual deve colocar o papel e abaixar a tampa para que o desenho ou texto seja então transferido para o computador. Eles fazem a leitura a partir dispositivos de carga dupla.

O digitalizador cilíndrico é o mais utilizado para trabalhos profissionais. Ele faz a leitura a partir de fotomultiplicadores. Sua maior limitação reside no fato de não poderem receber originais não flexíveis e somente digitalizarem imagens e traços horizontais e verticais. Ele tem a capacidade de identificar um maior número de variações tonais nas áreas de máxima e de mínima. Devido aos avanços recentes na área da fotografia digital, já começam a ser usadas câmeras digitais para capturar imagens e texto de livros.



Figura 20 – Scanner de Mesa

2.1.18- Kit Multimídia: O kit multimídia nada mais é que o nome dado ao conjunto de microfone e caixas de som, já que o drive de CD-ROM passou a ser algo obrigatório nos microcomputadores. O microfone é um dispositivo de entrada na qual transforma o som analógico em digital com o auxílio da placa de som, enquanto as caixas de som fazem o trabalho inverso, ou seja, trabalham como dispositivos de saída, levando o conteúdo do computador (digital) para ondas sonoras (analógica) para que o usuário possa escutar.



Figura 21 – Caixas de Som

2.1.19- Webcam: Webcam ou câmera web, é uma câmera de vídeo de baixo custo que capta imagens e as transfere para um computador. Pode ser usada para videoconferência, monitoramento de ambientes, produção de vídeo e imagens para edição, entre outras aplicações. Atualmente existem webcams de baixa ou de alta resolução (acima de 2.0 *megapixels*) e com ou sem microfones acoplados. Algumas webcams vêm com *leds* (diodos emissores de luz), que iluminam o ambiente quando há pouca ou nenhuma luz externa.

A maioria das webcams é ligada ao computador por conexões USB, e a captura de imagem é realizada por um componente eletrônico denominado CCD (*charge-coupled device*), ou dispositivo de carga acoplada (traduzindo para o português).



Figura 22 – WebCam

3- COMPONENTES DA PLACA-MÃE

3.1- BIOS, CMOS e POST

3.1.1- BIOS: O BIOS - *Basic Input Output System* (sistema básico de entrada e saída) contém todo o software básico, necessário para inicializar a placa-mãe, checar os dispositivos instalados e carregar o sistema operacional, o que pode ser feito a partir do HD, CD-ROM, pendrive, ou qualquer outra mídia disponível. O BIOS inclui também o Setup, o software que permite configurar as diversas opções oferecidas pela placa. O processador é programado para procurar e executar o BIOS sempre que o micro é ligado, processando-o da mesma forma que outro software qualquer. É por isso que a placa-mãe não funciona "sozinha", você precisa ter instalado o processador e os pentes de memória para conseguir acessar o Setup.

Por definição, o BIOS é um software, mas, como de praxe, ele fica gravado em um chip alocado na placa-mãe. Na grande maioria dos casos, o chip combina uma pequena quantidade de memória Flash (256, 512 ou 1024 KB), o CMOS, que é composto por de 128 a 256 bytes de memória volátil e o relógio de tempo real. Nas placas antigas era utilizado um chip DIP, enquanto nas atuais é utilizado um chip PLCC (*plastic lead chip carrier*), que é bem mais compacto.

Chip PLCC: Armazena o BIOS da placa-mãe.

3.1.2- CMOS: serve para armazenar as configurações do setup. Como elas representam um pequeno volume de informações, ele é bem pequeno em capacidade. Assim como a memória RAM principal, ele é volátil, de forma que as configurações são perdidas quando a alimentação elétrica é cortada. Por isso, toda placa-mãe inclui uma bateria, que mantém as configurações quando o micro é desligado. A mesma bateria alimenta também o relógio de tempo real (*real time clock*), que, apesar do nome pomposo, é um relógio digital comum, que é o responsável por manter atualizada a hora do sistema, mesmo quando o micro é desligado.

Um caso clássico é tentar fazer um *overclock* muito agressivo e o processador passar a travar logo no início do boot, sem que você tenha chance de entrar no setup e desfazer a alteração. Atualmente basta zerar o setup para que tudo volte ao normal, mas, se as configurações fossem armazenadas na memória Flash, a coisa seria mais complicada. Como todo software, o BIOS possui *bugs*, muitos por sinal. De tempos em tempos, os fabricantes disponibilizam versões atualizadas, corrigindo problemas, adicionando compatibilidade com novos processadores (e outros componentes) e, em alguns casos, adicionando novas opções de configuração no Setup.

Na grande maioria dos casos, o programa também oferece a opção de salvar um backup do BIOS atual antes de fazer a atualização. Esse é um passo importante, pois se algo sair errado, ou você tentar gravar uma atualização para um modelo de placa diferente, ainda restará a opção de reverter o upgrade, regravando o backup da versão antiga. A maioria das placas atuais incorpora sistemas de proteção, que protegem áreas essenciais do BIOS, de forma que, mesmo que acabe a energia no meio da atualização, ou você tente gravar o arquivo errado, a placa ainda preservará as funções necessárias para que você consiga reabrir o programa e gravação e terminar o serviço. Em alguns casos, a placa chega a vir com um

"BIOS de emergência", um chip extra, com uma cópia do BIOS original, que você pode instalar na placa em caso de problemas.

Um truque muito usado era utilizar uma placa-mãe igual, ou pelo menos de modelo similar, para regravar o BIOS da placa danificada. Nesses casos, você dava *boot* com o disquete ou CD de atualização (na placa boa), removia o chip com o BIOS e instalava no lugar o chip da placa danificada (com o micro ligado), dando prosseguimento ao processo de regravação. Dessa forma, você usava a placa "boa" para regravar o BIOS da placa "ruim". Naturalmente, a troca precisava ser feita com todo o cuidado, já que um curto nos contatos podia inutilizar a placa-mãe.

3.1.3- POST: (*Power on self test*, que em português é algo como "Auto-teste de inicialização") é uma sequência de testes ao hardware de um computador, realizada pelo BIOS, responsável por verificar preliminarmente se o sistema se encontra em estado operacional.[WIKIPÉDIA, 2011]. Se for detectado algum problema durante o POST o BIOS emite uma certa sequência de bips sonoros, que podem mudar de acordo com o fabricante da placa-mãe. É o primeiro passo de um processo mais abrangente designado IPL (*Initial Program Loading*), *booting* ou *bootstrapping*.

HD ou FDD Failure – quando esta mensagem ocorre, os motivos que levam o DRIVE a não funcionar podem ser três: o cabo de ligação do drive está mal conectado ou partido; o cabo de alimentação do drive, que sai da fonte, também possui defeito ou está mal conectado e, por último, seu drive não foi citado ou citado incorretamente no CMOS.

Parity erro – quando ocorre tal descrição, mais de uma memória apresenta defeito. Caso contrário, a placa-mãe poderá estar com sérios defeitos.

Non – system disk or disk error – é comum que exista no drive A: um disquete que não possua sistema operacional. Retire este disco e reinicialize a máquina. Se o erro persistir, você deverá transferir ou reinstalar o sistema operacional para o seu disco rígido.

Bad Address 1005B-B00CF H – mais uma vez o problema se refere à memória. A má conexão pode ser resolvida com um aperto no soquete.

No ROM Basic – este tipo de erro requer a utilização de um programa específico para instalação de drive. Neste caso, é comum o utilitário tentar resolver. Este tipo de programa pode ser indicado pelo fabricante do disco rígido.

Sector not found / Undercoverable error – o disco em uso apresenta problema que podem ser tornar irreversíveis. Para evitar maiores perdas, utilize programas como o NORTON UTILITIES para recuperar as informações.

Outra forma utilizada pelo computador para indicar possíveis defeitos é através de sinais sonoros ou BIPS. Cada fabricante de BIOS possui uma codificação própria.

Ao ligar a máquina, está previsto que se ouça um **BIP**. Se isto não acontecer, verifique a fonte de alimentação, a placa mãe ou o emissor de som.

UM BIP: indica que o conjunto está completo, mas será iniciado um teste dos componentes e se não aparecer a inicialização no vídeo, verifique a conexão dos cabos de alimentação e lógicos, placas mal encaixadas nos slots – principalmente a placa de vídeo.

DOIS, TRÊS E QUATRO BIPS: verifique a colocação dos chips de memória na placa. Caso isto não funcione, as memórias devem ser testadas uma a uma. Em última instância, deve ser substituída a placa mãe, caso as memórias estejam em perfeito funcionamento.

CINCO E SETE BIPS: semelhante ao anterior, porém dê ênfase em apertar o processador.

