

TEMA 1: INFORMÁTICA BÁSICA. REPRESENTACIÓN Y COMUNICACIÓN DE LA INFORMACIÓN: ELEMENTOS CONSTITUTIVOS DE UN SISTEMA DE INFORMACIÓN. CARACTERÍSTICAS Y FUNCIONES. ARQUITECTURA DE ORDENADORES. COMPONENTES INTERNOS DE LOS EQUIPOS MICROINFORMÁTICOS.

ÍNDICE:

1. Informática básica

- 1.1. Sistemas informáticos
- 1.2. Evolución histórica de los Sistemas de Información
 - 1.2.1. Evolución histórica desde un punto de vista físico
 - 1.2.2. Evolución histórica desde un punto de vista funcional
- 1.3. Estructura básica de un sistema informático
- 1.4. Funcionamiento básico de un sistema informático
 - 1.4.1. Unidades de medida
 - 1.4.2. Códigos de representación de la información

2. Arquitectura de ordenadores

- 2.1. Introducción
- 2.2. Componentes de un PC

3. La Placa Base

- 3.1. Introducción
- 3.2. BIOS y UEFI
- 3.3. Componentes
- 3.3. Formatos de placa base
- 3.4. Backplane

4. Unidad Central del Proceso (CPU)

- 4.1. Introducción
- 4.2. Componentes
- 4.3. Funcionamiento
- 4.4. Historia
- 4.5. Arquitecturas de procesador RISC y CISC
- 4.6. Segmentación de instrucciones
- 4.7. Vulnerabilidades de seguridad
- 4.8. Microprocesador (resumen)

5. La memoria

- 5.1. Introducción
- 5.2. Tipos de memoria
 - 5.2.1. Según su velocidad y coste
 - 5.2.2. Según la lectura y escritura en las mismas

6. Buses: arquitecturas y funcionamiento

- 6.1. Introducción
- 6.2. Subsistema de E/S controladores y periféricos
- 6.3. Tipos de buses físicos

7. Clasificación de los ordenadores

- 7.1. Introducción. Tipos de ordenadores
- 7.2. Equipos departamentales (miniordenador)
- 7.3. Estaciones de trabajo (Workstations)

- 7.4. Servidores
- 7.5. Computación en la nube
 - 7.5.1. Introducción
 - 7.5.2. Historia
 - 7.5.3. Características, ventajas e inconvenientes
 - 7.5.4. Servicios ofrecidos
 - 7.5.5. Tipos de nubes
 - 7.5.6. Ejemplos
 - 7.5.7. Cloud Privado Virtual
 - 7.5.8. Nuevas tecnologías en la nube
- 7.6. Thin Client
- 7.7. Sistemas Multiprocesador. Taxonomía de Flynn

***A DESTACAR:**

1. **Sistemas de numeración: unidades de información (al menos llegar hasta el Pijabyte!), conversión entre bases fundamental! (binario-decimal-octal-hexadecimal).**
2. **Esquema físico de una placa base con sus elementos, puertos y funciones.**
3. **Microprocesadores:**
 - a. **Arquitectura Von Neumann: Unidad Aritmético-Lógica (ALU), Unidad Control y Banco de Registros. Funciones de cada unidad.**
 - b. **Métricas: MIPS, MFLOPS, SPEC, SPECint, SPECfp.**
 - c. **Arquitecturas: CISC (Complex Instruction Set Computer), RISC (Reduced Instruction Set Computer): conceptos y diferencias.**
4. **Memoria:**
 - a. **Características.**
 - b. **Jerarquía: banco registros-caché L1- caché L2-memoria principal-memoria secundaria.**
 - c. **Tipos: sólo lectura(ROM, PROM, EPROM, EEPROM) o lectura/escritura (DRAM, SRAM).**
5. **Buses: concepto, tipos (internos/externos, paralelos/serie) y velocidades (sobretudo últimas versiones de las distintas tecnologías: USB Y THUNDERBOLT!).**
6. **Computación en la Nube: concepto y características, servicios ofrecidos (Saas, PaaS, IaaS), tipos de nubes (pública, privada, híbrida, comunitaria), ejemplos (Amazon, Azure, Salesforce...).**



1. Informática básica

1.1. Sistemas informáticos

//esta parte unicamente lectura para entenderlo//

Vivimos rodeados de sistemas, formando parte de muchos de ellos. En ocasiones lo hacemos inconscientemente y otras no (ejemplos como sistemas financieros, sistemas políticos, sistemas sanitarios son claras muestras de los mismos).

En su acepción más general, llamamos sistema a aquel conjunto ordenado de elementos que se relacionan entre sí y contribuyen a un determinado objetivo.

Es evidente que existen múltiples tipos de sistemas pero para lo que nos ocupa, tomamos como punto de partida la idea de los sistemas de comunicación entendidos como aquél conjunto de elementos que emiten, reciben e interpretan información.

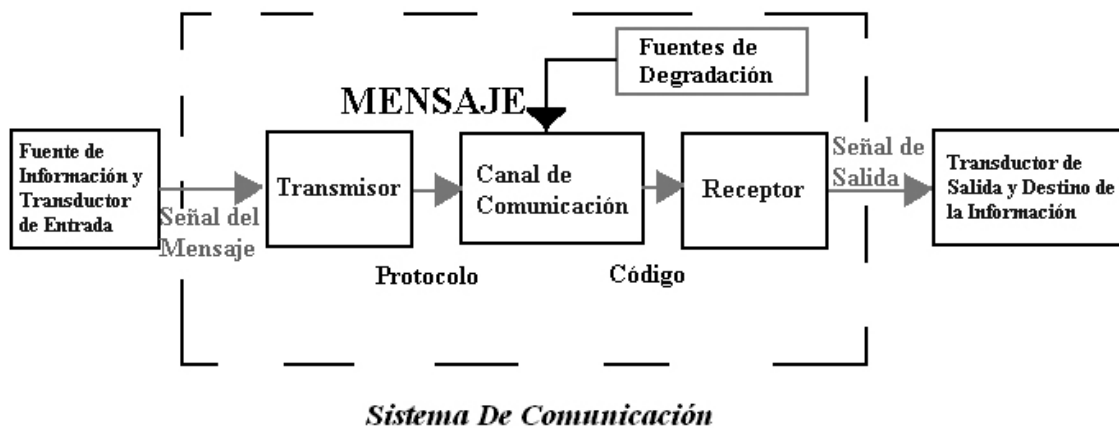


Figura 1. Esquema sistema de comunicación.

En la actualidad, debido al auge de las redes de ordenadores y la evolución de las nuevas tecnologías, el término sistema informático ha desplazado en el ámbito profesional, que no en el doméstico o popular, a otros términos como ordenador o computador.

Hoy día, la popular definición de un ordenador como máquina para el tratamiento automatizado o automático de la información, empleada hace ahora sólo unas décadas en puntuales sectores como el militar o el científico, ha dado paso a grandes sistemas informáticos empleados en casi todos los ámbitos y con los que se crea una excesiva dependencia. Existen sistemas informáticos en sectores tan dispares como el desarrollo industrial, la medicina, la enseñanza o la gestión empresarial sin olvidar el propio hogar.

Definición de sistema informático

Un sistema informático (SI) es un conjunto de dispositivos, con al menos una CPU o unidad central de proceso, que estarán física y lógicamente conectados entre sí a través de canales, lo que se denomina modo local, o se comunicarán por medio de diversos dispositivos o medios de transporte, en el llamado modo remoto. Dichos elementos se integran por medio de una serie de componentes lógicos o software con los que pueden llegar a interaccionar uno o varios agentes externos, entre ellos el hombre.

El objetivo de un sistema informático es el de dar soporte al procesado, almacenamiento, entrada y salida de datos que suelen formar parte de un sistema de información general o específico. Para tal fin es dotado de una serie de recursos que varían en función de la aplicación que se le da al mismo.

Elementos de un sistema informático

Todo SI debe disponer de dos elementos básicos: un sistema físico o hardware y un sistema lógico o software, a los que hay que añadirle un tercero que, sin pertenecer intrínsecamente, no se puede pensar funcionando sin él: los recursos humanos.

Tradicionalmente, los elementos que componen un SI son:

- **Hardware.** Formado por aquellos elementos físicos del SI, siendo elementos hardware el elemento terminal, los canales y los soportes de la información.
Lo constituyen dispositivos electrónicos y electromecánicos que proporcionan capacidad de captación de información, cálculos y presentación de información a través de dispositivos como sensores, unidades de procesado y almacenamiento, monitores, etc.
- **Software.** Aquellos elementos del sistema que no tienen naturaleza física y que se usan para el procesamiento de la información. Son programas de ordenador que suelen manejar estructuras de datos, entre las que destacan las bases de datos, entendidas como colecciones de información organizadas y que sirven de soporte al sistema.
- **Personal.** Entendido como el conjunto de usuarios finales u operadores del SI.
- **Documentación.** Son todo aquel conjunto de manuales impresos o en formato digital y cualquier otra información descriptiva que explican los procedimientos del sistema informático.

En un SI el software está condicionado por el hardware tanto en su uso como en su evolución.



Figura 2. Elementos de un sistema informático.

Todo sistema, y por supuesto un SI, se puede contemplar desde dos aspectos: su descripción física (cómo es físicamente, analizando los componentes que lo constituyen así como los elementos de interconexión) y su descripción funcional (funciones de sus componentes, cómo interactúan unos con otros, reglas o normas de comunicación, etc.)

A lo largo de este tema veremos con más detalle las características funcionales y físicas de un SI entendiendo por:

Estructura funcional del SI. Aquélla asociada al soporte físico o hardware que se encarga de estudiar las arquitecturas de organización y funcionamiento de los diversos componentes del mismo. Veremos la arquitectura tradicional o clásica de un ordenador personal o PC, que toma como punto de partida la histórica Arquitectura de Von Neumann.

Estructura física del SI. También asociada al hardware. En este caso estudiaremos lo que comúnmente se denomina hardware comercial. Veremos cómo son físicamente, para qué sirven y qué características tienen los diferentes componentes actuales que componen un PC, integrados a partir de la placa base y recogidos dentro de un chasis, comunicándose con distintos dispositivos de entrada, salida o entrada-salida.

1.2. Evolución histórica de los Sistemas de Información

//Esta parte unicamente lectura para entenderlo si teneis tiempo y como base de cultura informatica, no lo han preguntado nunca en TAI...//

Para hablar de la evolución histórica debemos analizarla desde las dos perspectivas propuestas, de forma que hablaremos de evolución histórica desde un punto de vista físico y funcional.

1.2.1. Evolución histórica desde un punto de vista físico

Las computadoras, entendidas como máquinas para procesar datos, no son un invento reciente ni mucho menos, sino que tienen detrás una larga historia y un interesante proceso evolutivo donde un grupo de personas, muchas caídas en el anonimato o el olvido, han contribuido aportando algún tipo de avance o mejora.

El hombre siempre ha buscado disponer de instrumentos al principio, y dispositivos después, capaces de efectuar cálculos precisos y rápidos. Hagamos un breve repaso de cómo hemos llegado a los sistemas informáticos actuales.

Hace más de 3.000 años a.c., ya los chinos, y posteriormente otras culturas, desarrollaron el ábaco, objeto creado a partir de un número de cuentas engarzadas en varillas que representan cifras numéricas y que permite realizar cálculos sencillos y operaciones aritméticas.

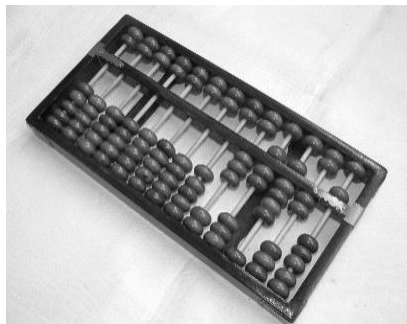


Figura 3. Ábaco.

En Europa, hacia el siglo XVII, el creciente interés por nuevas ciencias como la astronomía y la navegación impulsó el desarrollo de lo que se denominaron las calculadoras mecánicas.

En 1614, John Napier inventó las tablas logarítmicas que permitían efectuar complejas multiplicaciones como simples sumas.

Blaise Pascal en 1642 crea una máquina mecánica capaz de sumar con un sistema de ruedas dentadas que llamó la Pascalina. Posteriormente, Leibnitz en 1671 le agregó la posibilidad de restar, multiplicar y dividir.

Fue ya en el siglo XIX cuando se dio un nuevo empuje evolutivo por medio de Charles Babbage que diseñó la primera computadora de uso general, llamada "Máquina Diferencial" y posteriormente una segunda llamada "Máquina Analítica".

**Charles Babbage (1792-1871), matemático e ingeniero británico fue un hombre adelantado a su tiempo y tuvo muchas ideas que no pudo desarrollar e implementar porque no se lo permitieron los avances tecnológicos de la época. Considerado el inventor de las máquinas calculadoras programables, dedicó sus últimos años y recursos en crear una máquina infalible que fuese capaz de predecir los ganadores de las carreras de caballos.*

Un poco más tarde Lady Ada Byron se interesó por los descubrimientos de Babbage a quién ayudó e hizo una serie de aportaciones que la llevaron a ser considerada la primera mujer programadora.

En 1804, Joseph Jacquard inventó un telar que se servía de tarjetas perforadas para controlar la creación de complejos diseños textiles.

Una tarjeta perforada es una superficie de papel, cartón o plástico con unas perforaciones distribuidas de forma que representan información (en binario para las computadoras).

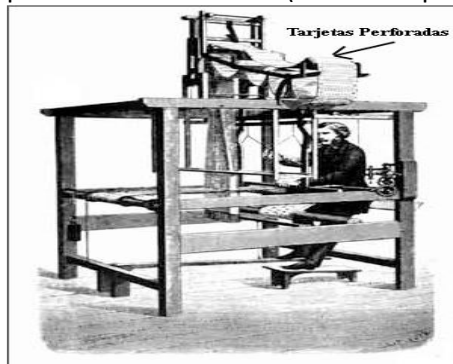


Figura 4. Telar de Jacquard.

Pero la primera operación real de procesamiento de datos la llevó a cabo Herman Hollerith en 1890 mediante el desarrollo de un sistema mecánico para llevar a cabo censos basado en tarjetas perforadas. Se usó en el censo de población en Estados Unidos donde se logró llevar a cabo el recuento en dos años y medio frente a los más de siete que se tardaba con anterioridad.

**Hollerith fundó la Tabulating Machine Company y vendió sus productos en todo el mundo. La demanda de sus máquinas se extendió incluso hasta Rusia. El primer censo llevado a cabo en Rusia en 1897, se registró con el Tabulador de Hollerith. En 1911, la Tabulating Machine Company, al unirse con otras Compañías, formó la Computing-Tabulating-Recording-Company. Para reflejar mejor el alcance de sus intereses comerciales, en 1924 la Compañía cambió el nombre por el de International Business Machines Corporation (IBM).*

Por entonces ya estaba muy claro que el sistema binario, basado en ceros y unos, es el que daba soporte al ordenador. Se hacían necesarios dispositivos electrónicos que permitiesen almacenar esa información. A ese tipo de dispositivos se les llamó dispositivos biestables y la evolución electrónica de los mismos fue determinante en los siguientes pasos que se dieron. Debido a los rápidos avances en el mundo de la electrónica, impulsados por la segunda guerra mundial, a partir de los años cuarenta, la historia de los ordenadores se clasifica por distintas etapas llamadas generaciones caracterizadas por los diferentes componentes que dan soporte a los biestables.

Podemos diferenciar las siguientes generaciones:

1ª Generación (1940-1956)

Comprende los primeros grandes ordenadores basados en la arquitectura Von Neumann, referente en el diseño actual de los ordenadores como veremos, y surgen por una necesidad vital al considerarse un instrumento armamentístico durante la Segunda Guerra Mundial.

Sus características principales son:

- Uso de la tecnología basada en válvulas de vacío, tecnología que sustituyó a los interruptores electromecánicos, para dar soporte a los biestables.
- Empleo de computadoras con fines militares y científicos.

- Máquinas muy grandes y pesadas, muy lentas en sus procesos, de tal forma que algunos programas largos implicaban días de espera. Aún así podían llegar a efectuar unos cinco mil cálculos por segundo.
- Destacan en esta época máquinas como el ENIAC o el EDVAC.

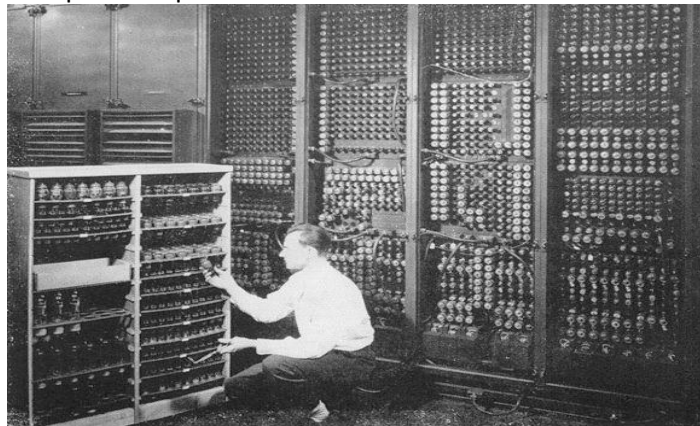


Figura 5. Máquina Eniac.

*Durante la Segunda Guerra Mundial, IBM comenzó a investigar en el campo de la informática. En 1944 terminó de construir el computador Automatic Sequence Controlled Calculator, también conocido como Mark I. El Mark I fue la primera máquina capaz de ejecutar cálculos complejos automáticamente, y estaba basada en interruptores electromecánicos.

En 1952, IBM creó el IBM 701, el primer gran computador basado en válvulas de vacío, tecnología que sustituyó a los interruptores electromecánicos. En 1954 introdujo la IBM 650.

2ª Generación (1956-1963)

Esta etapa coincide con la aparición del transistor (1956). Las funciones del transistor son similares a las de las válvulas de vacío pero con ahorro significativo en tamaño y consumo. Sus características principales son:

- Tecnología basada en transistores, con máquinas más pequeñas y de menor consumo energético.
- Sigue siendo el campo científico el de mayor aplicación aunque surgen computadoras con fines comerciales.
- Aparece la serie IBM 7090 que se empieza a comercializar en grandes empresas.
- Empleo de los primeros periféricos.
- Aparece el concepto de supercomputadora como aquel ordenador con capacidades de cálculo muy superiores a las comunes según la época.
- Aparecen los primeros lenguajes de programación y los famosos sistemas batch o de procesamiento por lotes, que permitían la ejecución de un programa sin el control o supervisión directa del usuario (frente a los sistemas interactivos que exigen el control o la presencia del usuario).

*Las supercomputadoras fueron introducidas en la década de los sesenta y fueron diseñadas principalmente por Seymour Cray en la compañía Control Data Corporation (CDC).

El término supercomputador o superordenador está en constante cambio de forma que los superordenadores de hoy tienden a convertirse en ordenadores ordinarios del mañana.

Los superordenadores actuales se usan para tareas de cálculos intensivos, tales como problemas que involucran física cuántica, predicción del clima, investigación de cambio climático, modelado de moléculas, simulaciones físicas tal como la simulación de aviones o automóviles en el viento, simulación de la detonación de armas nucleares e investigación en la fusión nuclear.

3ª Generación (1964-1971)

Se caracteriza por la aparición de los circuitos integrados. Se trata de integrar en un solo chip todos los transistores y circuitos analógicos que realizan las operaciones básicas de un ordenador. Este descubrimiento produjo grandes cambios en cuanto al tamaño de las computadoras, en velocidad, compatibilidad, etc.

Esta generación se caracteriza por:

- Uso de la tecnología basada primero en una escala de integración pequeña (SSI), con decenas de transistores, para luego pasar a una escala de integración media (MSI) que empleaba cientos de transistores en cada chip.
- Aparecen nuevos soportes de almacenamiento e interacción como los discos flexibles magnéticos creados por IBM o el monitor.
- Nuevas técnicas y lenguajes de programación.
- Aparecen los conceptos de miniordenador, computadora multiusuario de prestaciones intermedias, frente a los grandes sistemas o mainframe hasta entonces, y el concepto de estación de trabajo o computadora de altas prestaciones para determinadas tareas específicas.
- Se emplean por primera vez lenguajes de alto nivel no específicos, los lenguajes de programación de propósito general (C, Pascal, Basic, etc).

*A mediados de los 60 IBM lanzó el System/360, la primera arquitectura de computadores que permitía intercambiar los programas y periféricos entre los distintos equipos componentes de la arquitectura, al contrario de lo existente anteriormente, en que cada equipo era una caja cerrada incompatible con los demás. El desarrollo fue tan costoso que prácticamente llevó a la quiebra a IBM, pero tuvo tal éxito al lanzarse al mercado que los nuevos ingresos y el liderazgo que consiguió IBM respecto a sus competidores les resarcieron de todos los gastos.

Este éxito provocó que la empresa fuera investigada por monopolio. De hecho, tuvo un juicio, que comenzó en 1969, en el que fue acusada de intentar monopolizar el mercado de los computadores empresariales, que se prolongó hasta 1983.

4ª Generación (1971-1981)

Se caracteriza por la popularización del microordenador y de la computadora personal o doméstica. La tecnología permite integrar más circuitos en una sola pastilla. Esto reduce el espacio y el consumo haciendo asequible el ordenador a cualquier persona.

Sus características principales son:

- Tecnología de alta escala de integración (LSI) que empleaba miles de transistores (hasta 10.000).
- Aparece el microprocesador entendido como aquel circuito integrado que contiene algunos o todos los elementos hardware de la CPU.
- Proliferan los lenguajes de programación.
- Muchas familias comenzaron a tener computadoras en sus casas como las famosas Commodore 64 y 128, ZX Spectrum o Amstrad CPC.



5ª Generación (1982-1991)

El microprocesador sigue evolucionando reduciendo su tamaño y permitiendo muchas más operaciones y funcionalidades.

Sus características son:

- Tecnología VLSI o de muy alta escala de integración (de 10.000 a 100.000 transistores).
- Aparecen los microprocesadores de uso específico.
- Evolución muy rápida de la tecnología, lo que se fue a llamar la Ley de Moore.
 - *La Ley de Moore se trata de una ley empírica, basada en la observación, formulada por el co-fundador de Intel, Gordon E. Moore en 1965 que expresa que aproximadamente cada 18 meses se duplica el número de transistores en un circuito integrado. Desde hace unas décadas se aplica sobre ordenadores personales y su cumplimiento se ha podido constatar hasta hoy día.
- Desarrollo y expansión de la tecnología multimedia.
- En esta etapa la compañía Apple inventó la arquitectura abierta, que permitirá añadir, modernizar y cambiar sus componentes y la interfaz gráfica de usuario o GUI , que utiliza un conjunto de imágenes y objetos gráficos para representar información y acciones, en el cual se basan todos los sistemas operativos de la actualidad.
- Las computadoras bajaron sus precios, se hicieron accesibles a muchas más personas y se empezaron a utilizar en cada vez más diversos ámbitos.
- Aparecen microprocesadores trabajando en paralelo, y se extiende el uso generalizado de las redes.

6ª Generación (1992- actualidad)

Se emplean tecnologías superiores de integración como ULSI (Ultra Large Scale Integration) que empleaba entre 100.000 y 1.000.000 de transistores y la GLSI (Giga Large Scale Integration) con más de un millón de transistores.

En la actualidad la fabricación de las computadoras está basada en múltiples microprocesadores que trabajan al mismo tiempo de forma que algunas máquinas pueden llegar a realizar más de un billón de operaciones aritméticas por segundo.

Además, se ha extendido la conectividad de las computadoras mediante el empleo de redes, y cada vez crece más el uso de aplicaciones soportadas por la propia red como Internet.

Como supuestamente la sexta generación de computadoras está en marcha desde los años noventa, debemos por lo menos esbozar las características que deben tener las computadoras de esta generación. También mencionaremos algunos de los avances tecnológicos de la última década del siglo XX y lo que se espera lograr en el siglo XXI:

Características:

- Las computadoras de esta generación cuentan con arquitectura paralelo-vectorial, de tipo combinada que emplea cientos de microprocesadores vectoriales trabajando al mismo tiempo de forma que se han creado computadoras capaces de realizar más de un millón de millones de operaciones aritméticas de punto flotante por segundo (teraflops).
- Las redes de área mundial o WAN se han desarrollado pero seguirán creciendo desorbitadamente utilizando medios de comunicación como la fibra óptica y satélites, con anchos de banda impresionantes.
- Las tecnologías de esta generación ya han sido desarrolladas o están en ese proceso. Algunas de las más relevantes son la llamada inteligencia artificial distribuida, el empleo de la teoría del caos, uso de sistemas difusos, la holografía, transistores ópticos, nanotecnología , etc.
- Smartphones, IoT (Internet de las Cosas), Big Data, Blockchain, redes sociales, Smartcitys...



1.2.2. Evolución histórica desde un punto de vista funcional

//esta parte unicamente lectura para entenderlo si teneis tiempo, no lo han preguntado nunca//

Desde un punto de vista funcional los ordenadores aparentemente no han evolucionado mucho a grandes rasgos ya que todavía sigue vigente el esquema de funcionamiento de la Arquitectura Von Neumann, con una serie de módulos funcionales comunes (elementos de entrada y salida, memoria principal y secundaria, procesador y buses), aunque sí han existido grandes cambios en la forma de comunicarse y operar entre sí.

El sistema informático ha evolucionado desde una primera situación en que todos los componentes del mismo se encontraban centralizados en un mismo lugar (bien en forma de chasis en un sistema monopuesto o a través de sala de ordenadores en sistemas multipuesto) y nos encontrábamos con sistemas aislados, hasta la situación actual en la que los componentes del sistema se pueden encontrar repartidos en diferentes lugares físicos dando lugar a sistemas conectados o en red que pueden llegar a colaborar entre sí dando lugar a los llamados sistemas distribuidos.

Pero esta evolución e implantación de los sistemas distribuidos ha pasado por diferentes fases y no se puede dar aún por finalizada.

Los sistemas actuales distribuidos se pueden organizar de forma vertical o jerárquica y de forma horizontal.

a) **Distribución vertical.** Existen varios niveles como son el primer nivel o corporativo, segundo nivel o departamental y el tercer nivel o puesto personal y en cada uno se usan distintos tipos de equipos con configuraciones también características.

b) **Distribución Horizontal.** Todos los equipos tienen la misma categoría, por lo menos no existe un equipo central en el primer nivel de la jerarquía. Suelen existir un conjunto de ordenadores conectados que cooperan entre sí pero sin que ninguno de ellos centralice la información.

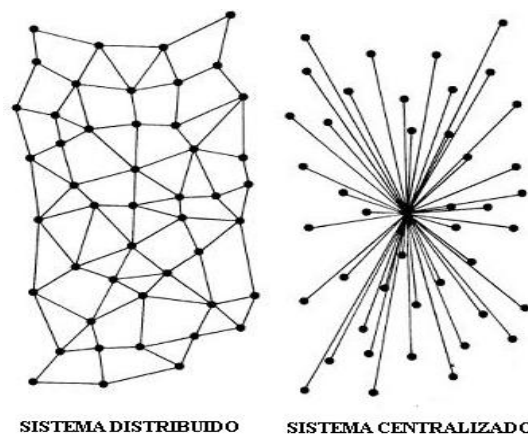


Figura 6. Sistemas distribuidos horizontal y vertical.

1.3. Estructura básica de un sistema informático

Como ya se adelantó con anterioridad, un sistema informático está formado por tres componentes básicos:

- **Componente Físico o Hardware.** Es el conjunto de elementos tangibles necesarios para el tratamiento eficaz de la información y está formado por varios elementos: la unidad central de proceso o CPU, la memoria auxiliar y los dispositivos de entrada-salida o periféricos.

- **Componente Lógico o Software.** Es el conjunto de recursos lógicos necesarios para que el SI se encuentre en condiciones de realizar todos los trabajos que se le vayan a encomendar.
- **Componente Humano o Humanware.** Variado colectivo de personas encargados de desarrollar o gestionar el sistema, conocidos como profesionales informáticos, o de aprovechar las prestaciones del mismo, lo que se conoce como usuarios, que pueden ser en distintos grados o a distinto nivel de conocimiento y manejo.



Figura 7. Esquema componentes de un sistema informático.

1.4. Funcionamiento básico de un sistema informático

1.4.1. Unidades de medida

// esto es importantísimo!!! Sobre todo la tabla (FIGURA 8) de unidades de medida tanto en sistema decimal (kilo,...) como binario (kibi,...)//

Un sistema informático maneja información de todo tipo (números, texto, imágenes, sonidos, video, etc.), dándole entrada, salida o procesándola. Para ello utilizará mecanismos de representación, almacenamiento y presentación como veremos a continuación.

Bajo nuestra perspectiva humana es fácil diferenciar rápidamente lo que son números de lo que es texto, imagen, sonido o video, utilizando para ello los órganos sensoriales y el cerebro. El ordenador en su funcionamiento trata de emular el comportamiento humano pero al ser una máquina digital, cuyo soporte es la electrónica, sólo es capaz de representar **información binaria** por lo que los ordenadores necesitan codificar la información del mundo real al equivalente binario y utilizar mecanismos para su presentación.

Puesto que toda la información de un ordenador se representa de forma binaria, se hizo indispensable el utilizar **unidades de medida** para poder indicar la capacidad de los dispositivos que manejaban dichos valores.

Las principales unidades de medida de bits son:

- Bit (de **Binary digit**, Dígito Binario). Es una unidad de medida de cantidad de información, equivalente a la elección entre dos posibilidades igualmente probables tomando valores 1 o 0.
- Nibble. Es el conjunto formado por cuatro bits.
- Byte (B). Es el conjunto formado por ocho bits.
- Kilobyte (kB). Son 1.024 bytes.
- Megabyte (MB). Son 1.024 Kilobytes.
- Gigabyte (GB). Son 1.024 Megabytes.
- Terabyte (TB). Son 1.024 Gigabytes.

- Petabyte (PB). Son 1.024 Terabytes.
- Exabyte (EB). Son 1.024 Petabytes.

*Existen unidades mayores al exabyte que con la evolución de la tecnología están empezando a emplearse como son: **Zettabyte** (1.024 Exabytes) y **Yottabyte** (1.024 Zettabytes). Recuerda también no confundir los bits (b) con los bytes (B) en las unidades de medida. **¡No es lo mismo kb que kB!** Las velocidades de transmisión de datos se miden normalmente en kilobits por segundo (kbps), mientras que el almacenamiento o tamaño de los archivos se especifica en bytes (Kilobytes, Megabytes, etc.)

Unidades de Medidas de Almacenamiento							
Sistema Internacional de Unidades Decimal				Sistema Binario ISO/IEC 80000-13			
NOMBRE	SIMBOLO	EQUIVALENCIA	PREFIJO	NOMBRE	SIMBOLO	EQUIVALENCIA	PREFIJO
bit	b	0 o 1				0 o 1	
Nibble	-	4 bits				4 bits	
Byte	B	8 b				8 b	
KiloByte	KB	1000 B	10 ³	KibiByte	KiB	1024 B	2 ¹⁰
MegaByte	MB	1000 KB	10 ⁶	MebiByte	MiB	1024 KB	2 ²⁰
GigaByte	GB	1000 MB	10 ⁹	GibiByte	GiB	1024 MB	2 ³⁰
TeraByte	TB	1000 GB	10 ¹²	TebiByte	TiB	1024 GB	2 ⁴⁰
PetaByte	PB	1000 TB	10 ¹⁵	PebiByte	PiB	1024 TB	2 ⁵⁰
ExaByte	EB	1000 PB	10 ¹⁸	ExbiByte	EiB	1024 PB	2 ⁶⁰
ZettaByte	ZB	1000 EB	10 ²¹	ZebiByte	ZiB	1024 EB	2 ⁷⁰
YottaByte	YB	1000 ZB	10 ²⁴	YobiByte	YiB	1024 ZB	2 ⁸⁰
RonnaByte	RB	1000 YB	10 ²⁷	RobiByte	RiB	1024 YB	2 ⁹⁰
QuettaByte	QB	1000 RB	10 ³⁰	QuebiByte	QiB	1024 RB	2 ¹⁰⁰

Figura 8. Unidades de medida.

1.4.2. Códigos de representación de la información

Desde los inicios de la informática la codificación ha sido problemática entre otras cosas por la falta de acuerdo en la representación de esa información. Hoy día existen numerosos **estándares** para tal fin. Un estándar es un conjunto de especificaciones que regulan procesos como la fabricación de componentes en nuestro caso para garantizar la interoperabilidad. Los números se almacenan dependiendo del tipo de valor que sea (natural, entero o real), utilizando diferentes sistemas de representación numérica como es el caso del **complemento a la base** (en sus diferentes versiones: C-1 y C-2) o el **exceso a la base** para números enteros o el **estándar IEEE 794** para números reales.

Para el caso del texto, lo que se hace es codificar cada carácter de la cadena a almacenar empleando una serie de valores binarios con los que se corresponde de acuerdo a un determinado código.

El código **ASCII** (American Standard Code for Information Interchange —Código Estándar Estadounidense para el Intercambio de Información) es un código de caracteres basado en el alfabeto latino, tal como se usa en inglés moderno. Fue creado en 1963 por el Comité Estadounidense de Estándares (ASA, conocido desde 1969 como el Instituto Estadounidense de Estándares Nacionales, o ANSI) como una refundición o evolución de los conjuntos de códigos utilizados entonces en telegrafía. Más tarde, en 1967, se incluyeron las minúsculas, y se redefinieron algunos códigos de control para formar el código conocido como US-ASCII. El código ASCII utiliza 7 bits para representar los caracteres, aunque inicialmente empleaba un bit adicional (bit de paridad) que se usaba para detectar errores en la transmisión. ASCII fue publicado como estándar por primera vez en 1967 y fue actualizado por última vez en 1986. En la actualidad define códigos para 32 caracteres no imprimibles, de los cuales la mayoría son

caracteres de control que tienen efecto sobre cómo se procesa el texto, más otros 95 caracteres imprimibles que les siguen en la numeración (empezando por el carácter espacio). Casi todos los sistemas informáticos actuales utilizan el código ASCII o una extensión compatible para representar textos y para el control de dispositivos que manejan texto como el teclado.

Poco después surgió un problema: este código es suficiente para los caracteres de la lengua inglesa, pero no para otras lenguas. Entonces se añadió el octavo bit para representar otros 128 caracteres que son distintos según los idiomas (Europa Occidental usa unos códigos que no utiliza Europa Oriental), llamándose así **ASCII Extendido**.

Otros códigos: *//esta clasificacion la considero importante//*

- **Unicode:** Una ampliación de ASCII es el código **Unicode** que puede utilizar hasta 4 bytes (32 bits) para representar cada carácter, con lo que es capaz de codificar cualquier símbolo en cualquier lengua del planeta utilizando el mismo conjunto de códigos, además de textos clásicos de lenguas muertas. Unicode proviene de los tres objetivos perseguidos: universalidad, uniformidad y unicidad.
- **BCD:** Binary-Coded Decimal (BCD) o Decimal codificado en binario es un estándar para representar números decimales en el sistema binario, en donde cada dígito decimal es codificado con una secuencia de 4 bits. Con esta codificación especial de los dígitos decimales en el sistema binario, se pueden realizar operaciones aritméticas como suma, resta, multiplicación y división.
- **EBCDIC:** (Extended Binary Coded Decimal Interchange Code - Código de intercambio decimal de código binario extendido), pronunciado como [ebquidic], es un código estándar de 8 bits usado por computadoras mainframe IBM. IBM adaptó el EBCDIC del código de tarjetas perforadas en los años 1960 y lo promulgó como una táctica customer-control cambiando el código estándar ASCII. EBCDIC es un código binario que representa caracteres alfanuméricos, controles y signos de puntuación. Cada carácter está compuesto por 8 bits = 1 byte, por eso EBCDIC define un total de 256 caracteres.
- **ISO/IEC 8859:** es un conjunto ISO y la IEC estándar de 8 bits para codificaciones de caracteres para su uso en computadoras.
- **UTF-8:** (8-bit Unicode Transformation Format) es un formato de codificación de caracteres Unicode e ISO 10646 utilizando símbolos de longitud variable.

***Arte ASCII** (pronunciado arte áski), es un medio artístico que utiliza recursos computarizados fundamentados en los caracteres de impresión ASCII. Hoy día puede crearse con cualquier editor de textos, aunque en la década previa al advenimiento del computador personal de escritorio (IBM PC, 1981), algunos artistas lo utilizaban de manera experimental (ver artículo: Angel Luis Arambilet Alvarez) y como medio alternativo de arte gráfico, utilizando tarjetas perforadas de 80 y 96 columnas, así como diversos programas compiladores o utilitarios (COBOL, RPG, IBM DITTO), combinado a impresoras de martillo de alta velocidad para fines de presentación.

En el caso del almacenamiento de otro tipo de datos como imágenes, vídeo o audio, se necesita una codificación mucho más compleja, no bastando una correlación símbolo-secuencia de bits. Además en este tipo de información no hay un patrón que se repita, por lo que hay decenas de formas de codificar.

