

Sistemas Distribuidos

- Un sistema distribuido es un conjunto de máquinas que coordinan sus acciones y que aparecen al usuario como un sistema único y coherente.
- emplea software sobre la red para lograr su cometido: middleware, componentes autónomos: computadores, impresoras, etc.
- paso de mensajes: medio único de interacción es el mensaje que se transmite por un canal de comunicación (red de computadores).
- usuarios: personas o programas que interactúan con un sistema único.
- porque distribuir:
 - confiabilidad: tolerancia a fallos / disponibilidad.
 - escalable: performance.
- Para aumentar la performance, incremento las réplicas.
- La consistencia puede ser a nivel de usuario o de datos.
- Un sistema debe ser **ESCALABLE**, es decir, que estos se puedan agrandar o achicarse de ser necesario.
 - horizontal: cobocar mas equipos.
 - vertical: es aumentar la RAM.
- Las réplicas ayudan a la performance y a la tolerancia a fallos.
- Técnicas de escalabilidad:
 - particionar y distribuir.
 - generar réplicas.
 - cache (memoria pequeña y de acceso rápido) / limitador.
 - agarrar lo relevante y ponerlo en memoria cache.
- Objetivos de diseño:
 1. Poner recursos a disposición:
 - facilitar a los usuarios y aplicaciones acceder recursos remotos de manera eficiente.
 2. Hacer transparente la distribución:
 - esconder el hecho que procesos y recursos estan físicamente distribuidos en multiples computadores.
 - transparencia:
 - a) acceso: o/s tienen diferentes nomenclatura para los archivos y como se manipulan.
 - b) ubicación: donde se encuentra el recurso.
 - c) migración: recurso puede ser movido y no acepta su localización.
 - d) Re-localización: recursos en uso de movido a otra ubicación.

e) **Replicación**: esconde el hecho que existan múltiples copias.

f) **Concurrencia**: dos usuarios no notan que ambos están accediendo a un recurso.

g) **Falla**: usuario no nota que el sistema falla, ni percibe que se recupera.

3. Ser abierto:

→ sistema distribuido abierto es aquel que ofrece servicios bajo reglas estándar (sintaxis, semántica).

- **Interoperabilidad**: permite que procesos se comuniquen y generen implementaciones diferentes.

- **Portabilidad**: extensiones pueden ser usados en otros sistemas.

- **fácil extensión**: agregar más componentes.

4. Ser escalable:

→ Habilidad del sistema para expandirse para satisfacer las necesidades de su negocio. Un sistema puede escalar agregando HW adicional o actualizando HW existente sin cambiar gran parte de la aplicación.

→ **escalabilidad horizontal**: aumentar la cantidad de los recursos, ej: más PCs.

→ **escalabilidad vertical**: aumentar las capacidades de un recurso, ej: más memoria.

3 técnicas para escalar:

- **Partición**: dividir un componente y distribuirlo.

- **Replicar**: hacer copias de los recursos en distintos lugares.

- **Cache**: subconjunto de recursos de acceso fácil y rápido.

- **Particionamiento**: término genérico que simplemente significa dividir sus entidades lógicas en diferentes entidades físicas para rendimiento, disponibilidad o algún otro propósito.

→ **particionamiento vertical**: me centro en un concepto.

→ **horizontal**: me centro en todos.

- **Replicación**: Compartir información para garantizar coherencia entre recursos redundantes, como SW y HW para mejorar la confiabilidad, la tolerancia a fallas o la accesibilidad.

→ Copia de un recurso que reside en un lugar diferente.

- **Cache**: espacio reducido de almacenamiento para la información de manera temporal y de acceso rápido.

→ dependiendo de lo que se guarde, influye en el rendimiento.

→ **políticas de remoción**: (LRU, FIFO, Random, MRU, etc).

→ **políticas de ingreso**: mecanismo de control de acceso, ley de potencias, analiza datos en frecuencia, etc.

PROARTE®

→ tiempo de vida (Time to live): depende de la aplicación y el tiempo para ver cuánto tiempo mantengo una entrada a caché.

Comunicación:

→ medio de comunicación es el paso de mensajes por medio de la red.

→ red saturada, comunicación fiable no y hace difícil el desarrollo de sistemas distribuidos.

→ fundamentos de comunicación:

* La red, tiene varios dispositivos HW, variados componentes SW, medio físico.

* performance:

- Latencia: tiempo de transmitir un mensaje vacío (TL).

- Tasa de transferencia: velocidad de transmisión de datos (ϵT).

- tiempo para transmitir un mensaje: $M(T, M) = TL + b / \epsilon T$.

- jitter: variación de tiempo que toma entregar una serie de mensajes.

→ LAN: Local Area Network: Alta velocidad

→ MAN: Metropolitan Area Network: Alta velocidad

→ WAN: Wide Area Network: Velocidad menor.

→ Wireless Networks: útiles para dispositivos móviles.

→ Internetworks: conexión de varios subsistemas de redes, útil para el desarrollo de SD extensibles.

Para comunicarse existen 2 protocolos:

→ orientados a conexión.

→ no orientados a conexión.

Modelo OSI: 7 capas

- capa física: transmitir 1 y 0.

- capa enlace de datos: detectar, corregir

- capa de red: protocolo IP, envío paquetes y rutas basadas en IP address.

- capa de transporte: red no fiable, se encarga de dar conexión end-to-end, particiona el mensaje y lo envía.

• middleware: SW informático que proporciona servicios a aplicaciones de SW más allá de los disponibles en el SO.

Para comunicarse en un SD existe entre interprocess communication, indirect communications, remote invocations.

1. Interprocess Communications: comunicación de bajo nivel entre procesos distribuidos, se utilizan primitivas, acceso directo a API entregadas entre protocolos internet.

2. Indirect Communications: emisor y receptor no necesitan existir al mismo tiempo, existe una tercera parte que interviene, e y r no están al tanto de lo que

reciben

3. Remote Invocation: técnicas basadas en intercambios de dos sentidos (emisor-receptor), request/reply patrón compuesto a un servicio basado en mensajes passing bajo un modelo C/S.

→ **Remote procedure call (RPC):** protocolo usado por un programa para solicitar un servicio de un programa ubicada en otra computadora en una red sin tener que comprender los detalles de la red.

- permite el llamado a procedimientos alojados en otra máquina.
- proceso A invoca un procedimiento en máquina B, define parámetros y entrega resultados.

• Propiedades:

→ fácil de entender, distribución transparente, garantía ante fallas,

• Fallas:

→ congestión de red o servidor, caídas de clientes o error de programador.

• desventajas:

→ usa solo C/S, cliente bloqueado debe esperar, no es orientado a objetos, no es full transparente, aumento de latencia overhead.

• Problemas:

→ diferentes espacios de direcciones, caída de máquinas.

• consideraciones:

→ paso de parámetros por valor o por referencia, ya que en ciertos C enteros y otros caracteres se pasan por valor, mientras que arreglos por referencia.

• Stub Cliente/Servidor:

Objetivo RPC generar llamadas a un procedimiento más parecido a una local, quien llama no nota que se ejecuta de manera remota, librería con llamadas específicas.

