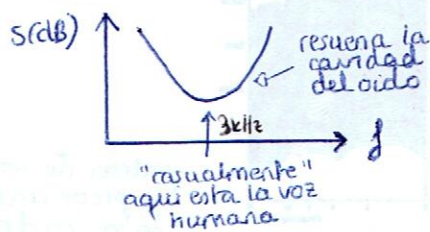


Sesión 1. Psicoacústica de la percepción espacial del sonido

Introducción

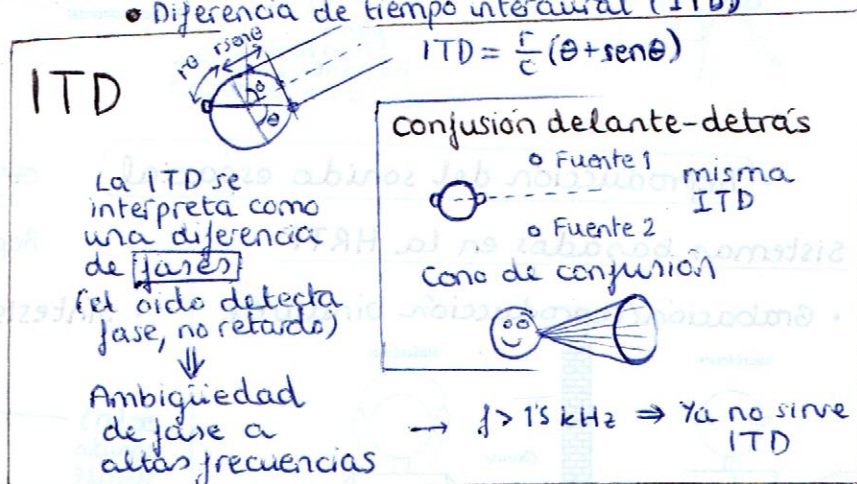
- El modelo cerebral del oído es mas complejo que el visual
→ localiza con rapidez y precisión el origen del sonido

- Sensibilidad del oído

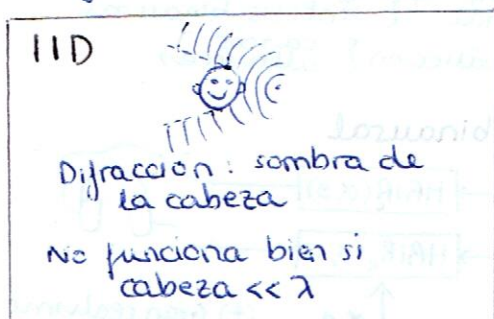


Localización espacial en el plano horizontal

- Diferencia de tiempo interaural (ITD)



- Diferencia de intensidad



Mecanismo combinado:

→ transición gradual

{	$< 1.5 \text{ kHz}$	→ ITD (detecta desfase)
	$> 1.5 \text{ kHz}$	→ IID (difracción cabeza)

Localización espacial en elevación

- Filtrado de las orejas y hombro según la dirección
→ ecos con retardos

Ambos oídos están en el mismo plano (no se puede detectar desfases en elevación)

Percepción de la distancia

→ Intensidad:

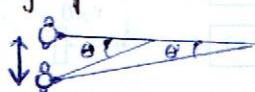
$$A_t = \frac{P_t}{P_r}$$

física: propagación atenua: el cerebro deduce la distancia
inteligencia: según el tipo de sonido (voz humana, pájaro, ...) el cerebro estima P_t

→ Atenuación de altas frecuencias:

- el aire tiene mas 'rozamiento' a frecuencias altas
- según la relación entre faltas y bajas el cerebro deduce la distancia

→ Paraleje por movimiento



se mueve ligeramente la cabeza y se compara el ángulo (típico para encontrar a un grillo)
(casi es como se mide la distancia de las estrellas)

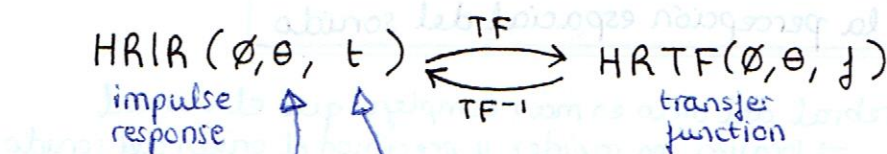
→ Exceso de diferencia de IID: para objetos muy cercanos → ej: mosquito en el oído