//esta parte unicamente saber los tipo de imágenes y de que van, solo los conceptos, los calculos no hay que saberlos//

Para las imágenes una forma básica de codificarlas en binario son las llamadas **imágenes rasterizadas, matriciales o de mapa de bits**, en las que se almacena la información de cada píxel (*cada uno de los puntos distinguibles en la imagen*) descomponiéndolo en tres colores primarios (R o nivel de rojo, G o nivel de verde, B o nivel de azul), con valores de tamaño

dependiendo del número de colores que admita la representación, lo que se conoce como **profundidad de color**.

Se establece un código formado por ceros y unos para cada color.

Por ejemplo, si tenemos una profundidad de color de cuatro colores (blanco, rojo, amarillo y negro), podemos establecer este código:

00-Blanco 01-Rojo 10-Amarillo 11-Negro

Hemos empleado dos bits para cada código y la imagen se representará mediante el código del color de cada punto de la imagen de forma ordenada.

					00	00	00	00	00
					00	01	01	01	01
					00	01	01	01	01
					00	10	10	10	10
					00	10	10	10	10
					00	01	01	01	01
					00	01	01	01	01
					00	11	00	00	00
					00	11	00	00	00

Figura 9. Imagen bitmap de bandera de España.

La imagen se representará dentro del ordenador por:

0000000000001010101000101010100101010100010101
0100001010101000101010100110000000011000000

La cantidad de información que supone la imagen viene dada por:

Tamaño Total = bits para cada color x resolución horizontal x resolución vertical

Para nuestro ejemplo, Tamaño Total = 2 x 5 x 9 = 90 bits.

Naturalmente en una imagen no sólo se graban los píxeles sino el tamaño de la imagen, el modelo de color, etc. De ahí que representar estos datos sea tan complejo para el ordenador (y tan complejo entenderlo para nosotros).

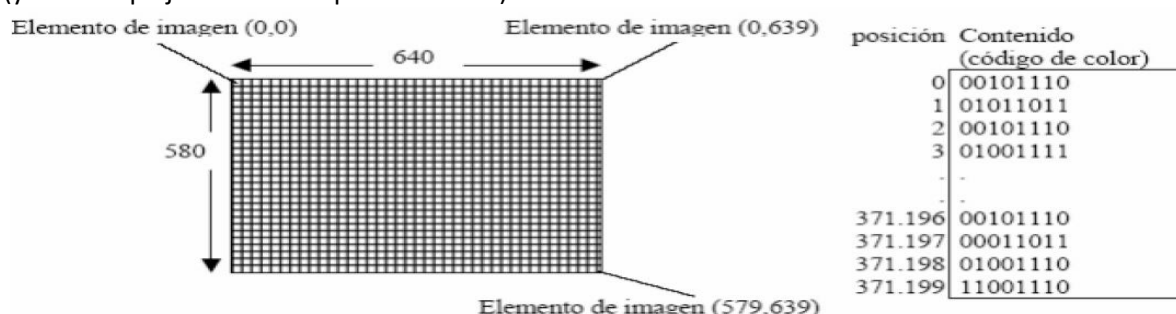


Figura 10. Imagen bitmap y almacenamiento de ésta.

Existen otro tipo de imágenes llamadas **imágenes vectoriales**, **vectorizadas** o **escalables** donde la imagen se construye a partir de vectores, que son objetos formados matemáticamente como segmentos, polígonos, arcos y otras figuras, almacenándose distintos atributos matemáticos de los mismos (por ejemplo, un círculo blanco se define por la posición de su centro, el tamaño de su radio, el grosor y color de la línea y el color de relleno).

En el caso del **audio** o **sonido**, que es por naturaleza información analógica o continua, es una onda que transcurre durante un tiempo. Para almacenar esesonido habrá que representar de

alguna forma esa onda para que después se pueda mandar la señal adecuada a dispositivos de salida de audio.

La onda de sonido suele tener un aspecto similar al siguiente, donde el eje horizontal representa el tiempo y el vertical la amplitud del sonido.

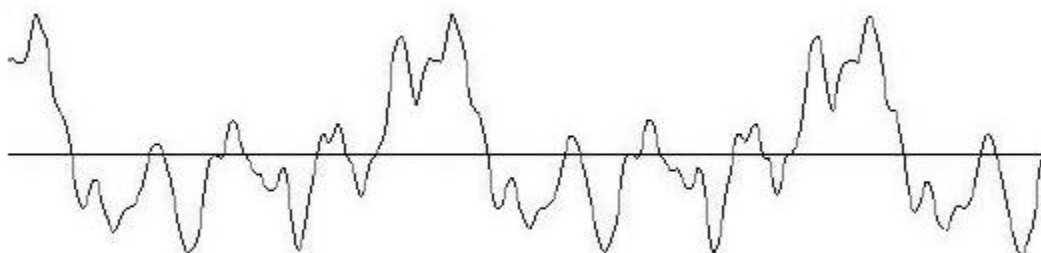


Figura 11. Onda de sonido.

Para guardar el sonido de esa onda se toma el valor de la amplitud (altura de la onda) en binario con un número de bits, llamado **calidad del muestreo**, que determinará la calidad del mismo, habitualmente 16 o 32 bits. Esta operación se hace cada cierto tiempo, tomándose un número de puntos por segundo a lo que se llama **frecuencia** que se mide en hercios (puntos por segundos). Son frecuencias habituales las de la telefonía (8 kHz, 8.000 puntos por segundo), la radio (22 kHz) o el CD (44,1 kHz).

Para reproducir el sonido se reconstruye la onda a partir de los valores almacenados. Es evidente que a mayor frecuencia (más cerca estén los puntos), la onda formada se parecerá más a la original.

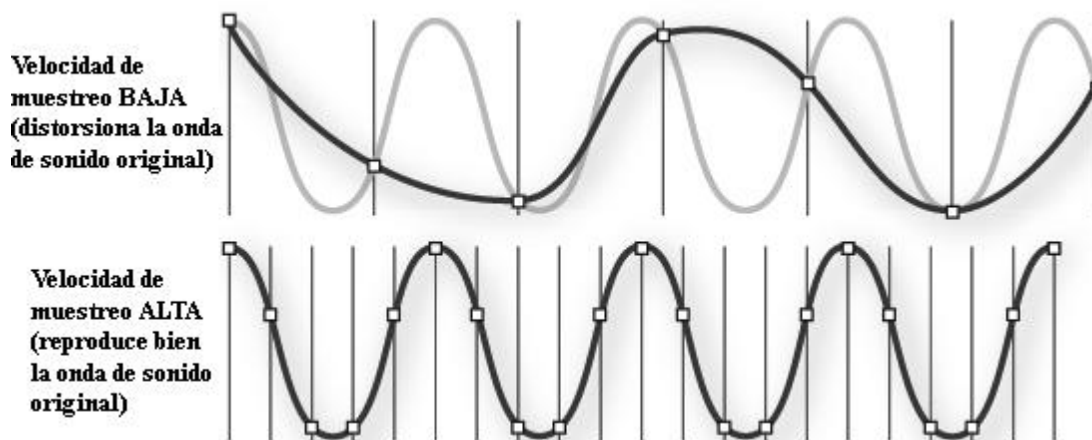


Figura 12. Onda de sonido digital con distintas frecuencias de muestreo.

Hay que tener también en cuenta que la música puede ser **sonido mono**, con un solo canal de sonido o **sonido estéreo** lo que implica dos ondas o dos canales, uno para cada altavoz.

El tamaño de un sonido almacenado vendría dado por:

Tamaño = número de canales x calidad de muestreo x frecuencia x duración (s)

Por ejemplo, si tenemos 30 segundos de sonido estéreo con una calidad de 32 bits y con una frecuencia de 22 kHz, el tamaño que ocupará será:

Tamaño = $2 \times 32 \times 22.000 \times 30 = 42.240.000$ bits = 5.280.000 bytes = 5.156,26 kB

Por último almacenar **video** es bastante sencillo si partimos de la base de que es una representación de imágenes o **frames** y sonido en el tiempo.

Una película no es más que una serie de cuadros desplegados unos tras otros para crear una ilusión de movimiento. El ritmo de imágenes por segundo es una característica del video llamada **frames por segundo (fps)**.

Vamos a ver un ejemplo con un video de 30 segundos grabado a una resolución 640x480 y 32 bits de profundidad de color a 30 fps con sonido estéreo de 32 bits de calidad con frecuencia de 22 kHz. ¿Cuánto ocupa todo el video?

Empezaremos calculando el tamaño de las imágenes:

Tamaño de imagen = $640 \times 480 \times 32 = 9.830.400$ bits = 1.200 kB.

El número total de imágenes será: 30 (fps) $\times$ 30 (segundos) = 900 imágenes.

La secuencia ocupará: 1.200 kB $\times$ 900 (imágenes) = $1.080.000$ kB = 1.054,68 MB

Por otro lado calculamos el tamaño del sonido:

Tamaño sonido = 2 (estéreo) $\times$ 32 (bits) $\times$ 22.000 (Hz) $\times$ 30 (segundos) = $42.240.000$ bits =

$5.280.000$ bytes = $5.156,25$ kB = 5,03 MB.

El tamaño total del video sería la suma de ambas cantidades:

$1.054,68$ MB + $5,03$ MB = $1.059,71$ MB = 1,034 GB.

1.4.3. Sistemas de numeración

//muy importante saber la mecanica de conversion entre bases, sobretodo decimal-binario y viceversa para el subnetting. Tenéis explicadas aquí las mecánicas de conversión y ejemplos explicados paso a paso, también en youtube hay bastantes tutoriales explicando cada caso. Aún así, si os quedan dudas me preguntáis//

Entendemos por sistema de numeración el conjunto de normas y símbolos utilizados para representar la información. Así pues diremos que el sistema de numeración del ser humano en la vida cotidiana es el sistema decimal, formado por diez símbolos, a través de los cuales se pueden conseguir todos los demás.

Informáticamente hablando nos podemos encontrar con el sistema binario, que utiliza dos símbolos para formar cualquier tipo de información. Pero también existen otros sistemas utilizados en la informática como pueden ser el sistema octal, formado por 8 componentes, y el sistema hexadecimal, formado por 16 componentes.

Se dice que se utiliza un sistema de numeración posicional en base b , lo que implica:

— Utilización de un alfabeto de b símbolos diferentes (o cifras).

— Representación de cualquier número como una secuencia de cifras, contribuyendo cada una de ellas con un valor que depende de la cifra en sí y de su posición en la secuencia en: ..., decenas de millar, unidades de millar, centenas, decenas, unidades, para la parte entera del número, y para la parte decimal: décimas, centésimas, milésimas...

$$N = n_2 \cdot b^2 + n_1 \cdot b^1 + n_0 \cdot b^0 + n_{-1} \cdot b^{-1} + n_{-2} \cdot b^{-2}$$

— Ejemplo: valor numérico del número 3479,52 interpretado en base 10 $\Rightarrow$ Conjunto de símbolos = $\{0, 1, 2, 3, 4, \dots, 9\}$

$$3479,52 = 3 \cdot 10^3 + 4 \cdot 10^2 + 7 \cdot 10^1 + 9 \cdot 10^0 + 5 \cdot 10^{-1} + 2 \cdot 10^{-2} = 3000 + 400 + 70 + 9 + 0,5 + 0,02$$

Si bien se puede representar la información en cualquier base b dentro del mundo de la informática, las más habituales son las siguientes:

- **Base 10 ($b = 10$): sistema decimal.** Alfabeto = $\{0, \dots, 9\}$. En el sistema decimal los símbolos válidos para construir números son $\{0, 1, \dots, 9\}$ (0 hasta 9, ambos incluidos), por tanto la base (el número de símbolos válidos en el sistema) es diez.
- **Base 2 ($b = 2$): sistema binario natural.** Alfabeto = $\{0, 1\}$. En este sistema los dígitos válidos son $\{0, 1\}$, y dos unidades forman una unidad de orden superior.
- **Base 8 ($b = 8$): sistema octal.** Alfabeto = $\{0, \dots, 7\}$. El sistema octal es el sistema de numeración posicional cuya base es igual 8, utilizando los dígitos indio árabigos: 0, 1, 2, 3, 4, 5, 6, 7.
- **Base 16 ($b = 16$): sistema hexadecimal.** Alfabeto = $\{0, \dots, 9, A, \dots, F\}$. El sistema hexadecimal (abreviado hex.) es el sistema de numeración posicional que tiene como base el número 16. En principio, dado que el sistema usual de numeración es



de base decimal y, por ello, solo se dispone de diez dígitos, se adoptó la convención de usar las seis primeras letras del alfabeto latino para suplir los dígitos que faltan. El conjunto de símbolos es el siguiente:

$$S = \{0, 1, 2, 3, 4, 5, 6, 7, 8, 9, A, B, C, D, E, F\}$$

La representación de un número en una base tiene su equivalencia en otra base mediante el procedimiento de conversión conveniente, por ejemplo, se muestran los siguientes números binarios de 3 cifras y sus valores en base decimal:

Base 2	Base 10
000	0
001	1
010	2
011	3
100	4
101	5
110	6
111	7

Procedimiento de conversión de base decimal a cualquier base

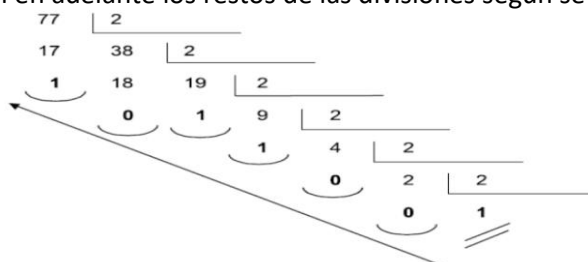
— Parte entera: ir dividiendo entre la base a la que queremos cambiar la parte entera del número decimal original y, sucesivamente, los cocientes que se van obteniendo. Los restos de estas divisiones y el último cociente son las cifras binarias. El último cociente es el bit más significativo y el primer resto el menos significativo.

— Parte fraccionaria: ir multiplicando por 2 la parte fraccionaria del número decimal original y, sucesivamente, las partes fraccionarias de los números obtenidos.

El número binario se forma con las partes enteras que se van obteniendo.

— Ejemplo de conversión de 77,1875 de base decimal a binario:

* La parte entera, que es 77 en decimal, en binario es 1001101, es decir, se toma el último cociente y de ahí en adelante los restos de las divisiones según se indica en la figura:



* La parte decimal que es 0,1875 en decimal, en binario es 0011 que se obtiene de ir tomando la cifra de la parte entera del resultado de cada multiplicación, que son las cifras que se encuentran en negrita de la figura que a continuación se muestra e iremos realizando esta tarea hasta que la parte decimal sea 0, tal y como se encuentra subrayado en la figura.

$$\begin{array}{r} 0,1875 \\ \times 2 \\ \hline 0,375 \end{array} \quad \begin{array}{r} 0,375 \\ \times 2 \\ \hline 0,75 \end{array} \quad \begin{array}{r} 0,75 \\ \times 2 \\ \hline 1,5 \end{array} \quad \begin{array}{r} 0,5 \\ \times 2 \\ \hline 1,0 \end{array}$$

* Por tanto, 77,1875 en base 10 es igual a 1001101,0011 en binario o base 2.

Procedimiento de conversión de cualquier base a base decimal

— Hay que adoptar la fórmula que ya vimos anteriormente en la que la base b es la base actual en la que se encuentra el número y según su posición estará elevada a la potencia correspondiente:

$$N = n_2 \cdot b^2 + n_1 \cdot b^1 + n_0 \cdot b^0 + n_{-1} \cdot b^{-1} + n_{-2} \cdot b^{-2}$$

— Para el ejemplo anterior, tenemos el siguiente número 1001101,0011 en binario y queremos convertirlo a base 10 (ya conocemos el resultado: 77,1875), veamos cómo se obtiene:

$$N = 1 \cdot 2^6 + 0 \cdot 2^5 + 0 \cdot 2^4 + 1 \cdot 2^3 + 1 \cdot 2^2 + 0 \cdot 2^1 + 1 \cdot 2^0 + 0 \cdot 2^{-1} + 0 \cdot 2^{-2} + 1 \cdot 2^{-3} + 1 \cdot 2^{-4} = 1 \cdot 64 + 0 \cdot 32 + 0 \cdot 16 + 1 \cdot 8 + 1 \cdot 4 + 0 \cdot 2 + 1 \cdot 1 + 0 \cdot 0,5 + 0 \cdot 0,25 + 1 \cdot 0,125 + 1 \cdot 0,0625 = 77,1875$$

DECIMAL A OCTAL:

Para convertir de decimal a octal lo que tenemos que hacer es **realizar divisiones sucesivas entre 8** hasta que el resto sea menor que ocho.

Posteriormente, **escribe el valor de los restos obtenidos en cada división** junto con el **del último cociente**. El último cociente será el primer dígito del número octal mientras que el primer resto será el último número.

EJEMPLO:

Número decimal 4248₁₀ en sistema octal:

Número decimal	Dividimos entre 8	
	Cociente	Resto
4248	531	0
531	66	3
66	8	2
8	1	0

Fijándonos en el último cociente y el valor de los residuos obtenidos, tenemos que:

$$4248_{10} = 10230_8$$

OCTAL A DECIMAL:

Para **convertir de octal a decimal** simplemente tienes que coger el número en octal de derecha a izquierda y asignar a cada uno la potencia en base ocho que le corresponde, siendo la primera de todas 8⁰, de tal forma que cuando tengas el resultado de cada una, los tienes que sumar entre sí. La suma de todos los términos te dará como resultado el número decimal que estás buscando.

EJEMPLOS:

Ejemplo 1: pasar 37 de octal a decimal

$$37_8 = 3 \times 8^1 + 7 \times 8^0 = 24 + 7 = 31_{10}$$

Ejemplo 2: convertir 7014 de octal decimal

$$7014_8 = 7 \times 8^3 + 0 \times 8^2 + 1 \times 8^1 + 4 \times 8^0 = 3584 + 0 + 8 + 4 = 3596_{10}$$

DECIMAL A HEXADECIMAL

- Ejemplos conversión de número decimal a hexadecimal

Convertiremos 51 decimal -----> Número Hexadecimal:

- 1.- Pasamos el número decimal a binario: se obtiene el número binario 110011
- 2.- Se separa la cifra binaria en grupos de 4, de derecha a izquierda: (11) (0011)
- 3.- Los números que no se completan en grupos de 4, se rellenan con ceros: (0011) (0011)

4.- Basándose en la tabla de equivalencia entre Binario y Hexadecimal, se buscan los números equivalentes: (0011) = 3 y (0011) = 3

5.- Se unen los números equivalentes en Hexadecimal: 33

DECIMAL	BINARIO	HEXADECIMAL
0	0000	0
1	0001	1
2	0010	2
3	0011	3
4	0100	4
5	0101	5
6	0110	6
7	0111	7
8	1000	8
9	1001	9
10	1010	A
11	1011	B
12	1100	C
13	1101	D
14	1110	E
15	1111	F

Figura 13. Tabla de equivalencia.

HEXADECIMAL A DECIMAL:

Método 1: para encontrar el equivalente decimal de un número hexadecimal, primero, convertir el número hexadecimal a binario, y después, el binario a decimal.

Ejemplo: Convertir a decimal los siguientes números hexadecimales:

(a) $1C_{16}$ (b) $A85_{16}$

Solución. Primero, hay que convertir a binario el número hexadecimal, y después a decimal:

(a)

$$\begin{array}{c} 1 \\ \downarrow \\ 0001 \end{array} \quad \begin{array}{c} C \\ \downarrow \\ 1100 \end{array} = 2^4 + 2^3 + 2^2 = 16 + 8 + 4 = 28_{10}$$

(b)

$$\begin{array}{c} A \\ \downarrow \\ 1010 \end{array} \quad \begin{array}{c} 8 \\ \downarrow \\ 1000 \end{array} \quad \begin{array}{c} 5 \\ \downarrow \\ 0101 \end{array} = 2^{11} + 2^9 + 2^7 + 2^2 + 2^0 = 2048 + 512 + 128 + 4 + 1 = 2693_{10}$$

Método 2: para convertir un número hexadecimal a su equivalente decimal, multiplicar el valor decimal de cada dígito hexadecimal por su peso, y luego realizar la suma de estos productos.

- Los pesos de un número hexadecimal crecen según las potencias de 16 (de derecha a izquierda).

- Para un número hexadecimal de 4 dígitos, los pesos son:

$$\begin{array}{cccc} 16^3 & 16^2 & 16^1 & 16^0 \\ 4096 & 256 & 16 & 1 \end{array}$$

Ejemplo: Convertir a decimal los siguientes números hexadecimales:

(a) $E5_{16}$ (b) $B2F8_{16}$

Solución. Las letras de la A hasta la F representan los números decimales de 10 hasta 15, respectivamente.

$$(a) E5_{16} = (E \times 16) + (5 \times 1) = (14 \times 16) + (5 \times 1) = 224 + 5 = 229_{10}$$

$$(b) B2F8_{16} = (B \times 4096) + (2 \times 256) + (F \times 16) + (8 \times 1)$$

$$= (11 \times 163) + (2 \times 162) + (15 \times 161) + (8 \times 160)$$

$$= (11 \times 4096) + (2 \times 256) + (15 \times 16) + (8 \times 1)$$

$$= 45056 + 512 + 240 + 8 = 45816_{10}$$

EJERCICIOS RESUELTOS:

1. Dado el número 450_{10} en decimal, convertirlo a base binaria, octal y hexadecimal:

Ir dividiendo entre la base a la que queremos cambiar la parte entera del número decimal original y, sucesivamente, los cocientes que se van obteniendo. Los restos de estas divisiones y el último cociente son las cifras binarias. El último cociente es el bit más significativo y el primer resto el menos significativo. Se toma el último cociente y de ahí en adelante los restos de las divisiones.

Binario:

$$450:2=225, \text{ resto } 0$$

$$225:2=112, \text{ resto } 1$$

$$112:2=56, \text{ resto } 0$$

$$56:2= 28, \text{ resto } 0$$

$$28:2=14, \text{ resto } 0$$

$$14:2=7, \text{ resto } 0$$

$$7:2=3, \text{ resto } 1$$

$$3:2=1, \text{ resto } 1$$

Por tanto: 111000010

Otra forma (la famosa tabla del subnetting...):

$$\begin{array}{cccccccc} 2^7 & 2^6 & 2^5 & 2^4 & 2^3 & 2^2 & 2^1 & 2^0 \\ \hline 128 & + & 64 & + & 32 & + & 16 & + & 8 & + & 4 & + & 2 & + & 1 & = & 255 \end{array}$$

$2^9=512$ (fuera de rango). Por tanto, a base de prueba y error cogemos el 2^8 :

$$256+128+64=448+2=450$$

2^8	2^7	2^6	2^5	2^4	2^3	2^2	2^1	2^0
-------	-------	-------	-------	-------	-------	-------	-------	-------



256	128	64	0	0	0	0	2	0
1	1	1	0	0	0	0	1	0

Octal:

$450:8= 56$, resto **2**

$56:8=7$, resto **0**

Por tanto: 702

Otro truco:

Que suele ser más rápido si tenéis solvencia en subnetting:

Pasar de decimal a binario, 450_{10} en binario es 111000010_2 , y una vez pasado a binario, como sabemos que los números octales van del 0 al 7, por tanto, un número octal son 3 bits:

$111_2=7_{10}$

Pues separamos el número binario anterior, 111000010_2 , en bloques de 3 bits empezando por la derecha, es decir: $111_2 \mid 000_2 \mid 010_2$.

Transformamos esos bloques individualmente de binario a decimal: $111_2=7_{10} \mid 000_2=0_{10} \mid 010_2=2_{10}$

Y ya tenemos el número octal correspondiente, uniendo los números decimales obtenidos anteriormente: 702_8

*Si por ejemplo nos dieran el siguiente número decimal 194, en binario sería 11000010_2 , siguiendo el truco anterior, quedarían estos bloques: $11_2 \mid 000_2 \mid 010_2$

Como podéis comprobar, el primer bloque solo tiene dos dígitos, habría que añadir un cero a la izquierda: $011_2 \mid 000_2 \mid 010_2$.

Que en octal nos quedaría: $011_2=3_{10} \mid 000_2=0_{10} \mid 010_2=2$, por tanto 302_8 .

Hexadecimal:

DECIMAL	HEXADECIMAL
0	0
1	1
2	2
3	3
4	4
5	5
6	6
7	7
8	8
9	9
10	A
11	B
12	C
13	D
14	E
15	F

$450:16=28$, resto **2**

$28:16=1$, resto **12(C)**

Por tanto: $1C2_{16}$

Otro truco:

Igual que el anterior para el octal, pero teniendo en cuenta que los números hexadecimales van del 0 al 15, por tanto, un número hexadecimal son 4 bits(un nibble o medio byte):

$1111_2=15_{10}$.

450_{10} en binario es 111000010_2 .

Separamos el número binario anterior, 111000010_2 , en bloques de 4 bits empezando por la derecha, es decir: $1_2 \mid 1100_2 \mid 0010_2$

En el primer bloque añadimos 3 “ceros” a la izquierda: $0001_2 \mid 1100_2 \mid 0010_2$

Transformamos esos bloques individualmente de binario a decimal: $0001_2=1_{10} \mid 1100_2=12_{10}$ (12 es C en hexadecimal) $\mid 0010_2=2$

Y ya tenemos el número hexadecimal correspondiente: $1C2_{16}$



2. Dado el número 0101001_2 en binario, convertirlo a base decimal, octal y hexadecimal:

Decimal:

Para convertir de binario a decimal simplemente tienes que coger el número en binario 0101001_2 de derecha a izquierda y asignar a cada uno la potencia en base dos que le corresponde, siendo la primera de todas 2^0 , de tal forma que cuando tengas el resultado de cada una, los tienes que sumar entre sí. La suma de todos los términos te dará como resultado el número decimal que estás buscando (igual que el anterior ejercicio)

$$0101001_2 = 0 \cdot 2^6 + 1 \cdot 2^5 + 0 \cdot 2^4 + 1 \cdot 2^3 + 0 \cdot 2^2 + 0 \cdot 2^1 + 1 \cdot 2^0 = 0 + 32 + 0 + 8 + 0 + 0 + 1 = 41_{10}$$

$$0101001_2 = 41_{10}$$

O el truco de la famosa tabla de subnetting:

Con $2^6 = 64$ (fuera de rango). Por tanto, a base de prueba y error cogemos el 2^5 : $32 + 8 + 1 = 41_{10}$

2^6	2^5	2^4	2^3	2^2	2^1	2^0
0	32	0	8	0	0	1
0	1	0	1	0	0	1

Octal:

Igual que anteriormente...

Separamos el número binario anterior, 0101001_2 , en bloques de 3 bits empezando por la derecha, es decir: $0_2 \mid 101_2 \mid 001_2$.

Rellenamos a la izquierda ceros... : $000_2 \mid 101_2 \mid 001_2$.

Transformamos esos bloques individualmente de binario a decimal: $000_2 = 0_{10} \mid 101_2 = 5_{10} \mid 001_2 = 1_{10}$

Y ya tenemos el número octal correspondiente: 51_8

Hexadecimal:

Igual que anteriormente...

Separamos el número binario anterior, 0101001_2 , en bloques de 4 bits empezando por la derecha, es decir: $010_2 \mid 1001_2$.

Rellenamos a la izquierda ceros... : $0010_2 \mid 1001_2$

Transformamos esos bloques individualmente de binario a decimal: $0010_2 = 2_{10} \mid 1001_2 = 9_{10}$

Y ya tenemos el número hexadecimal correspondiente: 29_{16}

3. Dado el número 533_8 en octal, convertirlos a base binaria, decimal y hexadecimal:

Binario:

Según lo visto en el ejercicio 1, cada dígito de un número octal son 3 bits, por tanto:

Transformamos cada dígito del número (a la inversa de antes), de decimal a binario, es decir:

$5_{10}=101_2$, $3_{10}=011_2$, $3_{10}=011_2$. Ahora juntamos el resultado y nos daría el número binario.

Por tanto: $533_8=101011011_2$

Decimal:

Para convertir de octal a decimal simplemente tienes que coger el número en octal 533_8 de derecha a izquierda y asignar a cada uno la potencia en base ocho que le corresponde, siendo la primera de todas 8^0 , de tal forma que cuando tengas el resultado de cada una, los tienes que sumar entre sí. La suma de todos los términos te dará como resultado el número decimal que estás buscando.

$$533_8=5*8^2+3*8^1+3*8^0=5*64+3*8+3*1=320+24+3=347; 533_8=347_{10}$$

Otro truco sería pasar primero de octal a binario y luego de binario a decimal (famosa tabla de subnetting):

$$533_8=101011011_2=347_{10}$$

$2^9=512$ (fuera de rango). Por tanto, a base de prueba y error cogemos el 2^8 :

$$256+64+16+8+2+1=347_{10}$$

2^8	2^7	2^6	2^5	2^4	2^3	2^2	2^1	2^0
256	0	64	0	16	8	0	2	1
1	0	1	0	1	1	0	1	1

Hexadecimal:

Se transforma cada dígito del número octal a binario:

$$533_8=101011011_2$$

Y ahora es sencillo, pues se trata de convertir el número binario obtenido a hexadecimal, igual que en el ejercicio 1:

Separamos el número binario anterior, 101011011_2 , en bloques de 4 bits empezando por la derecha, es decir: $1_2 | 0101_2 | 1011_2$

En el primer bloque añadimos 3 “ceros” a la izquierda: $0001_2 | 0101_2 | 1011_2$

Transformamos esos bloques individualmente de binario a decimal: $0001_2=1_{10}$ | $0101_2=5_{10}$ | $1011_2=11_{10}$ (11 es B en hexadecimal)

Y ya tenemos el número hexadecimal correspondiente: $15B_{16}$

4. Dado el número $3D_{16}$ en hexadecimal, convertirlo a base binaria, decimal y octal:

Binario:

Según lo visto en el ejercicio 1, cada dígito de un número hexadecimal son 4 bits, por tanto: Transformamos cada dígito del número (a la inversa de antes), de decimal a binario, es decir: $3_{10}=0011_2$, D (D hexadecimal a decimal es 13) $13_{10}=1101_2$. Ahora juntamos el resultado y nos daría el número binario.

Por tanto: $3D_{16}=00111101_2=111101_2$

Decimal:

Para convertir de hexadecimal a decimal simplemente tienes que coger el número en hexadecimal $3D_{16}$ de derecha a izquierda y asignar a cada uno la potencia en base dieciséis que le corresponde, siendo la primera de todas 16^0 , de tal forma que cuando tengas el resultado de cada una, los tienes que sumar entre sí. La suma de todos los términos te dará como resultado el número decimal que estás buscando.

$$3D_{16}=3*16^1+13*16^0=48+13=61_{10}$$

Otro truco sería pasar primero de hexadecimal a binario y luego de binario a decimal (famosa tabla de subnetting):

$$3D_{16}=111101_2=61_{10}$$

$2^6=64$ (fuera de rango). Por tanto, a base de prueba y error cogemos el 2^5 :

$$32+16+8+4+1=61_{10}$$

2^5	2^4	2^3	2^2	2^1	2^0
32	16	8	4	0	1
1	1	1	1	0	1

Octal:

Se transforma cada dígito del número hexadecimal a binario:

$$3D_{16}=111101_2$$

Y ahora es sencillo, pues se trata de convertir el número binario obtenido a octal, igual que en el ejercicio 1:

Separamos el número binario anterior, 111101_2 , en bloques de 3 bits empezando por la derecha, es decir: 111_2 | 101_2 .



Transformamos esos bloques individualmente de binario a decimal: $111_2=7_{10}$ | $101_2=5_{10}$

Y ya tenemos el número octal correspondiente: 75_8

*Herramienta online para comprobación de conversión entre todas las bases:

<https://www.rapidtables.com/convert/number/base-converter.html>

2. Arquitectura de ordenadores

2.1. Introducción

Un ordenador personal (PC) es un conjunto de hardware (componentes y dispositivos físicos) y software (programas lógicos que controlan el funcionamiento) utilizado para el tratamiento automatizado de la información (concepto directamente ligado con el de Informática). En la combinación de tareas de este tratamiento intervienen, dos procesos importantes, que caracterizan a todo sistema informático: el funcionamiento del ordenador y las técnicas de procesamiento de los datos (almacenamiento, proceso y control).

Por otra parte, un sistema Informático es el conjunto formado por uno o varios ordenadores y sus periféricos (componentes físicos o hardware), que ejecutan aplicaciones informáticas (componente lógico o software) y que son controlados por personal especializado (componente humano). Con más detalle se puede decir que:

La arquitectura de un Sistema Informático define el conjunto de reglas, normas y procedimientos que especifican las interrelaciones que deben existir entre los componentes del Sistema Informático y las características que deben cumplir cada uno de estos componentes.

2.2. Componentes de un PC

//importante esta parte, quedaros unicamente con los conceptos//

Un PC está formado por los siguientes componentes (sistemas) físicos y lógicos:

1. El sistema **Físico** consta de los siguientes subsistemas:

a) El subsistema **Central**: está constituido por la unidad de control, la unidad aritmética lógica, la memoria principal, y los buses. Al conjunto formado por la unidad aritmética lógica y la unidad de control se le denomina Unidad Central de Proceso (CPU).

i. Unidad de Control: dirige todas las actividades del ordenador, generándose señales de control para el resto de las unidades según el código de operación de la instrucción. La ejecución de la instrucción supone realizar una serie de operaciones elementales a través de la activación de distintas señales de control.

ii. Unidad Aritmética Lógica: ejecuta las operaciones aritméticas lógicas que le señala la instrucción residente en la unidad de control.

iii. Memoria Principal: almacena las instrucciones de los programas, los datos que ha de procesar y los resultados que obtiene al ejecutar los programas. Está constituida por celdas o elementos capaces de almacenar 1 bit de información (elemento de memoria que tiene el estado 0 o el estado 1). La memoria se organiza en conjuntos de elementos de un tamaño determinado y, a cada palabra, le corresponde una dirección única.

iv. Buses: son los canales de comunicación entre los distintos elementos del subsistema Central. Un bus es un conjunto de líneas conductoras que permiten transmitir direcciones, datos y señales de control entre los componentes de un sistema informático. Los buses pueden ser internos (como los que posee la CPU para interconectar sus componentes internos) o externos (como los que interconectan los



periféricos o la memoria con el microprocesador). Más específicamente se pueden mencionar los siguientes tipos de buses:

- I. De datos: son las líneas de comunicación por donde circulan los datos externos e internos del microprocesador.
- II. De dirección: línea de comunicación por donde viaja la información específica sobre la localización de la dirección de memoria del dato o dispositivo al que se hace referencia.
- III. De control: línea de comunicación por donde se controla el intercambio de información con un módulo de la unidad central y los periféricos.
- IV. De expansión (I/O): conjunto de líneas de comunicación encargadas de llevar el bus de datos, el bus de dirección y el de control a la tarjeta de interfaz (entrada, salida) que se agrega a la tarjeta principal.
- V. Del sistema: todos los componentes de la CPU se vinculan a través del bus de sistema, mediante distintos tipos de datos, el microprocesador y la memoria principal. La velocidad de transferencia del bus de sistema está determinada por la frecuencia del bus y el ancho del mínimo.

b) El subsistema de **Entrada/Salida**: constituido por procesadores de entrada/salida, controladores o procesadores de periféricos y los propios periféricos.

Este subsistema interactúa directamente con el subsistema Central a través de los procesadores de entrada/salida. El procesador de entrada/salida gestiona el acceso a la memoria del subsistema Central y la interacción con el controlador o procesador de periféricos. La conexión del procesador de entrada/salida y del controlador periférico se realiza a través de un canal por el que se envían datos, señales de control, etc. El procedimiento de intercambio de información y los protocolos utilizados están estandarizados, permitiendo que distintos controladores de periféricos utilicen el mismo interfaz. Una operación de entrada/salida es puesta en marcha por un componente software, denominado programa de canal.

Los sistemas operativos suelen tener rutinas para gestionar directamente los periféricos ya que sería extremadamente difícil realizarlo desde los lenguajes de alto nivel.

Los periféricos según su función pueden ser de entrada, de salida o que realizan funciones de almacenamiento auxiliar.

- Unidades de Entrada: son dispositivos por los que se introducen los datos al ordenador. En estas unidades se transforma la información codificada en señales eléctricas comprensibles por la CPU.
- Unidades de Salida: son los dispositivos en los que se obtienen los resultados de los programas ejecutados en el ordenador, normalmente transformando las señales eléctricas en caracteres escritos o visualizados.
- Unidades que realizan funciones de Almacenamiento Auxiliar: aunque la memoria principal es muy rápida, su capacidad es reducida y además la información contenida en ella se volatiliza en cuanto se apaga el ordenador, por ello, hay que recurrir a la memoria auxiliar mucho más lenta pero de mayor capacidad y de carácter no volátil. Los dispositivos típicos de memoria auxiliar son los discos y las cintas magnéticas. Son unidades de entrada/salida.

c) El subsistema de **Comunicaciones**: constituido también por procesadores de entrada/salida, procesadores de comunicaciones y por los elementos físicos de las redes de comunicaciones. Los componentes físicos del subsistema de comunicaciones son los procesadores y líneas de comunicaciones, los módems, los concentradores, los multiplexores, los terminales, etc. Los procesadores de comunicaciones realizan, en relación con el subsistema central, funciones similares a los controladores o procesadores de periféricos. El enlace es también a través del procesador de entrada/salida. Los procesadores de comunicaciones gestionan la red de comunicaciones.