• stub servidor:

recibe y transporta llamadas en locales; se bloquea al recibir una invocación y espera el mensaje bloqueado; llama al proceso local con los parámetros, los ejecuta y obtiene resultados.

• stub cliente:

recibe el mensaje y desempaqueta el resultado; se realiza una llamada a un procedimiento de manera tradicional y no hay invocación a send y receive.

• marshalling: transformación de los datos a un formato adecuado para transmisión de un mensaje.

→ **paso de parámetros por valor:** representar los datos (ASCII).

→ **paso de parámetros por referencia:** servidor recibe el arreglo y lo referencia en su espacio local permitiendo ver los cambios, este se transmite por la red hacia el cliente.

PROARTE®

Sincronización:

→ se necesita sincronización, cuando varios procesos deben decidir el orden de eventos o accesos a recursos compartidos.

- tiempos, relojes y sincronización de relojes:

Relojes físicos y lógicos:

→ relojes físicos:

- cristal + 2 registros (almacenador y contador).
- cada oscilación decremента el contador, al llegar a cero se genera una interrupción y reinicia el valor del registro almacenador, cada interrupción es llamado tic de reloj.
- cuando se inicia el computador se solicita la fecha y hora que se traduce a un número de tics.
- avanzar la hora es posible retrasarla no.

→ Algoritmo de Cristian: (UTC servidores sincronizados).

- Cada máquina pregunta el timeserver cada día 2 p segundos la hora, el servidor responde con su hora actual y cuando recibe el cliente debe actualizar su reloj. Si la hora actual es posterior a UTC se ralentiza el reloj, si la hora actual es anterior a lo recibido, lo adelanta al UTC o lo acelera progresivamente.

→ Algoritmo de Berkeley:

- Pensado en ambientes sin acceso a UTC, mantener sincronizadas las máquinas a una misma hora, servidor no es pasivo (pregunta la hora periódicamente).
- Se debe considerar el promedio tolerante a fallas, subutiliza un subconjunto de los tiempos recibidos para calcular el promedio.

→ NTP: Network Time Protocol

- Usa técnicas para filtrar los tiempos de calidad.
- Utiliza el algoritmo de Marquillo, el cual se utiliza para seleccionar fuentes de tiempo, relojes, para sincronizar tiempo entre varios relojes.
- servidores con estructura jerárquica y los niveles se conocen como stratum.
- 3 métodos de sincronización:
 1. Client-server: servidor chequea 1 o mas servidores para sincronizar.
 2. Symmetric: estratos bajo la arquitectura.
 3. Broadcast/multicast: se usa cuando se requiere poca precisión.

Conceptos:

- **delay**: tiempo de ida y vuelta desde que el cliente inicia hasta que recibe la respuesta.
- **offset**: diferencia de tiempo
- **filter**: error máximo del reloj del peer en relación al reloj local y la ruta de red que los separa.
- **poll**: indica el intervalo de contacto.
- **when**: indica el número transcurrido desde el último contacto.
- **type**: indica el tipo de sincronización.
- **Reach**: registro de cambio que se utiliza para determinar el estado de accesibilidad del peer.

→ Relojes lógicos (sincronización tiempo real).

• no es necesario sincronizar procesos que no interactúan, no es necesario que todos los procesos concuerden con el tiempo, sino que deben concordar con el orden en que ocurren los eventos.

• relojes lógicos de Lamport:

$a \sim b$ ocurre antes que b

1) Si a y b son eventos en el mismo proceso, y a ocurre antes que b , entonces $a \sim b$ es cierto.

La relación entre antes ($\sim$) es transitiva.

$a \sim b, b \sim c \rightarrow$ entonces $a \sim c$

2) Si x y y ocurren en procesos diferentes que no intercambian mensajes; $x \sim y$ no es cierto; $y \sim x$ tampoco, son concurrentes, es decir, no se puede saber el orden de como ocurrieron.

tiempo $C(a)$ manera de medir el tiempo, estos deben cumplir $a \sim b$ es decir, $C(a) < C(b)$.

El reloj C del proceso se incrementa en una unidad antes que se produzca el evento $C_p = C_p + 1$, cuando un proceso envía un mensaje, se le adjunta el valor actual de $t = C_p$.

• Multicast entre procesos

→ ubica mensajes en su cola de acuerdo al timestamp, receptor envía un ACK. luego en el caso Lamport, si el timestamp del mensaje recibido tiene que ser menor al ACK.

→ misma copia de cola

Reloj vectorial:

• $CV(a) < CV(b)$ entonces $a \sim b$

Reloj lógico

• $a \sim b$, entonces $C(a) < C(b)$, pero $C(a) < C(b)$ no implica que $a \sim b$.

Arquitecturas:

- tipo esta asociado a cual es la finalidad
- 8 tipos:
 - Computación distribuida
 - Sistema de informacion
 - sistemas distribuidos embebidos

Grid vs Cloud

- cloud publico (AWS, Azure, Google Cloud, etc.)
- no hay informacion/garantias en ubicacion, multi-tenancy, rendimiento real.
- cloud privado:
- permite monitoreo y toma de mediciones, no hay control sobre la configuracion de la infraestructura.
- Grid:
- se pueden modificar todos los niveles (VM, OS, Network).

Peer-to-Peer

- comunicacion directa, sistema abierto, completamente descentralizado.
- 3 tipos: estructurada, no estructurada, hibrida.
- ejecucion de programas que no se pueden ejecutar de manera centralizada, restricciones de computo.
- comparticion de recursos de calculo.
- descomposicion de la aplicacion en micro tareas.

Sistemas de informacion:

- sistemas principalmente para el manejo e integracion de funciones del negocio.
- EIS: cualquier sistema de informacion que mejora las funciones de un proceso de negocio.
- procesadores de transacciones: primitivas especiales, atomicas, consistentes, aisladas, durables (ACID).
- componentes o servicios SOA.

Sistemas Pervasivos.

- Incluyen la componente de movilidad.
- no hay control centralizado, caracteristicas especiales: energia, procesamiento, almacenamiento, etc.