SEIS BIPS: o chip da placa-mãe que controla o teclado está com problemas. Tente fixá-lo melhor, ou tente um novo teclado.

OITO BIPS: a sua placa de vídeo pode estar mal instalada. Caso contrário, deverá ser instalada uma nova.

NOVE BIPS: o BIOS está com sérios problemas e deve ser substituído.

ONZE BIPS: esse erro condiz com a memória CACHE e deve ser testado o processador, já que a memória cache está contida nele.

3.2- Interfaces

3.2.1- Interface Serial: Na comunicação entre um micro e um periférico que esteja distante ou na comunicação de um microcomputador com outro ou qualquer dispositivo digital que permita comunicação externa, como agenda eletrônicas, utilizamos comunicação em série através de uma interface serial.

As interfaces seriais são em números de quatro, sendo utilizado duas portas por método de compartilhamento. Desta forma, ao utilizar portas compartilhadas (mesma interrupção) não é possível o uso simultâneo das mesmas. Exemplo: Um mouse conectado na Serial 1 não pode ser utilizado ao mesmo tempo em que se utiliza um modem através da Serial 3 (ambos estarão no IRQ4).



Figura 23 – Porta Serial DB-9

3.2.2- Interface Paralela: Na interface paralela, basicamente destinada as impressoras e também utilizadas por outros periféricos como o ZIP Drive (drive de disco externo) e Scanner, os dados são transferidos byte a byte. No entanto, a interface tradicional é unidirecional, permitindo somente envio de dados do microcomputador ao periférico,

contando apenas com algumas linhas de controle que informam o estado do periférico ao microcomputador, como falta de papel, no caso de uma impressora.



Figura 24 – Porta Paralela DB-25

3.2.3- Interface IR (*infra red*): Interface utilizada pra conexão de dispositivos com tecnologia de infravermelho evitando a presença de cabos que limitam a distância entre o microcomputador e os mesmos.



Figura 25 – Placa IR

3.2.4- Interface Joystick ou Game: Interface utilizada pra conexão de Joystick para utilização em jogos. Sua presença está se tornando, cada vez mais, restrita às placas de multimídia ou semelhantes.



Figura 26 – Porta Joystick DB-15

3.2.5- Interface USB: Tecnologias equivalentes ao Firewire da Apple, é um padrão de barramento externo ao micro (veio para substituir todas as outras portas como seriais, paralelas e game), para a conexão de periféricos como teclados, monitores de vídeo, impressoras, mouses, joysticks, e outros. Através de um único plug padronizado é possível conectar de forma simultânea até 127 periféricos.



Figura 27 – Porta USB

3.2.6- Interface SCSI: A interface SCSI – *Small Computer System Interfaces* (Interface de Sistema de Computadores Pequenos), possuidora de grande taxa de transferência, não é apenas um padrão de discos rígidos. É um padrão de ligação de

periféricos em geral. A interface SCSI trabalha com grande folga, pois todo o controle de periféricos está no próprio periférico e isto é muito importante. Na verdade não existe uma controladora SCSI, mas apenas uma interface, uma servidora de SCSI, responsável somente pela troca de dados entre o microcomputador e o periférico.

Uma placa SCSI atende até oito ou dezesseis periféricos ao mesmo tempo numa mesma placa, se orientando pelo endereço SCSI que varia de 0 a 7 e determinado fisicamente nos próprios periféricos.



Figura 28 – Conector SCSI

3.2.7- Interface IDE: Para concorrer com o padrão SCSI, onde o controle é feito no próprio disco rígido evitando problemas com ruído, a Western Digital criou um padrão onde a controladora de dados estava integrada na mesma placa dos circuitos de controle do mecanismo do disco rígido. Com isto o problema de ruído, que antes fazia com que a interface pedisse diversas retransmissões de dados por divergência, foi simplesmente eliminado. Esta tecnologia recebeu o nome de IDE – *Integrated Drive Electronics* (Eletrônica de Drive Integrada), que passou a se tornar padrão devido ao relativo baixo custo.

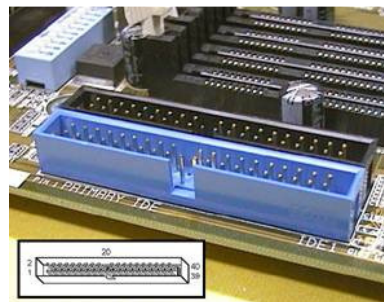


Figura 29 – Conector IDE

Discos com a tecnologia IDE, precisam apenas da conexão ao barramento do microcomputador através de um slot. Este tipo de conexão é a ATA – *AT Attachment* (Ligação AT – ou seja, ligação direto ao barramento ISA). Como muitas placas-mãe não apresentavam o conector para discos rígidos IDE, tornou-se necessário a criação da interface IDE para realizar esta conexão. Esta interface passou a integrar “Multi I/O” (Drive de disquetes, portas seriais, portas paralelas e joystick) já existente naquela época.

3.2.8- Master e Slave – Instalação de dois discos rígidos IDE

Ao instalar dois discos rígidos IDE num mesmo cabo teremos duas controladoras (lembre-se que as controladoras destes discos estão integradas aos mesmos) tentando controlar a interface IDE e disputando quem e quando a utilizará. Sendo assim, através de *jumpers*, desabilita-se a controladora de um dos discos rígidos IDE tornando-o escravo (*slave* – disco D:) do outro que é chamado de mestre (*master* – disco C:). E por não existir uma padronização a respeito, nem todos os discos rígidos IDE podem controlar ou ser controlado por outro disco rígido IDE. O HDD *Autodetect* (auto detecção), no Setup é a maneira mais confiável de testar se a conexão dois discos rígidos IDE estão funcionando corretamente. Como apenas uma controladora estará trabalhando com dois discos rígidos IDE, ela limitará

sua taxa de transferência pelo disco de menor performance. Por esta limitação, analise se compensa manter funcionando um disco antigo.

3.3- Slots de Barramento:

3.3.1- ISA: O ISA foi o primeiro barramento de expansão utilizado em micros PC. Existiram duas versões: os slots de 8 bits, que foram utilizados pelos primeiros PCs e os slots de 16 bits, introduzidos a partir dos micros 286. Embora fossem processadores de 16 bits, os 8088 comunicavam-se com os periféricos externos utilizando um barramento de 8 bits, daí o padrão ISA original também ser um barramento de 8 bits. Inicialmente, o barramento ISA operava a apenas 4.77 MHz, a frequência de *clock* do PC original, mas logo foi introduzido o PC XT, onde tanto o processador quanto o barramento ISA operavam a 8.33 MHz. Com a introdução dos micros 286, o barramento ISA foi atualizado, tornando-se o barramento de 16 bits que conhecemos. Na época, uma das prioridades foi preservar a compatibilidade com as placas antigas, de 8 bits, justamente por isso os pinos adicionais foram incluídos na forma de uma extensão para os já existentes.

Como podem notar na Figura 10, o slot ISA é dividido em duas partes. A primeira, maior, contém os pinos usados pelas placas de 8 bits, enquanto a segunda contém a extensão que adiciona os pinos extra.



Figura 30 – Slot ISA

Um fator que chama a atenção nos slots ISA é o grande número de contatos, totalizando nada menos que 98. Por serem slots de 16 bits, temos apenas 16 trilhas de dados, as demais são usadas para endereçamento, alimentação elétrica, sinal de *clock*, *refresh* e assim por diante. Este esquema mostra a função de cada um dos pinos em um slot ISA. Como podem ver, não é exatamente uma implementação "simples e elegante", mas enfim, funcionava e era o que estava disponível na época.

Cada um destes pinos podia ser controlado individualmente, via software e muitas placas não utilizavam todos os pinos do conector, por isso era comum que periféricos mais simples, como placas de som e modems viessem com alguns dos contatos "faltando". Outra curiosidade é que, justamente por serem fáceis de programar, as controladoras ISA foram as preferidas por programadores que trabalham com automatização e robótica durante muito tempo. Quando as placas-mãe com slots ISA começaram a desaparecer do mercado, alguns chegaram estocá-las. Apesar de toda a complexidade, o barramento ISA é incrivelmente lento. Além de operar a apenas 8.33 MHz, são necessários tempos de espera entre uma transferência e outra, de forma que, na prática, o barramento funciona a apenas metade da frequência nominal.