Los procesadores de entrada/salida aunque forman parte, según sus funciones, del subsistema de entrada/salida o del subsistema de comunicaciones físicamente, muchas veces están

incluidos en el procesador central. En cualquier caso, el esquema puede variar según el sistema físico concreto.

Arquitectura de un sistema físico

En una primera acepción se entiende que dos ordenadores tienen la misma arquitectura si tienen el mismo juego de instrucciones. Sin embargo, por arquitectura se entiende un término más amplio que incluye todos los aspectos relativos a la estructura, organización y funcionamiento del Sistema.

Se pueden mencionar las siguientes arquitecturas de Sistemas Físicos:

- Arquitectura de bus único: en la que todos los componentes del ordenador —unidad de control, unidad aritmética lógica, memoria, etc.— intercambian datos compartiendo el bus. Es una arquitectura típica de los mini y microordenadores.
- Arquitectura distribuida: sistemas multiprocesadores en los que las diferentes funciones del Sistema se encuentran distribuidas en procesadores especializados: de instrucciones, de cálculo, de direcciones, etc.
- Arquitectura paralela: varios procesadores idénticos trabajan en paralelo. Existen diferentes tipos, según la taxonomía de Flynn: SISD (*Single Instruction Single Data*), SIMD (*Single Instruction Multiple Data*), MISD (*Multiple Instruction Single Data*) y MIMD (*Multiple Instruction Multiple Data*).
- Arquitectura “Pipeline”: la ejecución de una instrucción se divide en fases. Existen componentes especializados dentro del procesador para la ejecución de cada una de las fases.
- Arquitectura “Fault Tolerant”: es la arquitectura en la que existe suficiente redundancia y, por tanto, permite asegurar un funcionamiento correcto del Sistema, aún en caso de fallo de algunos de sus elementos.

2. El sistema **Lógico** consta de:

a) Software de **sistema**: son programas imprescindibles para el funcionamiento de un ordenador, que administran los recursos hardware de éste y facilitan ciertas tareas básicas al usuario y a otros grupos de programas. El ejemplo más importante de este tipo de software es el sistema operativo.

Un sistema operativo es un programa, que se ejecuta en un procesador, encargado de llevar a cabo una serie de funciones y cuyo objetivo es simplificar el manejo y la utilización del ordenador de modo seguro y eficiente. Las funciones que lleva a cabo un sistema operativo son las siguientes:

- ☑ Permitir la ejecución de programas ofreciendo para ello una máquina extendida.
- ☑ Administración y gestión de los recursos del ordenador.
- ☑ Interfaz de usuario.

Igualmente, se encarga de ofrecer a las aplicaciones una serie de servicios, conocidos como llamadas al sistema. Entre ellas se pueden mencionar: ejecución de programas, operaciones de Entrada/Salida y detección y tratamiento de errores.

b) Software de **desarrollo**: también se les denomina lenguajes de programación o sistemas de desarrollo. Son programas que sirven para crear otros programas. Utilizan los servicios prestados por el software del sistema.

c) Software de **aplicación**: son el resultado de los desarrollos realizados con los lenguajes de programación del grupo anterior, que originan las aplicaciones que manejan los usuarios finales.

En este grupo se encuentran los procesadores de texto, editores gráficos, programas de diseño, bases de datos, etc. Hace uso de los recursos que ofrece el software de sistema.

3. La Placa Base



3.1. Introducción

La **placa base** (*mother board*) de un ordenador es el dispositivo sobre el que se insertan los demás componentes del PC, tales como el *microprocesador*, las diferentes *tarjetas de expansión* y la *memoria*. La función principal de la placa base es servir de vía de comunicación entre los citados componentes, proporcionando las líneas eléctricas necesarias y las señales de control para que todas las transferencias de datos se lleven a cabo de manera rápida y fiable. Se diseña básicamente para realizar tareas específicas vitales para el funcionamiento del computador, como por ejemplo:

- ☐ Conexión física.
- ☐ Administración, control y distribución de energía eléctrica.
- ☐ Comunicación de datos.
- ☐ Temporización.
- ☐ Sincronismo.
- ☐ Control y monitorización.

3.2. BIOS y UEFI

//vip! Mas de una vez lo han preguntado//

Para que la placa base cumpla con su cometido, tiene instalado un software muy básico denominado **BIOS** (*Basic Input Output System*). Este software es almacenado en un chip de memoria EPROM* (*Erasable Programmable Read-Only Memory*, memoria sólo lectura que puede ser borrada y programada), cuyo contenido permanece inalterable al apagar el PC. Se puede reprogramar y es el primer software en ejecutarse en el proceso de arranque de una placa base.

El software del BIOS es almacenado en un circuito integrado de memoria ROM no volátil en la placa base. Está específicamente diseñado para trabajar con cada modelo de computadora en particular, interconectando los diversos dispositivos que componen el conjunto de chips complementarios del sistema. En computadoras modernas, el BIOS está almacenado en una memoria flash, por lo que su contenido puede ser reescrito sin retirar el circuito integrado de la placa base. Esto permite que el BIOS sea fácil de actualizar para agregar nuevas características o corregir errores, pero puede hacer que la computadora sea vulnerable a los rootkit de BIOS.

La Interfaz de Firmware Extensible Unificada, Unified Extensible Firmware Interface (**UEFI**), es una especificación que define una interfaz entre el sistema operativo y el firmware. UEFI reemplaza la antigua interfaz del Sistema Básico de Entrada y Salida (BIOS) estándar presentado en las computadoras personales IBM PC como IBM PC ROM BIOS. La UEFI puede proporcionar menús gráficos adicionales e incluso proporcionar acceso remoto para la solución de problemas o mantenimiento.

La interfaz UEFI incluye bases de datos con información de la plataforma, inicio y tiempo de ejecución de los servicios disponibles listos para cargar el sistema operativo.

UEFI destaca principalmente por:

- Compatibilidad y emulación del BIOS para los sistemas operativos solo compatibles con esta última.
- Soporte completo para la Tabla de particiones GUID (GPT), se pueden crear hasta 128 particiones por disco, con una capacidad total de 8 ZB.1
- Capacidad de arranque desde unidades de almacenamiento grandes, dado que no sufren de las limitaciones del MBR.
- Independiente de la arquitectura y controladores de la CPU.
- Entorno amigable y flexible Pre-Sistema Operativo, incluyendo capacidades de red.
- Diseño modular.

La EFI hereda las nuevas características avanzadas del BIOS como ACPI (Interfaz Avanzada de Configuración y Energía) y el SMBIOS (Sistema de Gestión de BIOS), y se le pueden añadir muchas otras, ya que el entorno se ejecuta en 64 bits y no en 32 bits, como su predecesora. La EFI comunica el arranque además de con el ya clásico MBR (Master Boot Record), con el sistema GPT (tabla de particiones GUID) que solventa las limitaciones técnicas del MBR, a saber:

- Solo los procesadores little-endian pueden ser soportados.
- MBR soporta hasta 4 particiones Primarias por unidad física (si se desea realizar más particiones se tiene que convertir una o varias de esas particiones primarias a una o varias "particiones extendidas"). Todas las particiones independientemente de "su tipo" tienen un límite de 2,2 TB, es decir, un disco duro u otro dispositivo de almacenamiento de más de 2,2 TB no se podría aprovechar su capacidad al 100% mediante el uso de una única partición (sigue siendo posible particionar ese HDD/Dispositivo de almacenamiento en varias particiones de 2,2 TB).
- GPT soporta teóricamente hasta 9,4 ZB y no exige un sistema de archivos concreto para funcionar.
- Microsoft Windows soporta GPT a partir de las versiones de 64 bits de Windows Vista y posteriores.
- Algunos sistemas basados en Unix utilizan un híbrido entre MBR y GPT para arrancar.

3.3. Componentes

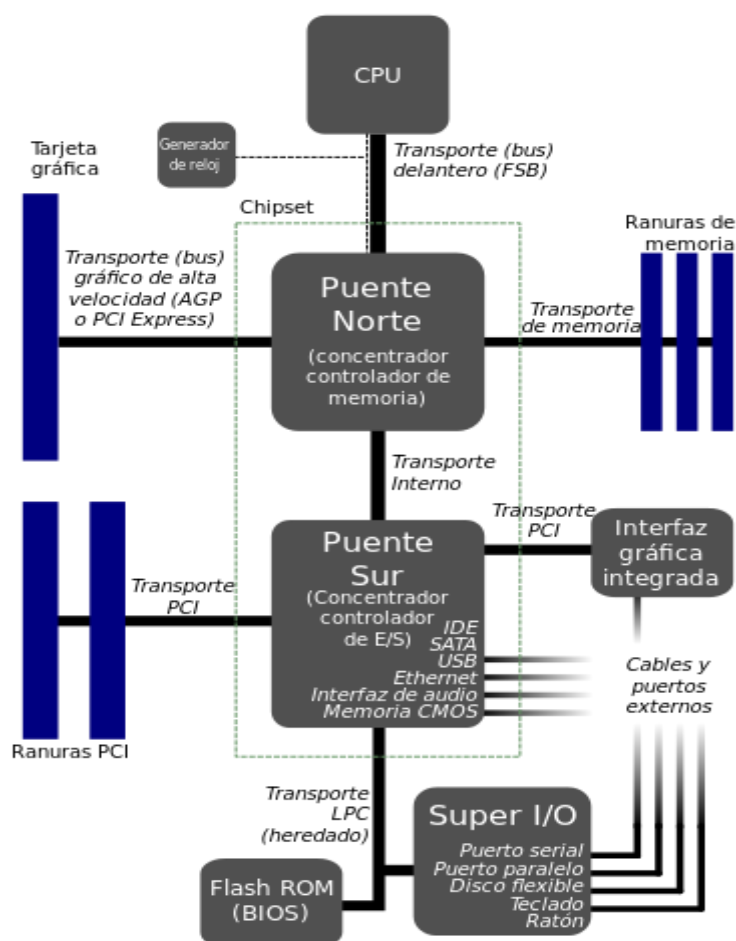


Figura 14. Diagrama de una placa base típica.

Los principales componentes de una placa base son los siguientes:

//VIP!!!, sobretodo los conectores de que tipo son cada uno//

a) Zócalos o ranuras: son elementos para la inserción de componentes hardware. Hay dos tipos de zócalos:

i. Zócalos para el procesador: los principales son Intel y AMD. Tienen dos denominaciones distintas:

- *Slot*: zócalo de microprocesador de formato longitudinal.
- *Socket*: zócalo de microprocesador cuadrangular o rectangular. El *socket* 478 contiene 478 pines y es de la marca de Intel. Fue reemplazado por el socket 775 que carece de pines. En la actualidad existen diferentes tipos de sockets como los de la serie 1xxx o el 2011 utilizado por los microprocesadores de gama alta.

La placa base necesita una configuración de zócalo de microprocesador en función del tipo de procesador.

ii. Zócalos para la memoria: sirven para la inserción de módulos de memoria y siempre son de tipo SLOT.

Los zócalos de la memoria pueden ser de tipo DIMM (*Dual Inline Memory Module*) que son reconocibles externamente por poseer sus contactos (o pines) separados en ambos lados. Los hay para módulos de 168, 184 y 240 contactos. Antecedentes de los DIMM están los módulos SIMM (*Single Inline Memory Module*), que poseen los contactos de modo que los de un lado están unidos con los del otro.

Un DIMM puede comunicarse con el PC a 64 bits (y algunos a 72 bits) en vez de los 32 bits de los SIMM. En DIMM se insertan módulos de memoria a ambas caras del módulo, en cambio en SIMM sólo se hace por una cara.

Hay casos en los que no sirve ninguno de los módulos anteriores, entonces se usa SO-DIMM (*Small Outline Dual Inline Memory Module*) que consisten en una versión compacta de los módulos DIMM convencionales. Se usan para dispositivos Tablet PC, impresoras, portátiles, algunas PDAs, etc.

Los módulos SO-DIMM tienen 100, 144 o 200 pines. Los de 100 pines soportan transferencias de datos de 32 bits, mientras que los de 144 y 200 lo hacen a 64 bits. Estas últimas se comparan con los DIMMs de 168 contactos (que también realizan transferencias de 64 bits). A simple vista, se diferencian porque las de 100 tienen 2 hendiduras guía, las de 144 una sola hendidura casi en el centro y las de 200 una hendidura parecida a la de 144 pero más desplazada hacia un extremo.

b) Voltajes: los valores de tensión están entre 3 y 12 Voltios.

c) Batería: es una pequeña pila que viene incorporada en todas las placas base y su función básica es mantener la alimentación eléctrica del reloj de tiempo real (RTC), así como diversos parámetros (sobre el disco duro y de configuración de usuario), que son utilizados por la BIOS en el arranque del computador.

d) Frecuencia (velocidad): es la que maneja el elemento de conexión único que existe en la placa base y el FSB (*Front Side Bus*). Este bus representa el camino por el cual es posible integrar en la placa base los distintos componentes hardware para el intercambio de información entre microprocesador, memoria y el subsistema de Entrada/ Salida (subsistema de E/S).

Para realizar intercambio de información con cualquiera de los elementos microprocesador, memoria y subsistema de E/S se debe hacer uso del FSB. La velocidad del FSB es siempre inferior a la velocidad interna del microprocesador.

En todo PC hay que diferenciar dos frecuencias distintas: la frecuencia interna o velocidad de ciclo del procesador y la frecuencia externa o velocidad del FSB. Hay que ajustar velocidades de los distintos componentes hardware. El elemento que introduce esa posibilidad de ajuste entre distintas velocidades es el *chipset*.

e) *Chipset*: es un conjunto de circuitos integrados con alta escala de integración diseñado con técnicas similares a las del microprocesador y que se especializan en la realización de dos técnicas fundamentales:

- i. Controlar los intercambios de información entre microprocesador y memoria principal usando el FSB.
- ii. Controlar los intercambios de información entre microprocesador y subsistema de E/S haciendo uso del FSB.

En las placas base antiguas no existían los *chipset*. En estas placas existía un bus (ISA o EISA) que conectaba todos los componentes: la memoria, las tarjetas de controladores, tarjetas de video, procesador, teclado, etc. Todos estos componentes (incluyendo la CPU) funcionaban a la misma velocidad. Cada nueva generación de procesadores requería unos componentes completamente distintos y el bus, cada vez, era más exigido para proporcionar mayor velocidad. Algunos fabricantes intentaron construir dispositivos que pudieran funcionar a distintas velocidades, aunque no lo consiguieron lograr con la suficiente calidad.

Para reducir los problemas de compatibilidad y evitar el consiguiente gasto en cada nueva generación de procesadores, Compaq (en 1985) separó la velocidad del bus de la velocidad del procesador. De esta forma, se evitaba que con cada nuevo procesador se tuvieran que construir nuevas versiones de los dispositivos conectados. Esta mejora fue muy importante, pero requería una nueva interfaz para que todos los dispositivos pudieran seguir funcionando a distintas velocidades.

Los fabricantes de placas base comenzaron a proporcionar más chips para componentes en vez de la tarjetas de expansión. Los primeros fueron para los controladores de teclado, a los que se añadieron controladores de discos flexibles (*diskettes*), puertos serie, y controladores de disco duro (año 1994).

Se había comenzado el cambio de los componentes integrados en las placas base.

Después de sufrir durante mucho tiempo los buses lentos, Intel introdujo (en 1994) el bus PCI pensado como un bus de expansión, más que propiamente como solo un bus. Con PCI, se comenzó el uso de un lenguaje propio y único, con un *bridge* incorporado que permitía la traducción de los comandos PCI al lenguaje de la CPU y de la memoria RAM. De esta forma, se podía trabajar tanto con el bus PCI como con la CPU y la memoria RAM de una forma sencilla. Este *bridge* fue llamado **North Bridge**.

Para conservar la compatibilidad con los buses anteriores (ISA o EISA) y los componentes no integrados de estos buses, se construyó otro *bridge* (**South Bridge**) que se comunicaba con ellos.

Todos estos componentes constituyeron el origen de los *chipset*'s.

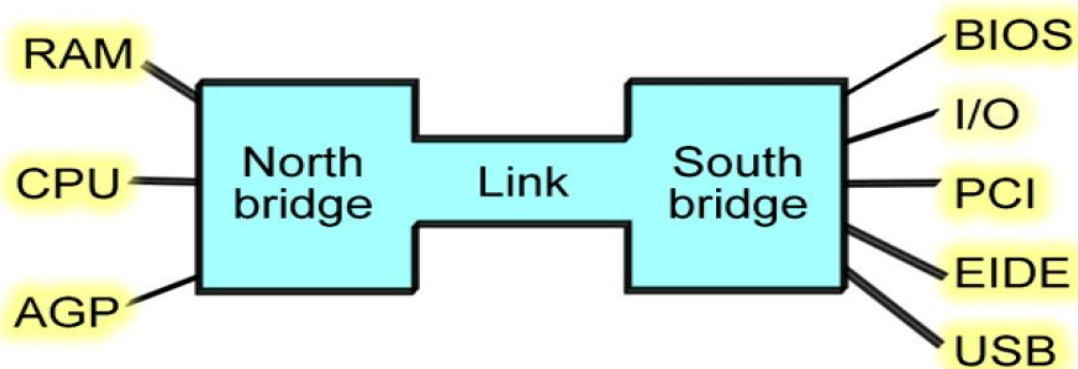


Figura 15. Comunicación entre North Bridge y South Bridge.

En la actualidad, cualquier placa base debería proporcionar, como mínimo, dos chipsets integrados a la placa (no insertados en zócalos)

- i. Chipset1, que vincule microprocesador y memoria principal.
- ii. Chipset2, que vincule microprocesador y subsistema de E/S.

El *chipset* es un elemento esencial de la placa base. Hace posible que la placa base funcione como eje del sistema, dando soporte a varios componentes e interconectándolos de forma que se comuniquen entre ellos a través de diversos buses. Es uno de los pocos elementos que tiene conexión directa con el procesador, gestiona la mayor parte de la información que entra y sale por el bus principal del procesador, del sistema de vídeo y muchas veces de la memoria RAM.

El primer grado de obsolescencia de una placa base depende del chipset, ya que no hay posibilidad de modificar el chipset pues va soldado a ella. El grado de actualización, en cuanto a conectores y buses de E/S, está condicionado por el chipset². Algunos fabricantes de chipsets son Intel, AMD, VIA Technologies, NVidia y ALI.

Las diferentes funciones lógicas que suele integrar cualquiera de los chipsets actuales son las siguientes:

- Soporte para el microprocesador. Una de las principales funciones del chipset es la detección correcta del microprocesador y el pleno soporte de todas sus funciones. Esta es una de las principales razones de la rápida evolución de los chipsets: nuevos procesadores, cada vez más rápidos, necesitan nuevos chipsets que les proporcionen un soporte completo. Cada chipset se diseña pensando en un procesador o familia de procesadores concretos. Igualmente, el chipset es el responsable directo de que la placa base soporte más de un microprocesador, en el caso de placas base duales o con más de dos microprocesadores.
- Controlador de memoria (MMU, *Memory Management Unit*). Gestiona la memoria RAM del sistema y en general todo el subsistema de memoria, incluidos los diferentes niveles de memoria caché.
- Controlador IDE/ATAPI para discos duros y otros dispositivos de almacenamiento que cumplan con el estándar IDE/ATAPI. Esta parte del chipset es la encargada de controlar los dos conectores IDE que habitualmente suelen integrar todas las placas bases actuales. Directamente relacionados con esta función, están los modos de transferencia.
- Control de los periféricos y del bus de E/S. Las placas base actuales disponen de una serie de buses, principalmente PCI y AGP. El chipset es el responsable de la gestión de los buses PCI y de ofrecer el soporte para el bus gráfico AGP. Algunas de las denominaciones de los chipsets históricos de Intel están en línea directa con esta función, con denominaciones como PCIsset y AGPset. Esta función también incluye el soporte para tecnologías más modernas como USB o HDMI.
- Controlador de Interrupciones. Gestiona todo el sistema de interrupciones del PC.
- Reloj de Tiempo Real (RTC). Mantiene la hora del sistema.
- El soporte para la gestión de energía. Todos los *chipsets* actuales soportan una serie de funciones para gestión y ahorro de energía eléctrica
- Controlador de Acceso Directo a Memoria (DMA, *Direct Memory Access*). Permite el acceso directo a la memoria a determinados dispositivos, sin pasar por el microprocesador, lo que agiliza el rendimiento de ciertas operaciones con dispositivos específicos como los discos duros. El DMA es controlado por una parte del chipset denominada controlador de DMA. Igualmente, el driver es el que soporta la función de arbitraje de bus (bus mastering), que es una mejora del DMA, que permite que un dispositivo tome directamente el control del bus del sistema para llevar a cabo las transferencias de datos.
- Controlador de infrarrojos (IrDA). Controla la conexión de dispositivos que funcionen mediante rayos infrarrojos.
- Controlador tipo PS/2. Controla toda la actividad y el funcionamiento del teclado y de ratones con este formato.



f) Elementos de expansión: la expansión está asociada con la funcionalidad del Chipset2 y se concreta con buses de expansión que pueden dar entrada a interfaces de E/S estandarizados y actualizados lo más posible.

g) Puertos y conectores:

Conectores externos

- Puertos serie controlados por un chip UART (*Universal Asynchronous Receiver-Transmitter, Transmisor Asíncrono Universal*) controlador serie de alta velocidad, compatible con el chip UART 16550 original de *National Semiconductor*. Estos dos puertos se suelen denominar COM1 y COM2 respectivamente. El COM1 se usa habitualmente para la conexión de ratones serie, mientras que el COM2 queda libre para la conexión de dispositivos serie tales como módems externos. Existen conectores serie externos de tipo DB9 y DB25 (de 9 y 25 patillas, respectivamente).
- Puerto paralelo multimodo para la conexión de dispositivos paralelos, generalmente impresoras y, en menor medida, escáneres y otros dispositivos de esta misma naturaleza. El término multimodo hace referencia a que el puerto es capaz de soportar tres modos de funcionamiento característicos: SPP (*Single Paralell Port*), ECP (*Extended Capabilities Port*) y EPP (*Enhanced Paralell Port*). Éstos suelen ser los tres modos habituales de transmisión del puerto paralelo, en orden creciente de prestaciones. El conector es de tipo hembra con 25 pines agrupados en dos filas.
- Puertos USB (*Universal Serial Bus*). Tienen forma estrecha y rectangular y permiten la conexión en caliente de dispositivos que cumplan este estándar. Actualmente, la mayoría de placas base soporta la especificación USB 3.1.
- Puerto IEEE 1394 (*Firewire*). Permiten la conexión en caliente de dispositivos que cumplan este estándar de alta velocidad. Actualmente los principales dispositivos IEEE 1394 son sistemas de video digital (DV) y unidades de almacenamiento externas.
- Puertos PS/2 uno para teclado y otro para ratón. Ambos son conectores de tipo mini-DIN de seis patillas. Éste suele ser el tipo habitual de conectores para ratón y teclado en las actuales placas base ATX.
- Puerto para juegos en el que habitualmente se suelen conectar dispositivos como palancas o mandos de juegos (*Joysticks* y *Gamepads*), o dispositivos de audio, tales como teclados MIDI (*Musical Instrument Digital Interface*).
- Conectores de audio, generalmente para clavijas de tipo *jack* estéreo, siendo los más habituales los de entrada y/o salida de línea (*line in/line out*), entrada de micrófono (*mic in*) y salida de altavoces (*speaker out*).
- Conector VGA para la tarjeta gráfica: es un conector estándar para tarjeta gráfica. Consta de 15 pines agrupados en tres filas.
- Conector RJ45 para conectar redes de ordenadores con cableado estructurado.
- Conector DVI "Interfaz Visual Digital", interfaz de video diseñada para obtener la máxima calidad de visualización posible en pantallas digitales, tales como los monitores LCD de pantalla plana y los proyectores digitales.
- Conector HDMI para video de alta definición.

Conectores internos.

- Conector para disquetera. Donde se inserta la banda de cable para datos de la disquetera.
- Conectores IDE. Para la conexión de dispositivos IDE, principalmente discos duros y lectores de CD- ROM.
- Conectores para el refrigerador del microprocesador. Denominados generalmente *Fan Power*.
- Conector para arranque desde red (**WakeOn-LAN, WoL**). Permite encender y apagar remotamente nuestro equipo. Wake on LAN (WoL) es un protocolo que permite encender de forma remota tu ordenador cuando este esté apagado, suspendido o hibernando. Lo hace enviándole cuando tu quieras una señal de forma remota a través de Internet o

cualquier red, y cuando el ordenador la detecta, entonces pasa a encenderse. El protocolo WoL incluye un componente llamado ***Paquete mágico***, que es un paquete que se envía por la red local compuesto por una cadena de 6 bytes cuyo valor es 255 en valor hexadecimal, la cual es seguida por 16 repeticiones de la dirección MAC del equipo al que se le envía. ***DIRECCION MAC** (se estudiara en temas del bloque IV, pero se aclara el concepto): En las redes de computadoras, la dirección MAC (siglas en inglés de *Medium Access Control*) es un identificador de 48 bits (6 bloques de dos caracteres hexadecimales 8 bits) que corresponde de forma única a una tarjeta o dispositivo de red. Se la conoce también como dirección física, y es única para cada dispositivo.

- Conector para arranque desde línea telefónica mediante módem (*WakeOn-Ring*). Permite encender y apagar remotamente nuestro equipo a través de la línea telefónica mediante un módem que soporte esta función.
- Conector para módulo de infrarrojos (irDA).
- Controlador de Entrada/Salida: similar al chipset, es un chip que integra la mayoría de funciones de entrada y salida de información de la placa base. Dentro de este chip se suelen integrar, entre otros, los siguientes elementos: Incorporan una UART (controlador de puertos serie), el controlador del puerto paralelo, la controladora de la disquetera. Muchos chips de
- E/S incluyen además el reloj de tiempo real y el controlador de teclado.

h) Sistemas de monitorización de la placa base

Desde hace un tiempo se han hecho bastante frecuentes una serie de sistemas y estándares para la monitorización de determinadas características de la placa base, lo que permite disponer, en tiempo real, de un completo diagnóstico sobre el estado de salud de nuestro PC. Se pueden mencionar las siguientes tecnologías y estándares existentes relacionados con esta monitorización:

- **DMI (*Desktop Management Interface*)**: es un estándar que permite la monitorización de determinadas funciones del PC, así como parámetros que pueden servir como indicadores del buen o mal funcionamiento de los diferentes componentes, todo ello, de manera centralizada y única. El DMI es un estándar totalmente independiente del hardware y del sistema operativo y permite gestionar, tanto ordenadores aislados, como conectados en red.

En relación con el ahorro de energía, la agencia estadounidense de protección medioambiental (EPA) desarrolló (en 1992) el programa *Energy Star*, que certificaba a aquellos PC que se consideraban eficientes en cuanto al ahorro de energía. A mediados de 2007, se actualizó esta especificación.

Actualmente, este programa es voluntario para los integradores de sistemas. Solo es obligatorio para adquisiciones públicas y del gobierno federal en Estados Unidos y Europa (los 27 estados). Pero la demanda de sistemas certificados aumenta, a medida que crece el coste energético y la toma de conciencia sobre la conservación del medio ambiente.

- **APM (*Advanced Power Management*)**: es un componente software (desarrollado por Intel y Microsoft) que se integra en determinados sistemas operativos junto con la BIOS del sistema, y basándose en ésta, permite controlar las diferentes funciones de ahorro de energía (la velocidad de la CPU, apagar el disco duro o apagar el monitor después de un período de inactividad para conservar corriente eléctrica).
- **ACPI (*Advanced Configuration and Power Interface*)**: es un estándar de gestión de energía desarrollado por Intel, Microsoft y Toshiba en 1996 que permite pasar todo el control de la gestión de energía directamente al sistema operativo, lo que facilita funciones tales como apagar automáticamente el sistema o dejarlo en estado de hibernación. Todas las versiones actuales de Microsoft Windows lo soportan.

Tanto APM como ACPI son dos estándares de gestión de energía. La diferencia fundamental es que APM realiza la gestión desde la BIOS, a un nivel más bajo, mientras que con ACPI la



práctica totalidad de funciones de gestión de energía se encuentra en el propio sistema operativo, que debe incluir soporte para ACPI.

i) Las peticiones de interrupción (IRQ *Interrupt ReQuest*)

Los recursos del sistema cobran una gran importancia, puesto que son compartidos por diferentes dispositivos físicos. La cantidad de recursos del sistema es bastante limitada por razones derivadas del diseño del PC original, de tal forma, que a medida que se añaden dispositivos al sistema, la cantidad de recursos disponibles disminuye, hasta el punto de convertirse en un serio problema en sistemas con gran cantidad de dispositivos instalados. Esto puede llevar a conflictos entre dispositivos que intenten hacer uso de un mismo recurso, generando uno de los principales problemas en la configuración de un PC.

Básicamente, una interrupción es un mensaje enviado por algún componente del PC a otro componente, generalmente al microprocesador, que indica a éste que debe detener la ejecución de todo lo que esté haciendo, atender al dispositivo que envía la petición de interrupción y, posteriormente, continuar donde lo había dejado. Las señales enviadas se denominan peticiones de interrupción ó IRQ. **//no es necesario saberse los IRQ de memoria//**

Interrupción	Utilidad
IRQ0	Timer Output 0
IRQ1	Teclado
IRQ2	Reservada
IRQ3	Puerto serie COM 2
IRQ4	Puerto serie COM 1
IRQ5	Libre
IRQ6	Controlador unidad 0 de discos
IRQ7	Puerto paralelo LPT 1
IRQ8	Reloj-Calendario
IRQ9	Libre
IRQ10	Libre
IRQ11	Pantalla
IRQ12	Raton PS-2
IRQ13	Coprocesador matematico
IRQ14	Interfaz primaria para discos
IRQ15	Interfaz secundaria para discos

Figura 16. IRQ actuales.

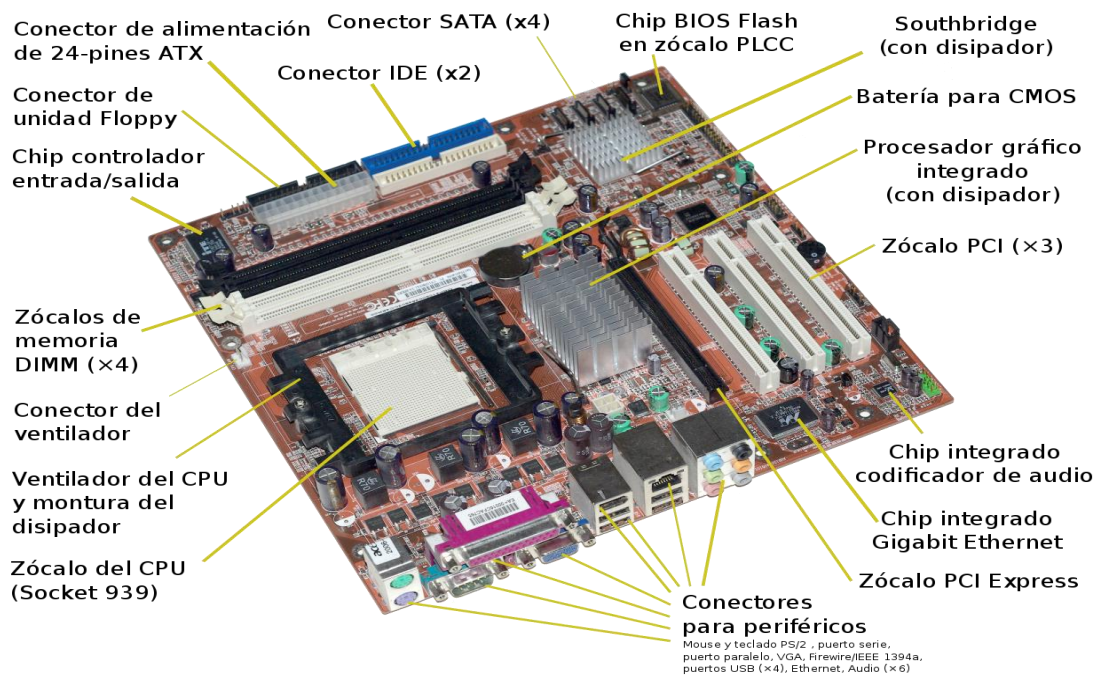


Figura 17. Placa Base ATX.

Con la evolución de las computadoras, más y más características se han integrado en la placa base, tales como circuitos electrónicos para la gestión del vídeo, de sonido o de redes, evitando así la adición de tarjetas de expansión:

- interfaz gráfica integrada o **unidad de procesamiento gráfico (GPU, Graphics Processing Unit, o IGP, Integrated Graphic Processor)**;
- interfaz integrada de audio o sonido;
- interfaz integrada **Ethernet** o puertos de red integrados ((10/100 Mbit/s)/(1 Gbit/s)).

En la placa también existen distintos conjuntos de pines, llamados **jumpers** o puentes, que sirven para configurar otros dispositivos:

- JMDM1: Sirve para conectar un módem por el cual se puede encender el sistema cuando este recibe una señal.
- JIR2: Este conector permite conectar módulos de infrarrojos IrDA, teniendo que configurar la BIOS.
- JBAT1: Se utiliza para poder borrar todas las configuraciones que como usuario podemos modificar y restablecer las configuraciones que vienen de fábrica.
- JP20: Permite conectar audio en el panel frontal.
- JFP1 Y JFP2: Se utiliza para la conexión de los interruptores del panel frontal y los ledes.
- JUSB1 Y JUSB3: Es para conectar puertos USB del panel frontal.

3.3. Formatos de placa base

//esto lo preguntaron un año//

Las placas base necesitan tener dimensiones compatibles con las cajas que las contienen, de manera que desde los primeros computadores personales se han establecido características mecánicas, llamadas **factor de forma**. Definen la distribución de diversos componentes y las dimensiones físicas, como por ejemplo el largo y ancho de la tarjeta, la posición de agujeros de sujeción y las características de los conectores.

Con los años, varias normas se fueron imponiendo.

XT

1983: XT (sigla en inglés de *eXtended Technology*, «tecnología extendida») es el formato de la placa base de la computadora IBM PC XT (modelo 5160), lanzado en 1983. En este factor de forma se definió un tamaño exactamente igual al de una hoja A4 y un único conector externo para el teclado.

AT

1984: AT (*Advanced Technology*, «tecnología avanzada») es uno de los formatos más grandes de toda la historia de la PC (305×279–330 mm), definió un conector de potencia formado por dos partes. Fue usado de manera extensa de 1985 a 1995.

- AT: 305×305 mm (IBM)
- Baby-AT: 216×330 mm

ATX

1995: ATX (*Advanced Technology eXtended*, «tecnología avanzada extendida») fue creado por un grupo liderado por Intel, en 1995 introdujo las conexiones exteriores en la forma de un panel E/S y definió un conector de 24 pines para la energía. Se usa en la actualidad en la forma de algunas variantes, que incluyen conectores de energía extra o reducciones en el tamaño.

- ATX: 305×244 mm (Intel)
- microATX: 244×244 mm
- FlexATX: 229×191 mm
- MiniATX: 284×208 mm

ITX

2001: ITX (*Integrated Technology eXtended*), con rasgos procedentes de las especificaciones microATX y FlexATX de Intel, el diseño de VIA se centra en la integración en placa base del mayor número posible de componentes, además de la inclusión del hardware gráfico en el propio chipset del equipo, siendo innecesaria la instalación de una tarjeta gráfica en la ranura AGP.

- ITX: 215×195 mm (VIA)
- Mini-ITX: 170×170 mm
- Nano-ITX: 120×120 mm
- Pico-ITX: 100×72 mm

BTX

2004: BTX fue retirada en muy poco tiempo por la falta de aceptación, resultó prácticamente incompatible con ATX, salvo en la fuente de alimentación. Fue creada para intentar solventar los problemas de ruido y refrigeración, como evolución de la ATX.

- BTX: 325×267 mm (Intel)
- Micro BTX: 264×267 mm
- Pico BTX: 203×267 mm
- Regular BTX: 325×267 mm

DTX

2007: DTX eran destinadas a las PC de pequeño formato. Hacen uso de un conector de energía de 24 pines y de un conector adicional de 2x2.