Arquitecturas Distribuidas

→ arquitectura del sistema:

Según coloquemos los componentes del SW:

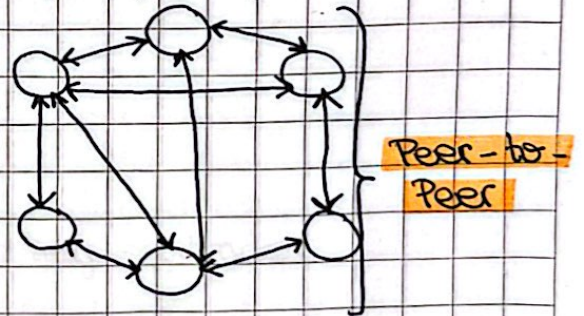
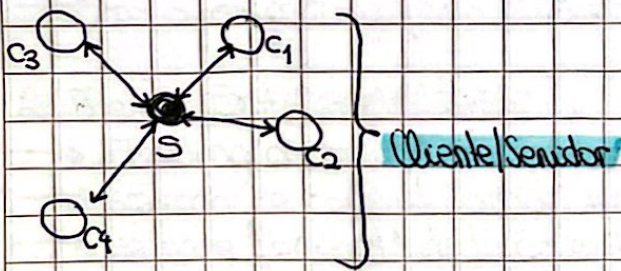
- arquitecturas centralizadas: Cliente/Servidor
- arquitecturas descentralizadas: Peer-to-Peer
- arquitecturas híbridas: Edge

• Cliente/Servidor

- 2 subsistemas o procesos
- clientes: proceso que solicita un servicio.
- servidores: proceso que implementa un servicio en específico.
- la app es modelada como un conjunto de servicios proveídos por los servidores y un conjunto de clientes que usan los servicios.
- escalamiento vertical, modelo seguro, cuellos de botella posible.
- mucha responsabilidad en el cliente
- cliente/servidor (3-tier), el mismo servidor puede actuar como cliente.

• Peer-to-Peer

- una arquitectura de red distribuida puede ser llamada P2P si los participantes comparten una parte de sus recursos. Ellos son accesibles por otros pares directamente, sin pasar por entidades intermedias.
- colaboración y compartición de recursos a gran escala.
- potencial de tener millones de nodos y sus participantes son simétricos



- P2P es altamente escalable
- 100% descentralizado, no hay cuellos de botella.
- Tolerancia a fallos (no hay punto único de fallo).
- conocimiento local.

• Overlay:

- red de computadores construida sobre otra red.

• Tipos de Overlay P2P

1. Estructurados
2. No estructurados
3. Híbridos.

• P2P estructurados

- estructura bien definida
- se utilizan tablas hash distribuidas para su generación.
- proveen búsquedas completas, eficientes y escalables
- mapeo a espacios lógicos por medio de un hash.
- Problemas: alta latencia de búsqueda, poca seguridad, alto dinamismo

P2P no estructurado

- Redes no estructuradas
- conmutación local, estructura aleatoria, gran autonomía, etc.
- para conectarse, se llama Bootstrap node, y lo localigo con listas pre-existentes, bootstrap cache.
- como buceo, son incompletas, poco eficientes, completamente descentralizadas
- difícil encontrar contenido no popular, no garantizan resultados.
- Recursos: se almacenan peers aleatorios, búsquedas más profundas.

• Tipos de búsquedas:

1. Flooding o inundación de mensajes.

- mensaje enviado a todos los participantes conocidos
- poca reactiva y poco eficiente

2. Breath First Search (BFS)

- Flooding controlado
- sistema de control basado en TTL.
- se hace "forward" de la consulta si no hay resultados.

3. Iterative Deeping

- múltiples y secuenciales BFS o expanding ring.
- comienza con búsquedas poco profundas, de no encontrar el recurso itera y prueba mayor profundidad.

4. Directed BFS:

- se elige un subconjunto de los vecinos y se envía el mensaje.
- trata de adivinar que vecino es más probable que responda (Random, resultados con menos saltos, más vecinos, etc.).

4. Local Indexes

- que archivos tienen mis vecinos y los vecinos de mis vecinos.
- todos los nodos conocen el radio.

5 Random Walks

- Diferentes a BFS, mensaje enviado a un vecino aleatorio
- Si se pilló el recurso retoma un hit, sino el proceso repite hasta TTL max.
- donde ir se evalúa de manera random en cada paso.
- basado en información es abundante y uniforme.

Super - Peers

- encontrar información en una red no estructurada es problemático.
- no hay proceso determinístico para encontrar un recurso en específico.
- Super-peers son peer con mayores capacidades o roles especiales. (procesamiento, ancho de banda, estabilidad).
- organización jerárquica.
- Peers regulares conectados a Super peer.
- Toda la comunicación va a través del super peer.

Edge Architectures

- soluciones en los bordes (edge) de la red.
- end users se conectan a edge servers.
- Permite optimizar la entrega de contenido.
- Se aplican técnicas similares a todo sistema distribuido.

Problema → resolver un nombre necesitamos un nodo directorio

— necesitamos seleccionar un contexto donde comenzar la resolución de nombres

Domain Name System

relaciona direcciones IPs con nombres

sino existiese habría que aprender las ips comprometiendo el desarrollo de internet ejemplo: publicidad

estructura jerárquica de árbol

nombres son llamados nombres de dominio ej: udp.el

contiene componentes de nombre separados por un '.'

no reconoce nombres relativos, concatena dominios por omisión

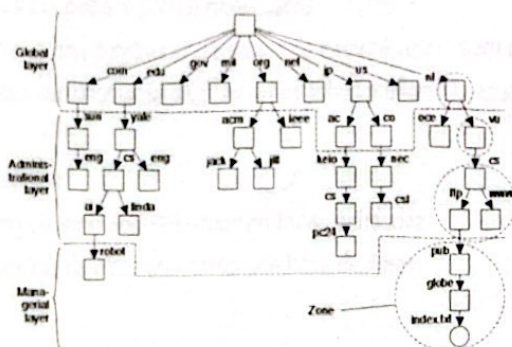
distribuir la resolución de nombres y el espacio de nombres de nombres sobre múltiples máquinas

tres niveles:

Global Level — directorios de alto nivel manejados por varias administraciones

Administrational Level — directorios de nivel intermedio

Managerial Level — directorios de bajo nivel manejados por una administración (local DNS)



¿Cómo funciona?

Participantes:

DNS Resolver — recibe consultas de los clientes por medio de aplicaciones como web browser — resolver es responsable de realizar requests adicionales para satisfacer la consulta del cliente

Root Nameserver — primer paso es la traducción de nombres a IP addresses — entrega información más específica de una localización

TDL Nameserver — paso siguiente en la búsqueda de una IP address y almacena la última porción de un hostname ejemplo example.com almacena '.com'

Authoritative Nameserver — ultima parada en la búsqueda — accede al registro y devuelve la ip address al host solicitado y lo envía de vuelta al DNS resolver

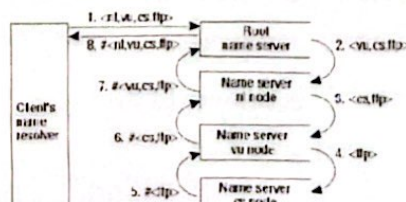
DNS Lookup

- usuario tipea 'example.com' en el web browser -- consulta (querrá) viajar por internet -- llega al DNS resolver

resolver envía la query a DNS root nameserver -- root server responde a el resolver con una dirección de un TLD DNS server -- almacena la información de estos dominios -- dada la query 'example.com' la request es redirigida a .com TDL -- resolver realiza la request a .com TLD -- TLD server responde con el IP address del dominio nameserver → example.com -- resolver envía una consulta al dominio nameserver -- ip address es devuelto al resolver desde el nameserver -- DNS resolver responde al web browser con la IP address del dominio solicitado inicialmente

Tipos de queries:

Recursivo — DNS resolver client se comunica con un DNS server típicamente un DNS recursive resolver — responde con el nombre o error en caso que no se encuentre

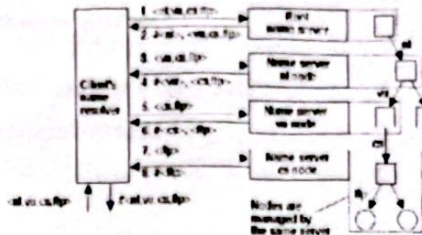


Recursive DNS resolver

- máquina que responde a una request recursiva y se encarga de gestionar las request para obtención el DNS record
- realiza una serie de request hasta encontrar el authoritative DNS nameserver para el record solicitado (o entrega el error si no lo encuentra)

Iterativo: — DNS resolver client permite al DNS server responde con la mejor respuesta que encuentre — sino encuentra el nombre, se retorna una referencia a un dominio más bajo

— proceso continua hasta tener una respuesta, error o timeout



No Recursivo — DNS resolver client se comunica con un servidor de nombre buscando un registro y este es Authoritative Nameserver o bien se encuentra en el cache de este

técnica ampliamente utilizada para reducir el consumo de ancho de banda

Resolución recursiva permite a cada servidor aprender acercada donde ubicar la consulta — ideal para cache

en approach iterativo el cache es necesario pero a nivel cliente

en general es de menor costo que una iterativa en términos de comunicación, pero es más costosa en computo para el servidor demanda más recursos, por eso servidores de nivel superior(global) prefieren búsqueda iterativa

Attribute-Based naming

servicios de directorio tradicionales

problema — costoso, matching de atributos -> buscar en todas entidades

solución— implementar un servicio de directorio como una base de datos

Relojes físicos

cristal +2 registros (contador y almacenado)

cada oscilación decremента el contador, al llegar a cero se genera una interrupción y reinicializa al valor del registro almacenador

cada interrupción es tic de reloj

al iniciar el computador se solicita fecha y hora que se traduce a un número de tics

proceso locales no les afecta reloj en cuanto avance normalmente

a pesar de que la frecuencia del cuarzo es estable, es imposible garantizar que todos los cristales de los computadores oscilen a la misma frecuencia

- intensidad de corriente y corte del cristal

problema —> gradualmente los relojes se desajustan, diferencia entre relojes de las máquinas es conocido como deriva del reloj

¿Cómo medimos el tiempo?

— intervalos entre 2 tránsitos del sol (punto más alto)

— intervalo rotación tierra no es constante

— rotación se hace más lenta —> día más largo

TAI — international atomic time

solución— introducir saltos

NIST —> provee UTC por medio de una estación de radio de onda corta (WWV)

Sincronización

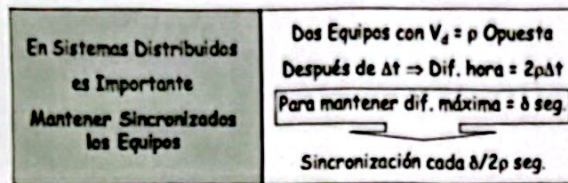
Algoritmo de Cristian

— pensado en entornos con servidor de tiempo sincronizado con UTC y las estaciones desean sincronizarse con el

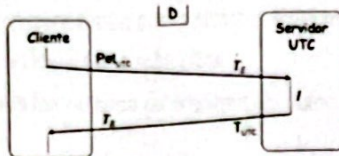
— cada máquina pregunta a el time server cada:

Segundos de diferencia
 2/2p segundos la hora
 máxima deriva

— servidor de tiempo responde con su hora actual Tuto— cuando recibe el cliente debe actualizar su reloj



si la hora actual es posterior a UTC, debe ralentizar su reloj, en cada interrupción se añaden 9 ms
 si la hora actual es anterior al recibido, lo adelanta al UTC o lo acelera progresivamente



Algoritmo de Cristian →

$$\text{Error promedio} = \frac{(T_i - T_0)}{2 + 1}$$

Problema → punto único de fallo — es centralizado

solución → varios servidores UTC, clientes considera la primera respuesta que llega

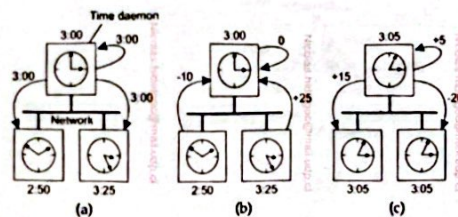
problema → relojes erróneos y propagación de error, aun se puede tener consenso si $N > 3f$

Algoritmo de Berkeley

— pensados para ambientes sin acceso a UTC

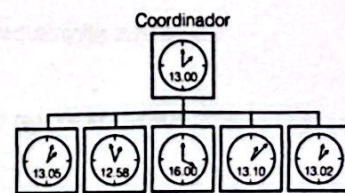
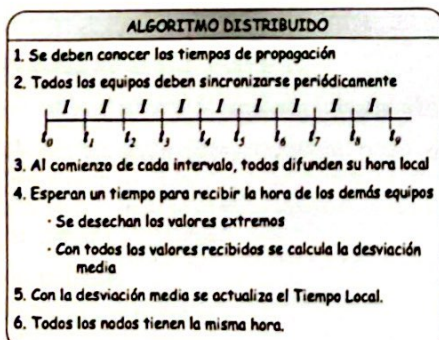
— principio radica en mantener sincronizadas las máquinas a una misma hora — no es necesario sincronizar a tiempo real

— selección de coordinador o servidor de tiempo



Promedio tolerante a fallos → solo utiliza un subconjunto de los tiempos recibidos para calcular el promedio — los que no varían en más de un valor Δ . El resto se consideran relojes fallados

si falla el coordinador se selecciona otro — no pregunta hora, no envía notificaciones, etc



Network time protocol (NTP)

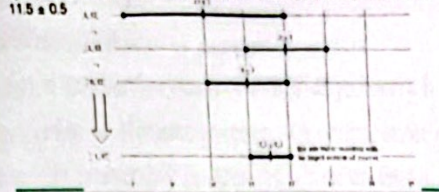
generado para proveer servicio de tiempo en internet

tiempos de transmisión en internet son variables

usa técnicas estadísticas para filtrar los tiempos de calidad

utiliza el algoritmo de marzullo → se utiliza par seleccionar fuentes de tiempo, relojes, para sincronizar tiempo entre varios mejores, basado en un origen o fuente único

Ejemplo: 10 ± 2 , 12 ± 1 y $11 \pm 1 \rightarrow [8, 12]$, $[11, 13]$ y $[10, 12] \rightarrow [11, 12]$ o 11.5 ± 0.5



servicio debe ser escalable para lidiar con constantes peticiones de sincronización

autenticación de servidores para proveer un servicio seguro (maliciosos o accidentales)

red de servidores en internet $\rightarrow$ servidores conectados en una estructura jerárquica o sub-red de sincronización

niveles de jerarquía son conocidos como strata o stratum

errores son mayores en los stratos más altos

se considera igualmente los tiempos de transmisión como métrica

de calidad

strata 2 sincroniza con strata 1, strata 3 sincroniza con strata

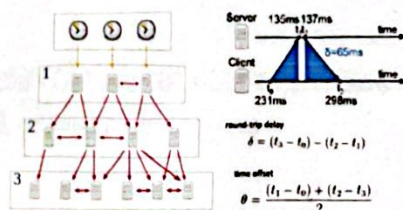
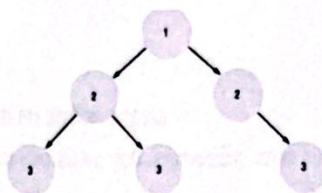
2, etc