→ Relación entre sonido directo y reverberación: (en salas)



la energía del sonido directo depende de la distancia
la energía TOTAL de los ecos es mas o menos igual en todos los puntos de la sala

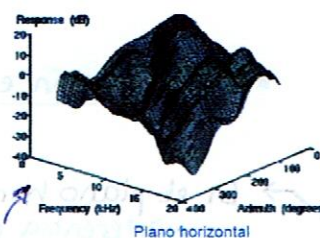
Función de transferencia relacionada con la cabeza (HRTF)



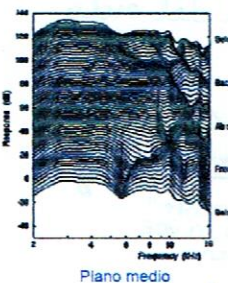
una distinta para cada oído!!
(nadie es simétrico!)

Medida:

Fuente en ϕ, θ que emita un pulso $\delta(t)$



plano horizontal bastante parecido entre personas



zona suave (solo hombres)

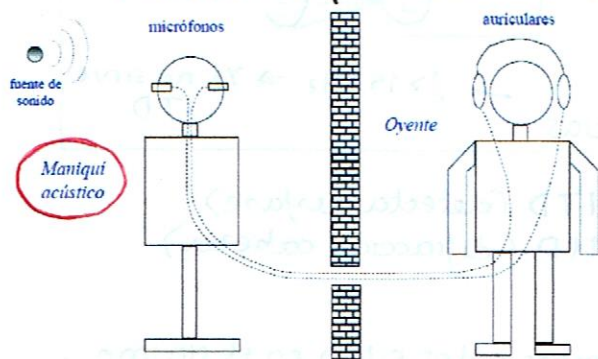
patrón de valles y picos único para cada persona

Reproducción del sonido espacial

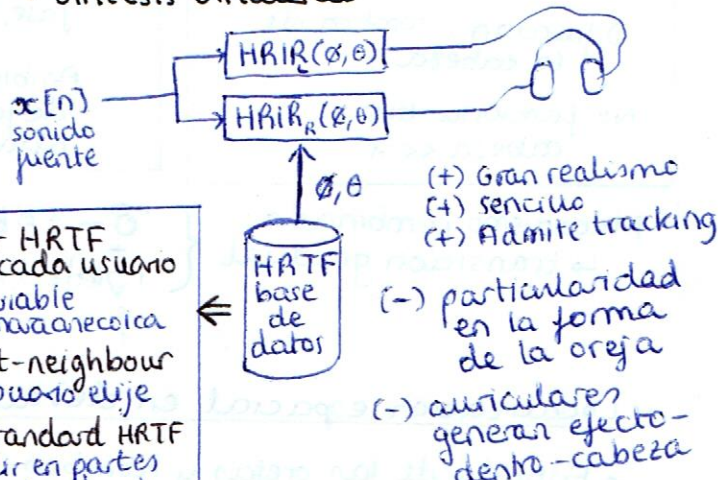
Obtener la señal {
- grabación binaural
- síntesis binaural
Reproducción {
- cascos
- altavoces

Sistemas basados en la HRTF

Grabación/reproducción binaural



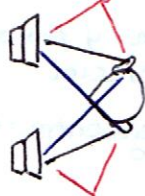
Síntesis binaural



- Medir HRTF para cada usuario
 - inviable
 - cansada
- Nearest-neighbour
 - el usuario elige
- Scale standard HRTF
 - dividir en partes (pina, head, ...)
 - para ir ajustando
- Adaptive model

Reproducción binaural con altavoces:

Problema
- crosstalk
- reflexiones

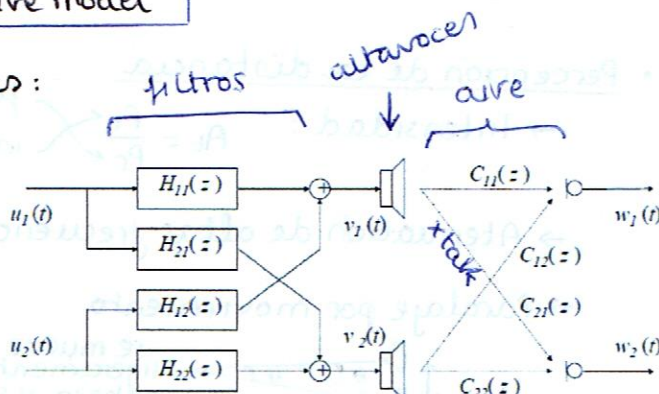


Solución: cancelar los caminos cruzados

→ sólo sirve si el usuario no se mueve del sweet spot:
estudios:

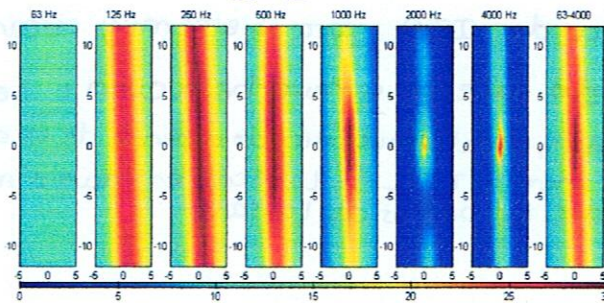
- teóricos en espacio libre
- medidos

altavoces



$$H = (C^T C)^{-1} C^T A$$

moverse delante/detrás no afecta mucho

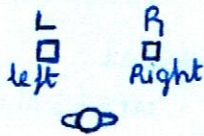


Degradación (relación directo/interferente) en función del desplazamiento

En 2kHz hay singularidad debido a distancia entre oídos → único sweet spot

EVOLUCIÓN DEL SONIDO ENVOLVENTE

Estéreo



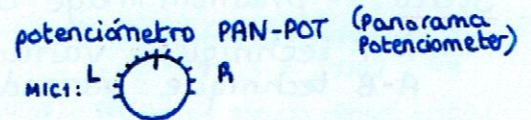
Efecto phantom:

• Por diferencia de volumen el cerebro interpreta una posición entre L y R

• Problema: no hay diferencia entre tiempos de llegada

• Solución: el cerebro se acostumbra

• Problema 2: sólo funciona si estás centrado

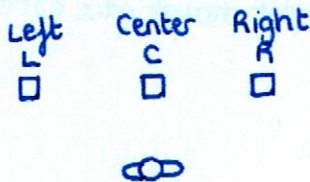


$$\frac{x}{D} = \frac{y}{D} \Rightarrow x=y$$

$$V_L = y \cdot V$$

$$V_R = x \cdot V$$

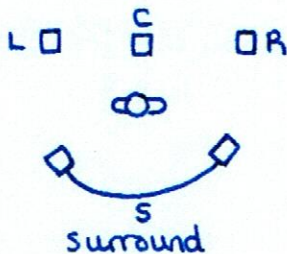
Canal central para diálogo



Los espectadores no centrados se desorientaban al oír el diálogo a los lados.

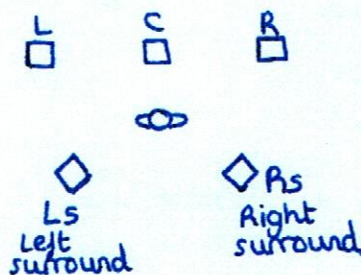
Canal central para el diálogo

Canal surround



- Llena el ambiente
- Simula reverberación
- Efectos/sustos

Desdoblar canal surround



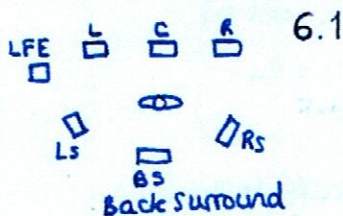
Permite el efecto phantom detrás

LFE: low frequency enhancement

usa gran altavoz (subwoofer) para frecuencias [20Hz, 120Hz]