3.3.2- PCI: O PCI opera nativamente a 33 MHz, o que resulta em uma taxa de transmissão teórica de 133 Mb/s. Entretanto, assim como em outros barramentos, a frequência do PCI está vinculada à frequência de operação da placa-mãe, de forma que, ao fazer *overclock* (ou *underclock*) a frequência do PCI acaba também sendo alterada. Conforme a frequência das placas foi subindo, passaram a ser utilizados divisores cada vez maiores, de forma a manter o PCI operando à sua frequência original. Em uma placa-mãe operando a 133 MHz, a frequência é dividida por 4 e, em uma de 200 MHz, é dividida por 6. Como você pode notar, o barramento PCI tem se tornado cada vez mais lento com relação ao processador e outros componentes, de forma que com o passar do tempo os periféricos mais rápidos migraram para outros barramentos, como o AGP e o PCI-Express. Ou seja, a história se repete, com o PCI lentamente se tornando obsoleto, assim como o ISA há uma década atrás.

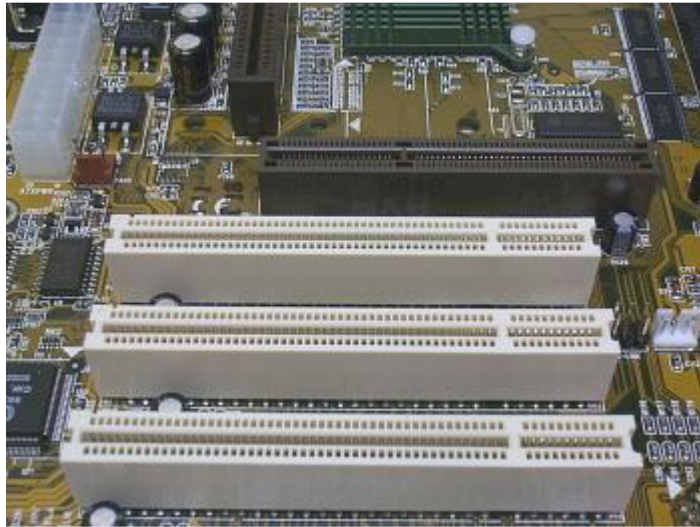


Figura 31 – Slot AGP acima e PCIs logo abaixo

Uma das principais vantagens do PCI sobre os barramentos anteriores é o suporte a *Bus Mastering*. Tanto o EISA quanto o VLB ofereciam um sistema de *Bus Mastering* rudimentar, mas o recurso acabou não sendo muito utilizado por um conjunto de fatores, incluindo as dificuldades no desenvolvimento dos drivers. Apenas com o PCI foi criado um padrão realmente confiável, que foi adotado em massa pelos fabricantes. O *Bus Mastering* é um sistema avançado de acesso direto à memória, que permite que HDs, placas de vídeo e outros periféricos leiam e gravem dados diretamente na memória RAM, deixando o processador livre.

Um dos melhores exemplos é quando o sistema está ocupado inicializando vários programas simultaneamente. O HD precisa transferir grandes quantidades de arquivos e bibliotecas para a memória, a placa de vídeo precisa exibir as telas de progresso e atualizar a tela, enquanto o processador fica ocupado processando as configurações e outras informações necessárias ao carregamento dos programas. Graças ao *Bus Mastering*, um micro atual ainda continua respondendo aos movimentos do mouse e às teclas digitadas no teclado, os downloads e transferências de arquivos através da rede não são interrompidos e assim por diante, muito diferente dos micros antigos que literalmente "paravam" durante transferências de arquivos e carregamento dos programas.

Complementando, temos a questão do *plug-and-play* (ligue e use). Atualmente, estamos acostumados a instalar o dispositivo, instalar os drivers e ver tudo funcionar, mas antigamente as coisas não eram assim tão simples, de forma que o *plug-and-play* foi tema de grande destaque. Tudo começa durante a inicialização do micro. O BIOS envia um sinal de requisição para todos os periféricos instalados no micro. Um periférico PnP é capaz de

responder ao chamado, permitindo ao BIOS reconhecer os periféricos PnP instalados. O passo seguinte é criar uma tabela com todas as interrupções disponíveis e atribuir cada uma a um dispositivo. O sistema operacional entra em cena logo em seguida, lendo as informações disponibilizadas pelo BIOS e inicializando os periféricos de acordo. As informações sobre a configuração atual da distribuição dos recursos entre os periféricos são gravadas em uma área do CMOS chamada de ESCD. Tanto o BIOS (durante o POST) quanto o sistema operacional (durante a inicialização) lêem essa lista e, caso não haja nenhuma mudança no hardware instalado, mantêm suas configurações. Isso permite que o sistema operacional (desde que seja compatível com o PnP) possa alterar as configurações caso necessário. Na maioria das placas-mãe, você encontra a opção "*Reset ESCD*" ou "*Reset Configuration Data*" que, quando ativada, força o BIOS a atualizar os dados da tabela, descartando as informações anteriores. Em muitos casos, isso soluciona problemas relacionados à detecção de periféricos, como, por exemplo, ao substituir a placa de som e perceber que a nova não foi detectada pelo sistema. Nos micros atuais, os conflitos de endereços são uma ocorrência relativamente rara. Na maioria dos casos, problemas de detecção de periféricos, sobretudo no Linux, estão mais relacionados a problemas no ACPI, à falta de drivers ou à falta de suporte por parte dos drivers existentes.

3.3.3- AGP: Embora seja mais recente que o PCI e tenha sido largamente utilizado, o AGP é atualmente um barramento em vias de extinção, devido à popularização do PCI-Express. Desde o final de 2006, placas novas com slots AGP são um item raro, com exceção de algumas placas da PC-Chips, ECS e Phitronics. A idéia central do AGP é ser um barramento rápido, feito sob medida para o uso das placas 3D de alto desempenho. A versão original do AGP foi finalizada em 1996, desenvolvida com base nas especificações do PCI 2.1. Ela operava a 66 MHz, permitindo uma taxa de transferência teórica de 266 Mb/s. Na época, as placas 3D ainda eram bastante primitivas, de forma que ainda não existia uma demanda tão grande por um barramento mais rápido. Por causa disso, o AGP demorou um pouco para se popularizar. O primeiro *chipset* com suporte a ele foi o Intel i440LX, lançado no final de 1997, e a adoção ocorreu de forma gradual durante 1998 e 1999. O padrão AGP inicial não chegou a ser muito usado, pois em 1998 surgiu o padrão AGP 2X, que mantém a frequência de 66 MHz, mas introduz o uso de duas transferências por ciclo (assim como nas memórias DDR), dobrando a taxa de transferência. Em seguida foi introduzido o AGP 4X e o 8X, que realizam, respectivamente, 4 e 8 transferências por ciclo, atingindo taxas de transferência teóricas de 1066 e 2133 Mb/s.

O desempenho de uma placa 3D é fortemente atrelado à velocidade de acesso à memória. Mais de 95% das informações que compõem uma cena 3D de um game atual são texturas e efeitos, que são aplicados sobre os polígonos. As texturas são imagens 2D, de resoluções variadas que são "moldadas" sobre objetos, paredes e outros objetos 3D, de forma a criar um aspecto mais parecido com uma cena real. A velocidade do barramento AGP é importante quando o processador precisa transferir grandes volumes de texturas e outros tipos de dados para a memória da placa de vídeo, quando a memória da placa se esgota e ela precisa utilizar parte da memória principal como complemento e também no caso de placas de vídeo onboard, que não possuem memória dedicada e, justamente por isso, precisam fazer todo o trabalho usando um trecho reservado da memória principal. Naturalmente, tudo isso também pode ser feito através do barramento PCI. O problema é que a baixa velocidade faz com que a queda no desempenho seja cada vez maior, conforme cresce o desempenho da placa de vídeo. Durante muito tempo, fabricantes como a nVidia e a ATI continuaram oferecendo suas placas também em versão PCI, mas a partir de certo ponto, a diferença de desempenho entre as duas versões passou a ser tamanha que, por mais que ainda existisse certa demanda, as placas PCI foram sumariamente descontinuadas.

Outra vantagem do AGP é que o barramento é reservado unicamente à placa de vídeo, enquanto os 133 Mb/s do barramento PCI são compartilhados por todas as placas PCI instaladas. Note que existe uma diferença entre barramento e slot. Uma placa de vídeo onboard é apenas um chip instalado na placa-mãe, ou mesmo um componente integrado diretamente ao *chipset* e não uma "placa" propriamente dita. Mesmo assim, ela pode ser ligada ao barramento AGP, utilizando uma conexão interna. É muito comum ainda que a placa-mãe inclua um *chipset* de vídeo *onboard* e, ao mesmo tempo, um slot AGP, que permite instalar uma placa *offboard*. Neste caso, entretanto, a placa *onboard* é desativada ao instalar uma placa *offboard*, já que o AGP não pode ser compartilhado pelas duas placas. Além da questão da velocidade, existe também a questão da tensão utilizada. O padrão AGP 1.0 previa placas AGP 1X e 2X, que utilizam tensão de 3.3V. O padrão AGP 2.0, finalizado em 1998, introduziu o AGP 4X e a tensão de 1.5V (utilizada pelas placas atuais), quebrando a compatibilidade com o padrão antigo. Placas de vídeo que utilizam sinalização de 3.3V (como a nVidia TNT2, à esquerda na Figura 12) possuem o chanfro do encaixe posicionado ao lado esquerdo, enquanto nas placas que utilizam 1.5V, ele é posicionado à direita. A maioria das placas AGP fabricadas de 2003 em diante são "universais" e podem ser utilizadas tanto nas placas-mãe antigas, com slots de 3.3V, quanto nas placas com slots de 1.5V. Elas possuem os dois chanfros (como a ATI Radeon, à direita na Figura 12), o que permite que sejam encaixadas em qualquer slot.