3 modos de sincronización:

Client/Server: servidor chequea uno o más servidores para sincronizar

Symmetric: para stratos bajos de la arquitectura— se utilizan entre ellos para generar backup mutuo de referencias

Broadcast/Multicast: se usa cuando se requiere poca precisión— uno sincroniza y notifica a todo el resto



se calculan los tiempos de propagación aplicando el concepto de dispersión (suma de dispersiones servidor-root)

se aplica una fase para la modificación de la frecuencia de los relojes en base a las derivas observadas

Delay — indica tiempo de ida y vuelta en segundos desde que el cliente inicia la comunicación hasta que recibe la respuesta

Offset — indica la diferencia de tiempo, en segundos entre el reloj del cliente y el de la fuente

Jitter — indica el error máximo en segundos, del reloj del peer en relación al reloj local y la ruta de red que los separa

Poll — indica el intervalo de contacto

When — indica el número de segundos transcurridos desde el último contacto

T(type) — indica el tipo de sincronización, puede ser:

l: local por ej. wwwb o gps

u: unicast

m: multicast

—: desconocido

Reach — registro de cambio que se utiliza para determinar el estado de accesibilidad del peer — un peer se considera accesible si al menos un bit en este registro está establecido en 1

Relojes lógicos

no es necesario sincronizar procesos que no interactúan, su falta de sincronización no será observable, deben concordar con el orden en que ocurren los eventos

A
B
C
D
E
F
G
H
I
J
K
L
M
N
Ñ
O
P
Q
R
S
T
U
V
W
X
Y
Z

Precedencia — ocurre antes de .

$a \sim b$ — a ocurre antes que b

todos los procesos concuerdan que a ocurre y luego ocurre b

1) si a y b son eventos en el mismo proceso, y a ocurre antes que b entonces $a \sim b$ es cierto

2) si a es un evento enviado por un proceso y b es un evento recibido por otro, entonces $a \sim b$ es cierto

la relación "ocurre antes que" — es transitiva

si $a \sim b$ y $b \sim c$, entonces $a \sim c$

Concurrentes — no se puede saber el orden de como ocurrieron ($\sim$ o $\parallel$)

evento a tiene un valor de tiempo $C(a)$ en el que todos los procesos concuerdan

estos tiempos deben cumplir $a \sim b$, es decir $C(a) < C(b)$

el tiempo de reloj C , siempre debe avanzar y nunca retroceder

problema —> actualizaciones deberían realizarse en el mismo orden en ambas copias

solución —> totally ordered multicast —> multicast donde todos los mensajes son entregados en el mismo orden, como relojes lógicos de lamport

Multicast entre procesos

cuando un proceso recibe un mensaje lo pone en su cola local ordenado de acuerdo a su timestamp

receptor envía un ACK a los otros procesos

siguiendo lamport, si el timestamp del mensaje recibido tiene que ser mejor al ACK

cada proceso tendrá la misma copia de la cola

proceso pasa el mensaje a la aplicación si: mensaje si esta en el reader de la cola, ACKs de todos los procesos han llegado para ese mensaje —

se asume comunicación confiable

Reloj lógico —> $a \sim b$, entonces $C(a) < C(b)$, pero $C(a) < C(b)$ no implica que $a \sim b$

Reloj de vector —> $CV(a) < CV(b)$, entonces $a \sim b$

Vectores:

cada proceso P_i tiene un arreglo $VC_i[n]$ con n numero de procesos involucrados

$VC_i[j]$ es el número de eventos que el proceso P_i sabe que se han llevado a cabo en P_j

cuando un proceso i envía un mensaje comparte su VC_i

Algoritmo:

- Inicialmente, $V[j] = 0, \forall j$
- Evento local o envío de mensaje en p_i : $V[j] = V[j] + 1$
- Cuando p_i envía un mensaje m a p_j , incluye el valor de su vector de tiempos, V_{p_i} . Al recibir dicho mensaje, p_j actualiza su vector de tiempos de la siguiente manera:
(1) $V[k] = \max(V[k], V_{p_i}[k])$. (2) $V[j] = V[j] + 1$

Reloj Vectorial

existirá una relación de causalidad entre 2 eventos a y b si solo si $VC(a) < VC(b)$ o $VC(b) < VC(a)$, sino son concurrentes

$VC_i < VC_j$ si solo si —> $VC_i[j] < VC_j[j]$ para todo j y existe un j / $VC_i[j] < VC_j[j]$

Exclusión Mutua

procesos deben cooperar de manera concurrente

recursos compartidos pueden ser accedido de manera concurrente

para asegurar inconsistencias o corrupción de información

Requisitos:

Seguridad — proceso puede ejecutar el recurso compartido (región crítica)

Viveza — se debe garantizar que un proceso que desea acceder a el recurso haga ingreso en un momento dado

Orden — garantizar orden causal, si un proceso P_1 envía un mensaje a P_2 y ambos requieren un recurso P_1 deberá accederlo antes que P_2

garantizan múltiples procesos distribuidos no se interfieran y logren cooperar

2 tipos: Token-Based y Permission-Based

Token-Based

exclusión mutua se realiza por medio de un token — token único quien lo tenga puede hacer uso del recurso compartido — cuando termina, el token es pasado al siguiente proceso, si el proceso siguiente no está interesado en usar el recurso, se pasa inmediato al siguiente

propiedades:

aseguran de manera fácil que los procesos van acceder el recurso compartido, no hay "starvation" o inanición

procesos se esperan unos a otros por el token por lo que no hay deadlocks

problema —> pérdida del token (caída del nodo), se debe iniciar un proceso complejo de creación de un nuevo token

Permission-Based

proceso solicita permiso a otro dos procesos — 2 mensajes: mensaje solicitud y mensaje autorización

Algoritmos de exclusión mutua:

Centralizados — servidor único

- único proceso coordinador

- si algún proceso desea acceder al recurso compartido, envía mensaje a coordinador — recurso está libre, entrega acceso respondiendo con el permiso — al recibir el mensaje con el permiso, proceso puede acceder al recurso

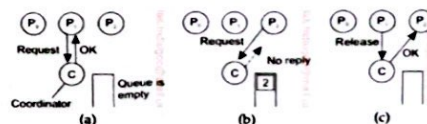
- solo un proceso a la vez puede acceder el recurso compartido

- se respeta el orden de las solicitudes

- ningún proceso espera por siempre

solo 3 mensajes por recurso — request, grant, release

problemas: punto único de fallo, si falla todo el sistema también — luego de realizar una petición, el proceso no puede identificar un coordinador caído de un permiso denegado — cuello de botella en un sistema de gran escala