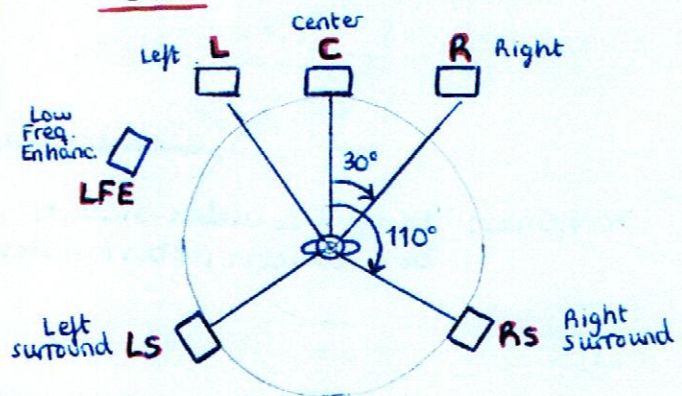
No requiere posicionamiento ya que apenas hay diferencia de fase y volumen entre los oídos a bajas frecuencias ($\lambda \gg d \rightarrow$ difracción de cabeza $\leftrightarrow$)

Tercer canal de surround

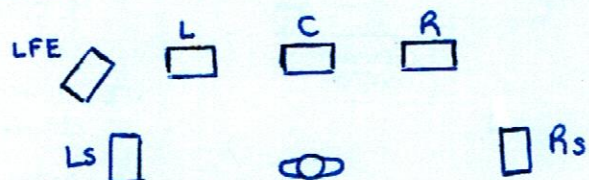


Desdoblar el Back surround → 7.1

5.1 (estándar ITU)



7.1



Otro 7.1 con 5 delante



- más tolerante a la posición del usuario
- Poco exito
- compatible con 5.1 (matriz de compatibilidad)



Stereo: - phantom image created

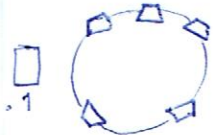
X-Y technique: variando intensidad ← Pan pot: el más usado

A-B technique: variando time-of-arrival

Stereo Sweet Spot:



Surround 5.1



→ pan-pot delante con 3 altavoces → sweet spot mayor que estéreo

→ localización lateral y detrás es pobre
7.1 lo mejora, pero no lo suficiente

Solución → Wave Field Synthesis



LFE: low frequency enhancement

no requiere posicionamiento
por que opera en una frecuencia
baja frecuencia [20Hz - 120Hz]
usa gran altavoz (subwoofer)
(LFE - diferencia de canales 4+)

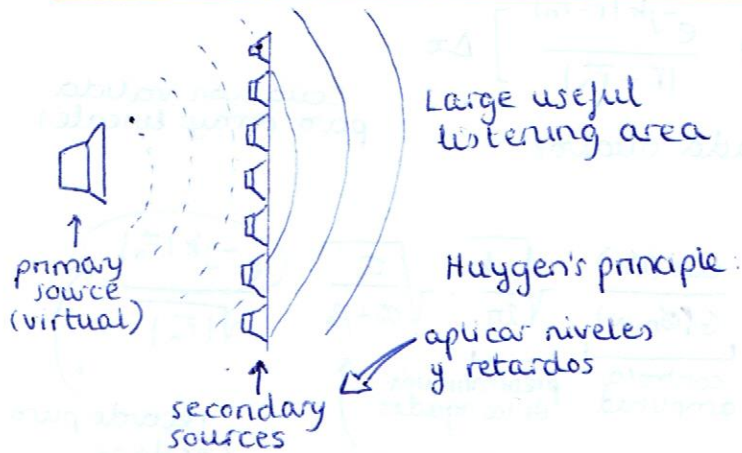


Desplazar el Back Surround →

Otro 7.1 con 5 delante
- más potente a la posición
- poco efecto
- compatible con 5.1
(más de compatibilidad)