- DTX: 248×203 mm (AMD)
- Mini DTX: 170×203 mm
- Full DTX: 243×203 mm

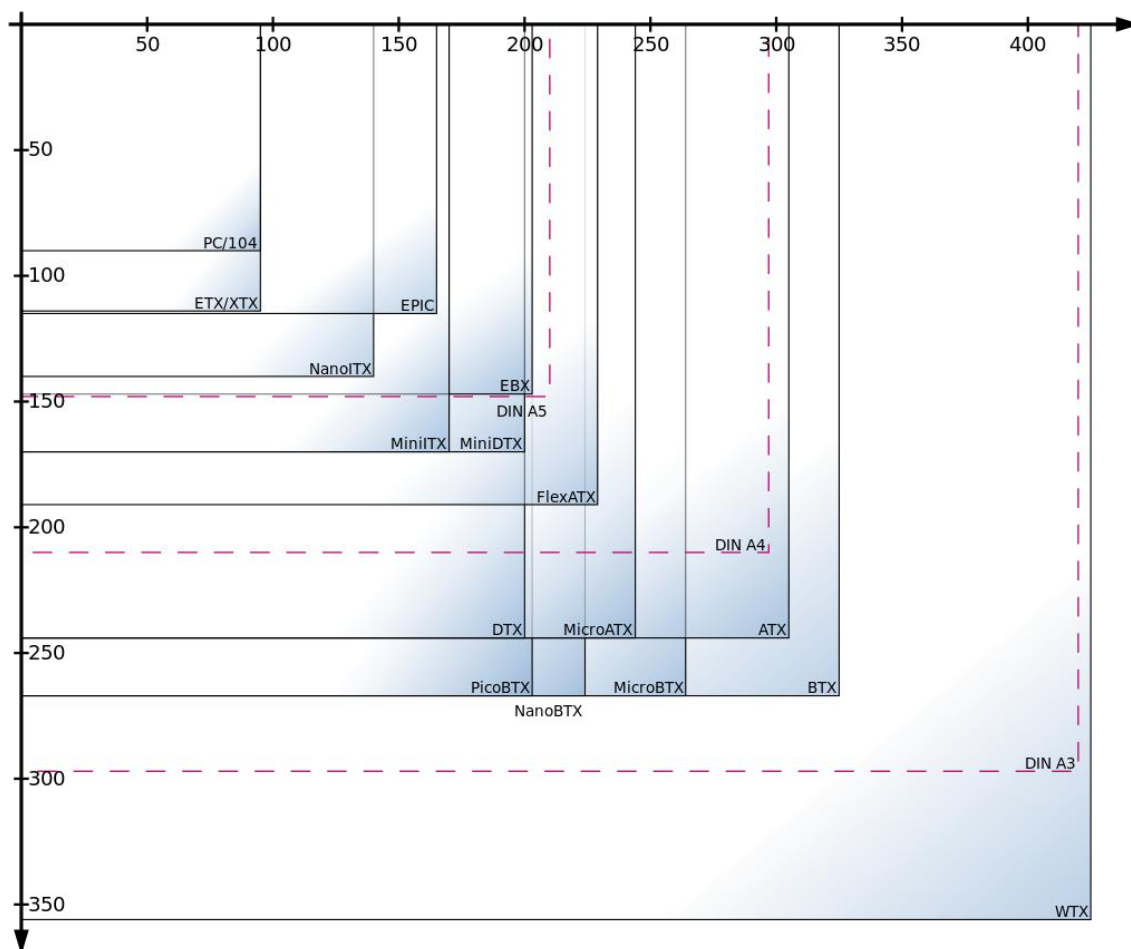


Figura 18. Formatos y tamaños de Placas Bases.

Formatos propietarios

Durante la existencia del PC, muchas marcas han intentado mantener un esquema cerrado de hardware, denominado formato propietario, fabricando placas base incompatibles físicamente con los factores de forma con dimensiones, distribución de elementos o conectores que son atípicos. Entre las marcas más persistentes está Dell, que rara vez fabrica equipos diseñados con factores de forma de la industria.

Fabricantes de placa base

Varios fabricantes se reparten el mercado de placas base, tales como: Advantech, Albatron, Aopen, ASUS, AsRock, Biostar, Chaintech, Dell, DFI, ECS EliteGroup, FIC, Foxconn, Gigabyte Technology, iBase, iEi, Intel, Lenovo, MSI, Pc Chips, Sapphire Technology, Super Micro, Tyan, VIA, XFx, Zotac.

Algunos diseñan y fabrican uno o más componentes de la placa base, mientras que otros ensamblan los componentes que terceros han diseñado y fabricado.

3.4. Backplane

//lectura, mínimo saber lo que es//

Un backplane es una placa de circuito (por lo general, una placa de circuito impreso) que conecta varios conectores en paralelo uno con otro, de tal modo que cada pin de un conector esté conectado al mismo pin relativo del resto de conectores, formando un bus de ordenador. Se utiliza como columna vertebral para conectar varias placas de circuito impreso (tarjetas) que juntas forman una computadora. Uno de los primeros sistemas en utilizar este enfoque fue el Bus S-100, llamado así porque los conectores tenían 100 pines, que fue muy popular en

los primeros ordenadores personales como el Altair 8800. Tanto el Apple II como el IBM PC integraban un backplane en la placa madre para tarjetas de expansión.

Mientras que una placa madre puede incluir un backplane, el backplane es en realidad una entidad separada. Un backplane se diferencia generalmente por la falta de CPUs, constando como mucho de los chips necesarios para manejar el bus, mientras que la CPU y el chipset residen en una tarjeta Single Board Computer.

Se utilizan preferentemente Backplanes en lugar de cables por su mayor fiabilidad. En un sistema con cables, estos se ven flexionados cada vez que se inserta o remueve una tarjeta, lo que causa eventualmente fallos mecánicos. Un backplane no se ve afectado por ese problema, por lo que su vida operativa está solo limitada por la longevidad de sus conectores.

Un backplane proporciona una funcionalidad mínima sin un Computador en una tarjeta instalado que le proporcione la CPU y otras funcionalidades del ordenador. Un Computador en una tarjeta (Single Board Computer) que cumple con la especificación PICMG 1.3 y compatible con un backplane PICMG 1.3 es denominado System Host Board (SHB).

Un backplane puede utilizarse sin un Computador en una tarjeta para proporcionar alimentación eléctrica a las tarjetas conectadas en él. Es de uso común en las empresas que fabrican tarjetas de ampliación para testearlas y grabarlas.

Además, existen cables de bus de expansión que permiten extenderlo a un backplane externo, generalmente situado en una carcasa auxiliar, para proporcionar más o diferentes ranuras de expansión de las que el ordenador principal proporciona. Estos grupos de cables tienen un circuito transmisor situado en el ordenador, una tarjeta de expansión situada en el backplane externo, y un cable entre ambos. Estos cables no necesitan de un Computador en una tarjeta en el backplane remoto para controlar las tarjetas de entrada/salida, y este sólo proporciona la electrónica de expansión necesaria.

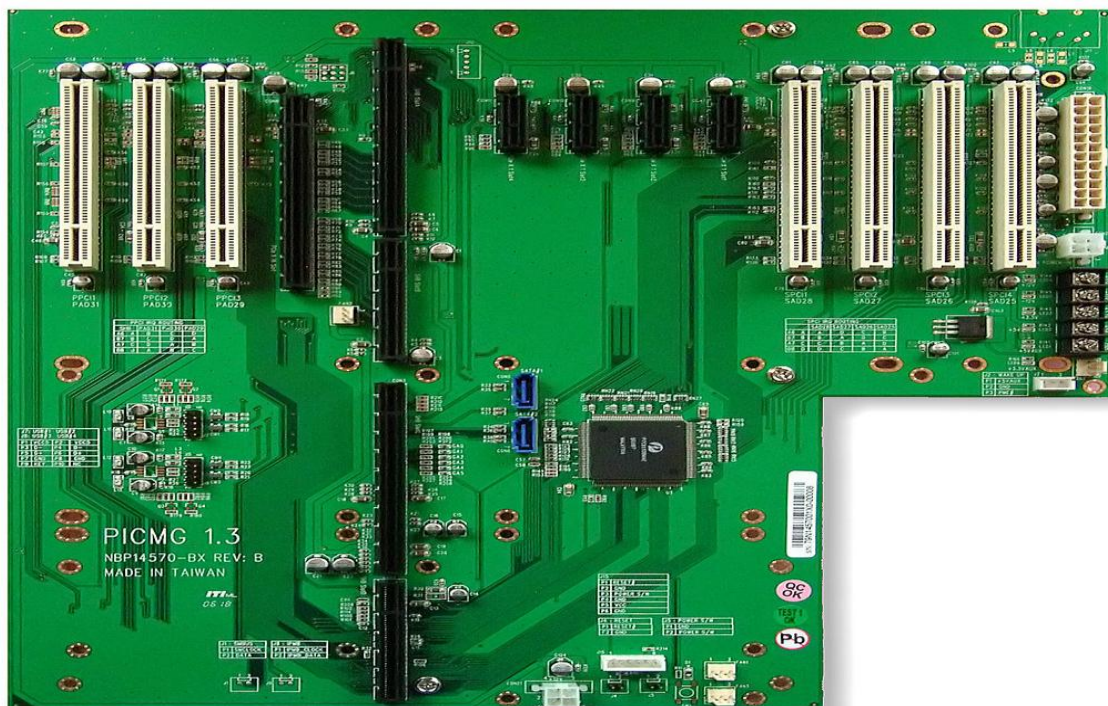


Figura 19. Backplane con ranuras PCI, PCIe x1 y PCIe x16.

Tipos de backplanes

Los Backplanes han crecido en complejidad desde los simples ISA (utilizado en el IBM PC) o Bus S-100 donde todos los conectores estaban conectados a un bus común. Debido a las limitaciones inherentes a las especificaciones PCI para las ranuras de ampliación, los backplanes se dividen ahora en **activos** y **pasivos**

- Pasivos: aquellos que consisten únicamente en conexiones eléctricas para la intercomunicación.
- Activos: aquellos que incorporan circuitería y lógica para el encaminamiento de dicha intercomunicación.

Los Backplanes pasivos no ofrecen circuitos de gobierno del bus. Cualquier lógica de arbitraje deseado se ponga en las tarjetas hijas. Los Backplanes activos incluyen chips con buffer de las distintas señales en las ranuras.

La distinción entre los dos no siempre es muy claro, pero puede convertirse en un problema importante si el sistema en su conjunto se espera que no tenga un SPOF (punto único de fallo). Un backplane pasivo, incluso si es único, no se considera un SPOF. Los Backplanes activos son más complicados y por lo tanto no tienen cero riesgos de mal funcionamiento.

Backplanes vs Placas madres

Cuando un backplane se utiliza conectado a un computador en una tarjeta (SBC) o system host board (SHB), la combinación proporciona la misma funcionalidad que una placa madre proporcionando potencia de procesamiento, memoria, E/S y ranuras para conectar tarjetas. Si bien hay algunas placas base que ofrecen más de 8 ranuras, ese es el límite tradicional.

Además, a medida que avanza la tecnología, la disponibilidad y el número de un tipo de ranura en particular puede ser limitado en términos de lo que actualmente ofrecen los fabricantes de placa base.

Por otro lado, la arquitectura de placa madre es algo relacionado con la tecnología del SBC que está conectado a ella. Existen algunas limitaciones a lo que se puede construir en el conjunto de chips de un SBC y el procesador tiene que proporcionar la capacidad de soportar el tipo de ranura. Además, prácticamente un número ilimitado de ranuras se puede proporcionar con 20, incluyendo la ranura SBC, como un práctico, aunque no absoluto límite. Por lo tanto, un backplane PICMG puede proporcionar cualquier número y cualquier combinación de ranuras ISA, PCI, PCI-X y PCI-e limitado solamente por la capacidad de la SBC de direccionar y gobernar esas ranuras. Por ejemplo, un SBC con el último procesador Intel Core i7 (Sandy bridge) podría interactuar con una placa base que proporciona hasta 19 ranuras ISA para manejar viejas tarjetas de entrada/salida.

Backplanes Mariposa

Algunas backplanes se construyen con ranuras en ambos lados. No son lo mismo que un backplane de medio plano. Un backplane mariposa se construye para maximizar el número de ranuras con la mínima altura vertical. El backplane se monta verticalmente en un chasis orientado de frente hacia atrás y la SBC y tarjetas conectadas se montan tumbadas, pudiendo sobresalir por ambos lados del backplane. Esto, por ejemplo, permite el uso de hasta cuatro tarjetas de altura completa en un chasis 2U.

Backplanes en almacenamiento

Los servidores suelen tener un backplane para conectar discos duros intercambiables en caliente, los pines del backplane pasan directamente a las tomas del disco duro sin cables. Pueden tener un solo conector para conectar un controlador de matriz de discos o varios conectores que se pueden conectar a uno o más controladores de forma arbitraria. Se encuentran backplanes comúnmente en discos externos, matrices de discos, y servidores. Backplanes para discos duros SAS y SATA HDD utilizan con mayor frecuencia el protocolo SGPIO como medio de comunicación entre el HBA y el backplane. Alternativamente se puede utilizar SCSI Enclosure Services. Con subsistemas Parallel SCSI, SAF-TE se utiliza en computadoras, principalmente en Servidores blade, donde los servidores blade residen en un lado y los periféricos (alimentación, redes, y otras entradas/salidas) y módulos de servicios residen en el otro. Midplanes son también populares en la creación de redes y equipos de

telecomunicaciones donde un lateral del chasis acepta tarjetas de sistema de procesamiento y el otro lado del chasis acepta tarjetas de interfaz de red.

PICMG



Figura 20. PICMG 1.3 Single Board Computer (SHB) en un Backplane activo.

Un computador en una tarjeta que cumpla la normativa PICMG 1.3 y un backplane compatible que también cumpla con PICMG 1.3 backplane se conoce como un System Host Board.

En el mundo de las Single Board Computer basadas en Intel, PICMG establece normas para la interfaz del backplane: PICMG 1.0, 1.1 y 1.2 proporcionan sucesivamente soporte para ISA y PCI y PCIX. PICMG 1.3 proporciona soporte para PCI-Express.

4. Unidad Central del Proceso (CPU)

4.1. Introducción

El microprocesador o Unidad Central de Procesamiento (CPU) es un circuito integrado que sirve como cerebro del computador. En el interior de este componente electrónico existen millones de transistores integrados.

Internamente, un microprocesador está compuesto por millones de pequeños componentes electrónicos llamados transistores. Actualmente, los microprocesadores se fabrican con unas distancias entre transistores del orden de 14 nm. El hecho de disminuir la separación entre transistores permite, que en el mismo espacio, se puedan concentrar más transistores, aumentando así la potencia del procesador.

Fabricantes como Intel han adoptado desde 2007 un modelo denominado Tick-Tock (tictac en español) para seguir cada cambio de la microarquitectura de una contracción de tecnología del proceso de fabricación anterior. Cada tic es una contracción de la tecnología de proceso de la microarquitectura anterior y cada tac es una nueva microarquitectura. Cada año se espera que sea un tic o un tac aunque actualmente Intel ya no está siguiendo este modelo debido a los problemas para mejorar los nodos de fabricación para disminuir el nivel de integración por debajo de los 14 nm. Ahora es más bien un tictac seguido de un refresco de procesadores: Process-Architecture-Optimization (PAO).

La litografía de 10 nm no está prevista hasta 2017 con el microprocesador Cannonlake de Intel. Otro parámetro que mide la potencia de un procesador es su frecuencia de reloj. La frecuencia de reloj de un procesador se mide en Megahercios (MHz) o bien en Gigahercios (GHz). A cada herzio también se le suele denominar ciclo. En cada ciclo un procesador puede realizar una o varias microoperaciones. Cada microoperación es una parte de una operación. De este modo, cuanto mayor sea esta frecuencia de reloj, mayor número de microoperaciones podrá realizar. Por último, hay que destacar que la potencia de un microprocesador no se debe medir sólo por su frecuencia de reloj, sino también por su número de transistores, por su tecnología de fabricación, por sus niveles internos de memoria caché, etc. Existen además diversos índices mediante los cuales se puede medir la potencia de un ordenador: **//muy importante esto//**

- **MIPS:** mide cuántos millones de instrucciones por segundo es capaz de ejecutar un procesador. Este índice no tiene en cuenta que hay procesadores con instrucciones más potentes que las de otros, por lo que se puede considerar como una medida simplista.
- **MFLOPS:** mide cuantos millones de instrucciones de coma flotante por segundo es capaz de ejecutar el procesador. Análogamente, esta medida no tiene en cuenta las diferencias entre las instrucciones de procesadores diferentes. Normalmente, los fabricantes de procesadores suelen proporcionar para los índices MIPS y MFLOPS picos de rendimiento.
- **SPEC:** ejecuta en el procesador un conjunto de programas y combina las medidas de prestaciones de éstos con la media aritmética o geométrica.
- **SPECint y SPECfp:** son unos índices que miden las velocidades en operaciones con enteros y con coma flotante. La medida resultante se denomina SPECmarks.

Los microprocesadores suelen tener forma de prisma chato, y se instalan sobre un elemento llamado zócalo de la placa base. También, en modelos antiguos, se solía soldar directamente a la placa madre.

Aparecieron algunos modelos donde se adoptó el formato de cartucho, sin embargo no tuvo mucho éxito.

Actualmente, se dispone de un zócalo especial para alojar el microprocesador y el sistema de enfriamiento, que comúnmente es un ventilador (*cooler*).

4.2. Componentes

//vip!!! En la pagina 50 (Punto 4.8.) teneis un resumen con lo principal que hay que saber. Al menos saber cuales son estos componentes y su funcion principal//

El microprocesador está compuesto por:

- a) Banco de Registros
- b) Unidad de control
- c) Unidad aritmético-lógica
- d) Unidad de Coma Flotante (dependiendo del microprocesador)
- e) Bus de control

A continuación se explica con más detalle cada uno de estos elementos:

a) Banco de Registros

Un registro es una memoria de alta velocidad y poca capacidad, integrada en el microprocesador, que permite guardar y acceder a valores muy usados, generalmente en operaciones matemáticas.

Los registros están en el punto más alto de la jerarquía de memoria, y son la manera más rápida que tiene el sistema de almacenar datos. Los registros se miden, generalmente, por el número de bits que almacenan; por ejemplo, un "registro de 8 bits" o un "registro de 32 bits". Por otro lado, un procesador puede tener registros de estos tres tipos:

- Registros de propósito general: son aquellos que pueden ser usados explícitamente por el programador mediante el lenguaje ensamblador.
- Registros de propósito específico: son los que su valor sólo puede modificarse mediante instrucciones específicas. En este grupo, se encuentra un registro llamado contador de programa (PC, *Program Counter*) cuya función es apuntar, en todo momento, a la siguiente instrucción a ejecutar. Dentro de este grupo de registros también se encuentran otros como el registro de estado (almacena cierta información sobre las operaciones que tienen lugar en la ALU) y el registro de puntero de pila (contiene una dirección de memoria a partir de la cual se encuentran una serie de datos).



- Registros transparentes: son registros sobre los que el programador no puede ni actuar ni referenciar. Son utilizados internamente por el procesador. Entre ellos se encuentra un registro muy importante, el denominado registro de instrucción, que contiene la instrucción que se está ejecutando.

b) Unidad de Control

La Unidad de control es el "cerebro del microprocesador". Es la encargada de activar o desactivar los diversos componentes del microprocesador en función de la instrucción que el microprocesador esté ejecutando y de la etapa de dicha instrucción en que se encuentre.

La unidad de control (UC) interpreta y ejecuta las instrucciones almacenadas en la memoria principal y genera las señales de control necesarias para ello.

Existen dos tipos de diseños para construir una Unidad de Control:

- Unidad de control **cableada**: se construye, en su mayor parte, mediante puertas lógicas. Una puerta lógica es un pequeño circuito electrónico que se comporta como alguna de las operaciones lógicas vistas en el tema anterior, con la diferencia de que los valores verdadero y falso son ahora señales eléctricas con tensión alta o baja. El conjunto de puertas lógicas actuará en función del tipo de instrucción leída.
- Unidad de control **microprogramada**: está formada por puertas lógicas y por una memoria que almacena, para cada instrucción, qué señales de control internas se deben activar y desactivar en cada instante con el fin de realizar la función de dicha instrucción. La información referente a la señales de control se almacena en esa memoria mediante secuencias de unos y ceros. A cada palabra de esta memoria formada por unos y ceros se le denomina microinstrucción, la cual define uno de los pasos por los que está compuesta una instrucción. Cada bit de ésta se referirá a una señal interna específica de la unidad de control. Al conjunto ordenado de microinstrucciones necesarias para la ejecución de una se le denomina microprograma.

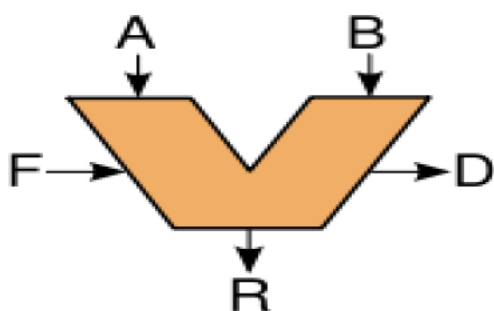
c) Unidad aritmético-lógica

La Unidad Aritmético Lógica (UAL), o *Arithmetic Logic Unit* (ALU), es un circuito digital que calcula operaciones aritméticas (como suma, resta, etc.) y operaciones lógicas (como OR, NOT, XOR, etc.) entre dos números.

Muchos tipos de circuitos electrónicos necesitan realizar algún tipo de operación aritmética, así que incluso el circuito dentro de un reloj digital tendrá una ALU minúscula que se mantiene sumando 1 al tiempo actual y comprobando si debe activar el pitido del temporizador, etc.

Los circuitos electrónicos más complejos son los que están contruidos dentro de los chips de microprocesadores modernos. Por lo tanto, estos procesadores tienen en su interior una ALU más compleja y poderosa. De hecho, un microprocesador moderno (y los mainframes) pueden tener múltiples núcleos, cada uno con múltiples unidades de ejecución, cada una de ellas con múltiples ALU.

La ilustración muestra un típico símbolo esquemático para una ALU:



☐ A y B son operandos.

☐ R es la salida.

☐ F es la entrada de la unidad de control.

☐ D es un estado de la salida.

Imagen 21. Unidad Aritmético-lógica.

d) Unidad de Coma Flotante

Una Unidad de Punto Flotante o Unidad de Coma Flotante (*Floating Point Unit, FPU*) o, más comúnmente conocido como, coprocesador matemático, es un componente de la CPU especializado en el cálculo de operaciones en coma flotante. Las operaciones básicas que toda FPU puede realizar son las aritméticas (suma y multiplicación), si bien algunos sistemas más complejos son capaces también de realizar cálculos trigonométricos o exponenciales.

No todas las CPUs tienen una FPU dedicada. En ausencia de FPU, la CPU puede utilizar programas en microcódigo para emular una función en coma flotante a través de la unidad aritmético-lógica (ALU), la cual reduce el coste del hardware a cambio de una sensible pérdida de velocidad.

En algunas arquitecturas, las operaciones en coma flotante se tratan de forma completamente distinta a las operaciones enteras, con registros dedicados y tiempo de ciclo diferentes. Incluso para operaciones complejas, como la división, podrían tener un circuito dedicado a dicha operación.

e) Bus de control

Generalmente, este bus interno es supervisado y controlado por el microprocesador debido a que es una vía de comunicación compartida por varios elementos del computador.

Un bus tiene multitud de características, entre las cuales se destacan las siguientes:

- Ancho de banda: número máximo de elementos de información que se pueden transmitir por unidad de tiempo (por ejemplo, MB/s).
- Frecuencia: los buses también poseen una frecuencia de trabajo. Cuanto mayor sea esta frecuencia más cantidad de información por unidad de tiempo será capaz de transmitir.
- Grado de paralelismo: indica cuántas unidades de información, usualmente bits, es capaz de transmitir el bus simultáneamente. Aquí se encuentran los buses de tipo serie que sólo pueden transmitir bit a bit, y los buses paralelos que son capaces de transmitir n bits en un determinado instante. A esta característica también se denomina ancho de bus.
- Función: el bus puede ser de uso general o tener una función específica, como utilizarse solamente para la transmisión de datos, direcciones o información de control.
- Longitud: dependiendo del ámbito de aplicación del bus, éste deberá tener una longitud máxima establecida.
- Físicas y eléctricas: el bus, al estar fabricado con alguno de los materiales conductores, tendrá unas propiedades eléctricas determinadas.

4.3. Funcionamiento

El microprocesador ejecuta instrucciones almacenadas como números binarios en la memoria principal. La ejecución de las instrucciones se puede realizar en varias fases:

- *PreFetch*, Pre lectura de la instrucción desde la memoria principal.
- *Fetch*, ordenamiento de los datos necesarios para la realización de la operación.
- Decodificación de la instrucción, es decir, determinar qué instrucción es y por tanto qué se debe hacer.
- Ejecución.
- Escritura de los resultados en la memoria principal o en los registros.

Cada una de estas fases se realiza en uno o varios ciclos de CPU, dependiendo de la estructura del procesador, y concretamente de su grado de supersegmentación. La duración de estos ciclos viene determinada por la frecuencia de reloj, y nunca podrá ser inferior al tiempo requerido para realizar la tarea individual (realizada en un solo ciclo) de mayor coste temporal. El microprocesador dispone de un oscilador de cuarzo capaz de producir pulsos a un ritmo constante, de modo que genera varios ciclos (o pulsos) en un segundo.

Velocidad

Actualmente se habla de frecuencias de Gigahercios (GHz), o de Megahercios (MHz). Lo que supone miles de millones o millones, respectivamente, de ciclos por segundo. El indicador de la frecuencia de un microprocesador es un buen referente de la velocidad de proceso del mismo, pero no el único. La cantidad de instrucciones necesarias para llevar a cabo una tarea concreta, así como la cantidad de instrucciones ejecutadas por ciclo ICP (*Instructions per cycle*), son los otros dos factores que determinan la velocidad de la CPU. La cantidad de instrucciones necesarias para realizar una tarea depende directamente del juego de instrucciones disponible, mientras que ICP depende de varios factores, como el grado de supersegmentación y la cantidad de unidades de proceso o "*pipelines*" disponibles, entre otros. La cantidad de instrucciones necesarias para realizar una tarea depende directamente del juego de instrucciones.

Bus de datos

Los modelos de la familia x86 (a partir del 386) trabajan con datos de 32 bits, al igual que muchos otros modelos en la actualidad. Pero los microprocesadores de las tarjetas gráficas, que tienen un mayor volumen de procesamiento por segundo, se ven obligados a aumentar este tamaño, y así hoy en día, existen microprocesadores gráficos que trabajan con datos de 128, 256, 384 y 512 bits. Estos dos tipos de microprocesadores no son comparables, ya que ni su juego de instrucciones ni su tamaño de datos son parecidos y, por tanto, el rendimiento de ambos no es comparable en el mismo ámbito.

La arquitectura x86 se ha ido ampliando a lo largo del tiempo a través de conjuntos de operaciones especializadas denominadas "extensiones", las cuales han permitido mejoras en el procesamiento de tipos de información específica. Este es el caso de las extensiones MMX y SSE de Intel, y sus contrapartes, las extensiones 3DNow! de AMD. A partir de 2003, el procesamiento de 64 bits fue incorporado en los procesadores de arquitectura x86 (x86-64) a través de la extensión AMD64 y, posteriormente, con la extensión EM64T (Intel 64) en los procesadores AMD e Intel respectivamente.

4.4. Historia

//lectura//

El primer procesador comercial, el Intel 4004, fue presentado el 15 de noviembre de 1971. Los diseñadores fueron Ted Hoff y Federico Faggin de Intel, y Masatoshi Shima de Busicom (más tarde ZiLOG).

A lo largo de la historia y desde su desarrollo inicial, los microprocesadores han mejorado enormemente su capacidad de tratamiento de instrucciones, evolucionando desde 4, 8, 16, 32, 64, 128, 256... bits.

Los microprocesadores modernos están integrados por millones de transistores y otros componentes empaquetados en una cápsula cuyo tamaño varía según las necesidades de las aplicaciones a las que van dirigidas. Las partes lógicas que componen un microprocesador son, entre otras: unidad aritmético-lógica, registros de almacenamiento, unidad de control, unidad de ejecución, memoria caché y buses de datos control y dirección.

Hay que destacar que los grandes avances en la construcción de microprocesadores se deben más a la Arquitectura de Computadoras que a la miniaturización electrónica. El microprocesador se compone de muchos componentes. En los primeros procesadores gran parte de éstos estaban ociosos el 90% del tiempo. Sin embargo, hoy en día, los componentes están repetidos una o más veces en el mismo microprocesador, y los buses están hechos de forma que siempre trabajan todos los componentes. Por eso, los microprocesadores son tan rápidos y tan productivos.

La fuerte competencia en el mundo de los procesadores, especialmente entre Intel y AMD, ha producido que la tecnología actual de fabricación de procesadores esté llegando a sus límites. Cada vez, la miniaturización de los componentes del procesador es más difícil (el límite de construcción del silicio se estima en los 10 nm, donde a la falta de consistencia de este elemento se suman los posibles efectos cuánticos por las distancias manejadas. Hoy en día se están fabricando chips en los 14 nm), el problema de la generación de calor ha aumentado, originando una mayor dificultad para aumentar la frecuencia principal del procesador. Todos estos problemas dificultan el incremento de rendimiento de los procesadores.

Con la mencionada tecnología, los procesadores actuales pueden alcanzar la velocidad de los 4.2 GHz pero necesitan grandes disipadores y ventiladores porque generan mucho calor. No se podía continuar fabricando procesadores de la misma manera, se estaba llegando a un estancamiento; era necesario tomar otro camino, utilizar otra variable que hiciera que el rendimiento del procesador aumentará. Entonces, se empezaron a construir los procesadores multinúcleo (a partir del año 2001).

Los procesadores multinúcleo contienen dentro de su empaque varios núcleos o "cerebros". Se basan en los sistemas distribuidos, la computación paralela, y las tecnologías como el *Hyperthreading* (tecnología creada por Intel para los procesadores de la familia Pentium 4), que mostraban cómo dividir el trabajo entre varias unidades de ejecución.

4.5. Arquitecturas de procesador RISC y CISC

//VIP, muy preguntado en GSI y en TAI ya lo han preguntado alguna vez ultimamente//

El procesador es todo un complejo universo en sí mismo y, aunque los primeros modelos eran comparativamente muy similares, con su evolución se han ido desarrollando distintos diseños que han afectado a numerosos elementos siendo de destacar las diferentes tendencias desarrolladas asociadas al juego de instrucciones que empleaban.

Podemos decir que frente a esta cuestión caben dos filosofías de diseño: las denominadas arquitecturas CISC y RISC.

La arquitectura **CISC** (*complex instruction set computer*), que ya se daba en los primeros diseños de UCP, se caracterizaba por disponer de un grupo amplio de instrucciones complejas y potentes. El ordenador era más potente a medida que era más amplio su repertorio de instrucciones.

Toman como principio la microprogramación, que significa que cada instrucción de máquina es interpretada, empleando un microprograma localizado en una memoria en el circuito integrado del procesador. Las instrucciones compuestas son codificadas internamente y ejecutadas con una serie de microinstrucciones que se almacenan en una memoria de control. Esto era efectivo y muy práctico porque al principio la memoria principal era más lenta que la UCP y el tiempo de una instrucción podía ser de varios ciclos de reloj ya que cuando una



instrucción era procesada en un único ciclo de reloj, no se podía continuar con la siguiente instrucción inmediatamente ya que todavía no estaba lista (al ser la memoria principal mucho más lenta que la de control).

Buscando aumentar la velocidad de procesamiento se descubrió que con una determinada arquitectura, la ejecución de programas compilados directamente con microinstrucciones estando residentes en memoria externa al circuito resultaba más eficiente.

A finales de los setenta, al aumentar las prestaciones de la memoria principal la consecuencia inmediata fue que ya no tenía que esperar la UC a ésta, lo que permitió trabajar con instrucciones mucho más simples que se completasen en un ciclo de reloj y acelerando la ejecución de instrucciones.

Esta arquitectura es conocida como **RISC** (*reduced instruction set computer*) y está formada por un juego de instrucciones lo más reducido posible, la mayoría completadas en un ciclo de reloj.

Debido a que se tiene un conjunto de instrucciones simplificado, éstas se pueden implantar por hardware directamente en la CPU, lo que elimina el microcódigo y la necesidad de decodificar instrucciones complejas. **//esta tabla si lo considero VIP saberlo//**

ARQUITECTURA CISC (pocas instrucciones complejas)	ARQUITECTURA RISC (muchas instrucciones pequeñas)
- Formatos de instrucción de varios tamaños - Interpreta microinstrucciones - Muchos modos de direccionamiento - Pocos registros de propósito general - Repertorio de instrucciones flexible - Son lentas. Ejecución por software - UC microprogramada.	- Formatos de instrucción de pocos tamaños - Interpreta microoperaciones - Pocos modos de direccionamiento - Muchos registros de propósito general - Repertorio de instrucciones rígido. - Son rápidas. Ejecución directa por Hardware - UC cableada.

Figura 22. Arquitecturas RISC y CISC.

Para ejecutar una tarea se necesitan más instrucciones en RISC que en CISC, ya que en RISC las instrucciones son más elementales. Pero el hecho de disponer hoy de memorias tan rápidas y buses de alta velocidad hace que existan diversos mecanismos como el llamado pipeline , para que se lleguen incluso a solapar varias instrucciones en un mismo ciclo.

Realmente las diferencias son cada vez más borrosas entre las arquitecturas CISC y RISC. Las CPU actuales combinan elementos de ambas y no son fáciles de encasillar ya que desde hace tiempo se ha empezado a investigar sobre microprocesadores “híbridos”, con el propósito de obtener ventajas procedentes de ambas tecnologías.

Los procesadores ARM, muy utilizados en Smartphones y pequeños dispositivos (IoT), utilizan la Tecnología RISC.

4.6. Segmentación de instrucciones

La segmentación de instrucciones es una técnica que permite implementar el paralelismo a nivel de instrucción en un único procesador. La segmentación intenta tener ocupadas con instrucciones todas las partes del procesador dividiendo las instrucciones en una serie de pasos secuenciales que efectuarán distintas unidades de la CPU, tratando en paralelo diferentes partes de las instrucciones. Permite una mayor tasa de transferencia efectiva por

parte de la CPU que la que sería posible a una determinada frecuencia de reloj, pero puede aumentar la latencia debido al trabajo adicional que supone el propio proceso de la segmentación.

Los modernos microprocesadores que incluyen segmentación utilizan **ejecución especulativa*** para reducir el coste de las instrucciones que dependen de la predicción de saltos, y lo hacen mediante técnicas que predicen la ruta de ejecución de un programa basándose en el historial de los saltos realizados anteriormente. Con el fin de mejorar el rendimiento y optimizar la utilización de los recursos del sistema informático, las instrucciones pueden planificarse cuando todavía no se ha determinado si será o no necesario ejecutarlas, antes de un salto.

***Ejecución especulativa:** es una forma de optimización en la que un sistema informático realiza una tarea que podría no ser necesaria; la idea consiste en llevar a cabo un trabajo antes de saber si será realmente necesario con la intención de evitar el retraso que supondría realizarlo *después* de saber que sí es necesario. Si el trabajo en cuestión resulta ser innecesario, la mayoría de los cambios realizados por ese trabajo se revierten y los resultados se ignoran. El objetivo de esta técnica es proporcionar una mayor conurrencia en caso de disponer de más recursos.

4.7. Vulnerabilidades de seguridad

//Vip!!! Saber al menos el nombre, el resto lectura//

En junio de 2017, un equipo formado por investigadores independientes, laboratorios universitarios de investigación, algunos miembros del Proyecto Cero de Google y Cyberus Technology descubrió dos vulnerabilidades de seguridad que eran posibles gracias al uso ampliamente extendido de la ejecución especulativa. El problema fue descubierto también de forma independiente por otros investigadores más o menos al mismo tiempo. Las vulnerabilidades terminaron haciéndose públicas en enero de 2018 y fueron denominadas **Meltdown** y **Spectre**. Ambas tienen el potencial de permitir que software malicioso pueda leer de la memoria protegida de un sistema informático, obteniendo acceso a información sensible como pueden ser contraseñas y claves de encriptación.

Para hacer que la ejecución especulativa ofrezca la máxima eficacia, algunas combinaciones de hardware y sistemas operativos permiten que ésta *toque* datos de la memoria privada del sistema operativo antes de que sea necesario. Las vulnerabilidades parten de la capacidad que tiene el programa malicioso de deducir en qué consiste esa información que por lo demás le resulta inaccesible.

La vulnerabilidad más ampliamente comentada de las dos, Meltdown, se basa en ciertas decisiones del hardware de leer la información sensible del usuario, pero puede solventarse mediante actualizaciones del software que se ejecuta en el sistema en cuestión. Spectre, la vulnerabilidad menos conocida y también la más difícil de aplicar de las dos, hace posible que un programa acceda a información correspondiente a otro programa que se ejecuta en el chip, y es mucho más difícil de solventar con una única solución. Aunque no es lo ideal, las mejores defensas generales sugeridas hasta el momento incluyen por un lado la detección, y por el otro la separación de la actividad referente al caché de los procesos concurrentes. La mitigación del segundo caso puede acarrear una pérdida significativa de rendimiento en ciertas arquitecturas, dependiendo también del tipo de tarea que ejecute el sistema.