Algoritmo Descentralizado — usando un sistema P2P

—> solución: descentralización — algoritmo de voto sobre un DHT, recurso replicado n veces cada uno con su coordinador, se requiere mayoría de $m > n/2$ coordinadores, m números de mensajes mayoría y k intentos (back-off)

cundo el coordinador no concede acceso, el coordinador notifica

elimina cuello de botella y punto único de fallo

Algoritmos Distribuidos — sin topología o usando anillo lógico

algoritmos de marcas de tiempo — mejora el algoritmo de lamport — se requiere ordenamiento total

Algoritmo —>

proceso que desea un recurso compartido genera un mensaje que contiene —> nombre del recurso; número del proceso; reloj lógico actual

envía mensajes a todos los procesos incluido el

no hay pérdida de mensajes

Proceso recibe un mensaje —> receptor no está accediendo al recurso y no quiere accederlo —> envía permiso ok

receptor ya accedió al recurso —> no responde y lo pone en cola

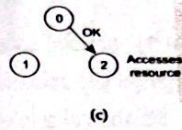
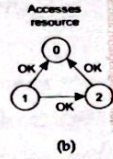
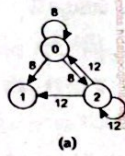
proceso también quiere acceder al recurso y aun no lo ha hecho, compara timestamp del mensaje con el suyo — si es menor el del mensaje — envía permiso ok, sino pone en cola el mensaje y no responde

debe esperar todos los permisos para acceder a el recurso

cundo termina envía ok a todos los procesos en su cola y los elimina

garantiza que todo accedan al recurso y que no se produzcan deadlocks

as híbridas.



problemas : caída de nodos, la no respuesta es considerada como una negación de acceso

solución— responder siempre o mayoría de votos

— manejo de multicast en cada proceso, mantener la lista de membership

— más complejo que el centralizado, n puntos fallo y con más tráfico

Token-Ring

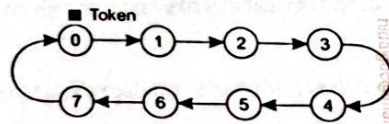
anillo lógico de equipos, cada uno con una posición dada según ip

cada nodo conoce la dirección de su vecino siguiente

al arrancar, primer proceso recibe el token el cual ira circulando por el anillo

cuando proceso k tiene el token, debe transferirlo al proceso k+1 por un mensaje

si el proceso desea hacer uso del recurso y tiene token, puede acceder a el reteniendo el token — cuando libera el recurso pasa el token



problemas: perdida de token, cada de nodos

Repaso Solemne 2

Propiedades para hacer escalar al DNS

- particionar : particiona por dominios y subdominios → guardamos información relativa a una URL hace una relación texto a una dirección IP
- replicas : niveles superiores → globales y administrativos → niveles que agregan mayor cantidad de carga
- hash/cache : nivel cliente → resolver → recibe la consulta y resuelve la consulta

tablas de hash distribuidas (DHT) → P2P estructurado

como logra:

- alta disponibilidad → cada par se alojaba en un espacio numérico → porque al aplicar un hash sobre la dirección ip: port
- se logra por el proceso de mapeo donde rompe localidad y eso evita particionamiento
- escalabilidad → explota la calidad en las entradas de la tabla → número reducido de entradas → se puede localizar a cualquiera
- balance de carga → hash → se distribuyen uniformemente los recursos indexados → ejemplo hash sha-1

Problema de sincronización— se puede abordar con relojes físicos y lógicos

¿por que al utilizar un reloj físico hay problemas de sync?

- corriente y calidad cuarzo → difieren de par a par siempre hay un efecto de deriva

para determinar los tiempos— los tiempos divergen por eso no se pueden usar en gran escala porque al divergir los tiempos el promedio del tiempo no será el mismo para cada uno

solución → NTP a gran escala se tiene mediciones por cada nodo y peticiones de los peer que están participando dentro de lo mismo

metodos recursivos DNS → se elige el iterativo sobre el recursivo para evitar la sobre carga del TTL

A

B

C

D

E

F

G

H

I

J

K

L

M

N

Ñ

O

P

Q

R

S

T

U

V

W

X

Y

Z

Orientada a la computación de alto rendimiento (HPC)

Clusters

- > nodos homogéneos conectados en una red rápida
- > Hardware homogéneo
- > utilizados tradicionalmente para aplicaciones paralelas, programa único corriendo en paralelo en varias máquinas

Grid

- > nodos heterogéneos conectados a una red compartida bajo varias administraciones
- > organización virtual con recursos para compartir
- > computación y almacenamiento
- > colaboración de diferentes recursos

Cloud público (AWS, azure, google cloud, etc)

- no hay información/garantías en ubicación, multi-tenancy, rendimiento real

Cloud Privado

- permite monitoreo y toma de mediciones
- no hay control sobre la configuración de la infraestructura

Grid -> completamente reconfigurable

- se pueden modificar todos los niveles (VM, OS, Network...)

Peer-to-peer -> nodos heterogéneos conectados a una red compartida completamente abierta

- comunicación directa
- sistema abierto
- completamente descentralizado
- tipos: estructurada, no estructurada, híbrida
- ejecución de programas que no se pueden ejecutar de manera centralizada, restricciones de cómputo
- descomposición de la aplicación en micro-tareas

Sistemas de información

- sistemas principalmente para el manejo e integración de funciones del negocio
- Enterprise information system (EIS) : cualquier sistema de información que mejora las funciones de un proceso de negocio
- procesadores de transacciones : -> primitivas especiales -> atómicas, consistentes, aisladas, durables (ACID)
- componentes o servicios (SOA)

Sistemas Pervasivos

- no hay control centralizado
- incluyen la componente la movilidad
- características especiales : energía, procesamiento, almacenamiento, etc
 - home systems, health , smartcities

Arquitecturas distribuidas -> una buena organización es clave para alcanzar las propiedades deseadas

Arquitectura del sistema -> si consideramos donde colocamos los componentes de software tenemos:

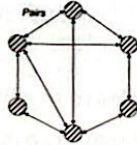
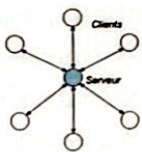
- arquitecturas centralizadas : cliente/servidor
- arquitectura descentralizada: peer-to-peer
- arquitecturas híbridas : edge

Cliente/Servidor

- dos sub-sistemas o procesos (2-tier)- cliente/servidor simple
- clientes: proceso que solicita un servicio
- servidores: proceso que implementa un servicio en específico
- aplicación es modelada como un conjunto de servicios proveídos por los servidores y un conjunto de clientes que usan los servidores
- paso de mensajes (interacción request-reply)
- escalamiento vertical
- cuello de botellas posibles
- modelo seguro
- dos máquinas (2-tier) y 3 capas (3-levels): interface application y database
- tendencia a mantener el procesamiento en el lado del servidor
- cliente/servidor (3-tier) -> multi-tier: mismo servidor puede actuar como un cliente