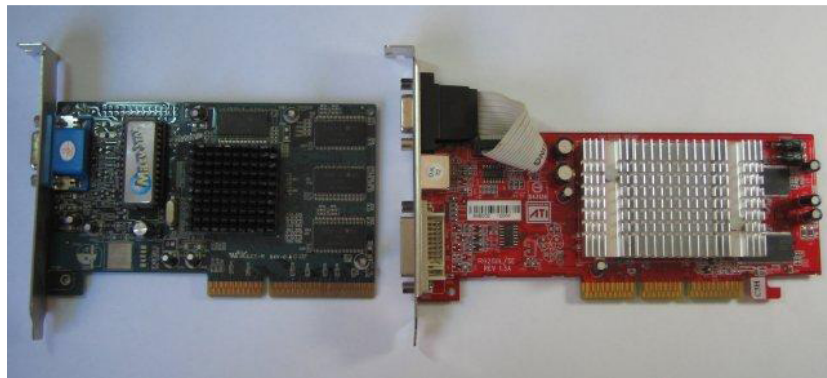


Figura 32 – Placa AGP de 3.3V e placa AGP universal

A mesma distinção existe no caso das placas-mãe. Placas antigas, que utilizam slots de 3.3V possuem o chanfro à esquerda, enquanto as placas com slots de 1.5V utilizam o chanfro posicionado à direita, como nestes dois exemplos: Placa com slot AGP de 3.3V e placa com slot de 1.5V. Existem ainda placas com slots AGP universais, em que o slot não possui chanfro algum e por isso permite a instalação de qualquer placa. Nesse caso a placa-mãe é capaz de detectar a tensão utilizada pela placa e fornecer a tensão adequada. Elas são mais raras, pois a necessidade de instalar tanto os circuitos reguladores para 1.5V quanto para 3.3V, encarece a produção, pois existem muitos casos de incompatibilidades entre placas de vídeo AGP de fabricação mais recente e placas-mãe antigas (e vice-versa), mesmo em casos em que os encaixes são compatíveis. Além dos problemas relacionados a deficiências nos drivers e incompatibilidade por parte do BIOS, existem problemas relacionados à alimentação elétrica, onde a placa de vídeo não indica corretamente qual é a tensão utilizada (fazendo com que a placa-mãe utilize 1.5V para uma placa que trabalhe com 3.3V, por exemplo) ou que a placa-mãe não seja capaz de alimentar a placa de vídeo com energia suficiente. Esse último caso é o mais comum, já que as placas AGP mais recentes consomem muito mais energia que as antigas.

O padrão AGP 3.0 inclui como pré-requisito que a placa-mãe seja capaz de fornecer 41 watts de energia para a placa de vídeo. O padrão AGP 2.0 fala em 25 watts,

enquanto muitas placas antigas fornecem ainda menos. Com a corrida armamentista, entre a nVidia e a ATI, o *clock* e, conseqüentemente o consumo elétrico das placas de vídeo cresceu de forma exponencial. Já se foi o tempo em que a placa de vídeo utilizava um simples dissipador passivo e consumia menos de 10 watts. Muitas das placas atuais superam a marca dos 100 watts de consumo e algumas chegam a ocupar o espaço equivalente a dois slots da placa-mãe devido ao tamanho do cooler, como no caso da ATI X850.

3.3.4- PCI Express: Ao longo da história da plataforma PC, tivemos uma longa lista de barramentos, começando com o ISA de 8 bits, usado nos primeiros PCs, passando pelo ISA de 16 bits, MCA, EISA, e VLB, até finalmente chegar no barramento PCI, que sobrevive até os dias de hoje. O PCI trouxe recursos inovadores (para a época), como o suporte a *plug-and-play* e *bus mastering*. Comparado com os barramentos antigos, o PCI é bastante rápido. O problema é que ele surgiu no começo da era Pentium, quando os processadores ainda trabalhavam a 100 MHz. Hoje em dia temos processadores na casa dos 3 GHz e ele continua sendo usado, com poucas melhorias. Por ser compartilhado entre todos os dispositivos ligados a ele, o barramento PCI pode ser rapidamente saturado, com alguns dispositivos rápidos disputando toda a banda disponível.

O barramento se torna então um gargalo, que limita o desempenho global do PC. A fase mais negra da história do barramento PCI foi durante a época das placas soquete 7 (processadores Pentium, Pentium MMX, K6 e 6x86), quando o barramento PCI era o responsável por praticamente toda a comunicação entre os componentes do micro, incluindo todos os periféricos, a comunicação entre a ponte norte e ponte sul do *chipset*, as interfaces IDE, etc. Até mesmo o antigo barramento ISA era ligado ao PCI através do PCI-to-ISA *bridge* (ponte PCI-ISA), um controlador usado nos *chipsets* da época. Isso fazia com que o barramento ficasse incrivelmente saturado, limitando severamente o desempenho do micro.

Eram comuns situações onde o desempenho do HD era limitado ao rodar games 3D, pois a placa de vídeo saturava o barramento, não deixando espaço suficiente para os demais componentes. A história começou a mudar com o aparecimento do barramento AGP. Ele desafogou o PCI, permitindo que a placa de vídeo tivesse seu próprio barramento rápido de comunicação com o *chipset*. O AGP matou dois coelhos com uma cajadada só, pois permitiu o aparecimento de placas 3D absurdamente mais rápidas e desafogou a comunicação entre os demais componentes.

Rapidamente todas as placas de vídeo passaram a utilizá-lo, com os fabricantes oferecendo versões PCI apenas dos modelos mais simples. O passo seguinte foi a criação de barramentos dedicados pra a comunicação entre os diversos componentes do *chipset* (como o DMI, usado em *chipsets* Intel, e o *HyperTransport*), fazendo com que as interfaces IDE ou SATA e outros componentes também ganhassem seu canal exclusivo. O PCI passou então a ser exclusividade das próprias placas PCI.

O problema é que, mesmo desafogado, o PCI é muito lento para diversas aplicações. É lento demais para ser utilizado por placas de rede Gigabit Ethernet (embora seja suficiente na teoria, na prática a história é um pouco diferente, devido ao compartilhamento da banda), por placas SCSI modernas, ou mesmo por placas RAID e controladoras SATA. Além disso, os slots PCI utilizam um número muito grande de trilhas na placa-mãe, o que é dispendioso para os fabricantes. Existiram tentativas de atualização do PCI, como o PCI de 64 bits, o PCI de 66 MHz e o PCI-X, que além de ser um barramento de 64 bits, trabalha a 133 MHz, resultando num barramento de 1024 Mb/s. Em termos de velocidade, o PCI-X supriria as necessidades dos periféricos atuais, o problema é que, devido ao grande número de contatos e ao tamanho físico dos slots, ele acaba sendo um barramento muito dispendioso e imprático, que ficou relegado aos servidores topo de linha. O PCI Express, ou PCIe, é um

barramento serial, que conserva pouco em comum com os barramentos anteriores. Graças a isso, ele acabou se tornando o sucessor não apenas do PCI, mas também do AGP.

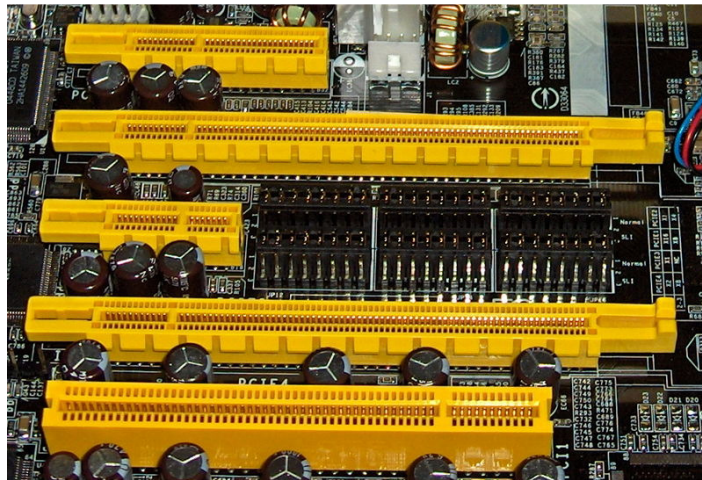


Figura 33 – Slots PCI-e

3.3.5- Chipsets e Placas: Nos primeiros PCs, os chips controladores da placa-mãe ficavam espalhados em diversos pontos da placa. Não é preciso dizer que este design não era muito eficiente, já que mais componentes significam mais custos. Com o avanço da tecnologia, os circuitos passaram a ser integrados em alguns poucos chips. Isso trouxe duas grandes vantagens: a primeira é que, estando mais próximos, os componentes podem se comunicar a uma velocidade maior, permitindo que a placa-mãe seja capaz de operar a frequências mais altas. A segunda é a questão do custo, já que produzir dois chips (mesmo que mais complexos) sai mais barato do que produzir vinte.