Ambas vulnerabilidades funcionan realizando una carga indexada desde la memoria. Durante esta carga se lee de la memoria un dato A (que supuestamente está "fuera de límites") y entonces este dato se emplea para calcular la dirección de otro dato B_A (accesible) para ser leído también de la memoria. Puesto que A es inaccesible, el procesador termina cancelando cualquier efecto *directo* de la operación sobre los registros y la memoria una vez que se da cuenta de que la lectura no está autorizada. Sin embargo, B_A sigue estando presente en el caché, y esta condición puede ser detectada por el atacante leyendo B_A para todos los posibles

valores de A y observando qué operación de lectura tarda bastante menos tiempo en completarse que las demás.

La diferencia entre ambas vulnerabilidades radica en la condición *no permitido* que se está soslayando y en los medios para conseguirlo:

- En Meltdown, el atacante hace que la ejecución especulativa rompa la barrera de protección situada entre el espacio de memoria de un programa de usuario y una página protegida del kernel, haciendo para ello que la ejecución especulativa lea una dirección de memoria y coloque su contenido en el caché. Los modernos procesadores de Intel ejecutan el código y luego toman y almacenan los resultados especulativos, dejándolos en el caché. Posteriormente, el código usado para explotar la vulnerabilidad puede tomar notas de lo ocurrido, ya sea después de un fallo, o después de preparar el código original como parte de una transacción de memoria que atenuará el fallo. Se cree que algunos procesadores, como los de AMD, son inmunes a esto, ya que realizan la comprobación de permiso de acceso a la página *antes* de ejecutar la lectura especulativa.
- En Spectre, el atacante no depende de estos mecanismos de fallo, sino que va a por otro proceso del usuario de una forma más generalista. Spectre depende de que un fallo en la predicción de saltos realice de forma especulativa la lectura de la celda de una matriz, aunque en el salto precedente se percibiese que la lectura iría más allá del final de la matriz. Empieza por entrenar la maquinaria de predicción de saltos del procesador para forzar un fallo en la predicción que se sale los límites de un proceso y entonces manipula al proceso objetivo para que ejecute parte de su propio código, que en realidad toca la rama especulativa. De nuevo, lo que ha tocado se revela mediante un canal lateral basado en medir los tiempos del caché. En este caso se muestra un ejemplo que combinado con un fragmento de código JavaScript pueden emplearse para leer todo el espacio de direcciones del proceso al que se ha engañado para ejecutar el código en cuestión.

Se cree que sólo los navegadores web que realizan compilación JIT (Just In Time) de JavaScript son vulnerables a esta forma de usar Spectre. Sin embargo, todos los navegadores modernos – incluyendo Chrome, Firefox, Safari y Edge– utilizan un motor JIT. Es más, Spectre puede emplearse también para manipular el predictor de saltos de varias formas, aunque no intervenga un lenguaje de scripts.

En principio Google tenía prevista una publicación coordinada del informe al completo de Proyecto Cero el martes 9 de enero de 2018, y dijo que había estado trabajando con empresas tanto de hardware como de software a lo largo de varios meses para mitigar los riesgos. Sin embargo, la escalada de rumores al respecto parece haber forzado un adelanto en la revelación de ambas vulnerabilidades.

4.8. Microprocesador (resumen)

//resumen de CPU//

La Unidad Central de Proceso (CPU), también denominada procesador, es el elemento encargado del control y ejecución de las operaciones que se efectúan dentro del ordenador con el fin de realizar el tratamiento automático de la información.

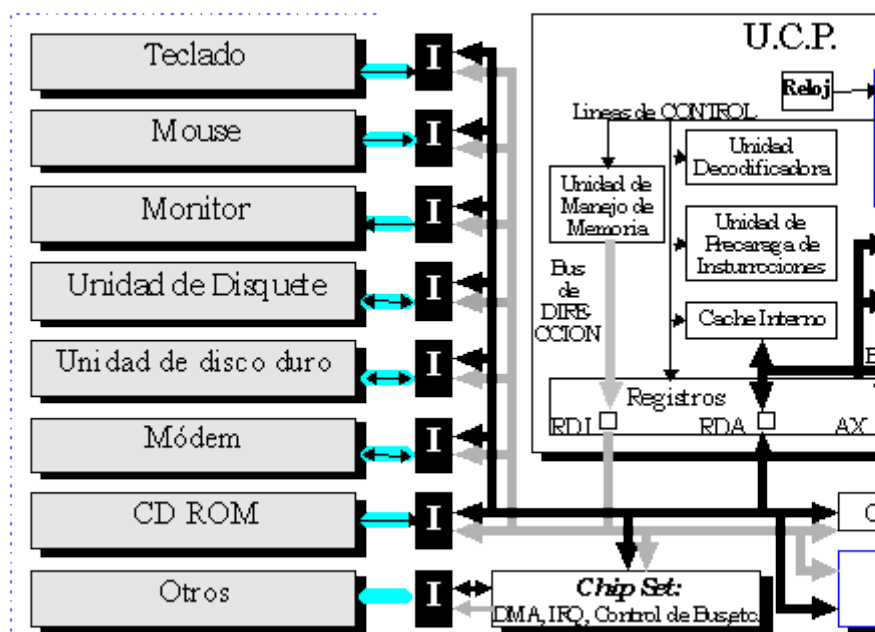
El procesador es la parte fundamental del ordenador; se encarga de controlar todas las tareas y procesos que se realizan dentro del él. Está formado por la unidad de control (UC), la unidad aritmético-lógica (ALU) y

su propia memoria interna, integrada en él. El procesado es la parte que gobierna el ordenador

Para que el procesador pueda trabajar necesita utilizar la memoria principal o central del ordenador.

En la mayoría de los casos también será necesario la intervención de la unidad de entrada-salida y los periféricos de entrada-salida.

El procesador gestiona lo que recibe y envía la memoria desde y hacia los periféricos mediante la unidad de entrada salida, los buses y los controladores del sistema.



Componentes de la CPU:

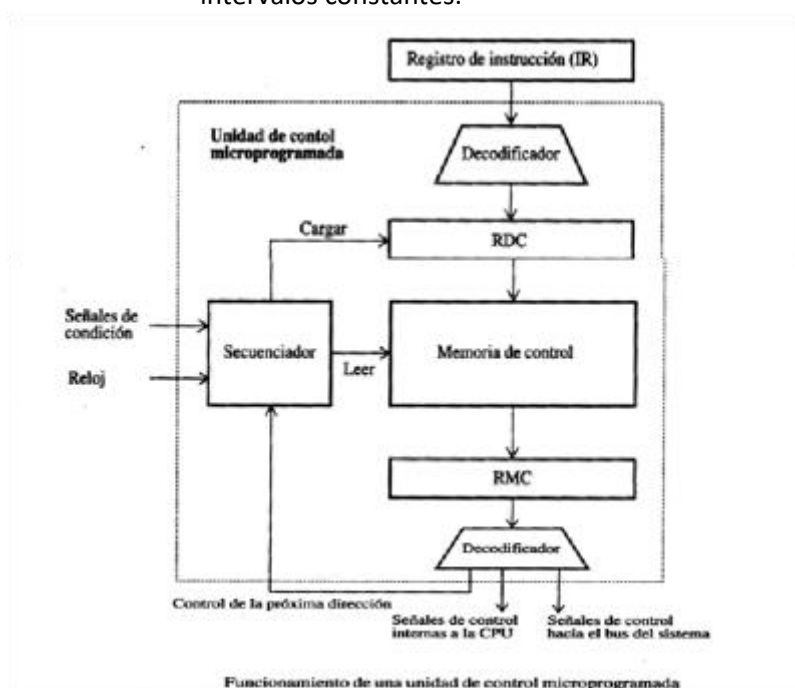
- UC: es la parte pensante del ordenador, es la que se encarga del gobierno y funcionamiento del ordenador. La tarea fundamental de la UC es recibir información para interpretarla y procesarla después mediante las órdenes que envía a los otros componentes del ordenador.

Se encarga de traer a la memoria interna o central del ordenador (memoria RAM) las instrucciones necesarias para la ejecución de los programas y el procesamiento de los datos. Estas instrucciones y datos se extraen, normalmente, de los soportes de almacenamiento externo. Además, la UC interpreta y ejecuta las instrucciones en el orden adecuado para que cada una de ellas se procese en el debido instante y de forma correcta.

Para realizar todas estas operaciones, la UC dispone de pequeños espacios de almacenamiento, denominados registros. Además de los registros, tiene otros componentes. Todos ellos se detallan a continuación:

- Registro de instrucción: Contiene la instrucción que se está ejecutando. Consta de diferentes campos:
 - CO: Código de la operación que se va a realizar.
 - MD: Modo de direccionamiento de la memoria para acceder a la información que se va a procesar.

- CDE: Campo de dirección efectiva de la información.
- Registro contador de programas: Contiene la dirección de memoria de la siguiente instrucción a ejecutar.
- Controlador y decodificador: Controla el flujo de instrucciones de la CPU e interpreta la instrucción para su posterior procesamiento. Se encarga de extraer el código de operación de la instrucción en curso.
- Secuenciador: Genera las microórdenes para ejecutar la instrucción
- Reloj: Proporciona una sucesión de impulsos eléctricos a intervalos constantes.

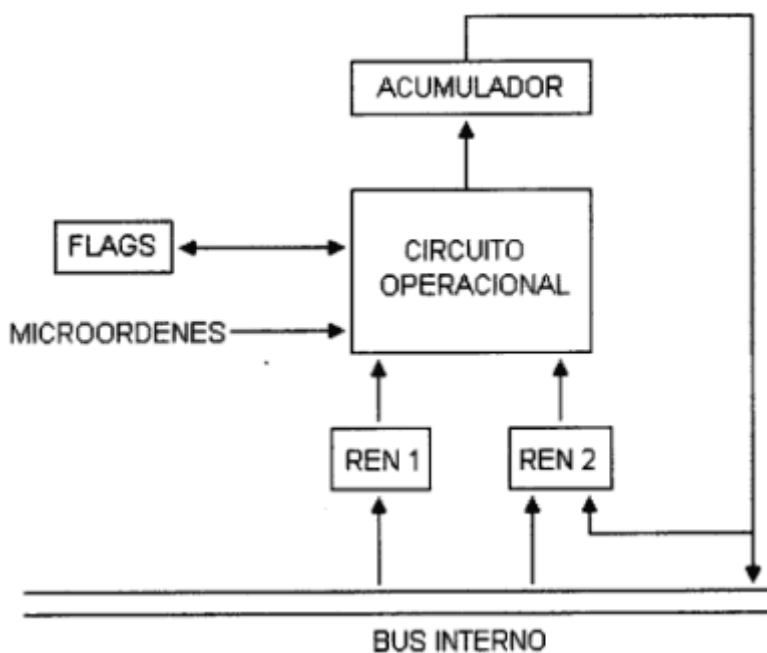


- ALU: es la parte de la CPU encargada de realizar las operaciones de tipo aritmético (suma, multiplicación, etc.), así como las de tipo lógico (comparación, etc.). Algunas de estas operaciones se detallan en la siguiente tabla:

Operación	Operador
Mayor que	>
Menor que	<
Mayor o igual	>=
No mayor	NOT > (<=)
Y lógico	AND
O lógico	OR

Los elementos que componen la ALU son los que se indican a continuación:

- Circuito combinacional u operacional. Realiza las operaciones con los datos de los *registros de entrada*
- Registros de entrada. Contienen los operandos de la operación.
- Registro acumulador. Almacena los resultados de las operaciones.
- Registro de estado. Registra las condiciones de la operación anterior.



5. La memoria

//VIP al menos los conceptos en negrita//

5.1. Introducción

La memoria es aquel elemento o unidad encargado de almacenar la información que necesita el ordenador, por tanto, las instrucciones que forman los programas y los datos que se emplean en su ejecución.

Se encuentra dividida en **celdas o palabras** que se identifican mediante una dirección y sobre las que se llevan a cabo operaciones de lectura y/o escritura.

El elemento básico de la memoria digital es el **biestable**, dispositivo electrónico capaz de almacenar un bit.

Mediante agrupamientos de estos dispositivos en distintas variantes tecnológicas que determinan las características de las memorias (el coste por bit, el tiempo de acceso y la capacidad o tamaño), se establece lo que se ha dado en llamar una **jerarquía de memorias**.



Figura 23. Esquema jerarquías de memorias.

5.2. Tipos de memoria

5.2.1. Según su velocidad y coste

Históricamente han existido dos tipos de memorias que se diferencian, principalmente, por su velocidad y coste, la **memoria interna** y la **memoria externa o secundaria**.

1. Memoria principal o interna

Se compone de los tres escalones superiores de la figura representada en el gráfico: un conjunto de **registros**, la **memoria caché** y la **memoria principal**.

Cabe recordar que el procesador es el elemento principal del ordenador y que interesa que las instrucciones y los datos con los que va a trabajar estén lo más próximos a él.

Cuando la CPU no encuentra un dato en alguno de los niveles de la memoria interna se obtiene del nivel inmediatamente inferior.

A medida que descendemos de nivel, el coste de adquisición se reduce, la velocidad disminuye, el tiempo de acceso aumenta y la capacidad será mayor.

Los **registros**, integrados en la CPU, están formados por un conjunto de biestables y almacenan bloques de bits que habitualmente se denominan **palabras**. Son capaces de realizar operaciones a la misma frecuencia que el procesador y su capacidad es muy pequeña.

La **memoria caché** es un tipo de memoria intermedia entre el procesador y la memoria principal. Este tipo de memoria suele estar formada por circuitos integrados **SRAM** o **RAM estáticos** que suelen ser más rápidos que los circuitos **DRAM** o **RAM dinámicos** empleados en la memoria principal. Al ser más rápidos, son más caros y voluminosos, de menor capacidad y mayor consumo energético. Este tipo de memorias se emplean para mantener la información más comúnmente usada por el procesador, evitando accesos continuos y más lentos a memoria principal. Dicha memoria sólo se usa como caché, debido a que su fabricación resulta cara, y por esta razón muchos comerciantes abaratan el coste de los microprocesadores quitando alguna caché o reduciendo el tamaño. Ejemplos son los microprocesadores CELERON de INTEL y DURON de AMD.

Los microprocesadores actuales incluyen en su propio chip total o parcialmente su caché.

Existen tres tipos diferentes de memoria caché para procesadores:

- **Caché de Primer nivel o L1:** Integrada en el núcleo del procesador, trabajando a la misma velocidad que éste. Su tamaño varía de un procesador a otro y suele estar dividida en dos partes dedicadas, una para instrucciones y otra para datos.

- **Caché de Segundo nivel o L2:** Integrada también en el procesador, aunque no directamente en el núcleo, tiene las mismas ventajas que la caché L1, aunque es algo más lenta que ésta. La caché L2 suele ser mayor que la caché L1. A diferencia de la caché L1, esta no está dividida, y su utilización está más encaminada a programas que al sistema.
- **Caché de Tercer nivel o L3:** Es un tipo de memoria caché más lenta que la L2, muy poco utilizada en la actualidad, incorporada a la placa base.

La memoria principal de un ordenador está organizada en grupos de celdas de memoria llamados **palabras de memoria**. Una palabra es el conjunto de bits que se pueden leer o memorizar en un instante dado y al número de bits se le denomina ancho de memoria o longitud de palabra.

2. Memoria secundaria o auxiliar

El gran inconveniente de la memoria principal, a pesar de ser muy rápida, es su baja capacidad de almacenamiento por lo que para guardar información de forma masiva se usan otros tipos de memorias.

La información guardada en este tipo de memorias permanece indefinidamente hasta que el usuario la borre de manera expresa (es lo que se denomina un almacenamiento **no volátil**). Estos dispositivos tienen mucha más capacidad que las memorias internas pero no podemos ejecutar programas desde esta memoria, es necesario pasar el programa completo o parte de éste a la memoria RAM para su ejecución.

Además pueden llegar a ser intercambiables, pudiéndose cambiar el soporte de almacenamiento de la información sin necesidad de cambiar la unidad lectora/ grabadora (cd, dvd o disco flexible).

Se puede llevar a cabo una **clasificación de este tipo de memorias** según diferentes criterios.

☐ Según la **tecnología empleada**:

- **Tecnología Magnética.** Emplean sustrato de plástico o aluminio cubierto de material magnetizable (óxido férrico o de cromo). La información se graba en celdas que forman líneas o pistas. Cada celda puede estar sin magnetizar o magnetizada con dos posibles valores: norte (0) y sur (1).
- **Tecnología Óptica.** Usan energía lumínica (como el rayo láser) para almacenar o leer información. Los ceros o unos se representan por la presencia o ausencia de señal luminosa.
- **Tecnología Magneto-Óptica.** Los soportes originarios poseen una magnetización previa (norte o sur) y mediante láser es posible cambiar la magnetización de las celdas.
- **Tecnología Flash-USB.** Usan memorias semiconductoras de tipo flash nand, con la característica de que no necesitan refresco al usar tecnología de puerta flotante.

☐ Según el **tipo de operaciones** que se pueda realizar sobre los mismos:

- **Reutilizables** (cinta magnética, disco cd-rw, etc.).
- **No Reutilizables** (disco cd-rom, etc.).

☐ Según la **forma de acceder** a la información:

- **Secuencial** (cinta magnética).
- **Directo** (cd-rom, disco duro).

☐ Según la **ubicación física de dicha unidad**:

- **Interna** (disco duro, unidad flexible).

- **Externa** (memoria USB, discos duros externos).

☐ Según la relación existente entre el soporte de almacenamiento y el elemento que lleva a cabo la lectura escritura:

- **Removibles** (discos flexibles).

- **No removibles** (discos duros).

Recordemos que este tipo de memoria aparece recogida en el esquema de Von Neumann* como formando parte del subsistema de entrada-salida siendo periféricos de almacenamiento.

***John Von Neumann** (1903-1957), nacido en Budapest, fue un niño prodigio con gran memoria fotográfica y talento para las matemáticas aunque su padre le prohibió dichos estudios porque pensaba que no era una carrera con la que ganar dinero así que lo engañó y estudiaba química en Berlín aunque estaba matriculado en Budapest en Matemáticas y sólo iba a los exámenes. En 1930 viaja a EEUU y al comenzar la Segunda Guerra Mundial trabaja para el gobierno americano y viendo las limitaciones del ENIAC y otras máquinas de computación definió un nuevo sistema lógico de computación.

El modelo básico de arquitectura empleado en los computadores digitales fue establecido en 1946 por **John Von Neumann**. Su aportación más significativa fue la de construir una computadora con programa almacenado ya que los computadores existentes hasta entonces trabajaban con programas cableados que se introducían estableciendo manualmente las conexiones entre las distintas unidades. La idea de Von Neumann consistió en conectar permanentemente las unidades de las computadoras, siendo coordinado su funcionamiento por un elemento de control. Esta tecnología sigue estando vigente en la actualidad aunque con pequeñas modificaciones y sigue siendo empleada por la mayoría de los fabricantes.

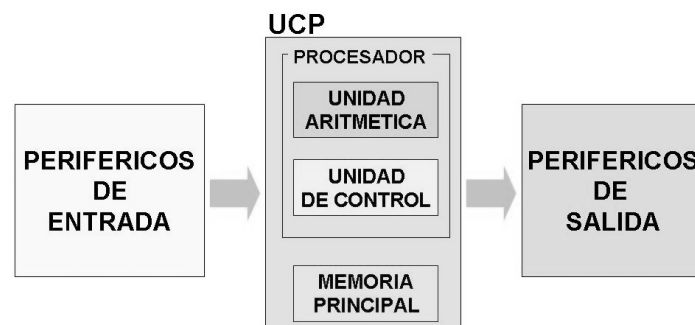


Figura 24. Arquitectura de Von Neumann.

Arquitectura Harvard (//saber como se estructura, el grafico que teneis mas abajo//)

Es una arquitectura de computadora con pistas de almacenamiento y de señal físicamente separadas para las instrucciones y para los datos. El término proviene de la computadora Harvard Mark I basada en relés, que almacenaba las instrucciones sobre cintas perforadas (de 24 bits de ancho) y los datos en interruptores electromecánicos. Estas primeras máquinas tenían almacenamiento de datos totalmente contenido dentro la unidad central de proceso, y no proporcionaban acceso al almacenamiento de instrucciones como datos. Los programas necesitaban ser cargados por un operador; el procesador no podría arrancar por sí mismo. Hoy en día (2018), la mayoría de los procesadores implementan dichas vías de señales separadas por motivos de rendimiento, pero en realidad implementan una arquitectura Harvard modificada, para que puedan soportar tareas tales como la carga de un programa desde una unidad de disco como datos para su posterior ejecución.

En la arquitectura Harvard, no hay necesidad de hacer que las dos memorias compartan características. En particular, pueden diferir la anchura de palabra, el momento, la tecnología de implementación y la estructura de dirección de memoria. En algunos sistemas, se pueden almacenar instrucciones en memoria de solo lectura mientras que, en general, la memoria de datos requiere memoria de lectura-escritura. En algunos sistemas, hay mucha más memoria de instrucciones que memoria de datos así que las direcciones de instrucción son más anchas que las direcciones de datos.

Contraste con arquitecturas von Neumann

Bajo arquitectura de von Neumann pura, la CPU puede estar bien leyendo una instrucción o leyendo/escribiendo datos desde/hacia la memoria pero ambos procesos no pueden ocurrir al mismo tiempo, ya que las instrucciones y datos usan el mismo sistema de buses. En una computadora que utiliza la arquitectura Harvard, la CPU puede tanto leer una instrucción como realizar un acceso a la memoria de datos al mismo tiempo, incluso sin una memoria caché. En consecuencia, una arquitectura de computadores Harvard puede ser más rápida para un circuito complejo, debido a que la instrucción obtiene acceso a datos y no compite por una única vía de memoria.

Además, una máquina de arquitectura Harvard tiene distintos código y espacios de dirección de datos: dirección de instrucción cero y dirección de datos cero son cosas distintas. La instrucción cero dirección podría identificar un valor de veinticuatro bits, mientras que dirección de datos cero podría indicar un byte de ocho bits que no forma parte de ese valor de veinticuatro bits.

En Contraste con la arquitectura Harvard modificada:

Una máquina de arquitectura Harvard modificada es muy similar a una máquina de arquitectura Harvard, pero relaja la estricta separación entre la instrucción y los datos, al mismo tiempo que deja que la CPU acceda simultáneamente a dos (o más) memorias de buses. La modificación más común incluye cachés de instrucciones y datos independientes, respaldados por un espacio de direcciones en común. Si bien la CPU ejecuta desde la memoria caché, también actúa como una máquina de Harvard pura. Cuando se accede a la memoria de respaldo, actúa como una máquina de von Neumann pura (donde el código puede moverse alrededor como datos, que es una técnica poderosa). Esta modificación se ha generalizado en modernos procesadores, tales como la arquitectura ARM y los procesadores x86. A veces se llama vagamente arquitectura Harvard, con vistas al hecho de que en realidad está "modificada".

Otra modificación proporciona un camino entre la memoria de instrucciones (como ROM o flash) y la CPU para permitir que las palabras de la memoria de instrucciones sean tratados como datos de solo lectura. Esta técnica es utilizada en algunos micro controladores, incluyendo el Atmel AVR. Esto permite datos constantes, tales como cadenas de texto o tablas de funciones, que puede acceder sin necesidad de ser previamente copiadas en datos de memoria, preservando memoria de datos escasa (y hambrienta de poder) de lectura / escritura de variables. Las instrucciones especiales de lenguaje de máquina se proporcionan para leer datos desde la memoria de instrucciones. (Esto es diferente a las instrucciones que a sí mismos embebiendo datos constantes, aunque para las constantes individuales de los dos mecanismos pueden sustituir unos por otros.)



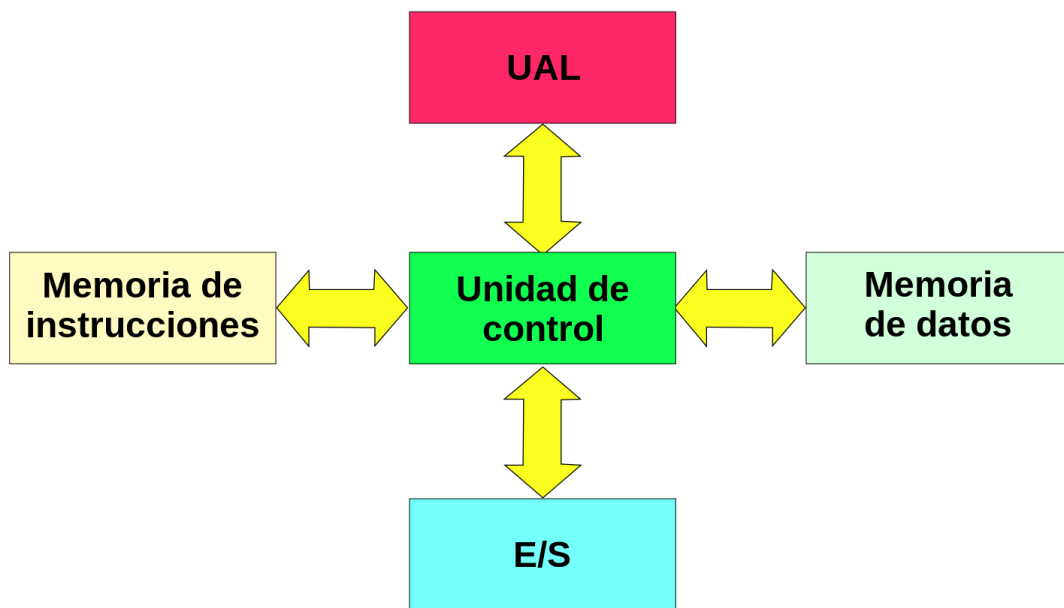


Figura Arquitectura Harvard.

5.2.2. Según la lectura y escritura en las mismas

//vip!!/

Dentro de las memorias internas, éstas se pueden clasificar atendiendo a la posibilidad de lectura o escritura en las mismas. Hablamos de:

- **Memorias de sólo lectura o programables.** No volátiles, no pierden la información en ausencia de alimentación.

Son memorias de este tipo: **ROM, PROM, EPROM, EEPROM y Flash.**

TIPOS DE MEMORIA ROM

PROM	EPROM (Erasable Programmable Read-Only Memory)	EEPROM (Electrically Erasable Programmable Read-Only Memory)	MEMORIA FLASH
Memoria inalterable programable. Un PROM es un chip de memoria en el que se puede salvar un programa, pero una vez que se haya utilizado el PROM no se puede reutilizarlo para salvar algo más. Como la ROM la PROM es permanente 	Se utiliza para corregir errores de última hora en la ROM, el usuario no la puede modificar y puede ser borrada exponiendo la ROM a una luz ultravioleta. 	Puede ser borrada y volver a ser programada por medio de una carga eléctrica, pero sólo se puede cambiar un byte de información a la vez. 	Tipo de memoria EEPROM que es reprogramable, su utilización por lo regular es en BIOS de ahí su nombre. 

Figura 25. Tipos Memoria ROM.

- **Memorias de lectura y Escritura.** La llamada memoria RAM, son memorias volátiles, que pierden la información en ausencia de alimentación.

Son memorias de este tipo: **SRAM, DRAM.**

1. **SRAM** (Static Random Access Memory), RAM estática, memoria estática de acceso aleatorio. Memoria que no necesita refresco. Se graba una vez la información y perdura a lo largo del tiempo. El coste es elevado. Su diseño es más complejo y la velocidad más alta. El tiempo de acceso es notablemente inferior.
 - volátiles.
 - no volátiles:
 - **NVRAM** (non-volatile random access memory), memoria de acceso aleatorio no volátil
 - **MRAM** (magnetoresistive random-access memory), memoria de acceso aleatorio magnetorresistiva o magnética
2. **DRAM** (Dynamic Random Access Memory), RAM dinámica, memoria dinámica de acceso aleatorio. Necesita refresco, es decir, que se realimente con tensión su contenido para que no se pierda. Son memorias más lentas y más sencillas de construir (consisten en una circuitería y elementos más fáciles de integrar) y menos costosa.
 - **DRAM Asíncrona** (Asynchronous Dynamic Random Access Memory, memoria de acceso aleatorio dinámica asíncrona*)
 - **FPM RAM** (Fast Page Mode RAM)
 - **EDO RAM** (Extended Data Output RAM)
 - **SDRAM** (Synchronous Dynamic Random-Access Memory, memoria de acceso aleatorio dinámica sincrónica)
 - Rambus:
 - a. **RDRAM** (Rambus Dynamic Random Access Memory)
 - b. **XDR DRAM** (eXtreme Data Rate Dynamic Random Access Memory)
 - c. **XDR2 DRAM** (eXtreme Data Rate two Dynamic Random Access Memory)
 - **SDR SDRAM** (Single Data Rate Synchronous Dynamic Random-Access Memory, SDRAM de tasa de datos simple)
 - **DDR SDRAM** (Double Data Rate Synchronous Dynamic Random-Access Memory, SDRAM de tasa de datos doble)
 - **DDR2 SDRAM** (Double Data Rate type two SDRAM, SDRAM de tasa de datos doble de tipo dos)
 - **DDR3 SDRAM** (Double Data Rate type three SDRAM, SDRAM de tasa de datos doble de tipo tres)
 - **DDR4 SDRAM** (Double Data Rate type four SDRAM, SDRAM de tasa de datos doble de tipo cuatro)

* **Síncrono/Asíncrono:** Unos conceptos particulares asociados a las memorias son el de memoria síncrona y asíncrona. Se dice que una memoria es síncrona cuando, el intercambio de datos entre la CPU y la memoria, se realiza mediante una señal de reloj externa a ésta que le indica en qué momentos puede recibir y devolver las peticiones de lectura o escritura, con lo que se evitan así ciclos de espera. La memoria es asíncrona cuando depende de unos relojes internos que sincronizan el microprocesador y los buses con el resto de elementos.

Todo sistema asíncrono tiene asignado unos tiempos mínimos para cada una de las operaciones internas, por lo que se pueden tener en accesos diferentes, distintos tiempos de respuesta del sistema. Esta es una de las desventajas de estos sistemas, pues aunque acaben en medio de un ciclo su operación, hasta el siguiente no pueden empezar otra, dado que han de resincronizarse con el resto de componentes.

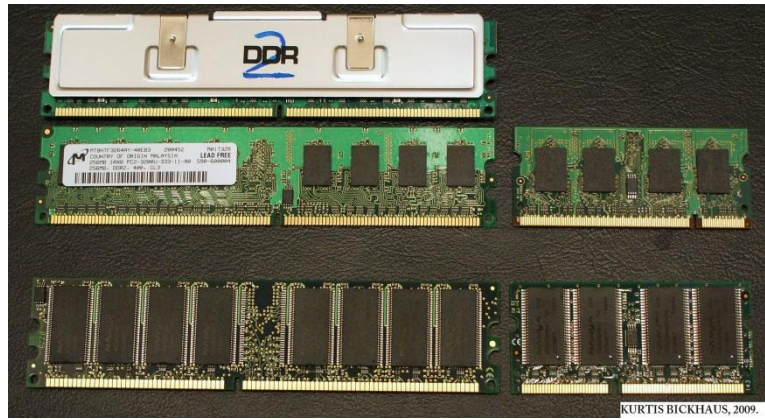


Figura 26. Memorias RAM.

MEMORIA CACHÉ:

En el diseño de la memoria caché se deben considerar varios factores que influyen directamente en el rendimiento de la memoria y por lo tanto en su objetivo de aumentar la velocidad de respuesta de la jerarquía de memoria. Estos factores son las políticas de ubicación, extracción, reemplazo y escritura.

Política de ubicación

Decide dónde debe colocarse un bloque de memoria principal que entra en la memoria caché. Las más utilizadas son:

- Directa: al bloque i -ésimo de memoria principal le corresponde la posición $i \bmod n$, donde n es el número de bloques de la memoria caché. Cada bloque de la memoria principal tiene su posición en la caché y siempre en el mismo sitio. Su inconveniente es que cada bloque tiene asignada una posición fija en la memoria caché y ante continuas referencias a palabras de dos bloques con la misma localización en caché, hay continuos fallos habiendo sitio libre en la caché.
- Asociativa: Los bloques de la memoria principal se alojan en cualquier bloque de la memoria caché, comprobando solamente la etiqueta de todos y cada uno de los bloques para verificar acierto. Su principal inconveniente es la cantidad de comparaciones que realiza.
- Asociativa por conjuntos: Cada bloque de la memoria principal tiene asignado un conjunto de la caché, pero se puede ubicar en cualquiera de los bloques que pertenecen a dicho conjunto. Ello permite mayor flexibilidad que la correspondencia directa y menor cantidad de comparaciones que la totalmente asociativa.

Política de extracción

La política de extracción determina cuándo y qué bloque de memoria principal hay que traer a memoria caché. Existen dos políticas muy extendidas:

- Por demanda: un bloque sólo se trae a memoria caché cuando ha sido referenciado y no se encuentre en memoria caché.
- Con prebúsqueda: cuando se referencia el bloque i -ésimo de memoria principal, se trae además el bloque $(i+1)$ -ésimo. Esta política se basa en la propiedad de localidad espacial de los programas.

Política de reemplazo

Determina qué bloque de memoria caché debe abandonarla cuando no existe espacio disponible para un bloque entrante. Básicamente hay cuatro políticas:

- Aleatoria: el bloque es reemplazado de forma [aleatoria](#).
- *FIFO*: se usa el algoritmo *First In First Out* ([FIFO](#)) (primero en entrar primero en salir) para determinar qué bloque debe abandonar la caché. Este algoritmo generalmente es poco eficiente.
- Usado menos recientemente (*LRU*): Sustituye el bloque que hace más tiempo que no se ha usado en la caché, traemos a caché el bloque en cuestión y lo modificaremos ahí.
- Usado con menor frecuencia (*LFU*): Sustituye el bloque que ha experimentado menos referencias.

Política de Actualización o Escritura

Determinan el instante en que se actualiza la información en memoria principal cuando se hace una escritura en la memoria es ejecutada. Existen 2 casos diferentes:

Sobre la memoria caché

- Escritura Inmediata: Se escribe a la vez en Memoria caché y Memoria principal para mantener una coherencia en todo momento.
- Postescritura: El bloque donde se escribió queda marcado con un bit llamado bit de basura, cuando se reemplaza por la política de re-emplazamiento se comprueba si el bit se encuentra activado, si lo está, se escribe la información de dicho bloque en la memoria principal.

Sobre la memoria principal

- Asignación en escritura: El bloque referenciado se copia de la memoria principal a la cache y después el bloque se envía a la CPU.
- No asignación en escritura: Envía el bloque directamente de la memoria principal a la CPU y al mismo tiempo la carga en la cache.

6. Buses: arquitecturas y funcionamiento

//VIP todo lo que esta en negrita, los buses fisicos saber que es cada uno sin entrar en mucho detalle//

6.1. Introducción

La interconexión de todas las unidades estudiadas se lleva a cabo a través de una serie de canales de conexión denominados **buses** que, físicamente, son un conjunto de líneas por las que se transmite la información binaria (sea de una instrucción, un dato, o una dirección, en un instante dado).

Se denomina **ancho de bus** al tamaño de ese número de hilos o bits que se transmiten simultáneamente por uno de esos canales.

Se pueden distinguir tres tipos de buses:

Bus de datos (bidireccional). Transporta datos procedentes o con destino a la memoria principal y las unidades de entrada-salida. Cabe destacar cómo la velocidad de este bus en su conexión con la memoria RAM es un factor determinante en el rendimiento del sistema.

Bus de direcciones (unidireccional). Transporta las direcciones de la unidad de control a la memoria principal o a los periféricos.

Bus de control (bidireccional). Transporta las señales de control (microórdenes) generadas por la unidad de control.



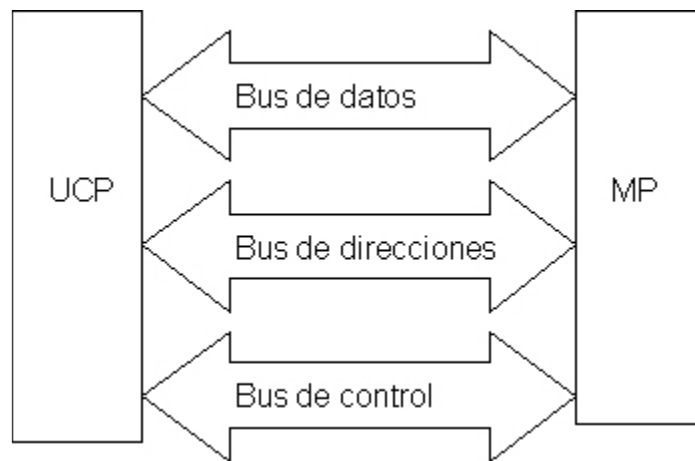


Figura 27. Tipos de buses en un ordenador.