Peer-to-Peer

- arquitectura de red distribuida puede ser llamada P2P si los participantes comparten una parte de sus recursos. Ellos son accesibles por otros pares directamente, sin pasar por entidades intermediarias
- colaboración y compartición de recursos a gran escala
- participantes son simétricos



C/S vs P2P

- altamente escalables
- 100% descentralizado, no hay cuellos de botella
- auto-organizados, participantes autónomos
- tolerancia a fallos (no hay un único punto de falla)
- conexión y participantes no confiables (abierto y anónimo)

Overlay-> red de computadores construida sobre otra red

P2P Estructuradas

- construidas con proceso determinístico
- tienen estructura bien definida
- se utilizan tablas de hash distribuidas para su generación
- capaces de escalar a los cientos de millones de participantes
- proveen de búsquedas completas
- búsqueda eficiente y escalable orden logarítmico ($\log N$)
- rapeo a espacio lógico por medio de hash tipo sha-1 o md5
- rutas diferentes dependiendo de la geometría del overlay
- Key Based Routing protocols (KBR)
- > tree, hypercube, butterfly, platoon trees

Chord

- organización tipo anillo $\rightarrow O \dots 2^n$
- cada peer posee un identificador generado con un SHA-1 (consistent hash)
 - $\rightarrow$ node id: SHA(IP Address + puerto)
 - $\rightarrow$ recurso id (key): SHA(recurso)
- auto-organizado
- ruteo de mensajes de tipo árbol
- asignación de recursos $\rightarrow$ sucesor de la key
- asignación en sentido del reloj (uni-direccional)
- peer entra (join) $\rightarrow$ recibe recursos $\rightarrow$ se hace responsable de un grupo de recursos (sucesor)
- peer deja la red (leave) $\rightarrow$ entrega recursos al nuevo sucesor

* Mantenimiento caro ($\log N^2$)

Chord vs Churn

Churn $\rightarrow$ join/leave crean inconsistencias en tablas

- reparación es costosa, protocolo de estabilización
- se detectan caídas por hear beats messages
- guarda replicas en S sucesores

¿Cómo participamos?

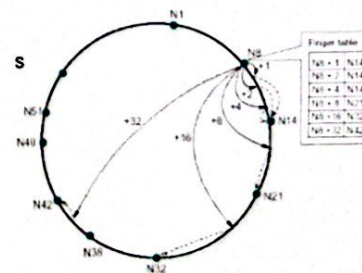
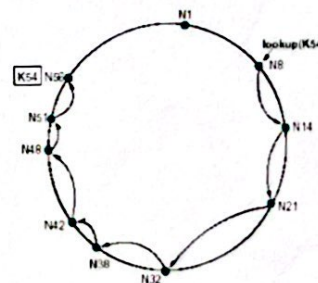
- identificador en el mismo espacio, se rutea identificador
- contacta al sucesor con su identificador
- consistencia de tablas
 - $\rightarrow$ sucesor entrega información de su antecesor (entradas tabla)
 - $\rightarrow$ sucesor agrega a un nuevo peer como antecesor
- distribuye los recursos de los cuales es responsable el nuevo peer

¿Cómo se busca?

- cada peer conoce su sucesor en el anillo
- modelo simple de búsqueda $\rightarrow$ seguir links a sucesores

¿Cómo mejorarla?

- agregar información
- fingertable(DTH) $\rightarrow$ tabla que contiene información acerca de $\log m$ peer en la red $\rightarrow$ permite saltar peers y acercarse más rápido al peer responsable
- objetivo $\rightarrow$ encontrar el sucesor de la clave $\rightarrow$ métrica: cercanía entre identificadores



Pastry

- organizado en anillo
- topología ruteo Plaxon-Tree
- ruteo basado en prefijo
- identificador de 128 bits generado con hash SHA-1
- espacio $O(2^{128})$
- ruteo $\log N$ saltos al destino

Pastry Leafset

- pastry guarda L peers llamados leafsets nodes
- para un peer X , su leafset esta compuesto $L/2$ y $L/2$ peers en el anillo
- para balancear cargas, tolerancia a fallo, replicas, multi-path search

Tabla de ruta Pastry

- organizada por prefijo
- fila i contiene aquellos ids de peer que comparten un prefijo de tamaño i
- tuteo por prefijo asegura acercarse sucesivamente a la clave buscada
- mas prefijo en común, más cerca del destino

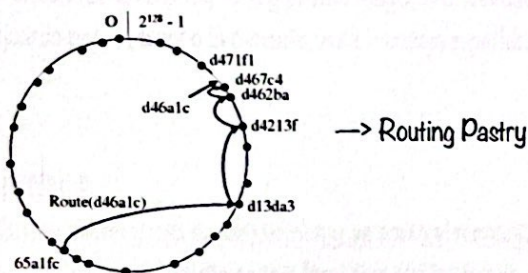


Tabla de vecinos

- peer cercanos en términos de latencia
- útiles para generar búsquedas diversas en caminos
- se genera -> ping para estimar latencia o simplemente IP

Overlays estructurados

- búsqueda puntual es un problema para las aplicaciones
- se resuelve con Prefix-tree (PHT), multianillos, etc
- se puede generar un overlay que preserve la localidad -> LSH

P2P no estructurado

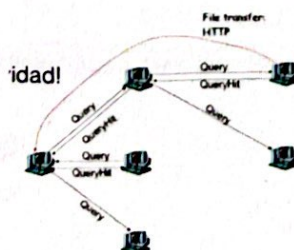
- estructura aleatoria
- conocimiento local
- métodos de búsqueda basado en inundación de mensajes o flooding
- gran autonomía

¿Cómo me conecto?

- participante llamado Bootstrap node
- lo localizo a través de listas pre-existentes, bootstrap cache
- problema para limitar conexiones con el fin de balancear la carga

¿Cómo busco?

- incompletas
- poco eficientes, muchos mensajes
- comparación de string
- completamente descentralizada
- efectivo en contenido popular
- búsqueda no garantiza resultado



Recursos

- se almacenan en peers aleatorios
- recursos pueden quedar lejos de los peers -> max TTL=7 en gnutella 0.4
- búsquedas más profundas, más mensajes
- caching -> problema de investigación



Tipos de búsqueda

Flooding o inundación de mensajes

- mensaje enviado a todos los participantes conocidos
- poco eficiente
- calidad resultados vs cantidad mensajes

Breadth First Search (BFS)

- flooding controlado
- sistema de control basado en TTL o routing hops
- si no hay resultados -> forward
- proceso se repite TTL veces
- mensajes duplicados, gran overhead

Iterative Deeping

- multiple y secuenciales BFS o expanding ring
- comienza con búsquedas poco profundas
- al no encontrar el recurso, itera y prueba mayor profundidad
- > guarda por un tiempo la consulta evita mensaje repetidos

Directed BFS

- se elige un subconjunto de los vecinos y se envía el mensaje
- trata de adivinar que vecino es más probable que responda —> random, resultados con menos saltos, + mensajes recibidos, + vecinos,
- latencia, etc