Muitas vezes, temos a impressão de que novas tecnologias (sobretudo componentes miniaturizados) são mais caras, mas, na maior parte dos casos, o que acontece é justamente o contrário. Produzir chips utilizando uma técnica de 45 nanômetros é mais barato do que produzir utilizando uma técnica antiga, de 90 ou 180 nanômetros, pois transistores menores permitem produzir mais chips por *wafer*, o que reduz o custo unitário.

Usando uma técnica de 180 nanômetros (0.18 micron), temos transistores 16 vezes maiores que ao utilizar uma técnica de 45 nanômetros. Isso significa que, utilizando aproximadamente o mesmo volume de matéria-prima e mão de obra, é possível produzir quase que 16 vezes mais chips. É bem verdade que migrar para novas tecnologias implica um grande custo inicial, já que a maior parte do maquinário precisa ser substituído. Os fabricantes aproveitam o impulso consumista do público mais entusiasta para vender as primeiras unidades muito mais caro (o que cria a impressão de que a nova tecnologia é mais cara), mas, uma vez que os custos iniciais são amortizados, os produtos da nova geração acabam custando o mesmo, ou menos que os anteriores, mesmo incluindo mais funções.

A grande maioria dos *chipsets* segue o projeto tradicional, onde as funções são divididas em dois chips, chamados de ponte norte (*north bridge*) e ponte sul (*south bridge*). Nos últimos anos essa designação anda um pouco fora de moda, com os fabricantes adotando nomes pomposos, mas ainda pode ser utilizada como uma definição genérica. A ponte norte é o chip mais complexo, que fica fisicamente mais próximo do processador. Ele incorpora os barramentos "rápidos" e as funções mais complexas, incluindo o controlador de memória, as linhas do barramento PCI Express, ou o barramento AGP, além do *chipset* de vídeo onboard, quando presente. As placas para processadores AMD de 64 bits não possuem o controlador de memória, já que ele foi movido para dentro do processador [TORRES, 2010]. Nas placas atuais, a ponte norte do *chipset* é sempre coberta por um dissipador metálico, já que o chip

responde pela maior parte do consumo elétrico e, consequentemente, da dissipação de calor da placa-mãe. Em alguns casos, os fabricantes chegam a utilizar coolers ou até mesmo *heat-pipes* para refrigerá-lo.



Figura 34 – Ponte norte do chipset (com o dissipador removido)

A ponte sul é invariavelmente um chip menor e mais simples que o primeiro. Nas placas atuais ela incorpora os barramentos mais lentos, como o barramento PCI, portas USB, SATA e IDE, controladores de som e rede e também o controlador Super I/O, que agrupa portas "de legado", como as portas seriais e paralelas, porta para o drive de disquete e portas do teclado e mouse (PS/2).



Figura 35 – Ponte sul

É comum que os fabricantes adicionem funções adicionais ou substituam componentes disponíveis na ponte sul, incluindo controladores externos. Com isso, podem ser adicionadas portas SATA ou IDE adicionais, o controlador de áudio pode ser substituído por outro de melhor qualidade ou com mais recursos, uma segunda placa de rede *onboard* pode ser adicionada e assim por diante. Entretanto, com pouquíssimas exceções, as funções da ponte norte do *chipset* não podem ser alteradas. Não é possível adicionar suporte a mais linhas PCI Express ou aumentar a quantidade de memória RAM suportada (por exemplo) adicionando um chip externo. Estas características são definidas ao escolher o *chipset* no qual a placa será baseada. Embora incorpore mais funções (em número) as tarefas executadas pela ponte sul são muito mais simples e os barramentos ligados a ela utilizam menos trilhas de dados. Normalmente os fabricantes utilizam as tecnologias de produção mais recente para

produzir a ponte norte, passando a produzir a ponte sul utilizando máquinas ou fábricas mais antigas. No caso de um fabricante que produz de tudo, como a Intel ou a AMD, é normal que existam três divisões.

Novas técnicas de produção são usadas para produzir processadores, a geração anterior passa a produzir *chipsets* e chips de memória, enquanto uma terceira continua na ativa, produzindo chips menos importantes e controladores diversos. Isso faz com que o preço dos equipamentos seja mais bem amortizado. No final, o maquinário obsoleto (a quarta divisão) ainda acaba sendo vendido para fabricantes menores, de forma que nada seja desperdiçado. Por exemplo, o chip MCH (ponte norte) do *chipset* P35, lançado pela Intel em julho de 2007, é ainda produzido em uma técnica de 0.09 micron, a mesma utilizada na produção do Pentium 4 com *core Prescott*, cuja produção foi encerrada mais de um ano antes.

O chip ICH9 (ponte sul), por sua vez, é ainda produzido utilizando uma técnica de 0.13 micron, a mesma usada no Pentium 4 com *core Northwood* e no Pentium III com *core Tualatin*. A diferença na técnica de produção é justificável pela diferença de complexidade entre os dois chips. Enquanto o MCH do P35 possui 45 milhões de transistores (mais que um Pentium 4 *Willamette*, que possui apenas 42 milhões), o ICH9 possui apenas 4.6 milhões, quase 10 vezes menos. Nos antigos *chipsets* para placas soquete 7 e slot 1, como o Intel i440BX e o VIA Apollo Pro, a ligação entre a ponte norte e ponte sul do *chipset* era feita através do barramento PCI. Isso criava um grande gargalo, já que ele também era utilizado pelas portas IDE e quase todos os demais periféricos.

Nas placas atuais, a ligação é feita através de algum barramento rápido (muitas vezes proprietário) que permite que a troca de informações seja feita sem gargalos. Não existe uma padronização para a comunicação entre os dois chips, de forma que (com poucas exceções) os fabricantes de placas-mãe não podem utilizar a ponte norte de um *chipset* em conjunto com a ponte sul de outro, mesmo que ele seja mais barato ou ofereça mais recursos. O *chipset* é de longe o componente mais importante da placa-mãe.

Excluindo o *chipset*, a placa-mãe não passa de um emaranhado de trilhas, conectores, reguladores de tensão e controladores diversos. Placas que utilizam o mesmo *chipset* tendem a ser muito semelhantes em recursos, mesmo quando fabricadas por fabricantes diferentes. Devido a diferenças no barramento e outras funções, o *chipset* é sempre atrelado a uma família de processadores específica. Não é possível desenvolver uma placa-mãe com um *chipset* AMD que seja também compatível com processadores Intel, por exemplo.

3.4- Tipos de Memórias

3.4.1- Memórias SDRAM: As memórias SDRAM (*Synchronous Dynamic RAM*) por sua vez, são capazes de trabalhar sincronizadas com os ciclos da placa-mãe, sem tempos de espera. Isso significa que a temporização das memórias SDRAM é sempre de uma leitura por ciclo. Veja que o primeiro acesso continua tomando vários ciclos, pois nele é necessário realizar o acesso padrão, ativando a linha (RAS) e depois a coluna (CAS). Apenas a partir do segundo acesso é que as otimizações entram em ação e a memória consegue realizar uma leitura por ciclo, até o final da leitura.

O *burst* de leitura pode ser de 2, 4 ou 8 endereços e existe também o modo "*full page*" (disponível apenas nos módulos SDRAM), onde o controlador pode especificar um número qualquer de endereços a serem lidos sequencialmente, até um máximo de 512, ou seja, em situações ideais, pode ser possível realizar a leitura de 256 setores em 260 ciclos!

Só para efeito de comparação, se fossem usadas memórias regulares, com tempos de acesso similares, a mesma tarefa tomaria pelo menos 1280 ciclos. Outra característica que

ajuda as memórias SDRAM a serem mais rápidas que as EDO e FPM é a divisão dos módulos de memória em vários bancos. Um módulo DIMM pode ser formado por 2, 4, ou mesmo 8 bancos de memória, cada um englobando parte dos endereços disponíveis. Apenas um dos bancos pode ser acessado de cada vez, mas o controlador de memória pode aproveitar o tempo de ociosidade para fazer algumas operações nos demais, como executar os ciclos de *refresh* e também a pré-carga dos bancos que serão acessados em seguida.