Se suele denominar **ruta de datos** o **datapath** a todos aquellos elementos (buses, registros y circuitos lógicos) interconectados y encargados de transferir, memorizar y procesar las informaciones (instrucciones, direcciones y operandos) en su conexión con la memoria principal y los dispositivos periféricos.

6.2. Subsistema de E/S controladores y periféricos

Un ordenador tendría una utilidad nula sin la presencia de algún medio que permitiese realizar las entradas y salidas de datos para poder interactuar con el medio. El concepto de entrada y salida hace referencia a toda comunicación o intercambio de información entre la CPU o la memoria central con el exterior.

Estas operaciones se suelen llevar a cabo a través de una cada vez más amplia gama de dispositivos externos llamados **periféricos** que proporcionan al ordenador las vías para intercambiar datos con el exterior.

La parte del equipo que permite esta comunicación es la **unidad de entrada-salida**, entendiendo por tal concepto en arquitecturas reales a un conjunto de módulos o **canales de entrada-salida** encargados de gobernar uno o más periféricos asociados a los que suministra la inteligencia necesaria para su funcionamiento coordinado con el ordenador.

Estos **módulos de entrada-salida** estarían formados por los **controladores de periféricos** (circuitos de interfaz), de forma que cada periférico necesita su propio controlador para comunicarse con la CPU y los **puertos de entrada-salida**, que son registros que se conectan directamente a uno de los buses del ordenador. Cada puerto tiene asociada una dirección o código, de forma que el procesador ve al periférico como un puerto o un conjunto de puertos. Básicamente estos módulos o canales se usan para resolver las diferencias (velocidad de transmisión, formato de datos, etc.) que pueden existir entre el procesador y dichos periféricos. Sus funciones fundamentales son direccionamiento, transferencia y sincronización.

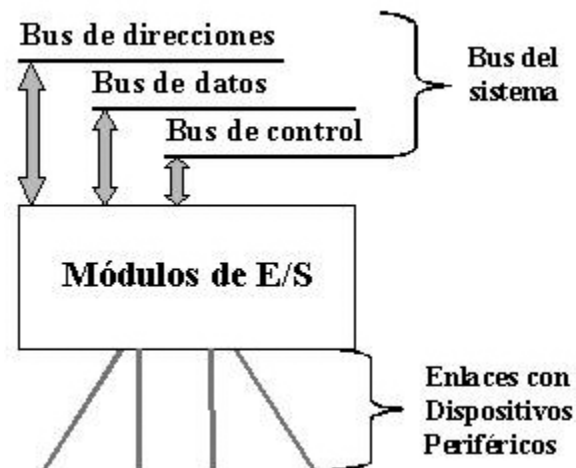


Figura 28. Módulos o canales de entrada-salida.

Existen distintas arquitecturas de ordenador según el direccionamiento de los dispositivos de entrada-salida.

- **Buses separados de memoria y entrada-salida.**
- **Entrada-salida mapeada en memoria o máquina de bus único.**

Los puertos de entrada-salida se tratan como si fuesen direcciones de memoria.

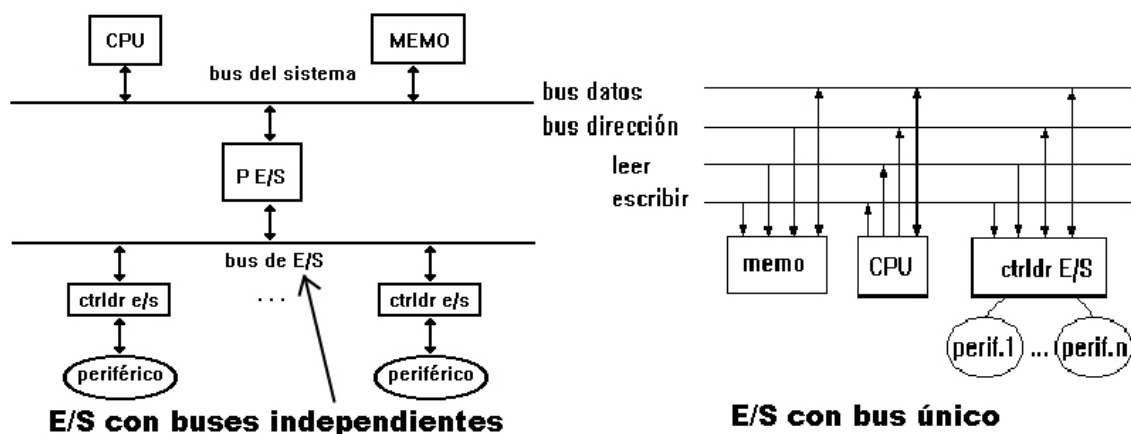


Figura 29. Entrada-salida mapeada a memoria o de bus único.

También existen distintas arquitecturas de ordenador según como se establece el **control del tránsito de datos**:

- **Entrada-salida controlada por programa.** Mediante la ejecución de unas instrucciones especiales si es de bus separado (in, out, etc.) e instrucciones de almacenamiento si es entrada-salida mapeada a memoria. Usada en periféricos con velocidades menores que la CPU.
- **Entrada-salida por acceso directo a memoria.** Todas las funciones se implementan mediante un circuito controlador llamado **controlador de DMA** (*direct access memory*).

Por último, también existen diferentes arquitecturas según el modo en que se produce la **sincronización de la CPU con los periféricos**:

- **Entrada-salida con sincronización por sondeo y selección.** La CPU hace encuesta o consulta a los dispositivos de su situación y los va atendiendo.

- **Entrada-salida con sincronización por interrupciones.** Los dispositivos son los que interrumpen la ejecución del programa en CPU cuando están en disposición de realizar una operación de entrada-salida.

6.3. Tipos de buses físicos

Los buses se pueden clasificar en dos tipos: internos (los que posee la CPU para interconectar sus componentes internos) o externos (los que interconectan los periféricos o la memoria con el microprocesador). Dentro de esta clasificación se puede especificar aún más, distinguiendo entre:

- ☐ Buses paralelos: los datos son enviados por *bytes*, al mismo tiempo, con la ayuda de varias líneas que tienen funciones fijas. La cantidad de datos enviada es bastante grande con una frecuencia moderada y es igual al ancho de los datos por la frecuencia de funcionamiento.
- ☐ Buses en serie: los datos son enviados, bit a bit, y se reconstruyen por medio de registros o rutinas de software. Está formado por pocos conductores y su ancho de banda depende de la frecuencia.

Partiendo de esta doble clasificación, se pueden mencionar los siguientes desarrollos comerciales:

☐ Buses Internos

- i. Paralelos: EV6 (de Athlon y Alpha), GTL+/AGTL+ de Intel ISA, EISA, VESA, MCA, PCI, AGP
- ii. Serie: PCI Express (PCIe), I2C, *HyperTransport*

☐ Buses Externos

- i. Paralelos: ATA (IDE, EIDE, ATAPI), SCSI, PCMCIA
- ii. Serie: SATA, USB, IEEE 1394 (FireWire)

A continuación se describen con más detalle alguno de ellos:

a) Bus EV6 de AMD y GTL+ (*Gunning Transceiver Logic*) /AGTL+ (*Assisted Gunning Transceiver Logic*) de Intel

Ambos son buses *Front Side Bus* (FSB), bus interfaz entre la CPU y la placa base.

El bus EV6 (creado por *Digital Equipment Corporation*) es el usado por el sistema AMD Athlon. Funciona a 200 MHz, y puede proporcionar un ancho de banda de datos de 1,6 GB/s, escalable hasta un ancho de banda de datos de 3,2 GB/s a 400 MHz.

Para alcanzar los 200 Mhz de frecuencia, en el bus se utiliza la tecnología DDR (*Double Data Rate*), que consiste en transferir datos en ambos sentidos del bus por cada ciclo, duplicando así el rendimiento. Esta tecnología es similar a la utilizada en DDR-SDRAM y en la interfaz Ultra ATA 66.

El EV6, más que un bus, se podría entender como un *switch* (conmutador) de red, ya que cada procesador tiene conexión completa con el *North Bridge* ejecutándose a una velocidad efectiva de 200Mhz (un 50% más rápido que el FSB de Intel).

El bus GTL+ es usado por Intel en sus procesadores Pentium II (100Mhz) y Pentium III (133Mhz).

Proporciona una única conexión al *North Bridge* que es compartida por todas la CPU's del ordenador, cuyo ancho de banda se divide entre todas ellas. El principio del que parte es que casi nunca una CPU necesita el ancho de banda total. Desafortunadamente, en esta situación de todo o nada de bus compartido, cada CPU debe esperar su turno para el uso del mismo (apareciendo la latencia de uso del bus). Este problema se agrava con la situación actual de CPU's con 800Mhz de velocidad y memorias a 133MHz, donde una CPU tiene que esperar seis ciclos de procesador para acceder a los datos.

Técnicamente el bus EV6 aporta una serie de ventajas respecto al GTL+, entre ellas, se pueden destacar:

- ☐ Una topología punto-a-punto (cada procesador es un sistema multiproceso dispone de su propio bus a 200Mhz), en el GTL+ los procesadores comparten este bus a 100Mhz por lo que cuantos más unidades menos ancho de banda).
- ☐ Una interfaz de memoria asíncrona (la velocidad de la memoria es independiente a la velocidad del bus y es determinada solamente por el *chipset*). Como curiosidad, se puede citar que el bus EV6 soporta sistemas multiprocesadores de hasta 14 procesadores. El espacio de direcciones del EV6 es también mayor que el del bus P6, es de 43 bits de profundidad contra los 36-bits del bus P6.
- ☐ Además, el EV6 permite transferencias de 64 bytes en ráfagas, en vez de los 32 soportados por el GTL+, por lo que el EV6 es capaz de mandar más datos a la vez a través del bus que el GTL+.
- ☐ El bus EV6 también añade seguridad, ya que añade un byte de ECC (*Error Checking Correction*) en todas las transferencias, una característica implementada en el bus GTL+ de Intel, pero que está ausente en todas las plataformas Super7.

b) Bus ISA (*Industrial Standard Architecture*, Arquitectura Industrial Estandarizada)

Podía utilizar entre 8 bits a 4,77Mhz (año 1980) y 16 bits a 8Mhz (año 1984). Los usuarios de máquinas basadas en ISA tenían que disponer de información especial sobre el hardware que iban a añadir al sistema. Aunque algunas tarjetas eran esencialmente *Plug-and-play* (enchufar y listo), no era lo habitual. Frecuentemente había que configurar varias cosas al añadir un nuevo dispositivo, como la IRQ, las direcciones de entrada/salida, o el canal DMA.

Este bus da conexión a las siguientes interfaces:

- ☐ Puertos paralelos: LPT1, LPT2
- ☐ Puertos serie: COM1, COM2
- ☐ Teclado y ratón: conector PS2
- ☐ FD o disquetera

Las evoluciones del ISA son el bus EISA con 32 bits (año 1987) que resuelve la integración de los periféricos y el bus VESA Local Bus (año 1992) dedicado a sistemas gráficos.

c) Bus MCA (*Micro Channel Architecture*)

Bus exclusivo mejorado diseñado por IBM en 1987 para utilizar en su línea de equipos PS/2. Este bus de 16 a 32 bits no era compatible con el bus ISA y podía alcanzar un rendimiento de 20 MB/s.

d) Bus PCI (*Peripheral Component Interconnect*)

Es un bus de ordenador estándar para conectar dispositivos periféricos directamente a su placa base.

A diferencia del bus ISA, el bus PCI permite configuración dinámica de un dispositivo periférico. En el tiempo de arranque del sistema, las tarjetas PCI y la BIOS interactúan y negocian los recursos solicitados por la tarjeta PCI. Su velocidad se mueve de los 66Mhz hasta los 133Mhz y da entrada a los conectores o interfaces de Entrada/Salida siguientes:

- ☐ Controladora IDE, para el disco duro
- ☐ Controladora SCSI, para el disco duro
- ☐ Conector/interfaz Firewire, para video en tiempo real
- ☐ Tarjeta de red, tarjeta de sonido, etc.

Ha evolucionado a distintas versiones:

- ☐ PCI-X (*PCI eXtended*): supera al bus PCI por su mayor ancho de banda, que suele ser exigido por servidores (en requerimientos de banda ancha: tarjetas *Ethernet Gigabit*, canal de Fibra, SCSI). Es una versión con el doble de ancho del PCI, ejecutándose hasta cuadruplicar la velocidad de reloj, estrategia similar en otras implementaciones eléctricas que usan el mismo protocolo. Fue desarrollado conjuntamente con IBM, HP y Compaq y presentado para su aprobación en 1998.



☐ PCIe (*PCI Express*): es un nuevo desarrollo del bus PCI que usa los conceptos de programación y los estándares de comunicación existentes, pero se basa en un sistema de comunicación serie mucho más rápido. Este sistema es apoyado principalmente por Intel. Aunque ambos (PCI-X y PCIe) son buses de alta velocidad para ordenadores y periféricos internos, se diferencian en muchas cosas. La primera es que PCI-X es un interfaz paralelo que es compatible con versiones anteriores de los dispositivos PCI a excepción de los que trabajan con 5 voltios. Por el contrario, PCIe es un bus serie con una interfaz física distinta que fue diseñada para superar a PCI y PCI-X.

Los buses PCI-X y PCI estándar pueden funcionar en un puente PCIe, de forma similar a como lo hacían los buses ISA, que funcionan en buses PCI normales en algunos ordenadores. PCIe también supera a PCI-X y PCI-X 2.0 en ancho de banda máximo: PCI-E 1.0 x1 ofrece 250 MB/s en cada sentido, y hasta en 32 líneas, otorgando un máximo de 8 GB/s en cada sentido.

PCI-X sufre una serie de desventajas tecnológicas y económicas respecto a PCIe:

☐ El interfaz paralelo de 64 bits requiere de por sí un difícil enrutado, ya que como todos los interfaces paralelos, las señales del bus deben llegar simultáneamente o en un margen muy corto de tiempo, y el ruido de las ranuras adyacentes puede causar interferencias. Por el contrario, el interfaz serie de PCI-E sufre muchos menos problemas y, por lo tanto, requiere diseños mucho menos complejos y más económicos.

☐ Los buses PCI-X, como los PCI estándar, son bidireccionales *half-duplex*, mientras PCIe son full duplex.

☐ Los buses PCI-X sólo son tan rápidos como el dispositivo más lento de su bus, y PCIe puede negociar las velocidades de transferencia dispositivo por dispositivo.

☐ Las ranuras (*slots*) PCI-X ocupan más espacio en las placas, lo cual puede llegar a ser un problema en caso de las ATX.

e) Bus USB (*Universal Serial Bus*)

// Todo tema de velocidades de estos buses lo preguntan mucho, sobretodo USB y thunderbolt. Aquí recomiendo IMPRESCINDIBLE estudiar por wikipedia la parte de USB: https://es.wikipedia.org/wiki/Universal_Serial_Bus//

Es un bus serie, a diferencia de los anteriores, que son paralelos y por tanto la información se transfiere bit a bit de forma consecutiva. El diseño del USB quería eliminar la necesidad de adquirir tarjetas separadas para poner en los puertos bus ISA o PCI, y mejorar las capacidades *plug-and-play*, permitiendo a esos dispositivos ser conectados o desconectados al sistema sin necesidad de reiniciar.

Sin embargo, en aplicaciones donde se necesita ancho de banda para grandes transferencias de datos, o si se necesita una latencia baja, los buses PCI o PCIe salen ganando. Igualmente sucede si la aplicación requiere de robustez industrial. A favor del bus USB, cabe decir que cuando se conecta un nuevo dispositivo, el servidor lo enumera y agrega el software necesario para que pueda funcionar (esto dependerá ciertamente del sistema operativo que esté usando el ordenador).

USB puede manejar 128 dispositivos compartiendo entre ellos un ancho de banda de 1.5Mbps.

USB está constituido por un conector de 4 hilos con las siguientes características:

☐ 2 hilos para entrada balanceada, es decir para datos que duplican la información por dos hilos. Se lee alternativamente de cada hilo, no de los dos hilos a la vez.

☐ 1 hilo para la tensión

☐ 1 hilo para masa o toma de tierra

Los USB tipo Mini (cámaras digitales) y Micro (dispositivos móviles) incluyen un hilo adicional que permite la distinción entre los tipo A y tipo B.

Respecto a la velocidad de transferencia, en la versión de USB 2.0, la velocidad de transferencia máxima es de 60 Megabytes por segundo.

En la versión 3.0 se multiplica por 10, pasando a 600 Megabytes por segundo.

La principal diferencia apreciable entre las versiones 2 y 3 es la velocidad de transferencia de datos, que es muy superior en el estándar USB 3.0. El soporte de formato HD es casi nulo en USB 2.0, pero ampliamente soportado por USB 3.0. Los dispositivos USB 3.0 se pueden conectar en conectores USB 2.0 y viceversa, si es de tipo A. Si es de tipo B o micro-B, los dispositivos USB 2.0 se pueden conectar en conectores USB 3.0, pero no al revés.

El problema es que en la simplicidad, que sin duda ha sido una ayuda durante muchos años, el USB 3.0 encontraba limitaciones importantes como no poder provisionar múltiples conexiones de forma simultánea. Esto, solo esto, ya hacía que las velocidades del interfaz fueran muy limitadas en la práctica.

Surge entonces el UAS (USB Attached SCSI) y en consecuencia el UASP (USB Attached SCSI Protocol). Este estándar o protocolo de transporte de datos por USB fue desarrollado inicialmente por el ANSI (Instituto Nacional Estadounidense de Estándares) pero es realmente desarrollado por el USB-IF (USB implementers fórum). Básicamente introduce técnicas de manejo de datos ya resueltas en un estándar tan veterano y avanzado como es SCSI. Esto incluye la autogestión del mejor modo de manejar múltiples conexiones, el ordenamiento de las mismas, etc. UASP desarrolla verdaderamente todo el potencial de ancho de banda de USB 3.0.

Tras la presentación de los nuevos macbook en 2015, Apple incorpora el estándar USB 3.1 con conector USB-C en sus equipos portátiles. El conector es similar al tamaño del microUSB, pero su forma es distinta ya que es reversible y retrocompatible con USB 3.0 y USB 2.0. USB-C permite cargar el portátil pero también conectarlo (de momento, vía un adaptador) a dispositivos USB convencionales, o bien a salidas de vídeo HDMI y D-SUB (VGA).

Esto posibilita prescindir del resto de puertos, a excepción del conector de 3,5 mm para auriculares.

USB 3.1 es una actualización de la anterior en la que se ha logrado doblar algunos datos. Por ejemplo, permite transferencias dos veces más rápidas, exactamente de 10Gbps. Además, puede manejar mayores cantidades de energía. Esto le permite cargar dispositivos que sólo requieren 5V o bien otros más demandantes gracias a que llegará hasta una potencia máxima de 100W. Es por eso que portátiles como el nuevo MacBook de Apple se puede cargar directamente a través de su conector USB-C o el último disco de Seagate se pueda alimentar solamente con un cable USB-C. Esto es posible gracias a la implementación de la especificación USB Power Delivery (USB PD).

USB-C es un conector compatible con los nuevos estándares USB y también retrocompatible con los anteriores. Gracias al conector USB-C y el uso conjunto del estándar USB 3.1 se puede contar de todas las ventajas de ambos. Pero también se puede implementar USB 3.1 en conectores como USB A, USB B, USB Micro B, etc.

Por tanto, existen cables con conectores USB-C pero que no obtienen ningún beneficio adicional a nivel de rendimiento respecto a cables microUSB. Esto se debe a que hay fabricantes que adoptan el conector por la ventaja de ser reversible pero implementan la especificación USB 2.0.

Además de la carga, USB 3.1 también ofrece novedades como su capacidad de transmitir vídeo con resolución hasta 4K o la integración con Thunderbolt 3.

A continuación se presenta una imagen comparativa de los tipos de conectores USB:



Figura 30. Comparativa de conectores USB.

USB 3.2: El 25 de julio de 2017, un comunicado de prensa del Grupo de Promotores USB 3.0 detalla una actualización pendiente de la especificación USB Tipo-C, que define la duplicación del ancho de banda para los cables USB-C existentes. Según la especificación USB 3.2, los cables existentes USB-C 3.1 Gen 1 con certificación SuperSpeed podrán funcionar a 10 Gbit / s (hasta 5 Gbit / s), y los cables USB-C 3.1 Gen 2 con certificación SuperSpeed + podrán operar a 20 Gbit / s (por encima de 10 Gbit / s). El aumento en el ancho de banda es el resultado de la operación de varios carriles sobre cables existentes que estaban destinados a las capacidades de flip-flop del conector tipo C.

El estándar USB 3.2 es compatible con USB 3.1 / 3.0 y USB 2.0. Define los siguientes modos de transferencia:

→USB 3.2 Gen 1x1 - SuperSpeed, 5 Gbit / s (0.625 GB / s) de velocidad de señalización de datos en 1 carril con codificación 8b / 10b, lo mismo que USB 3.1 Gen 1 y USB 3.0.

→USB 3.2 Gen 1x2 - SuperSpeed +, nueva velocidad de datos de 10 Gbit / s (1.25 GB / s) en 2 carriles con codificación 8b / 10b.

→USB 3.2 Gen 2x1 - SuperSpeed +, velocidad de datos de 10 Gbit / s (1.25 GB / s) en 1 carril con codificación 128b / 132b, lo mismo que USB 3.1 Gen 2.

→USB 3.2 Gen 2x2 - SuperSpeed +, nueva velocidad de datos de 20 Gbit / s (2.5 GB / s) en 2 carriles con codificación 128b / 132b.

En mayo de 2018, Synopsys demostró la primera conexión USB 3.2, donde una PC con Windows estaba conectada a un dispositivo de almacenamiento, alcanzando una velocidad de 1.6 GB / s promedio.

USB 3.2 es compatible con los controladores USB predeterminados de Windows 10 y en Linux Kernel 4.18.

f) Bus Thunderbolt

[//vip https://es.wikipedia.org/wiki/Thunderbolt_\(bus\) //](https://es.wikipedia.org/wiki/Thunderbolt_(bus))

La tecnología Thunderbolt llegó al mercado a principios de 2011 poniendo nombre hasta lo que hasta entonces se conocía como Light Peak. Es una interfaz de datos de gran ancho de banda que aún interfaz DisplayPort y datos en un único conector.

El ancho de banda de esta tecnología es de 10 Gbps (hasta 20 Gbps para Thunderbolt2 y 40 Gbps para Thunderbolt3).

La tecnología Thunderbolt se desarrolló por Intel y Apple bajo dos protocolos: PCI-Express y DisplayPort. Gracias a ello se pueden mover datos y vídeo por el mismo canal. Fue concebido para reemplazar a algunos buses actuales, tales como FireWire y HDMI.

A nivel físico, el ancho de banda de Thunderbolt 1 y Thunderbolt 2 es idéntico y por tanto el cable de la versión 1 es compatible con los interfaces de la versión 2. A nivel lógico, Thunderbolt 2 habilita la agregación de dos canales separados de 10 Gbps que pueden combinarse en uno único lógico de 20 Gbps.

La versión Thunderbot 3 utiliza conectores USB-C. En comparación con su versión anterior, Thunderbolt 3 reduce a la mitad el consumo de energía y acciona simultáneamente dos pantallas externas 4K a 60 Hz (o una sola pantalla externa 4K a 120 Hz) en lugar de sólo la única pantalla que los controladores anteriores podían manejar.

El nuevo controlador es compatible con PCIe 3.0 y otros protocolos, incluyendo HDMI 2.0 y DisplayPort 1.2 (resolución 4k a 60 Hz).

En virtud de ser un modo alternativo de USB-C, Thunderbolt 3 implementa USB Power Delivery, permitiendo hasta 100W de potencia, lo que permite a eliminar el cable de alimentación separado de algunos dispositivos.

Thunderbolt 3 permite la compatibilidad hacia atrás con las dos primeras versiones mediante el uso de adaptadores.

Los periféricos con thunderbolt son caros lo que hace que la tendencia del mercado vaya hacia USB 3.1 que además es retrocompatible con los dispositivos que usan versiones anteriores del estándar USB.

//VIP! Lo mas importante y preguntado historicamente de todo este punto son las siguientes tres tablas, sobre todo USB y Thunderbolt:

USB

Especificación USB	Velocidad de Transferencia Máxima	Característica especial
USB 1.0	1.5 Mbps	
USB 2.0	480 Mbps	Introdujo la capacidad de conectar dispositivos en caliente.
USB 3.0	5 Gbps	
USB 3.1	10 Gbps	Se añadió el Type C
USB 3.2 Gen 1x1	10 Gbps	
USB 3.2 Gen 1x2 2x1	10 Gbps	
USB 3.2 Gen 2x2	20Gbps	
USB 4.0	Hasta 40 Gbps	
Usb 4.0 V2	80 Gbps o 120 en asimétrico.	Tiene un modo en el que permite alcanzar 120 Gbps

THUNDERBOLT

	THUNDERBOLT 5	THUNDERBOLT 4	THUNDERBOLT 3	THUNDERBOLT 2	THUNDERBOLT 1
Velocidad	80Gbps & 120 Gbps con Bandwith Boost	40Gbps	40Gb/s	20Gb/s	10Gb/s
Video	Dos pantallas 8K - Tres pantallas 4k 144hz	Dos pantallas 4K	Una pantalla 4K	Pantallas 4K	Una pantalla 4K a 30hz
Carga	240w	100w	100w	No soporta	No soporta

DISPLAYPORT

	DisplayPort 1.2	DisplayPort 1.3	DisplayPort 1.4	DisplayPort 2.0	DisplayPort 2.1
Ancho de banda	17.28 Gbps	32.4 Gbit/s Gbps (4.32 Gbps por carril)	32.4 Gbps	77.4 Gbps	77.4 Gbps
Resoluciones	4K a 60Hz	4K a 120Hz / 5K a 60Hz / 8K a 30Hz	8K a 60Hz	15360 x 8640 (16K) a 60 Hz	15360 x 8640 (16K) a 60 Hz

Apunte: DisplayPort 2.1 se centra en mejorar la compatibilidad con el USB4 y el USB de tipo C.

Figura 31. TABLAS USB/THUNDERBOLT/DISPLAYPORT.

g) Bus AMR (Audio Modem Rise)

Es un bus para dispositivos de audio (como tarjetas de sonido) o módems lanzado en 1998 y presente en placas de Intel Pentium III, Intel Pentium IV y AMD Athlon. Fue diseñada por Intel

como una interfaz con los diversos chipsets para proporcionar funcionalidad analógica de Entrada/Salida, permitiendo que esos componentes fueran reutilizados en placas posteriores sin tener que pasar por un nuevo proceso de certificación. Fue sustituido por CNR, aunque actualmente ninguno de ambos es utilizado en las placas bases.

h) Bus CNR (*Communication and Network Riser*)

Se trata de una bus para dispositivos de comunicaciones como módems, tarjetas de red o USB. Un poco más grande que la AMR, CNR fue introducida en febrero de 2000 por Intel para procesadores Pentium y se trataba de un diseño propietario por lo que no se extendió más allá de las placas que incluían los *chipsets* de Intel, que más tarde fue implementada otros *chipset*.

i) Bus AGP (*Accelerated Graphics Port*)

Es un bus especializado para la generación de video, desarrollado por Intel en 1996. Todo tratamiento de señal de video necesita un gran ancho de banda, para ello, AGP constituye un bus dedicado que aísla el sistema de video.

El puerto AGP es de 32 bit como PCI, pero cuenta con diferencias tales como contar con 8 canales adicionales para acceso a la memoria RAM. Además, puede acceder directamente a ésta a través del puente norte, pudiendo emular así memoria de vídeo en la RAM. La velocidad del bus AGP 1X es de 66 MHz con una tasa de transferencia de 266 MB/s.

La versión AGP 8X tiene una velocidad de 533 MHz con una tasa de transferencia de 2 GB/s. Hoy en día se hace la integración del sistema de video con el bus PCI-X, es decir, las tarjetas gráficas se insertan en el bus PCI-X. Tiene mayor velocidad, aunque no sea dedicado exclusivamente y sea compartido. Por lo tanto, AGP tiende a desaparecer, siendo la tendencia el uso de PCI-X para todo.

j) PCMCIA (*Personal Computer Memory Card International Association*)

PCMCIA es una asociación Internacional centrada en el desarrollo de tarjetas de memoria para ordenadores personales que permiten añadir nuevas funciones.

Se fundó en 1989 con el objetivo de establecer estándares para circuitos integrados y promover la compatibilidad dentro de los ordenadores portátiles, donde solidez, bajo consumo eléctrico y tamaño pequeño son los factores más importantes. A medida que las necesidades de los usuarios de ordenadores portátiles han ido cambiando, lo ha hecho también el estándar de las tarjetas de PC.

En 1991 PCMCIA había definido una interfaz I/O (entrada/salida) para el mismo conector de 68 pines que inicialmente se usaba en las tarjetas de memoria. A medida que los diseñadores se daban cuenta de la necesidad de un software común para aumentar la compatibilidad, se fueron añadiendo primero las especificaciones de servicios de *socket*, seguidos de los servicios de especificación de tarjeta.

Los interfaces desarrollados bajo este estándar son los siguientes:

☐ **PCCard:** (originalmente PCMCIA), es un periférico diseñado para ordenadores portátiles. En un principio era usado para expandir la memoria, pero luego, se extendió a diversos usos como disco duro, tarjeta de red, tarjeta sintonizadora de TV, puerto paralelo, puerto serial, módem, puerto USB, etc. Fue diseñado por la industria estadounidense para competir con las tarjetas japonesas JEIDA.

☐ **CardBus:** es el nombre que reciben las tarjetas pertenecientes al estándar PCMCIA 5.0 o posteriores (JEIDA 4.2 o posteriores). Todas ellas son dispositivos de 32 bits y están basadas en el bus PCI de 33 MHz (a diferencia de las PC Card que pueden ser de 16 o 32 bits).

Incluye Bus Mastering, que admite la comunicación entre su controlador y los diversos dispositivos conectados a él sin que intervenga el CPU. Muchos *chipsets* están disponibles tanto para PCI y CardBus, como los que soportan Wi-Fi.



❑ **ExpressCard:** es un estándar de hardware que reemplaza a las tarjetas PCMCIA/PC Card/CardBus. El dispositivo host soporta conectividad PCI Express, USB 2.0 (incluyendo Hi-Speed) y USB3.0 (Superspeed) (sólo con ExpressCard 2.0) a través del *slot* ExpressCard, y cada tarjeta utiliza el modo de conectividad que el diseñador considere más apropiado para la tarea. Las tarjetas pueden conectarse al ordenador encendido sin necesidad de reiniciarlo (soportan *hot swap* o conexión en caliente). Es un estándar abierto de la organización internacional ITU-T.

El beneficio principal de la tecnología *ExpressCard* sobre CardBus es un gran aumento del ancho de banda, causado porque *ExpressCard* tiene una conexión directa al bus del sistema sobre una conexión PCI-Express x1 o USB, mientras que CardBus utiliza un controlador de interface que sólo interactúa con PCI.

ExpressCard tiene un rendimiento de procesamiento máximo de 2,5 Gb/s sobre PCIe ó 480 Megabits/segundo sobre USB 2.0 dedicado para cada ranura, mientras que CardBus debe compartir un ancho de banda de 1.066 Megabits/segundo.

La tecnología *ExpressCard* no es compatible con los dispositivos CardBus, lo cual puede representar un problema para quienes compren un nuevo sistema y quieran mantener sus dispositivos anteriores.

k) Bus I2C (Inter Integrated Circuit)

Es un bus de comunicaciones en serie. La versión 1.0 data del año 1992 y la versión 2.1 del año 2000, su diseñador es Philips. La velocidad es de 100Kbits por segundo en el modo estándar, aunque también permite velocidades de 3.4 Mbit/s. Es un bus muy usado en la industria, principalmente para comunicar microcontroladores y sus periféricos en sistemas integrados (Embedded Systems) y generalizando más para comunicar circuitos integrados entre sí que normalmente residen en un mismo circuito impreso.

Su principal característica es que utiliza dos líneas para transmitir la información: una para los datos y por otra la señal de reloj. También es necesaria una tercera línea, pero ésta sólo es la referencia (masa). Como suelen comunicarse circuitos en una misma placa que comparten una misma masa, esta tercera línea no suele ser necesaria.

Los dispositivos conectados al bus tienen una dirección única para cada uno. También pueden ser maestros o esclavos. El dispositivo maestro inicia la transferencia de datos y, además, genera la señal de reloj, pero no es necesario que el maestro sea siempre el mismo dispositivo. Esta característica se la pueden ir pasando los dispositivos que tengan esa capacidad y confiere a este bus la denominación de multimaestro.

l) HyperTransport (HT), también conocido como Lightning Data Transport (LDT).

Es una tecnología de comunicaciones bidireccional, que funciona tanto en serie como en paralelo, y ofrece un gran ancho de banda en conexiones punto a punto de baja latencia. Se publicó el 2 de abril de 2001. Esta tecnología se aplica en la comunicación entre chips de un circuito integrado, ofreciendo un enlace (ó bus) de conexión universal de alta velocidad y alto rendimiento para aplicaciones incorporadas y facilitando sistemas de multiprocesamiento altamente escalables. Esta conexión universal está diseñada para reducir el número de buses dentro de un sistema.

Existen tres versiones de HyperTransport (1.0, 2.0 y 3.0) que pueden funcionar desde los 200MHz hasta 2.6GHz (mientras el bus PCI corre a 33 o 66 MHz). También soporta tecnología DDR (*Double Data Rate*), lo cual permite alcanzar un máximo de 5200 MT/s (2600MHz hacia cada dirección: entrada y salida) funcionando a su máxima velocidad (2.6GHz).

m) Bus IEEE 1394 (conocido como FireWire por Apple Inc. y como i.Link por Sony)

Es un estándar multiplataforma para entrada/salida de datos en serie a gran velocidad. Se suele utilizar para la interconexión de dispositivos digitales como cámaras digitales y

videocámaras a ordenadores. Tiene una velocidad máxima de 800 Megabits por segundo (IEEE 1394b).

7. Clasificación de los ordenadores

//lectura, me parece importante que sepáis que es cada tipo como mínimo para tener cultura y comprender temas que se verán más adelante, tanto esto como lo de los servidores del siguiente punto//

7.1. Introducción. Tipos de ordenadores

Se puede distinguir entre dos tipos principales de ordenadores según el tipo de procesador:

- **Ordenadores analógicos:** Un ordenador analógico u ordenador real es un tipo de ordenador que utiliza dispositivos electrónicos o mecánicos para modelar el problema a resolver utilizando un tipo de cantidad física para representar otra. Los ordenadores analógicos ideales operan con números reales y son diferenciales. En el típico ordenador analógico electrónico, las entradas se convierten en tensiones que pueden sumarse o multiplicarse empleando elementos de circuito de diseño especial. Las respuestas se generan continuamente para su visualización o para su conversión en otra forma deseada.
- **Ordenadores digitales:** Un ordenador digital dispone de un sistema digital con tecnología microelectrónica y es capaz de realizar el procesamiento de datos a partir de un programa. Los ordenadores digitales se limitan a números computables y son algebraicos.

Hay un grupo intermedio, los **ordenadores híbridos**, en los que un ordenador digital es utilizado para controlar y organizar entradas y salidas hacia y desde dispositivos analógicos anexos.

Los ordenadores analógicos tienen una tasa de dimensión de la información, o potencial de dominio informático más grande que los ordenadores digitales. Esto, en teoría, permite a los ordenadores analógicos resolver problemas que son indecifrables con ordenadores digitales.

Dentro de los ordenadores digitales se definen los siguientes tipos principales según su tamaño:

- **Miniordenador:** Ordenador de nivel medio diseñado para realizar cálculos complejos y gestionar eficientemente una gran cantidad de entradas y salidas de usuarios conectados a través de un terminal. Normalmente, los miniordenadores se conectan mediante una red con otros miniordenadores, y distribuyen los procesos entre todos los equipos conectados. En la actualidad, el término miniordenador se sustituye con frecuencia por el de servidor.
- **Microordenador:** Ordenador de sobremesa o portátil, que utiliza un microprocesador como su unidad central de procesamiento o CPU. Los microordenadores más comunes son los ordenadores personales o PC. Se incluyen dentro de esta categoría los ordenadores de sobremesa, los ordenadores portátiles, los tablet pc, consolas de videojuegos, dispositivos móviles, etc. El desarrollo de los microordenadores fue posible gracias a dos innovaciones tecnológicas en el campo de la microelectrónica: el circuito integrado y el microprocesador.
- **Estación de trabajo (Workstation):** Las estaciones de trabajo son ordenadores que tienen unas prestaciones especiales (altas capacidades de rendimiento y fiabilidad) y se dedican a un fin específico, como puede ser el diseño asistido por ordenador. En una red de ordenadores, una estación de trabajo es un ordenador que facilita a los usuarios el acceso a los servidores y periféricos de la red.
- **Superordenador:** Ordenador de gran capacidad, gran potencia y de coste elevado, utilizado en cálculos complejos o tareas muy especiales. Normalmente se trata de una máquina capaz de distribuir el procesamiento de instrucciones y que puede utilizar instrucciones vectoriales. Los superordenadores se usan, por ejemplo, para hacer las



previsiones meteorológicas, para construir modelos científicos a gran escala y en los cálculos de las prospecciones petrolíferas.