Local Indexes

- que archivos tienen mis vecinos y los vecinos de mis vecinos
- todos los nodos conocen el radio
- problemas de memoria para mantener listas
- construir la lista, extra overhead

K-Random Walks

- multiples random walks paralelos
- disminuye latencia
- variantes : —> APS

—> id- based , enviable solo al que no ha ruteado consulta

—> two-level random walk, k1 random walkers por un TTL T1, luego k2 random walkers por un TTL T2

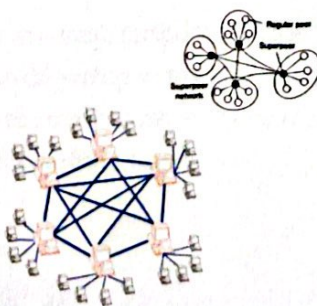
Condiciones de termino :

TTL-> difícil setear el parámetro , red dinámica no ayuda

Chequeo resultados -> se chequea con el nodo que consulta frecuentemente , normalmente cada 4 saltos

Super-Peers

- super peers son peer con mayores capacidades o roles especiales -> procesamiento, ancho de banda, estabilidad
- organización jerárquica
- peers regulares conectados a super peer (SP)
- toda comunicación va a través del super peer
- relación estable con super peer
- SP larga sesión, alta disponibilidad



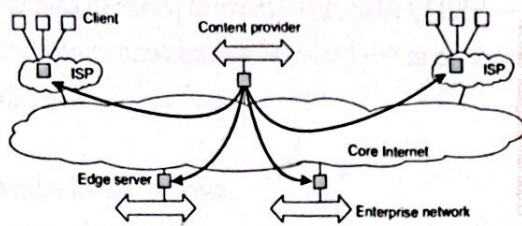
Random Walks

- mensajes enviado a un vecino aleatorio
- si se pilla el recurso retorna un hit, si no proceso se repite hasta TTL max
- se evalúa donde ir de manera random en cada paso
- ciego y sin memoria
- basado en que la información es abundante y uniforme
- utiliza un walker
- disminuye el numero de mensajes
- aumenta la latencia al esperar respuestas

Edge Architectures

- soluciones en los bordes (edges) de la red
- end users se conectan a edge servers
- permite optimizar la entrega de contenido

- se aplican técnicas similares a todos sistema distribuido -> replicas en otros edge servers
ejemplo: web pages



procesos no pueden compartir recursos a menos que los nombren consistentemente

Local Vs Distribuido

Foco en lo distribuido -> implementación del servicio de nombres es distribuido sobre múltiples máquinas — juega un rol importante en la eficiencia y escalabilidad del sistema

Naming— identificar

cadena de letras (string) que hace referencia a una entidad ejemplo: hosts, impresoras, discos, archivos, procesos, web pages, etc.

entidades pueden ser operadas ej: imprimir

primero hay que accederla -> dirección y mecanismo para localizar un recurso

se pueden tener múltiples direcciones ejemplo: combinación ip/puerto

Identificador vs Nombres

requiere que sean únicos, identifica a lo más una entidad

cada entidad es referenciada por a lo más un identificador

identificador siempre hace referencia a la misma entidad
reduce la ambigüedad para acceder a una entidad

problemas -> entidad cambian , nombres no son únicos

URI: Uniform Resource Identifiers

son uniformes y coherentes— permite el procesamiento de los nombres por softwares comunes ejemplo: browser

compuesta por 3 partes : esquema, nombre de máquina donde se aloja el recurso, nombre del recursos en forma de ruta(path o ruta)

subtipos: URL (uniform resource locators) y URN(uniform resource names)

URL:

provee información acerca de un recurso y como localizarlo — ejemplo: <http://www.google.cl>

eficiente para acceder recursos , pero no es eficiente ante eliminación o cambio del recursos

URN:

nombres de recursos puros, no para localizar — identifica un mensaje particular usando el "Message-id"

distingue este mensaje de otros, pero no provee la dirección del mensaje en ningún almacenamiento

proceso o mecanismo permite la asociación de un nombre con una dirección -> name-to-address

3 tipos de sistemas : Flat Naming, Structured Naming, Attribute-based naming

Flat Naming

soluciones simples → broadcast and multicast, forwarding pointers

home based solutions

DHT's

Identificador Random → Cadena de storing generada (Hash, SHA-1, MD5) — sin estructura

identificador no provee información de como localizar la entidad → se localiza por Lookup

espacio de búsqueda limitado

Broadcast y multicast

→ se envía un mensaje a todos los participantes

→ no escala más allá de una red local (LAN)

→ todos los procesos deben estar escuchando

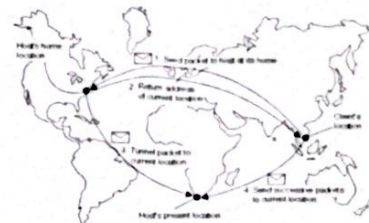
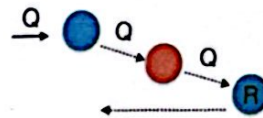
Multicast reduce el número de mensajes y mejora la escalabilidad sin embargo sigue teniendo escalabilidad limitada

Forwarding Pointers → cuando una entidad se mueve dejando una referencia hacia su nueva ubicación, cuando se localiza se hace el update de la ubicación

Home-Based Approaches: Mobile IP

- home mantiene donde la entidad se encuentra
- home address de la entidad es registrado en un servicio de nombres
- clientes contactan el servicio de nombres y luego van a la ubicación de la entidad
- chequea registro local — va al home address si falla el local

Problemas? — home address es fijo → sobrecarga cuando la entidad cambia frecuentemente de su lugar, escalabilidad limitada



DHT's — Distributed Hash Tables

Chord → nodos en espacio lógico 2^m

cada nodo tiene id random único (m-bits) $m=160 \rightarrow 1,46+48$ ids

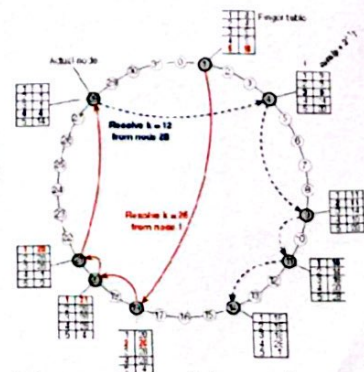
cada entidad tiene asignada una key de m-bits

entidad id es responsable de la key k si — $id \geq k$ (sucesor)

problemas? → proximidad semantica, latencia

soluciones → Proximity Aware Overlay, Proximity Routing

como buscamos?



Structured Naming — sistema cuyo identificador posee una estructura

"human-friendly" son estructurados y compuestos de varias partes → nombres, apellido paterno, apellido materno — subdominio.dominio

espacio de nombre representado como un grafo ciclico dirigido (DAG)

estructura jerárquica — existe un único primer nivel que se expande abarcando

toda la red, el root

red es dividida en una colección de dominios — cada dominio es dividido

en sub dominios de menor tamaño

hojas → dominios locales

puede crecer de manera infinita ejemplo: /etc/passwd