Nos módulos EDO e FPM, todas essas operações precisam ser feitas entre os ciclos de leitura, o que toma tempo e reduz a frequência das operações de leitura.

A partir da memória SDRAM, tornou-se desnecessário falar em tempos de acesso, já que a memória trabalha de forma sincronizada em relação aos ciclos da placa-mãe. As memórias passaram então a ser rotuladas de acordo com a frequência em que são capazes de operar. No caso das memórias SDRAM temos as memórias PC-66, PC-100 e PC-133, no caso das DDR temos as PC-200, PC-266, PC-333, PC-400 (e assim por diante), enquanto nas DDR2 temos as PC-533, PC-666, PC-800, PC-933, PC-1066 e PC-1200.

Um módulo de memória PC-133 deve ser capaz de operar a 133 MHz, fornecendo 133 milhões de leituras por segundo. Entretanto, essa velocidade é atingida apenas quando o módulo realiza um *burst* de várias leituras. O primeiro acesso continua levando 5, 6 ou mesmo 7 ciclos da placa-mãe, como nas memórias antigas, ou seja, o fato de ser um módulo PC-100 não indica que o módulo possui um tempo de acesso de 10 ns ou menos (nem mesmo os módulos DDR2 atuais atingem essa marca). Pelo contrário, a maioria dos módulos PC-100 trabalhavam com tempos de acesso de 40 ns. Graças as otimizações citadas, as leituras podiam ser paralelizadas, de forma que no final o módulo suportasse *bursts* de leitura onde, depois de um lento ciclo inicial, o módulo conseguia realmente entregar 64 bits de dados a cada 10 ns.

Independentemente da frequência de operação, temos também os módulos CL2 e CL3, onde o "CL" é abreviação de "*CAS latency*", ou seja, o tempo de latência relacionado ao envio do valor CAS, durante o primeiro acesso de cada *burst*. Em módulos CL2, o envio do valor CAS toma 2 ciclos, enquanto nos CL3 toma 3 ciclos. A eles, somamos um ciclo inicial e mais dois ciclos relacionados ao envio do valor RAS, totalizando 5 (nos módulos CL2) ou 6 (nos CL3) ciclos para o acesso inicial. A diferença acaba sendo pequena, pois os acessos seguintes demoram sempre apenas um ciclo.

Apesar disso, os módulos CL2 trabalham com tempos de acesso um pouco mais baixos e por isso suportam melhor o uso de frequências mais altas que o especificado, dando mais margem para *overclock*. Então desde as memórias regulares, até as SDRAM, foi possível multiplicar a velocidade das memórias sem fazer alterações fundamentais nas células, que continuam seguindo o mesmo projeto básico, com um transistor e um capacitor para cada bit armazenado. Desde a década de 80, as reduções nos tempos de acesso foram apenas incrementais, acompanhando as melhorias nas técnicas de fabricação.

O que realmente evoluiu com o passar do tempo foram os circuitos em torno dos módulos, que otimizaram o processo de leitura, extraindo mais e mais performance. Chegando então às memórias DDR, DDR2 e DDR3 usadas atualmente, que levam este processo crescente de otimização a um novo nível.

3.4.2- Memórias DDR2: Em 2005, quando os primeiros módulos DDR2-533 chegaram ao mercado, eles rapidamente ganharam a fama de "lentos", pois eram comparados a módulos DDR-400 ou DDR-466, que já estavam entrincheirados. Embora um módulo DDR2 ganhe de um DDR da mesma frequência em todos os quesitos (um DDR2-800 contra um DDR-400, por exemplo), o mesmo não acontece se comparamos módulos de frequências diferentes. Um DDR2-533 opera a apenas 133 MHz, por isso acaba realmente perdendo para um DDR-400 (200 MHz) na maioria das aplicações, pois a ganho de realizar 4 operações por

ciclo acaba não sendo suficiente para compensar a diferença na frequência de operação das células de memória. Vale lembrar que um módulo DDR2-533 trabalha com tempos de latência similares a um módulo DDR-266.

Dependendo da fonte, você pode ler tanto que as memórias DDR2 operam ao dobro da frequência que as DDR quanto que elas realizam quatro transferências por ciclo em vez de duas. Nenhuma das explicações está errada, mas ambas estão incompletas. Como se percebe, as células de memória continuam trabalhando na mesma frequência das memórias SDR e DDR, mas os *buffers* de entrada e saída, responsáveis por ler os dados, passaram a operar ao dobro da frequência. É justamente esta frequência que é "vista" pelo restante do sistema, de forma que a maioria dos programas de diagnóstico mostra a frequência dobrada usada pelos circuitos de entrada e não a frequência real das células de memória.

Presumindo que o módulo DDR2 do exemplo operasse a 100 MHz, teríamos as células de memória ainda operando na mesma frequência, mas agora entregando 4 leituras de setores sequenciais por ciclo. Os *buffers* e o barramento de dados operam agora a 200 MHz, de forma que as 4 leituras podem ser enviadas em 2 ciclos, com duas transferências por ciclo. Os dois ciclos do barramento são realizados no mesmo espaço de tempo que apenas um ciclo das células de memória.

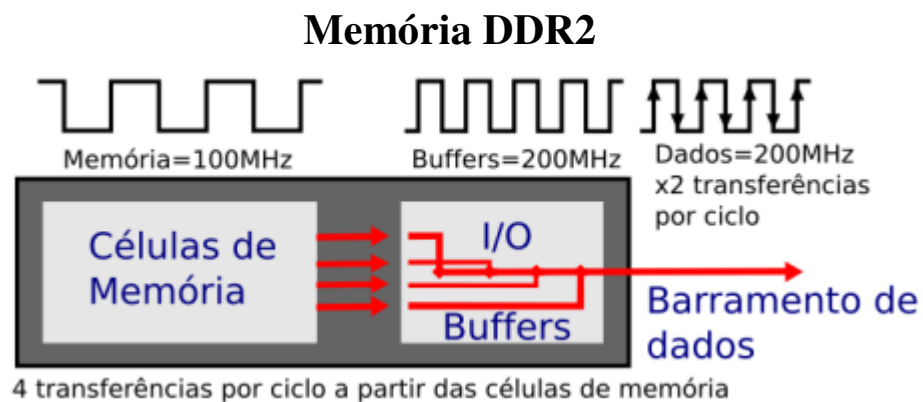


Figura 36 – Ciclos DDR2

As células de memória podem ser grosseiramente comparadas a uma planilha eletrônica, com inúmeras linhas e colunas. Não existe uma grande dificuldade em ler vários endereços diferentes simultaneamente, desde que o fabricante consiga desenvolver os circuitos de controle necessários. Graças a isso, o desenvolvimento das memórias tem sido focado em realizar mais leituras pro ciclo, combinada com aumentos graduais nas frequências de operação.

Quando as memórias DIMM surgiram, ainda na época do Pentium II, os módulos mais rápidos operavam a 100 MHz (os famosos módulos PC-100). Atualmente temos chips de memórias que ultrapassam os 300 MHz que, combinados com as 4 leituras por ciclo, resultam em módulos com transferência teórica de até 9.6 Gb/s: DDR2-533 (133 MHz) = PC2-4200 DDR2-667 (166 MHz) = PC2-5300 DDR2-800 (200 MHz) = PC2-6400 DDR2-933 (233 MHz) = PC2-7500 DDR2-1066 (266 MHz) = PC2-8500 DDR2-1200 (300 MHz) = PC2-9600.

3.4.3- Memórias DDR3: O lançamento das memórias DDR2 teve um impacto diferente para a Intel e a AMD. Para a Intel, a migração para as memórias DDR2 foi mais simples, já que o controlador de memória é incluído no *chipset*, de forma que aderir a uma nova tecnologia demanda apenas modificações nos *chipsets* e placas.

A Intel oferece suporte a memórias DDR2 em seus *chipsets* desde o i915P, lançado em 2004. Inicialmente, os *chipsets* ofereciam tanto suporte a memórias DDR quanto DDR2, de forma que ficava a cargo do fabricante escolher qual padrão utilizar. Existem inclusive placas híbridas, que suportam ambos os padrões, como a ECS 915P-A, que possuem dois slots de cada tipo, permitindo que você escolha qual padrão utilizar.