- **Ordenador central (Mainframe):** Ordenador grande, potente y costoso usado principalmente por una gran compañía para el procesamiento de una gran cantidad de datos; por ejemplo, para el procesamiento de transacciones bancarias. Los superordenadores se centran en los problemas limitados por la velocidad de cálculo mientras que los ordenadores centrales se centran en problemas limitados por los dispositivos de E/S y la fiabilidad.
- **Servidores.** Los servidores se pueden ejecutar en cualquier tipo de ordenador, incluso en equipos dedicados a los que se les conoce como “el servidor”. Un servidor es una aplicación en ejecución capaz de atender las peticiones de un cliente y devolverle una respuesta en concordancia. En la mayoría de los casos un mismo equipo puede proveer múltiples servicios y tener varios servidores en funcionamiento. La ventaja de montar un servidor en ordenadores dedicados es la seguridad y gestión, motivo por el que la mayoría de los servidores son procesos diseñados de forma que puedan funcionar en ordenadores de propósito específico.

Clasificación de ordenadores digitales en función de sus prestaciones:

Microordenador < Estación de trabajo < Miniordenador < Ordenador central $\cong$ Superordenador

7.2. Equipos departamentales (miniordenador)

Características

Hasta finales de la década de 1960, todos los ordenadores eran ordenadores centrales y eran excesivamente caros, en consecuencia sólo las grandes empresas podían permitirse adquirir estas máquinas. Sobre esa época, aparecieron los miniordenadores (o simplemente minis), más pequeños, ligeros y asequibles para pequeñas compañías.

Los miniordenadores están a medio camino entre el ordenador central y el microordenador. Eran bastante comunes en las empresas, en los departamentos de finanzas, personal y contabilidad para compartir recursos mediante un miniordenador y un ordenador central. Los miniordenadores eran usados en conjunción con terminales tontos sin capacidad de cálculo propio.

Además de las diferencias obvias en velocidad de proceso de datos, la mayor diferencia entre miniordenadores y ordenadores centrales es el número de estaciones de trabajo remotas a las que pueden dar servicio. En general, cualquier ordenador que da servicio a más de 100 estaciones de trabajo remotas no debe ser llamado miniordenador. Los ordenadores centrales más rápidos y potentes dan servicio a más de 10.000 estaciones de trabajo remotas.

Funcionalidad

Las principales aplicaciones de este tipo de máquinas se detallan a continuación:

- **Servidor principal** de una empresa pequeña o mediana, permite tener múltiples aplicaciones de gestión estandarizadas a un precio asequible. También se pueden utilizar en un proceso de downsizing para hacer una transacción progresiva: rescata procesos de un ordenador central y los delega en los otros servidores.
- **Servidor asociado** o incluso departamental instalado en una dependencia separada y dependiente de la sede central de la organización. Se concebía la máquina departamental como una herramienta de productividad con un precio bajo que evitaba el empleo masivo de líneas de teleproceso hacia los ordenadores centrales corporativos. No se asocian en proceso sino en acceso a datos.

En la actualidad esta aplicación se da en los Data Mining y Data Warehouse: se filtra la información para obtener los datos más significativos. De este modo no se reduce

capacidad del ordenador central. Por otra parte, el coste del software de Datamining para un ordenador central es muy superior al que precisa un miniordenador.

- **Servidor autónomo** para aplicaciones poco repetitivas o especializadas como las de 'Automatización de oficinas' con soporte directo del proveedor. Se utiliza cuando se requiere una capacidad transaccional en un servicio central del ordenador central y una capacidad transaccional distribuida y distante que recaee en el miniordenador. El miniordenador recoge los datos en un punto cercano a su origen, realiza la transacción y retransmite el resultado al ordenador central. Históricamente esto era lo que sucedía en la administración, recogía los datos en las delegaciones territoriales.

Tendencias

Actualmente, el término miniordenador se sustituye con frecuencia por el de servidor, como elemento central de una red de medianas dimensiones a la que se conectan ordenadores personales, hasta varios centenares de usuarios.

La tendencia generalizada es la descentralización en la capacidad de proceso, con lo cual proliferan las redes locales conectadas al ordenador central y se tiende a sustituir los terminales convencionales por ordenadores personales.

Aprovechando las cada vez más potentes prestaciones de las plataformas pequeñas y medias, se propone la migración de las tradicionales aplicaciones de mainframe a dichas plataformas (downsizing), utilizando como plataforma lógica los sistemas abiertos.

Se combina así un coste más favorable, proveniente de la utilización de máquinas más pequeñas, y la portabilidad e interoperabilidad de los sistemas abiertos.

Aunque se tiende a la implantación de sistemas abiertos en sistemas pequeños y medios, en sistemas grandes la migración no se ha producido tan rápidamente como se podía pensar. Las razones son dos: por un lado, las fuertes inversiones ya realizadas por los usuarios de grandes sistemas propietarios, tanto en equipos físicos como en el software de su propiedad, y por otro, las dificultades y riesgos involucrados en la migración de las aplicaciones.

La virtualización está provocando cambios en las operaciones e infraestructuras de TI. Esta tecnología está transformando el mercado y creando una competencia entre los fabricantes.

El aumento de la capacidad de proceso de los miniordenadores así como la mejora en las unidades de almacenamiento hace que la diferenciación entre miniordenadores y grandes ordenadores cada vez sea menor.

Las principales tendencias en estos sectores son:

- ☑ Aumento de la capacidad de proceso de los ordenadores que hacen posible el uso de bases de datos textuales y de formatos documentales multimedia.
- ☑ Disminución del coste de almacenamiento de datos.
- ☑ Bases de datos distribuidas
- ☑ Tolerancia a fallos y arquitecturas multiprocesador.

Soluciones comerciales

- ☑ FUJITSU Superordenador PRIMEHPC FX100 (SPARC64 Xlfx)
- ☑ Fujitsu M10 Servers de Oracle (SPARC SPARC64 X+)
- ☑ Oracle SPARC M7

7.3. Estaciones de trabajo (Workstations)

Características

Una estación de trabajo se puede definir como: Micro o miniordenador para un único usuario, de alto rendimiento, que ha sido especializada para gráficos, diseño asistido por ordenador, ingeniería asistida por ordenador o aplicaciones científicas.

Una estación de trabajo está optimizada para desplegar y manipular datos complejos como el diseño mecánico en 3D, la simulación de ingeniería, los diagramas matemáticos, etc. Las estaciones de trabajo fueron un tipo popular de ordenadores para ingeniería, ciencia y gráficos durante las décadas de 1980 y 1990. Últimamente se las asocia con CPUs RISC, pero inicialmente estaban basadas casi exclusivamente en la serie de procesadores Motorola 68000. Pueden disponer de más de una CPU.

Funcionalidad

Las principales aplicaciones de las estaciones de trabajo son las siguientes:

- **Aplicaciones técnicas**

- **CAD** (Computer Aided Design, Diseño Asistido por Ordenador), destinadas al diseño y análisis de sistemas de ingeniería y arquitectura.
- **AEC** (Architecture, Engineering and Construction, Sistemas para la arquitectura / ingeniería / construcción) aplicables a la edición de planos de construcción y arquitectura, elaboración de presupuestos y seguimiento de obras.
- **CAM** (Computer Aided Manufacturing, Fabricación Asistida por Ordenador), aplicables en la ingeniería de producción, control de procesos y gestión de recursos.
- **EDA** (Electronic Design Automation, Diseño Electrónico Automatizado), aplicables en el diseño, análisis y prueba de circuitos integrados, tarjetas y otros sistemas electrónicos.
- **CASE** (Computer Aided Software Engineering, Ingeniería de Software Asistida por Ordenador), que ayudan a la gestión completa de los ciclos de vida de los desarrollos de aplicaciones lógicas.

- **Aplicaciones científicas**

- **GIS** (Geographic Information System, Sistemas de Información Geográfica) para edición, captura y análisis de información sobre determinadas zonas geográficas, con base en referencias de mapas digitalizados.
- Aplicaciones orientadas a la industria química y farmacéutica, aplicaciones de laboratorio tales como instrumentación analítica, modelado de experimentos químicos y sistemas de información de laboratorios químicos.
- Aplicaciones dentro del campo de la medicina para la captura, manipulación y presentación de imágenes de rayos X, ultrasonidos y escáneres, así como sistemas de información propios de hospitales.
- Sistemas de análisis de los recursos de la Tierra para análisis geológicos y sísmicos sobre mapas.
- Sistemas expertos, basados en técnicas de programación de inteligencia artificial, para aplicaciones tales como detección electrónica de errores, funciones de diagnóstico o configuración de ordenadores.

- **Aplicaciones comerciales**

El desarrollo de aplicaciones comerciales para estaciones de trabajo se ha visto potenciado por la tendencia actual a migrar desde entornos de miniordenadores y grandes ordenadores a arquitecturas cliente-servidor y equipos de menor rango, fenómeno denominado downsizing.

Dentro del segmento de las aplicaciones comerciales se pueden englobar:

- **Aplicaciones empresariales:** Aplicaciones de investigación cuantitativa, seguridad, bases de datos, simulación de análisis reales.
- **Postproducción Digital:** Elaboración y tratamiento de imágenes de calidad cinematográfica.



- **Sistemas de Gestión Documental:** Creación y organización de documentos para su posterior publicación, como pueden ser: periódicos, revistas, presentaciones y documentación en general.
- **Telecomunicaciones:** Gestión de redes, desarrollo de aplicaciones de telecomunicaciones basadas en inteligencia artificial, aplicaciones de apoyo a la investigación y desarrollo (I+D). Las estaciones de trabajo también pueden ser utilizadas como pasarelas (gateways), para acceder a grandes ordenadores, así como servidores de Correo electrónico, Web, etc.

Tendencias

Se puede afirmar que las estaciones de trabajo encuentran su segmento del mercado donde la potencia de cálculo y la capacidad de tratamiento gráfico de los PCs no son suficientes y el precio de los grandes ordenadores no puede ser financiado o no compensa respecto a las necesidades que se pretenden cubrir.

Siguiendo las tendencias de rendimiento de los ordenadores en general, los ordenadores promedio de hoy en día son más potentes que las mejores estaciones de trabajo de una generación atrás. Como resultado, el mercado de las estaciones de trabajo se está volviendo cada vez más especializado, ya que muchas operaciones complejas que antes requerían sistemas de alto rendimiento pueden ser ahora dirigidas a ordenadores de propósito general. Las estaciones de trabajo están experimentando un espectacular incremento en cuanto a prestaciones. Entre sus características figuran estructuras propietarias de bus ultrarrápido, y tecnología de disco tolerante a fallos como las matrices de discos redundantes (RAID). Incluyen facilidades de administración del equipo lógico que permiten la monitorización remota de todos los componentes esenciales (memoria, espacio en disco, etc.).

Muchas de ellas incorporan también protección ante cortes de fluido eléctrico y de bajadas de tensión.

Hoy en día la mayoría de las estaciones de trabajo se basan en tecnologías RISC, aunque todavía se encuentran instaladas numerosas estaciones de trabajo basadas en tecnología CISC y existen algunos fabricantes que todavía ofrecen productos basados en esta última arquitectura. No obstante, el futuro parece dirigirse hacia arquitecturas RISC. Esta tecnología es la que en este momento ofrece mejores relaciones rendimiento/precio, planes de diseño más cortos, mayor grado de fiabilidad y mayor potencia de proceso con menores costes.

Las principales líneas de evolución de las estaciones de trabajo son:

- ☑ Empleo de la tecnología ULSI (Ultra Large Scale Integration) lo que disminuye los costes asociados a los circuitos electrónicos.
- ☑ Mejora de la portabilidad e interoperabilidad del equipo lógico mediante el soporte de interfaces estándares (POSIX) y de protocolos (TCP/IP, ISO/OSI).
- ☑ Redundancia de componentes para conseguir sistemas tolerantes a fallos que aumentan la disponibilidad.
- ☑ Gestión de memoria basada en caché, combinando chips de distintas potencias y velocidades.
- ☑ Mejoras en las prestaciones de las arquitecturas multiproceso y proceso paralelo.

El futuro próximo de las estaciones de trabajo y su éxito dependerán del número de aplicaciones que puedan soportar en diferentes entornos: técnicos, científicos, comerciales, comunicaciones, educación, financiero, administración, medicina, etc. En este sentido, todos los proveedores de equipos lógicos están preocupados por adaptar desarrollos, concebidos inicialmente para el mercado de los PCs, a plataformas RISC, más propias de las estaciones de trabajo. Este dato permite augurar que las inversiones que se realicen en este tipo de tecnologías quedarán aseguradas durante varios años.



Las estaciones de trabajo también podrían introducirse en el mercado de los PCs conforme vayan ganando la batalla de los costes, ofreciendo facilidades de red que permiten compartir discos, impresoras, etc.

Lo mismo parece suceder en el entorno de los servidores, donde el incremento en la potencia de proceso los convertiría en el ordenador ideal para situarlo como servidor de grupos de trabajo, con una potencia de cálculo próxima a la de los miniordenadores y con unos costes que es posible ajustar mucho más respecto a las necesidades de proceso y de cálculo, precisamente debido a sus características de escalabilidad.

Una importante característica que predice el éxito de las estaciones de trabajo es su diseño conforme a los estándares lo que asegura, por un lado, el aprovechamiento de las inversiones realizadas por los usuarios, y por otro, la interoperabilidad con productos de otros fabricantes. En este mismo sentido se está trabajando en el desarrollo de plataformas que aseguren la portabilidad de aplicaciones ya compiladas, desde una estación antigua a una más moderna sin necesidad de ninguna modificación sobre el código (compatibilidad a nivel binario).

Dentro de las líneas de actuación destaca la estrategia de Sistemas Abiertos de la Administración del Estado.

Las administraciones europeas, tanto las comunitarias como las nacionales, vienen impulsando desde hace años la implantación de Sistemas Abiertos.

Soluciones comerciales

Los principales fabricantes que ofrecen soluciones comerciales para estaciones de trabajo son:

☐ **Dell:** Series Precision

☐ **IBM:** Series IntelliStation

☐ **Oracle** Sun Fire basados en Opteron AMD

☐ **HP:** ISV y HP Z

7.4. Servidores

Características

Un servidor a nivel **software** es una aplicación en ejecución capaz de atender las peticiones de un cliente y devolverle una respuesta en concordancia. Los servidores se pueden ejecutar en cualquier tipo de ordenador o incluso en ordenadores dedicados.

Por ejemplo, un servidor web como Apache es multiplataforma y puede ejecutarse en cualquier equipo que cumpla con unos requerimientos mínimos.

A nivel **hardware**, un servidor (host o anfitrión) es un equipo informático especializado con muy altas capacidades de proceso, encargado de proveer diferentes servicios a las redes de datos o conjunto de ordenadores interconectados entre sí, tanto inalámbricas como cableadas.

Los servidores hardware se suelen montar en cabinas especiales denominados racks, donde es posible colocar varios servidores en los compartimentos especiales y ahorrar espacio. Tienen sistemas que les permiten resolver ciertas averías de manera automática así como sistemas de alerta para evitar fallas en operaciones de datos críticos, ya que deben estar encendidos los 365 días del año las 24 horas del día.

Funcionalidad //vip//

Comúnmente los servidores proveen servicios esenciales dentro de una red, ya sea para usuarios privados dentro de una organización o compañía, o para usuarios públicos a través de Internet. Los tipos de servidores más comunes son servidor de base de datos, servidor de



archivos, servidor de correo, servidor de impresión, servidor web, servidor de juego, y servidor de aplicaciones.

A continuación se describen algunos tipos comunes de servidores:

- **Servidor de archivos:** es el que almacena varios tipos de archivos y los distribuye a otros clientes en la red.
- **Servidores Storage:** se trata de servidores que tienen como fin principal el almacenamiento de grandes cantidades de información, por lo que su característica principal es contar con una conexión de red de alta velocidad, así como de dispositivos de almacenamiento de muy alta capacidad y velocidad como discos duros SATA, SCSI o SAS.
- **Servidor de impresión:** controla una o más impresoras y acepta trabajos de impresión de otros clientes de la red, poniendo en cola los trabajos de impresión (aunque también puede cambiar la prioridad de las diferentes impresiones), y realizando la mayoría o todas las otras funciones que en un sitio de trabajo se realizaría para lograr una tarea de impresión si la impresora fuera conectada directamente con el puerto de impresora del sitio de trabajo.
- **Servidor de correo:** almacena, envía, recibe, enruta y realiza otras operaciones relacionadas con el correo electrónico para los clientes de la red.
- **Servidor de telefonía:** realiza funciones relacionadas con la telefonía, como es la de centralita telefónica o PBX, contestador automático, realizando las funciones de un sistema interactivo para la respuesta de la voz, almacenando los mensajes de voz, encaminando las llamadas y controlando también la red o Internet, por ejemplo, la entrada excesiva de la voz sobre IP (VoIP), etc.
- **Servidor de actualizaciones:** por cuestiones de seguridad, la mayor parte de las empresas no permiten que los equipos se encuentren con acceso a Internet, ya que ello puede generar fuga de información o accesos no autorizados a la red. El problema que esto conlleva es que las actualizaciones que los sistemas operativos ponen a disposición requieren de acceso a Internet, por lo que un servidor de actualizaciones se encarga de descargar las actualizaciones y permitir que los equipos se conecten a él para descargarlas. El más popular es WSUS (Windows Server Update Services) utilizado por Microsoft Windows Server, actualmente integrado en su herramienta SCCM (System Center Configuration Manager).
- **Servidor VPN** para la gestión de las comunicaciones externas a nuestra red corporativa.
- **Servidor web** que distribuyen contenido web a los clientes que lo solicitan.
- **Servidor de base de datos:** proveen servicios de base de datos a otros programas o equipos.
- **Servidor de Seguridad** como Firewalls, IPS, Sistemas Antispam, antivirus, etc.
- **Servidor dedicado:** son aquellos que le dedican toda su potencia a administrar los recursos de la red, es decir, a atender las solicitudes de procesamiento de los clientes.
- **Servidor no dedicado:** son aquellos que no dedican toda su potencia a los clientes, sino también pueden jugar el rol de estaciones de trabajo al procesar solicitudes de un usuario local.

En cuanto al tamaño, se pueden clasificar de la siguiente forma:

- **Servidores rack** o armarios metálicos destinados al alojamiento de servidores y con dimensiones de anchura y altura (Unidades Rack o simplemente U) normalizadas. Una U equivale 4,445 cm de alto.
Ofrecen mayor capacidad de administración, consolidación, seguridad, ampliación y modularidad, contribuyendo así a reducir el coste de su implantación.
- **Servidor de torre** consistentes en un chasis o carcasa.
- **Servidores de miniatura**, consistente en equipos comunes reducidos que no necesitan tanta capacidad como un servidor normal. Están destinados a lugares pequeños o negocios con poco flujo de información.



- **Miniservidores rack.** Se trata de un rack para pocos equipos.
- **Servidor blade** es un tipo de computadora para los centros de proceso de datos (CPD). Los servidores blade están diseñados para su montaje en bastidores al igual que otros servidores. La novedad estriba en que pueden compactarse en un espacio más pequeño gracias a sus principios de diseño. Cada servidor blade es una delgada "tarjeta" que contiene únicamente microprocesador, memoria y buses. Es decir, no son directamente utilizables ya que no disponen de fuente de alimentación ni tarjetas de comunicaciones. Específicamente diseñado para aprovechar el espacio, reducir el consumo y simplificar su explotación. La densidad de un servidor blade puede ser seis veces mayor que la de los servidores normales.



Figura 32. Chasis de Servidores Blade.

- **Servidores móviles** accesibles desde cualquier ubicación.
- **Servidores ultra-densos** con grandes especificaciones, cientos de cores y varios TB de memoria RAM.
- **Superservidores** con alto poder de cómputo, diseñados para trabajos pesados. Cuentan con gran capacidad de almacenamiento en disco y con funciones de respaldo y expandibilidad.

Tendencias

Las soluciones en la nube o cloud computing están desplazando a los servidores tradicionales. Los diferentes informes hablan de caídas en la venta de servidores y las empresas, sobre todo las más pequeñas, ya no apuestan en adquirir nuevas máquinas.

El segmento de pequeños servidores (low-end) se está mostrando muy débil, el de servidores medianos (midrange) algo mejor, pero el segmento de grandes servidores, dedicados principalmente a prestar servicios de cloud, está mostrando un crecimiento significativo. Los centros de datos están experimentando un vertiginoso crecimiento. Y es aquí donde se está centrando el mercado de los servidores.

Se está viviendo una época de grandes cambios. Por ejemplo, el denominado Big Data ha incrementado notablemente la necesidad de contar con sistemas optimizados que permitan gestionar grandes volúmenes de información generados por múltiples dispositivos. Por ello, los fabricantes están trabajando por soluciones que sean capaces de integrarse fácilmente en las actuales infraestructuras tecnológicas, que manejan exigentes cargas de trabajo de virtualización y cloud; todo ello con una gestión avanzada del consumo energético a un precio más asequible.

En el mercado servidor, las organizaciones se están centrando en iniciativas de consolidación, migración y virtualización de servidores, con el fin de reducir los costes de infraestructura, incrementar su eficiencia y adoptar el nuevo modelo de servicio cloud.

Los fabricantes se están adaptando a las nuevas tendencias y ofrecen elementos como el consumo por nodo, la tasa de fallos y su integración en soluciones más complejas a nivel de estructura pero no de gestión.

Por otro lado, cada vez cobra más importancia la computación unificada, una innovadora aproximación que propone la unificación de los elementos dispersos del data center, servidores, redes y sistemas de almacenamiento, bajo una misma arquitectura que se apoya en la virtualización y en la gestión unificada.

Fabricantes como Cisco han sido pioneros en la informática unificada, presentando hace pocos años Cisco UCS (Unified Computing System), la primera plataforma de la industria capaz de unificar los servidores, las redes, el acceso al almacenamiento, la virtualización y toda la gestión en un mismo sistema escalable y modular.

En lo que se refiere a las arquitecturas, x86 se está convirtiendo en el estándar de facto en el mercado de servidores, desplazando a los sistemas RISC al poder soportar aplicaciones críticas y ofrecer un mejor precio-rendimiento, mayor flexibilidad y capacidades de unificación y virtualización.

7.5. Computación en la nube

//VIP!!! Al menos lectura. Como basico saber los tipos de servicios que se prestan: los clasicos IaaS, PaaS y SaaS y las nuevas modalidades que teneis en anexos (Edge y fog). Tambien los tipos de nube (publica, privada, hibrida) y soluciones comerciales (tipica pregunta de test)//

7.5.1. Introducción

La computación en la nube (del inglés cloud computing), conocida también como servicios en la nube, informática en la nube, nube de cómputo, nube de conceptos o simplemente "la nube", es un paradigma que permite ofrecer servicios de computación a través de una red, que usualmente es Internet.

En este tipo de computación todo lo que puede ofrecer un sistema informático se ofrece como servicio, de modo que los usuarios puedan acceder a los servicios disponibles "en la nube de Internet" sin conocimientos (o, al menos sin ser expertos) en la gestión de los recursos que usan.

La computación en la nube son servidores desde Internet encargados de atender las peticiones en cualquier momento. Se puede tener acceso a su información o servicio, mediante una conexión a Internet desde cualquier dispositivo móvil o fijo ubicado en cualquier lugar. Sirven a sus usuarios desde varios proveedores de alojamiento repartidos frecuentemente por todo el mundo. Esta medida reduce los costos, garantiza un mejor tiempo de actividad y que los sitios web sean invulnerables a los delincuentes informáticos, a los gobiernos locales y a sus redadas policiales pertenecientes.

Cloud computing es un nuevo modelo de prestación de servicios de negocio y tecnología, que permite incluso al usuario acceder a un catálogo de servicios estandarizados y responder con ellos a las necesidades de su negocio, de forma flexible y adaptativa, en caso de demandas no previsibles o de picos de trabajo, pagando únicamente por el consumo efectuado, o incluso gratuitamente en caso de proveedores que se financian mediante publicidad o de organizaciones sin ánimo de lucro.

El cambio que ofrece la computación desde la nube es que permite aumentar el número de servicios basados en la red. Esto genera beneficios tanto para los proveedores, que pueden ofrecer, de forma más rápida y eficiente, un mayor número de servicios, como para los usuarios que tienen la posibilidad de acceder a ellos, disfrutando de la 'transparencia' e inmediatez del sistema y de un modelo de pago por consumo. Así mismo, el consumidor ahorra los costes salariales o los costes en inversión económica (locales, material especializado, etc.).

Computación en nube consigue aportar estas ventajas, apoyándose sobre una infraestructura tecnológica dinámica que se caracteriza, entre otros factores, por un alto grado de automatización, una rápida movilización de los recursos, una elevada capacidad de adaptación



para atender a una demanda variable, así como virtualización avanzada y un precio flexible en función del consumo realizado, evitando además el uso fraudulento del software y la piratería. El concepto de “nube informática” es muy amplio, y abarca casi todos los posibles tipos de servicio en línea, pero cuando las empresas predican ofrecer un utilitario alojado en la nube, por lo general se refieren a alguna de estas tres modalidades: el software como servicio (por sus siglas en inglés SaaS —Software as a Service—), Plataforma como Servicio (PaaS) e Infraestructura como Servicio (IaaS).

El software como servicio (SaaS) es un modelo de distribución de software en el que las aplicaciones están alojadas por una compañía o proveedor de servicio y puestas a disposición de los usuarios a través de una red, generalmente la Internet. Plataforma como servicio (PaaS) es un conjunto de utilitarios para abastecer al usuario de sistemas operativos y servicios asociados a través de Internet sin necesidad de descargas o instalación alguna. Infraestructura como Servicio (IaaS) se refiere a la tercerización de los equipos utilizados para apoyar las operaciones, incluido el almacenamiento, hardware, servidores y componentes de red.

7.5.2. Historia

El concepto de la computación en la nube empezó en proveedores de servicio de Internet a gran escala, como Google (Google Cloud Services), Amazon AWS (2006), Microsoft (Microsoft Azure) o Alibaba Cloud y otros que construyeron su propia infraestructura. De entre todos ellos emergió una arquitectura: un sistema de recursos distribuidos horizontalmente, introducidos como servicios virtuales de TI escalados masivamente y manejados como recursos configurados y mancomunados de manera continua.

El concepto fundamental de la entrega de los recursos informáticos a través de una red global tiene sus raíces en los años sesenta. La idea de una “red de computadoras intergaláctica” la introdujo en los años sesenta JCR Licklider, cuya visión era que todo el mundo pudiese estar interconectado y poder acceder a los programas y datos desde cualquier lugar, según Margaret Lewis, directora de mercadotecnia de producto de AMD. “Es una visión que se parece mucho a lo que llamamos *cloud computing*”.

Otros expertos atribuyen el concepto científico de la computación en nube a John McCarthy, quien propuso la idea de la computación como un servicio público, de forma similar a las empresas de servicios que se remontan a los años sesenta. En 1960 dijo: “Algún día la computación podrá ser organizada como un servicio público”.

Desde los años sesenta, la computación en nube se ha desarrollado a lo largo de una serie de líneas. La Web 2.0 es la evolución más reciente. Sin embargo, como Internet no empezó a ofrecer ancho de banda significativo hasta los años noventa, la computación en la nube ha sufrido algo así como un desarrollo tardío. Uno de los primeros hitos de la computación en nube es la llegada de Salesforce.com en 1999, que fue pionero en el concepto de la entrega de aplicaciones empresariales a través de una página web simple. La firma de servicios allanó el camino para que tanto especialistas como empresas tradicionales de software pudiesen publicar sus aplicaciones a través de Internet.

El siguiente desarrollo fue Amazon Web Services en 2002, que prevé un conjunto de servicios basados en la nube, incluyendo almacenamiento, computación e incluso la inteligencia humana a través del Amazon Mechanical Turk. Posteriormente en 2006, Amazon lanzó su Elastic Compute Cloud (EC2) como un servicio comercial que permite a las pequeñas empresas y los particulares alquilar equipos en los que se ejecuten sus propias aplicaciones informáticas. “Amazon EC2/S3 fue el que ofreció primero servicios de infraestructura en la nube totalmente accesibles”, según Jeremy Allaire, CEO de Brightcove, que proporciona su plataforma SaaS de vídeo en línea a las estaciones de televisión de Reino Unido y periódicos.

George Gilder dijo en 2006: “El PC de escritorio está muerto. Bienvenido a la nube de Internet, donde un número enorme de instalaciones en todo el planeta almacenarán todos los datos que usted podrá usar alguna vez en su vida”.

Otro hito importante se produjo en 2009, cuando Google y otros empezaron a ofrecer aplicaciones basadas en navegador. "La contribución más importante a la computación en nube ha sido la aparición de 'aplicaciones asesinas' de los gigantes de tecnología como Microsoft y Google. Cuando dichas compañías llevan a cabo sus servicios de una manera que resulta segura y sencilla para el consumidor, el efecto 'pasar la pelota' en sí crea un sentimiento de mayor aceptación de los servicios online", según Dan Germain, jefe de la oficina de tecnología en IT proveedor de servicios Cobweb Solutions.

Otro de los factores clave que han permitido evolucionar a la computación en la nube han sido, según el pionero en computación en la nube británico Jamie Turner, las tecnologías de virtualización, el desarrollo del universal de alta velocidad de ancho de banda y normas universales de interoperabilidad de software. Turner añadió: "A medida que la computación en nube se extiende, su alcance va más allá de un puñado de usuarios de Google Docs. Solo podemos empezar a imaginar su ámbito de aplicación y alcance. Casi cualquier cosa puede ser utilizado en la nube".

7.5.3. Características, ventajas e inconvenientes

Características

La computación en nube presenta las siguientes características clave:

- **Agilidad:** Capacidad de mejora para ofrecer recursos tecnológicos al usuario por parte del proveedor.
- **Costo:** los proveedores de computación en la nube afirman que los costos se reducen. Un modelo de prestación pública en la nube convierte los gastos de capital en gastos de funcionamiento. Ello reduce barreras de entrada, ya que la infraestructura se proporciona típicamente por una tercera parte y no tiene que ser adquirida por una sola vez o tareas informáticas intensivas infrecuentes.
- **Escalabilidad y elasticidad:** aprovisionamiento de recursos sobre una base de autoservicio casi en tiempo real, sin que los usuarios necesiten cargas de alta duración.
- **Independencia entre el dispositivo y la ubicación:** permite a los usuarios acceder a los sistemas utilizando un navegador web, independientemente de su ubicación o del dispositivo que utilice (por ejemplo, PC, teléfono móvil).
- La tecnología de virtualización permite compartir servidores y dispositivos de almacenamiento y una mayor utilización. Las aplicaciones pueden ser fácilmente migradas de un servidor físico a otro.
- **Rendimiento:** Los sistemas en la nube controlan y optimizan el uso de los recursos de manera automática, dicha característica permite un seguimiento, control y notificación del mismo. Esta capacidad aporta transparencia tanto para el consumidor o el proveedor de servicio.
- **Seguridad:** puede mejorar debido a la centralización de los datos. La seguridad es a menudo tan buena o mejor que otros sistemas tradicionales, en parte porque los proveedores son capaces de dedicar recursos a la solución de los problemas de seguridad que muchos clientes no pueden permitirse el lujo de abordar. El usuario de la nube es responsable de la seguridad a nivel de aplicación. El proveedor de la nube es responsable de la seguridad física.
- **Mantenimiento:** en el caso de las aplicaciones de computación en la nube, es más sencillo, ya que no necesitan ser instalados en el ordenador de cada usuario y se puede acceder desde diferentes lugares.

Ventajas

Las principales ventajas de la computación en la nube son:

- Integración probada de servicios Red. Por su naturaleza, la tecnología de *cloud computing* se puede integrar con mucha mayor facilidad y rapidez con el resto de las



aplicaciones empresariales (tanto software tradicional como Cloud Computing basado en infraestructuras), ya sean desarrolladas de manera interna o externa.

- Prestación de servicios a nivel mundial. Las infraestructuras de *cloud computing* proporcionan mayor capacidad de adaptación, recuperación completa de pérdida de datos (con copias de seguridad) y reducción al mínimo de los tiempos de inactividad.
- Una infraestructura 100% de *cloud computing* permite también al proveedor de contenidos o servicios en la nube prescindir de instalar cualquier tipo de software, ya que este es provisto por el proveedor de la infraestructura o la plataforma en la nube. Un gran beneficio del *cloud computing* es la simplicidad y el hecho de que requiera mucha menor inversión para empezar a trabajar.
- Implementación más rápida y con menos riesgos, ya que se comienza a trabajar más rápido y no es necesaria una gran inversión. Las aplicaciones del *cloud computing* suelen estar disponibles en cuestión de días u horas en lugar de semanas o meses, incluso con un nivel considerable de personalización o integración.
- Actualizaciones automáticas que no afectan negativamente a los recursos de TI. Al actualizar a la última versión de las aplicaciones, el usuario se ve obligado a dedicar tiempo y recursos para volver a personalizar e integrar la aplicación. Con el *cloud computing* no hay que decidir entre actualizar y conservar el trabajo, dado que esas personalizaciones e integraciones se conservan automáticamente durante la actualización.
- Contribuye al uso eficiente de la energía. En este caso, a la energía requerida para el funcionamiento de la infraestructura. En los datacenters tradicionales, los servidores consumen mucha más energía de la requerida realmente. En cambio, en las nubes, la energía consumida es solo la necesaria, reduciendo notablemente el desperdicio.

Inconvenientes

- La centralización de las aplicaciones y el almacenamiento de los datos origina una interdependencia de los proveedores de servicios.
- La disponibilidad de las aplicaciones está sujeta a la disponibilidad de acceso a Internet.
- La confiabilidad de los servicios depende de la "salud" tecnológica y financiera de los proveedores de servicios en nube. Empresas emergentes o alianzas entre empresas podrían crear un ambiente propicio para el monopolio y el crecimiento exagerado en los servicios.
- La disponibilidad de servicios altamente especializados podría tardar meses o incluso años para que sean factibles de ser desplegados en la red.
- La madurez funcional de las aplicaciones hace que continuamente estén modificando sus interfaces, por lo cual la curva de aprendizaje en empresas de orientación no tecnológica tenga unas pendientes significativas, así como su consumo automático por aplicaciones.
- Seguridad. La información de la empresa debe recorrer diferentes nodos para llegar a su destino, cada uno de ellos (y sus canales) son un foco de inseguridad. Si se utilizan protocolos seguros, HTTPS por ejemplo, la velocidad total disminuye debido a la sobrecarga que estos requieren.
- Escalabilidad a largo plazo. A medida que más usuarios empiecen a compartir la infraestructura de la nube, la sobrecarga en los servidores de los proveedores aumentará, si la empresa no posee un esquema de crecimiento óptimo puede llevar a degradaciones en el servicio o altos niveles de jitter.

7.5.4. Servicios ofrecidos

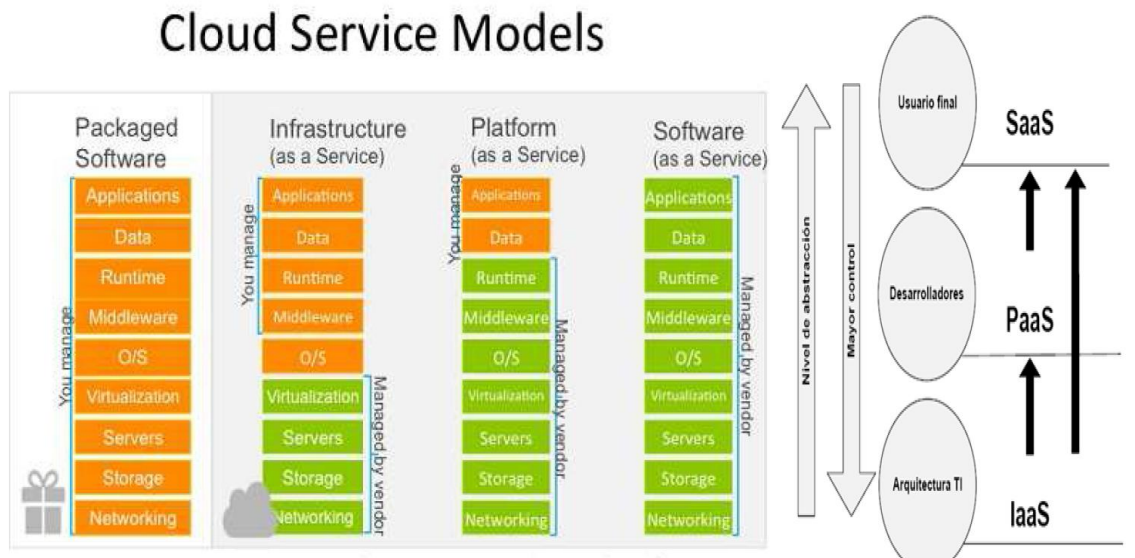


Figura 33. Modelos de servicio en la Nube.