A partir de certo ponto, entretanto, as memórias DDR2 caíram de preço e quase todas as placas com soquete 775 passaram a vir com suporte exclusivo a memórias DDR2. Para a AMD, a mudança foi mais tortuosa, já que o Athlon 64 e derivados utilizam um controlador de memória embutido diretamente no processador, desenvolvido de forma a minimizar os tempos de acesso. Por um lado isto é bom, pois oferece um ganho real de desempenho, mas por outro é ruim, pois qualquer mudança no tipo de memória usado demanda mudanças no processador e no soquete usado. Foi justamente isso que aconteceu quando a AMD decidiu fazer a migração das memórias DDR para as DDR2.

Além das mudanças internas no processador e controlador de memória, o soquete 754 foi substituído pelo soquete 939 e em seguida pelo AM2, quebrando a compatibilidade com as placas antigas. Com esta adoção por parte da AMD, a procura (e consequentemente a produção) das memórias DDR2 aumentou bastante, fazendo com que os preços passassem a cair rapidamente. A partir do final de 2006, os preços dos módulos de memória DDR2 (nos EUA) caíram a ponto de passarem a ser mais baratos que os módulos DDR regulares. Como sempre, as mudanças chegam ao Brasil com alguns meses de atraso, mas a partir do início de 2007 as memórias DDR2 passaram a ser encontradas por preços inferiores às DDR por aqui também. Como sempre, módulos de alto desempenho (como os *Corsair Dominator*) chegam a custar até duas vezes mais caro, mas em se tratando de memórias genéricas, os preços caíram muito.

Inicialmente, os módulos DDR3 foram lançados em versão DDR3-1066 (133 MHz x 8) e DDR3-1333 (166 MHz x 8), seguidos pelo padrão DDR3-1600 (200 MHz x 8). Os três padrões são também chamados de (respectivamente) PC3-8500, PC3-10667 e PC3-12800, nesse caso dando ênfase à taxa de transferência teórica. Para evitar isso, os módulos DDR3 incluem um sistema integrado de calibragem do sinal, que melhora de forma considerável a estabilidade dos sinais, possibilitando o uso de tempos de latência mais baixos, sem que a estabilidade seja comprometida.



Figura 37 – Memória DDR3 da Corsair

Os módulos DDR3 utilizam também 8 bancos em vez de 4, o que ajuda a reduzir o tempo de latência em módulos de grande capacidade. Elas também trouxeram uma nova redução na tensão usada, que caiu para apenas 1.5V, ao invés dos 1.8V usados pelas memórias DDR2. A redução na tensão faz com que o consumo elétrico dos módulos caia proporcionalmente, o que os torna mais atrativos para os fabricantes de notebooks.

As memórias DDR2 demoraram quase 3 anos para se popularizarem desde a introdução do *chipset* i915P, em 2004. As memórias DDR3 estão passando por um caminho similar, com os módulos inicialmente custando muito mais caro e caindo ao mesmo nível de

preço dos módulos DDR2. Não existe nada de fundamentalmente diferente nos módulos DDR3 que os torne mais caros de se produzir, o preço é determinado basicamente pelo volume de produção. Assim como no caso das memórias DDR2, a maior taxa de transferência oferecida pelas memórias DDR3 resulta em um ganho relativamente pequeno no desempenho global do sistema, de apenas 1 a 3% na maioria dos aplicativos, por isso não vale a pena pagar muito mais caro por módulos DDR3 ao comprar. Enquanto eles estiverem substancialmente mais caros, continue comprando (e indicando) placas com suporte a módulos DDR2.

4- DISCOS RAID

4.1- O que é RAID?

Como muitos já sabem a sigla RAID significa "*Redundant Array of Independent Disks*" (Conjunto Redundante de Discos Independentes). Mas você sabe como é que surgiu essa tecnologia???

Segundo o site Brasil Educador [EDUCADOR, 2011], uma equipe de pesquisadores da Universidade de Berkeley na Califórnia desenvolveram um estudo definindo o RAID e os seus níveis (inicialmente 5). O RAID surgiu como um método de substituir um único disco grande e muitíssimo caro na época por vários menores e com custo muito mais baixo. O princípio básico envolvido nessa tecnologia é uma teoria muito simples: através da combinação de uma matriz formada por discos "pequenos", um administrador poderá gravar dados com redundância para prover tolerância a falhas em um sistema ou dividi-los para aumentar a performance.

Existem basicamente dois tipos de RAID, um baseado em hardware e outro em software, o em hardware que é o mais difundido, este não depende de nenhum sistema operacional para ter um desempenho satisfatório.

É importante notar que o RAID foi desenvolvido há mais de 15 anos e que foi originalmente desenvolvido para discos rígidos SCSI e só recentemente “herdado” pelos discos rígidos SATA. Então não espere que todos os níveis do RAID sejam encontrados em discos SATA, mesmo porque muitos dos níveis abaixo explicados não são utilizados nem mesmo em discos SCSI, alguns desses níveis por não serem tão úteis e outros por serem economicamente inviáveis.

4.2- Os Níveis do Raid

O RAID pode trabalhar de varias maneiras distintas cada uma com funções diferentes. Essas “maneiras” que o RAID trabalha são conhecidas como Níveis de RAID.

4.2.1- RAID Nível 0: Esse nível também é conhecido como “Striping” ou “Fracionamento”. No RAID 0 os dados são divididos em pequenos segmentos e distribuídos entre os diversos discos disponíveis, o que proporciona alta performance na gravação e leitura de informações, porém não oferece redundância, ou seja, não é tolerante a falhas. O aumento da performance no RAID 0 é obtido porque se vários dados fossem gravados em um único disco esse processo aconteceria de forma “Sequencial” já nesse nível os dados são distribuídos entre os discos ao mesmo tempo.

O RAID 0 pode ser usado para estações de alta performance (CAD, tratamento de imagens e vídeos), porém não é indicado para sistemas de missão-crítica.

4.2.2- RAID Nível 1: O nível 1 também é conhecido como "Mirror", “Duplexing” ou “Espelhamento”.

No RAID 1 os dados são gravados em 2 ou mais discos ao mesmo tempo, oferecendo portanto redundância dos dados e fácil recuperação, com proteção contra falha em disco. Uma característica do RAID 1 é que a gravação de dados é mais lenta, pois é feita duas

ou mais vezes. No entanto a leitura é mais rápida, pois o sistema pode acessar duas fontes para a busca das informações.

O RAID 1 pode ser usado para Servidores pelas características de ter uma leitura muito rápida e tolerância à falhas.

4.2.3- RAID Nível 0+1: O RAID 0+1 é uma combinação dos níveis 0 (striping) e 1 (mirroring). No RAID 0+1 os dados são divididos entre os discos e duplicados para os demais discos. Assim temos uma combinação da performance do RAID 0 com a tolerância à falhas do RAID 1. Para a implantação do RAID 0+1 são necessários no mínimo 4 discos o que torna o sistema um pouco caro.

Ao contrario do que muitos pensam, o RAID 0+1 não é o mesmo do RAID 10. Quando um disco falha em RAID 0+1 o sistema se torna basicamente um RAID 0.

O RAID 0+1 pode ser utilizado em estações que necessitam de alta performance com redundância como aplicações CAD e edição de vídeo e áudio.

4.2.4- RAID Nível 2: O nível 2 também é conhecido como “Monitoring”. O RAID 2 é direcionado para uso em discos que não possuem detecção de erro de fábrica, pois “adapta” o mecanismo de detecção de falhas em discos rígidos para funcionar em memória. O RAID 2 é muito pouco usado uma vez que todos os discos modernos já possuem de fábrica a detecção de erro.

4.2.5- RAID Nível 3: No RAID 3 os dados são divididos (em nível de bytes) entre os discos enquanto a paridade é gravada em um disco exclusivo. Como todos os bytes tem a sua paridade (acréscimo de 1 bit para identificação de erros) gravada em um disco separado é possível assegurar a integridade dos dados para recuperações necessárias. O RAID 3 também pode ser utilizado para Servidores e sistemas de missão critica.

4.2.6- RAID Nível 4: O RAID 4 é muito parecido com o nível 3. A diferença é que além da divisão de dados (em blocos e não bytes) e gravação da paridade em um disco exclusivo nesse nível, permite que os dados sejam reconstruídos em tempo real utilizando a paridade calculada entre os discos. Além disso, a paridade é atualizada a cada gravação, tornado-a muito lenta. O RAID 4 pode ser utilizado para sistemas que geram arquivos muito grandes como edição de vídeo, porque a atualização da paridade a cada gravação proporciona maior confiabilidade no armazenamento.

4.2.7- RAID Nível 5: O RAID 5 é semelhante ao nível 4, exceto o fato de que a paridade não é gravada em um disco exclusivo para isso e sim distribuída por todos os discos da matriz. Isso faz com que a gravação de dados seja mais rápida, porque não existe um disco separado do sistema gerando um “Gargalo”, porém como a paridade tem que ser dividida entre os discos a performance é um pouco menor que no RAID 4. O RAID 5 é amplamente utilizado em Servidores de grandes corporações por oferecer uma performance e confiabilidade muito boa em aplicações não muito pesadas. E normalmente são utilizados 5 discos para aumento da performance.

4.2.8- RAID Nível 6: O RAID 6 é basicamente um RAID 5 porém com dupla paridade.

O RAID 6 pode ser utilizado para sistemas de missão critica aonde a confiabilidade dos dados é essencial.

4.2.9- RAID Nível 7: No RAID 7 as informações são transmitidas em modo assíncrono e são controladas e cacheadas de maneira independente, obtendo performance altíssima. O RAID 7 é raramente utilizado pelo custo do Hardware necessário.

4.2.10- RAID Nível 10: O RAID 10 precisa de no mínimo 4 discos rígidos para ser implantado. Os dois primeiros discos trabalham com striping enquanto os outros dois armazenam uma cópia exata dos dois primeiros, mantendo a tolerância a falhas. A diferença básica desse nível para o RAID 0+1 é que sobre certas circunstâncias o RAID 10 pode sustentar mais de uma falha simultânea e manter o sistema.

O RAID 10 pode ser utilizado em servidores de banco de dados que necessitem de alta performance e alta tolerância a falhas como em Sistemas Integrados e Bancos.

5- SISTEMAS NUMÉRICOS

5.1- Sistema Numérico Decimal

O sistema de numeração que se usa habitualmente é o sistema decimal, pois contamos em grupos de 10. A palavra decimal tem origem na palavra latina *decem*, que significa 10. Ele foi inventado pelos hindus, aperfeiçoado e levado para a Europa pelos árabes. Daí o nome indo-arábico.

Esse sistema de numeração apresenta algumas características:

Utiliza apenas os algarismos indoarábicos 0-1-2-3-4-5-6-7-8-9 para representar qualquer quantidade.

Cada 10 unidades de uma ordem formam uma unidade da ordem seguinte.

Observe:

10 unidades = 1 dezena = 10

10 dezenas = 1 centena = 100

10 centenas = 1 unidade de milhar = 1000

Outra característica é que ele segue o princípio do valor posicional do algarismo, isto é, cada algarismo tem um valor de acordo com a posição que ele ocupa na representação do numeral.

Temos, então, o seguinte quadro posicional (ou de ordens):

4ª ordem	3ª ordem	2ª ordem	1ª ordem
unidade de milhar	centena de unidades	dezena de unidades	unidades

Tabela 2 – Quadro de ordens

Observe:

Neste número: **649**, o algarismo 9 representa 9 unidades e vale 9 (1º ordem); o algarismo 4 representa 4 dezenas, ou seja, 4 grupos de 10 unidades e vale 40 (2º ordem); o algarismo 6 representa 6 centenas, ou seja, 6 grupos de 100 unidades e vale 600 (3º ordem). Ou seja, $600 + 40 + 9$ é igual a 649, que lemos seiscentos e trinta e dois.

Seria o mesmo que:

$$6 \times 10^2 + 4 \times 10^1 + 9 \times 10^0$$

Logo:

$$6 \times 100 + 4 \times 10 + 9 \times 1 = \mathbf{649}$$

Mesmo exemplo para: **121.977**

$$1 \times 10^5 + 2 \times 10^4 + 1 \times 10^3 + 9 \times 10^2 + 7 \times 10^1 + 7 \times 10^0$$

Logo:

$$1 \times 100.000 + 2 \times 10.000 + 1 \times 1.000 + 9 \times 100 + 7 \times 10 + 7 \times 1 = \mathbf{121.977}$$

5.2- Sistema Numérico Binário

O sistema numérico mais simples que usa notação posicional é o sistema numérico binário. Como o próprio nome diz, um sistema binário contém apenas dois elementos ou estados. Num sistema numérico isto é expresso como uma base dois, usando os

dígitos **0** e **1**. Esses dois dígitos têm o mesmo valor básico de 0 e 1 do sistema numérico decimal.

Devido a sua simplicidade, microprocessadores usam o sistema binário de numeração para manipular dados. Dados binários são representados por dígitos binários chamados "bits". O termo "bit" é derivado da contração de "*binary digit*". Microprocessadores operam com grupos de "bits" os quais são chamados de palavras. O número binário 11101101 contém **oito** "bits".

Notação Posicional:

Tal qual no sistema numérico decimal, cada posição de "bit" (dígito) de um número binário tem um peso particular o qual determina a magnitude daquele número. O peso de cada posição é determinado por alguma potência da base do sistema numérico.

Para calcular o valor total do número, considere os "bits" específicos e os pesos de suas posições Tabela 1. Por exemplo, o número binário 11101101 pode ser escrito com notação posicional como segue:

$$1 \times 2^7 + 1 \times 2^6 + 0 \times 2^5 + 1 \times 2^4 + 0 \times 2^3 + 1 \times 2^2 + 1 \times 2^1 + 1 \times 2^0$$

Como pode-se perceber a conversão neste caso é praticamente automática, vejamos:

$$1 \times 32 + 1 \times 16 + 0 \times 8 + 1 \times 4 + 0 \times 2 + 1 \times 1 = 53$$

Ou seja, o numeral binário **11101101** corresponde ao numeral decimal **53**.

5.3- Sistema Numérico Hexadecimal

O hexadecimal, segundo ICEA (2011), é outro sistema numérico que é normalmente usado com microprocessadores. Ele permite a fácil conversão ao sistema numérico binário. Devido a isso e também devido ao fato que a notação hexadecimal simplifica a manipulação de dados. Tal qual o nome diz, hexadecimal tem base 16. Ele usa os dígitos **0** até **9** e as letras **A** até **F**.

As letras são usadas para representar 16 valores diferentes com um simples dígito para cada valor. Portanto, as letras de A até F são usadas para representar os valores numéricos de 10 até 15.

Os números iniciais entre os sistemas decimal e hexadecimal são de valores iguais, $0^{10} = 0^{16}$; $3^{10} = 3^{16}$; $9^{10} = 9^{16}$.

Para números maiores que 9, as relações seguintes existem: $10^{10} = A^{16}$; $11^{10} = B^{16}$; $12^{10} = C^{16}$; $13^{10} = D^{16}$; $14^{10} = E^{16}$ e $15^{10} = F^{16}$.

Usar letras em contagem pode parecer grosseiro até a familiarização com o sistema. A tabela DEC e HEXA ilustra o relacionamento entre inteiros decimal e hexadecimal, enquanto tabela FRAÇÕES ilustra o relacionamento entre frações decimal e hexadecimal.

Como nos sistemas numéricos anteriores, cada posição dos dígitos de um número hexadecimal tem um peso posicional o qual determina a magnitude do número. O peso de cada posição é determinado por alguma potência do número base do sistema (neste caso, 16). O valor total do número pode ser calculado considerando os dígitos específicos e os pesos de suas posições. (a tabela mostra uma lista resumida das potências de 16). Por exemplo, o número hexadecimal E5D7,A3 pode ser escrito com notação posicional como se segue:

$$(E \times 16^5) + (5 \times 16^4) + (D \times 16^3) + (7 \times 16^2) + (A \times 16^1) + (3 \times 16^0)$$

REFERÊNCIAS BIBLIOGRÁFICAS:

- BATTISTI, Júlio. Apostila de Hardware.** Banco de Tutoriais do site Júlio Battisti, 2009. Disponível em: <<http://www.juliobattisti.com.br/tutoriais>>. Acessado em abril de 2011.
- CEFAS, Centro Franciscano de Assistência Social. Montagem e Configuração de Computadores.** Apostila de hardware do Cefas. Disponível em: <<http://www.cefaz.org.br>>. Acessado em março de 2011.
- EDUCADOR, Brasil. Discos Raid.** Apostila de Tecnologia Digital. Disponível em: <<https://sites.google.com/a/educadorbrasil.com/tecnologia/arquivos/DiscosRaid.doc?attredirects=0&d=1>>. Acessado em abril 2011.
- MORIMOTO, Carlos E. Hardware. O Guia Definitivo.** Disponível em: <<http://www.gdhpress.com.br/hardware/>>. Acessado em setembro. 2010.
- TORRES, Gabriel. Hardware e Redes.** Clube do Hardware, 2010. Disponível em: <<http://www.clubedohardware.com.br/.../Gabriel-Torres/1>>. Acessado em fevereiro de 2011.
- WIKIPEDIA, Enciclopédia Livre. Microsoft Windows.** Disponível em: <http://pt.wikipedia.org/wiki/Microsoft_Windows>. Acessado em março. 2011.