Software como servicio

El **software como servicio** (en inglés *software as a service*, SaaS) se encuentra en la capa más alta y caracteriza una aplicación completa ofrecida como un servicio, por-demanda, vía multitenencia —que significa una sola instancia del software que corre en la infraestructura del proveedor y sirve a múltiples organizaciones de clientes—. Las aplicaciones que suministran este modelo de servicio son accesibles a través de un navegador web —o de cualquier aplicación diseñada para tal efecto— y el usuario no tiene control sobre ellas, aunque en algunos casos se le permite realizar algunas configuraciones. Esto le elimina la necesidad al cliente de instalar la aplicación en sus propios computadores, evitando asumir los costos de soporte y el mantenimiento de hardware y software.

Plataforma como servicio

La capa del medio, que es la **plataforma como servicio** (en inglés *platform as a service*, PaaS), es la **encapsulación** de una abstracción de un ambiente de desarrollo y el empaquetamiento de una serie de módulos o complementos que proporcionan, normalmente, una funcionalidad horizontal (persistencia de datos, autenticación, mensajería, etc.). De esta forma, un arquetipo de plataforma como servicio podría consistir en un entorno conteniendo una pila básica de sistemas, componentes o APIs preconfiguradas y listas para integrarse sobre una tecnología concreta de desarrollo (por ejemplo, un sistema Linux, un servidor web, y un ambiente de programación como Perl o Ruby). Las ofertas de PaaS pueden dar servicio a todas las fases del ciclo de desarrollo y pruebas del software, o pueden estar especializadas en cualquier área en particular, tal como la administración del contenido.

Ejemplos comerciales son [Google App Engine](#), que sirve aplicaciones de la infraestructura [Google](#); [Microsoft Azure](#), una plataforma en la nube que permite el desarrollo y ejecución de aplicaciones codificadas en varios lenguajes y tecnologías como [.NET](#), [Java](#) y [PHP](#) o la [Plataforma G](#), desarrollada en [Perl](#). Servicios PaaS como estos permiten gran flexibilidad, pero puede ser restringida por las capacidades disponibles a través del proveedor.

En este modelo de servicio al usuario se le ofrece la plataforma de desarrollo y las herramientas de programación por lo que puede desarrollar aplicaciones propias y controlar la aplicación, pero no controla la infraestructura.

Infraestructura como servicio

La **infraestructura como servicio** (*infrastructure as a service*, IaaS) —también llamada en algunos casos *hardware as a service*, HaaS)— se encuentra en la capa inferior y es un medio de entregar almacenamiento básico y capacidades de cómputo como servicios estandarizados en la red. Servidores, sistemas de almacenamiento, conexiones, enrutadores, y otros sistemas se concentran (por ejemplo a través de la tecnología de virtualización) para manejar tipos específicos de cargas de trabajo —desde procesamiento en lotes (“batch”) hasta aumento de servidor/almacenamiento durante las cargas pico—. El ejemplo comercial mejor conocido es Amazon Web Services, cuyos servicios EC2 y S3 ofrecen cómputo y servicios de almacenamiento esenciales (respectivamente). Otro ejemplo es Joyent, cuyo producto principal es una línea de servidores virtualizados, que proveen una infraestructura en demanda altamente escalable para manejar sitios web, incluidas aplicaciones web complejas escritas en Python, Ruby, PHP y Java.

Otros:

- **SecaaS:** Security as a Service. Modelo de computación en la nube que ofrece servicios de seguridad gestionada a través de Internet. Está basado en el SaaS, pero limitado a los servicios especializados de seguridad de la información. (audiores, forenses, actuación frente a ataques, antivirus...)
- **NaaS:** Network as a Service. La capacidad proporcionada al usuario es utilizar la red/servicios de conectividad de transporte y/o servicios de conectividad de red entre las nubes. Los modelos de NaaS más comunes son:
 - VPN
 - BOD: Ancho de Banda bajo demanda. El ancho de banda se asigna en función de las necesidades de nodos o usuarios.
 - Red móvil Virtual. Operador de red móvil virtual (MVNO)

7.5.5. Tipos de nubes

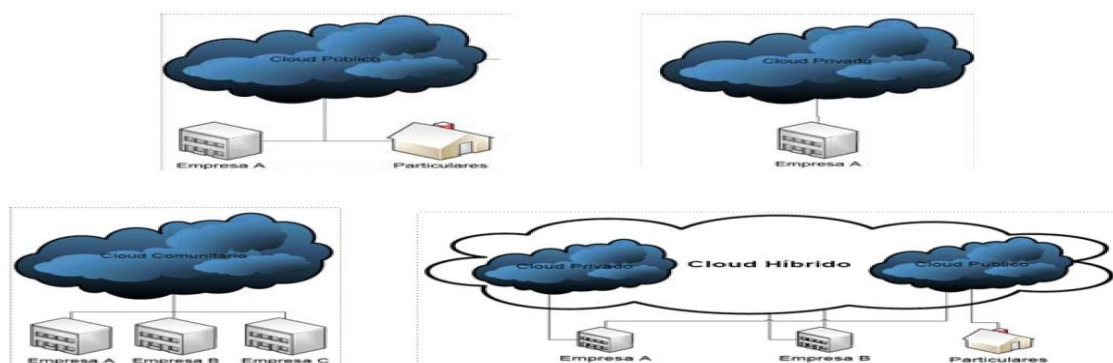


Figura 34. Tipos de Nube.

1. Una **nube pública** es una nube computacional mantenida y gestionada por terceras personas no vinculadas con la organización. En este tipo de nubes tanto los datos como los procesos de varios clientes se mezclan en los servidores, sistemas de almacenamiento y otras infraestructuras de la nube. Los usuarios finales de la nube no conocen qué trabajos de otros clientes pueden estar corriendo en el mismo servidor, red, sistemas de almacenamiento, etc. Aplicaciones, almacenamiento y otros recursos están disponibles al público a través del proveedor de servicios, que es propietario de toda la infraestructura en sus centros de datos; el acceso a los servicios solo se ofrece de manera remota, normalmente a través de internet.
Dentro de estas nubes podemos encontrar englobadas las Cloud Privado virtual, consideradas como nubes de dominio público pero que mejoran la seguridad de los datos. Estos datos se encriptan a través de la implantación de una VPN. Algunas grandes empresas como Amazon ya permiten aprovisionar una sección de su AWS con esta posibilidad VPC.
2. Las **nubes privadas** son una buena opción para las compañías que necesitan alta protección de datos y ediciones a nivel de servicio. Las nubes privadas están en una infraestructura bajo demanda, gestionada para un solo cliente que controla qué aplicaciones debe ejecutarse y dónde. Son propietarios del servidor, red, y disco y pueden decidir qué usuarios están autorizados a utilizar la infraestructura. Al administrar internamente estos servicios, las empresas tienen la ventaja de mantener la privacidad de su información y permitir unificar el acceso a las aplicaciones corporativas de sus usuarios.
3. Las **nubes híbridas** combinan los modelos de nubes públicas y privadas. Un usuario es propietario de unas partes y comparte otras, aunque de una manera controlada. Las nubes híbridas ofrecen la promesa del escalado, aprovisionada externamente, a demanda, pero añaden la complejidad de determinar cómo distribuir las aplicaciones a través de estos ambientes diferentes. Las empresas pueden sentir cierta atracción por la promesa de una nube híbrida, pero esta opción, al menos inicialmente, estará probablemente reservada a aplicaciones simples sin condicionantes, que no requieran de ninguna sincronización o necesiten bases de datos complejas. Se unen mediante la tecnología, pues permiten enviar datos o aplicaciones entre ellas. Un ejemplo son los sistemas de correo electrónico empresarial.
4. **Nube comunitaria**. De acuerdo con Joyanes Aguilar en 2012, el Instituto Nacional de Estándares y Tecnología (NITS, por sus siglas en inglés) define este modelo como aquel que se organiza con la finalidad de servir a una función o propósito común (seguridad, política...), las cuales son administradas por las organizaciones constituyentes o terceras partes.

7.5.6. Implementaciones

1. Amazon Web Services (AWS)

Este proveedor permite que sus usuarios creen una imagen de máquina virtual de Amazon (AMI), esto es, una máquina virtual con el sistema operativo Windows o Linux, en la que el usuario instala sus aplicaciones, librerías y datos que necesite. Posteriormente, Amazon ejecuta esa máquina en sus sistemas, y le asigna características físicas de acuerdo al contrato suscrito con el usuario.
Coste: se factura por hora de utilización y tipo de recursos asignados a cada máquina física.

Tipos de servicios de Amazon:

- EC2 (Amazon Elastic Compute Cloud): Servicio web que proporciona capacidad informática con tamaño modificable en la nube. Se ha diseñado con el fin de facilitar el trabajo de desarrolladores. Amazon EC2 reduce el tiempo necesario para obtener e iniciar nuevas instancias de servidor a cuestión de minutos, lo que permite escalar con rapidez su capacidad cuando cambien los requisitos informáticos.
Características:
 - Elástico
 - Control Total
 - Flexible: trabaja con Amazon S3, Amazon SimpleDB, Amazon Simple Queue Service (Amazon SQS).
 - Fiable
 - Seguro
- S3 (Amazon Simple Storage Service): Servicio de almacenamiento en la nube. Está creado intencionadamente con un conjunto de funciones mínimas. Los objetos almacenados en una Región nunca abandonan la misma (ej, objetos almacenados en la Región UE (Irlanda) nunca salen de la UE). Se incluyen mecanismos de autenticación. Utiliza interfaces REST y SOAP basadas en estándares. Está creado para ser flexible y permitir añadir fácilmente protocolos o capas funcionales.
El protocolo de descarga predeterminado es HTTP. Se proporciona un protocolo BitTorrent para reducir los costes de la distribución a gran escala.

2. Microsoft Windows Azure

Ofrece la creación de aplicaciones Web adaptadas a sus sistemas y su despliegue en los mismos con ciertas limitaciones de consumo. Admite varios lenguajes de programación. Windows Azure Platform ofrece a los desarrolladores un entorno para crear y ejecutar sus aplicaciones en los centros del proveedor. Dicho entorno proporciona las funcionalidades necesarias para que las aplicaciones creadas con él puedan realizar tareas de negocio, almacenar en Bds de la nube y comunicarse con otras apps. Permite integración con herramientas y aplicaciones ya existentes como Microsoft Office, Exchange (en modalidad SaaS) y CRM Microsoft Dynamics.

Servicios que ofrece Azure:

- WindowsAzure: S.O. Como un servicio en línea
- MicrosoftSQLAzure: solución completa de BD Cloud relacional
- Plataforma AppFabric de Windows Azure: conecta servicios cloud y aplicaciones internas.

Otras aplicaciones:

- Box - desarrollado por Box Inc.
- Campaign Cloud - desarrollado por ElectionMall Technologies (Cerrado)
- Doit!e ajaxplorer - desarrollado por Doit!e
- Dropbox - desarrollado por Dropbox
- Google Drive - desarrollado por Google
- Flokzu BPM en la nube - Desarrollado por Flokzu.
- iCloud - desarrollado por Apple
- OneDrive - desarrollado por Microsoft (antes SkyDrive)
- OwnCloud - desarrollado por OwnCloud Inc.
- Salesforce.com - desarrollado por Salesforce.com Inc.
- Scaleway - desarrollado por Online.net
- SugarSync - desarrollado por SugarSync
- Ubuntu One - desarrollado por Canonical (cerrado)
- Wuala - desarrollado por LaCie
- Datapius - desarrollado por Datapius
- NetQuatro Cloud - desarrollado por NetQuatro



7.5.7 Cloud Privado Virtual

Un **Cloud Privado Virtual (CPV)** es un conjunto de recursos computacionales configurables por demanda al interior de un ambiente de cloud público, el cual provee un cierto nivel de aislamiento entre las diferentes organizaciones (llamados *usuarios*) que utilizan dichos recursos. El aislamiento entre un CPV y los demás usuarios del cloud (tanto otros CPV como usuarios del cloud público) se logra normalmente a través de la utilización de una subred IP y un mecanismo de comunicación virtual (como una VLAN o un grupo de canales encriptados) para cada usuario. En un CPV el mecanismo descrito anteriormente para proveer aislamiento dentro del cloud, se complementa con un servicio de VPN (para cada usuario) el cual garantiza la seguridad del acceso remoto de un usuario a sus recursos a través de un sistema de autenticación y criptografía.

A través de estos mecanismos de aislamiento, la organización que utiliza el servicio está de hecho trabajando en un cloud '**virtualmente privado**' (es decir, como si la infraestructura no estuviese siendo compartida con otros usuarios), de allí el nombre CPV.

Los CPV se usan comúnmente en el contexto de plataforma como servicio. En dicho contexto, el proveedor de la infraestructura (cloud) y el proveedor del servicio de CPV sobre dicha infraestructura podrían ser compañías diferentes.

Implementaciones

Amazon Web Services lanzó su servicio en:Amazon Virtual Private Cloud el 26 de agosto de 2009, el cual permite que puedan crearse conexiones entre en:Amazon Elastic Compute Cloud y otras redes a través de una red privada virtual sobre protocolo IPsec.

En AWS, el CPV puede ser usado de manera gratuita, sin embargo, se cobra a los usuarios por las redes VPN que usen. Además hay algunas limitaciones en la cantidad de recursos que pueden garantizarse cuando se trabaja con Instancias reservadas.

Google App Engine soporta una funcionalidad similar a través de su producto *Secure Data Connector* lanzado el 7 de abril de 2009. Google dio de baja este servicio el 14 de marzo de 2013 y no acepta nuevos usuarios. Se espera que el servicio continúe funcionando para los usuarios actuales hasta (al menos) Abril de 2015.

HP ofrece un servicio Empresarial de Cloud que incluye cloud privado, cloud administrado y servicios de cloud público basados en OpenStack.

Microsoft Azure ofrece la posibilidad de crear un CPV utilizando redes virtuales.

Cloud-Bricks.net es un ejemplo de un proveedor de Cloud Privado sobre Hardware Dedicado en donde no se comparten recursos con otros usuarios.

FortyCloud es un ejemplo de un CPV que es ofrecido sobre la infraestructura de cloud público de un tercero como AWS EC2 y sobre ambientes híbridos de múltiples proveedores.

Host Virtual es un proveedor de infraestructura como servicio que incorpora CPV como una de sus características.

Existen también CPVs regionales como Cloud-A y cloud.ca, Plataformas Canadienses de cloud privado virtual.

7.5.8. Nuevas tecnologías en la nube

Edge Computing

Los millones de dispositivos de esa Internet de las Cosas que nos rodea tienen un problema: recolectan información, pero no hacen nada con ella. La envían a la nube, donde grandes centros de datos la procesan para obtener ciertas conclusiones o activar ciertos eventos.

Ese funcionamiento "pasivo" de todos esos dispositivos es lo que quiere cambiar la llamada Edge Computing, un tipo de filosofía aplicable especialmente en escenarios empresariales e industriales que aporta mucho más autonomía a todos esos dispositivos, haciendo que sean algo más "listos".

Hasta ahora en la mayoría de los casos las grandes plataformas de Cloud Computing se encargaban de hacer ese "trabajo sucio" de analizar los datos recolectados por los sensores y dispositivos IoT.

La eficiencia de este paradigma no es óptima en muchos casos en los que los propios nodos de la red pueden analizar esos datos para evitar ese paso por la nube.

Edge Computing "permite que los datos producidos por los dispositivos de la internet de las cosas se procesen más cerca de donde se crearon en lugar de enviarlos a través de largas recorridos para que lleguen a centros de datos y nubes de computación".

Eso tiene una ventaja fundamental, ya que "permite a las organizaciones analizar los datos importantes casi en tiempo real, algo que es una necesidad patente en muchas industrias tales como la fabricación, la salud, las telecomunicaciones o la industria financiera".

Esa definición es precisamente la que explica esa nueva tendencia que hace que esos dispositivos y sensores situados por todas partes se ocupen no solo de recolectar esos datos para enviarlos a la nube, sino de procesarlos directamente. Las aplicaciones industriales en este ámbito son diversas, y ese planteamiento puede hacer que efectivamente la mejora en muchos procesos sea sensible.

Lo cierto es que muchos de esos dispositivos conectados tendrán capacidad no solo de recolección, sino de proceso de datos. Y si no lo tienen ellos, lo tendrán nodos en esas redes empresariales o industriales. Es allí donde la filosofía Edge Computing toma forma, y donde entre otras cosas se producirá la mejora en eficiencia: no tener que transmitir todos esos datos a la nube ya supondrá un ahorro importante.

Fog computing

Hay otro término muy relacionado con Edge Computing que está usándose cada vez más en este ámbito, y es el de la llamada Fog Computing. Como revelaba un estudio de Cisco, esta plataforma permite *"extender la nube para que esté más cerca de las cosas que producen y se accionan mediante datos de dispositivos IoT"*. Los responsables del estudio añadían que "cualquier dispositivo con conectividad de red, capacidad de computación y almacenamiento puede ser un nodo de esa "niebla".

Esta filosofía podría decirse que permite que los grandes centros de datos de la nube "deleguen" parte de sus responsabilidades a dispositivos Edge Computing, y lo hagan a través de esa Fog Computing que define requisitos o necesidades en ese extremo de todo este ecosistema que como decimos tiene aplicaciones industriales claras.

La Edge Computing se refiere de forma específica a cómo los procesos computacionales se realizan en los "dispositivos edge", los dispositivos IoT con capacidad de análisis y procesos como routers o gateways de red, por ejemplo.

A diferencia de ese concepto, la Fog Computing se refiere a las conexiones de red entre los dispositivos edge y la nube.

Beneficios presentes y futuros

Hay además factores que harán que este tipo de paradigma lo tenga aún más fácil en el futuro: el coste cada vez más reducido de los dispositivos y los sensores se une a la cada vez mayor potencia que tienen incluso dispositivos modestos.

También hay necesidades industriales que contribuyen a apostar por el Edge Computing: en ciertos entornos la única forma de optimizar aún más los procesos es la de tratar de evitar la comunicación con la nube en la medida de lo posible. Eso permite reducir latencias, consumir

menos anchos de banda —no es necesario enviar todos los datos a la nube en todo momento— y acceder de forma inmediata a análisis y evaluación del estado de todos esos sensores y dispositivos.

El Edge Computing Consortium se centra precisamente en tratar de impulsar esta disciplina para que se consigan ventajas en forma de reducción de costes en mantenimiento e implantación, garantías en materia de seguridad, o menores consumos energéticos.

Hay otra ventaja interesante: la seguridad. Cuantos menos datos hay en un entorno cloud, menos vulnerable es ese entorno si se ve comprometido. Si la seguridad en esos "micro centros de datos" Edge Computing se cuida de la forma adecuada, este apartado podría ganar muchos enteros.

Eso no significa que la dependencia de la nube y de entornos Cloud Computing desaparezca: ambas tendencias deben aportar, y por ejemplo Edge Computing es más adecuado cuando sobre todo se necesita velocidad y baja latencia en esas transferencias de datos, mientras que la nube seguirá siendo protagonista para analizar y tratar grandes cantidades de datos que requieren una potencia de cálculo notable.

Los coches autónomos, el ejemplo perfecto

Si hay un campo en el que este tipo de filosofía tenga sentido, ese es el del coche autónomo. Estos "centros de datos sobre ruedas" no paran de recolectar información sobre sus sistemas y su entorno, y toda esa información debe ser procesada en tiempo real para que podamos disfrutar de una conducción autónoma óptima y segura.

Intel estima que un coche autónomo podría acabar generando 4 TB de datos al día: solo las cámaras del coche se encargarán de transferir al sistema entre 20 y 40 Mbits por segundo, a los que se sumarían entre otros 10 y 100 Kbits por segundo del radar. Esos son solo parte de los datos que tendrá que gestionar el coche para poder funcionar correctamente.

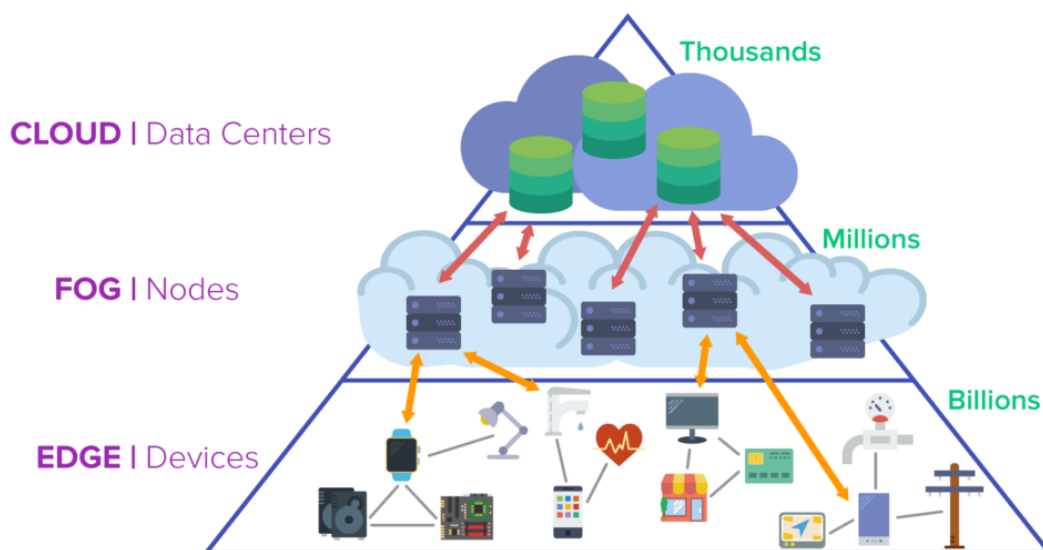
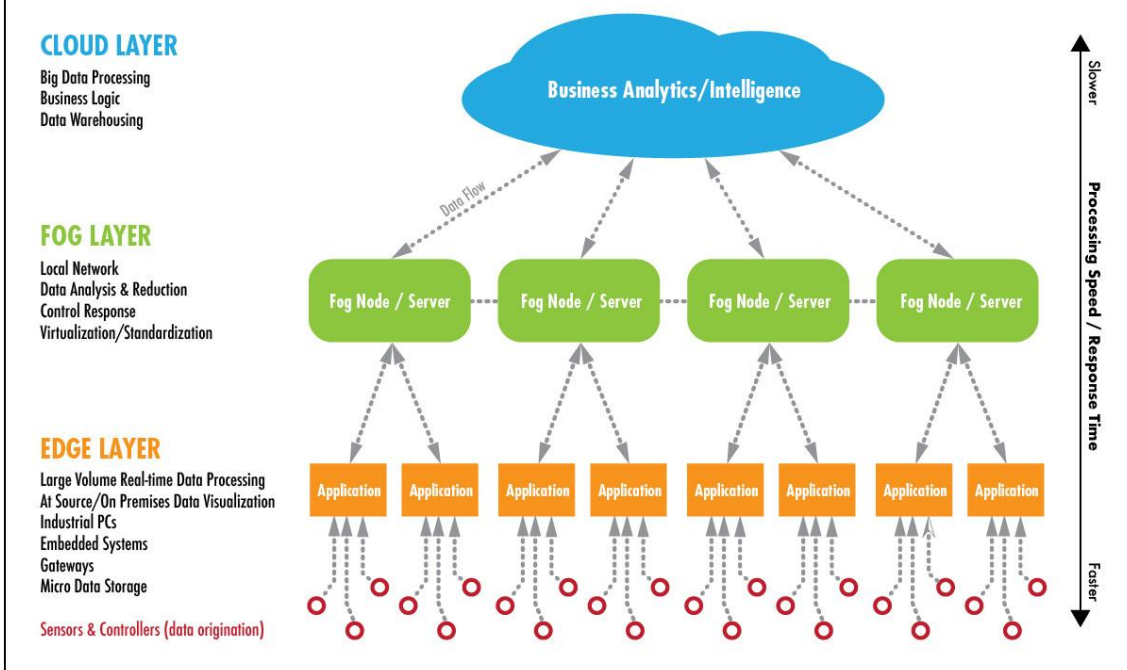
El coche autónomo no puede estar esperando a comunicarse con la nube y a esperar la respuesta: todo ese proceso y análisis de datos hay que hacerlo en tiempo real, y es ahí donde la Edge Computing entra en juego, confirmando el importante papel que el ordenador central del coche tiene para aglutinar, analizar y dar respuesta a las necesidades de la conducción autónoma en cada momento.

Esas necesidades en tiempo real se suman a otras igualmente interesantes que como decíamos hacen que no se pueda descartar la relevancia de la nube: todo lo "aprendido" por un coche autónomo se puede transferir a la nube para que el resto de modelos se beneficien de la experiencia y de la reacción más adecuada ante distintos eventos.

La idea final de esta tendencia cada vez más popular es en esencia bastante lógica: aprovechar los recursos en cada caso (en los centros de datos de la nube o en tu sistema repleto de sensores y componentes IoT) para optimizar esa capacidad de proceso y esa transferencia de datos. Ahora queda por ver si ese paradigma acaba imponiéndose.

<https://aunclicdelastic.blogthinkbig.com/cloud-edge-y-fog-computing/>

INDUSTRIAL IoT DATA PROCESSING LAYER STACK



Hyperscale Computing

Hyperscale se refiere a la combinación completa de hardware e instalaciones que puede escalar un entorno informático distribuido hasta miles de servidores.

La computación a hiperescala se usa generalmente en entornos como big data y computación en la nube. Además, generalmente está conectado a plataformas como Apache Hadoop.

El diseño estructural de la computación a hiperescala es a menudo diferente de la computación convencional. En el diseño de hiperescala, las construcciones informáticas de alto grado, como las que se encuentran normalmente en los sistemas blade, generalmente se abandonan. Hyperscale favorece el diseño de productos simplificados que es extremadamente rentable. Este nivel mínimo de inversión en hardware facilita la financiación de los requisitos de software del sistema.

La arquitectura informática de hiperescala está disponible en forma de una sola unidad, que



hace uso de redes convergentes, una combinación de almacenamiento conectado a la red y local, o como software de administración incorporado en un factor de forma modesto.

Containers as a Service (CaaS)

Es un modelo de servicio en la nube que permite a los usuarios administrar e implementar contenedores, aplicaciones y clústeres a través de la virtualización basada en contenedores. CaaS es muy útil para los departamentos de TI y los desarrolladores en la creación de aplicaciones en contenedores seguros y escalables. Con CaaS, esto se puede lograr utilizando centros de datos locales o en la nube.

En pocas palabras, CaaS proporciona una manera fácil de configurar un clúster contenedor. La orquestación, que esencialmente automatiza funciones clave de TI, es una calidad esencial de la tecnología CaaS. Google Kubernetes, Docker Swarm, Rackspace Carina y Apache Mesos son ejemplos de plataformas de orquestación CaaS. Algunos proveedores de CaaS en la nube pública incluyen Google, Amazon Web Services (AWS), Rackspace e IBM.

A menudo, se considera que CaaS es un subconjunto de IaaS (infraestructura como servicio), pero incluye los contenedores como su recurso fundamental, a diferencia de los sistemas básicos y las máquinas virtuales.

<https://www.suse.com/es-es/products/caas-platform/>

Function as a Service (FaaS) o Serverless

La función como servicio (FaaS) es una categoría de servicios de computación en la nube que proporciona una plataforma que permite a los clientes desarrollar, ejecutar y administrar las funcionalidades de la aplicación sin la complejidad de construir y mantener la infraestructura asociada típicamente al desarrollo y lanzamiento de una aplicación. La creación de una aplicación siguiendo este modelo es una forma de lograr una arquitectura "sin servidor", y se utiliza normalmente al crear aplicaciones de microservicios.

FaaS es un desarrollo extremadamente reciente en la computación en la nube, que se puso a disposición del mundo por hook.io en octubre de 2014, seguido de AWS Lambda, Google Cloud Functions, Microsoft Azure Functions, IBM / Apache's OpenWhisk* (código abierto) en 2016 y Oracle Cloud Fn (código abierto) en 2017, que están disponibles para uso público. Las capacidades de FaaS también existen en plataformas privadas, como lo demuestran los triggers Schemaless de Uber.

*Apache OpenWhisk (Incubating) is an open source, distributed Serverless platform that executes functions (fx) in response to events at any scale.

Diferencias con Paas:

Las arquitecturas de computación sin servidor en las que el cliente no tiene necesidad directa de administrar recursos también pueden lograrse utilizando los servicios de Plataforma como Servicio. Sin embargo, estos servicios suelen ser muy diferentes en su arquitectura de implementación, lo que tiene algunas implicaciones para el escalado. En la mayoría de los sistemas PaaS, el sistema ejecuta continuamente al menos un proceso de servidor y, incluso con el escalado automático, simplemente se agregan o eliminan una cantidad de procesos de ejecución más largos en la misma máquina. Esto significa que la escalabilidad es un problema más visible para el desarrollador.

En un sistema FaaS, se espera que las funciones se inicien en milisegundos para permitir el manejo de solicitudes individuales. En un sistema PaaS, por el contrario, normalmente hay un

subproceso de aplicación que se ejecuta durante un largo período de tiempo y maneja múltiples solicitudes. Esta diferencia es principalmente visible en el precio, donde los servicios FaaS se cobran por el tiempo de ejecución de la función, mientras que los servicios PaaS se cobran por el tiempo de ejecución del subproceso en el que se ejecuta la aplicación del servidor.

Enlace muy interesante para comprender estos conceptos:

<https://chernando.xyz/articulos/que-es-serverless-baas-y-faas/>

7.6. Thin Client

//vip!//

Un cliente liviano o cliente delgado (thin client) es una computadora cliente o un software de cliente en una arquitectura de red cliente-servidor que depende primariamente del servidor central para las tareas de procesamiento, y se enfoca principalmente en transportar la entrada y la salida entre el usuario y el servidor remoto. En contraste, un cliente pesado realiza tanto procesamiento como sea posible y transmite solamente los datos para las comunicaciones y el almacenamiento al servidor.

Los beneficios que representan los **Thin Client** para las organizaciones empresariales ha potenciado su demanda. Uno de las principales ventajas es el ahorro económico. Como tienen externalizada la computación en servidores tienen un consumo energético muy bajo y no hay que reponerlos en mucho tiempo, ya que no tienen que estar equipados con la última tecnología. También pueden evitar pérdidas de datos importantes para el correcto desarrollo de sus negocios.

Los **Thin Clients** son dispositivos que pesan muy poco, ya que no cuentan con apenas potencia de procesamiento. La computación la pueden realizar en servidores a través de sistemas de virtualización de cloud-computing. Normalmente, son ordenadores que se benefician de una infraestructura con un sistema de virtualización de escritorios **VDI**. Depende del tipo de **Thin Client**.

Tipos:

1. Basic ThinClient o Zero Client

Un terminal **Thin Client** de la tipología **Zero Client o Basic Thin Client** es un dispositivo que se conecta a un servidor de escritorio remoto o PC compartido. No tienen sistema operativo ni tampoco almacenamiento en local. Todas las aplicaciones son aprovisionadas por el servidor para arrancar más rápido y reducir el mantenimiento del mismo y la computación.

Estos tipos de dispositivos **Thin Client Zero Client** al no tener disco duro ni procesamiento tienen un nivel más alto de seguridad y consumen mucha menos energía que un ordenador con un escritorio físico. Incluso en las configuraciones Zero Client que se personalizan con un mínimo sistema operativo y disco duro las ventajas en la eficiencia energética y en la prevención de posibles infecciones se mantienen.

2. Browser ThinClient

Los terminales ThinClient se conectan a los servidores web para acceder a las aplicaciones basadas en el navegador. Estos se llaman también Browser Thin Clients y se benefician de todas las ventajas que los Zero client o clientes básicos, pero cuentan con funciones adicionales y más energía para arrancar procesos avanzados.

Para la eficiencia y el correcto comportamiento del navegador local, los ThinClient basados en el navegador o también denominados Cloud Client prosperan donde los trabajos basados en tareas o transacciones necesitan acceso web o aplicaciones basadas en web. Hay determinados negocios como los Call Centers o relacionadas con la atención al cliente que necesitan una solución de este tipo, ya que los usuarios tienen que acceder a información basada en web además de los datos del servidor central.

3. Flexible Thin Client

Este tipo de dispositivos arrancan el sistema operativo incrustado que aprovechan todas las ventajas de los **Thin Client** en las organizaciones, pero que son capaces de generar energía para procesar en local. Normalmente, arrancan en el dispositivo aplicaciones Java y .NET. Es para empresas que necesiten ejecutar alguna aplicación específica desde el propio ordenador. Este tipo de clientes Thin pueden ser configurados de manera personalizada para proporcionar la necesidad específica que requiera cada aplicación. Son una solución ideal para los profesionales de negocios de servicios de entrega de paquetes, operaciones minoristas y líneas áreas, así como en otras empresas que requieran utilizar los dispositivos como cajones de efectivo u otros terminales críticos para la empresa, ya que pueden funcionar incluso en el caso que haya un fallo de red.

7.7. Sistemas Multiprocesador. Taxonomía de Flynn

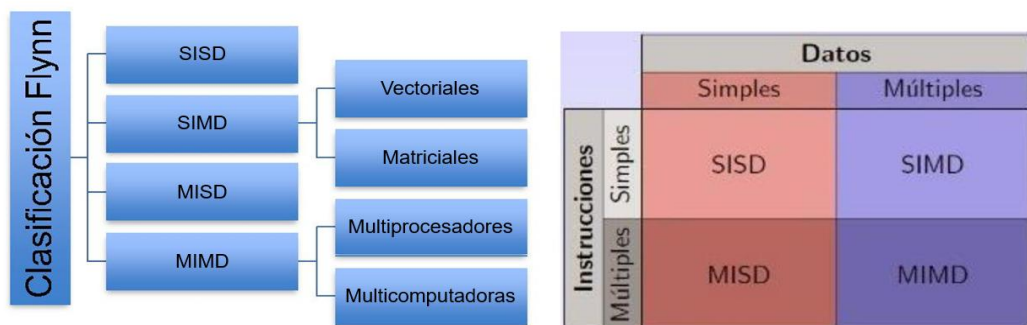
//Aunque es mas temario de GSI, lo de la taxonomia de Flynn ya lo han preguntado en TAI como minimo saber la tabla de las 4 categorias...//

Es una familia de sistemas de altas prestaciones que se basa en una sola máquina de mucha potencia de cálculo y elementos compartidos.

Taxonomía de Flynn

Se trata de una clasificación clásica de computadores en función de su arquitectura, publicada por Flynn por primera vez en 1966 y por segunda vez en 1970. Esta taxonomía se basa en el flujo que siguen los datos dentro de la máquina y de las instrucciones sobre esos datos. Se define como flujo de instrucciones al conjunto de instrucciones secuenciales que son ejecutadas por un único procesador y como flujo de datos al flujo secuencial de datos requeridos por el flujo de instrucciones.

Con estas consideraciones, Flynn clasifica los sistemas en **cuatro categorías**:



SISD (Single Instruction stream, Single Data stream)

Los sistemas de este tipo se caracterizan por tener un único flujo de instrucciones sobre un único flujo de datos, es decir, se ejecuta una instrucción detrás de otra. Este es el concepto de arquitectura serie de Von Neumann donde, en cualquier momento, sólo se ejecuta una única instrucción.

SIMD (Single Instruction stream, Multiple Data streams)

Estos sistemas tienen un único flujo de instrucciones que operan sobre múltiples flujos de datos. Ejemplos de estos sistemas los tenemos en las máquinas vectoriales con hardware escalar y vectorial.

El procesamiento es síncrono, la ejecución de las instrucciones sigue siendo secuencial como en el caso anterior, todos los elementos realizan una misma instrucción, pero sobre una gran cantidad de datos. Por este motivo existirá concurrencia real de operación, es decir, esta clasificación es el origen de la máquina paralela.

Ejemplos de este paradigma de arquitectura: ILLIAC IV, CRAY monoprocesador, CYBER 205, FUJITSU, NEC SUPERCOMPUTERS, IBM 390 VF, IBM 9000 VF, ALLIANT FX/1 Y CONVEX C-1

MISD (Multiple Instruction streams, Single Data stream)

Los sistemas MISD se contemplan de dos maneras distintas:

- Varias instrucciones operando simultáneamente sobre un único dato.
- Varias instrucciones operando sobre un dato que se va convirtiendo en un resultado que será la entrada para la siguiente etapa. Se trabaja de forma segmentada, todas las unidades de proceso pueden trabajar de forma concurrente.

Ejemplos de estos tipos de sistemas son los arrays sistólicos (el movimiento de avance del dato es el 'latido' del sistema) o arrays de procesadores. También podemos encontrar aplicaciones de redes neuronales en máquinas masivamente paralelas.

MIMD (Multiple Instruction streams, Multiple Data streams)

Sistemas con un flujo de múltiples instrucciones que operan sobre múltiples datos. Estos sistemas empezaron a utilizarse a principios de los 80.

Se trata de sistemas basados en un único bus multiplexado, de altas prestaciones, a través del cual podemos transmitir simultáneamente más de una instrucción y más de un dato.