SESIÓN 2. WAVE-FIELD SYNTHESIS

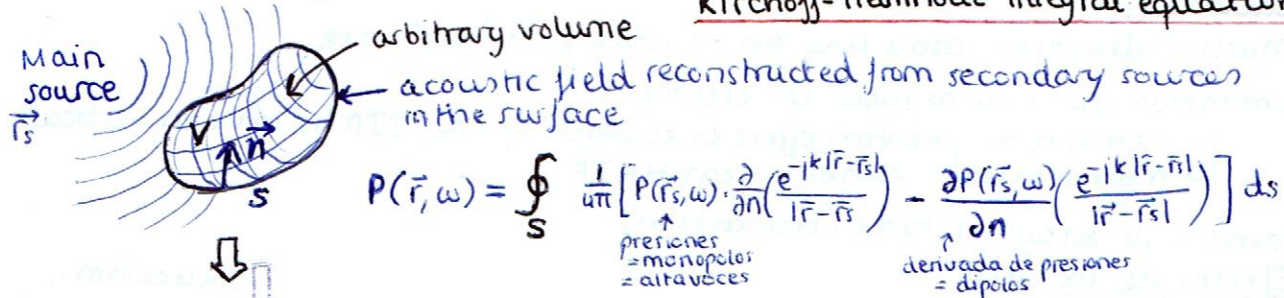
→ 2.5D (no tiene elevación)



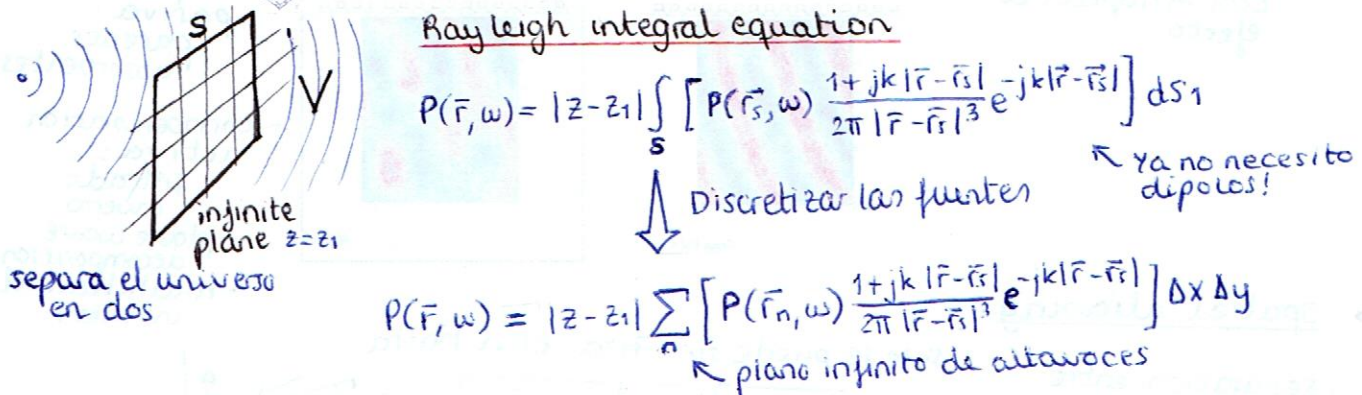
otra técnica aún mejor:
Ambisonics: descomposición del campo acústico en armónicos esféricos (con array de 32 micrófonos se capta hasta 4° orden)

WFS: Theory Background

Kirchoff-Helmholtz integral equation



Rayleigh integral equation



Simplification to a line array
 ↳ suponer que cada altavoz emite onda cilíndrica en lugar de esférica (no es cierto, pero el error es pequeño si estamos en el plano de los altavoces)

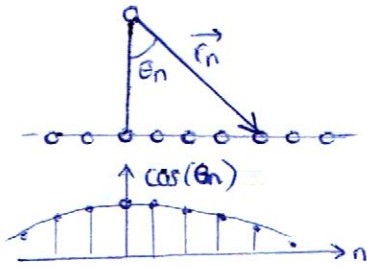
$$P(\vec{r}, \omega) = \frac{jk\rho c}{2\pi} \sum_n \left(u_n(\vec{r}_n, \omega) \frac{e^{-jk|\vec{r}-\vec{r}_n|}}{|\vec{r}-\vec{r}_n|} \right) \Delta x$$

Para N altavoces

$$P(\vec{r}, \omega) = \sum_{n=1}^N \left[Q(\vec{r}_n, \omega) \cdot G(\theta_n, \omega) \cdot \frac{e^{-jk|\vec{r}-\vec{r}_n|}}{|\vec{r}-\vec{r}_n|} \right] \Delta x$$

Excitación de cada altavoz

Ecuación válida para arrays lineales

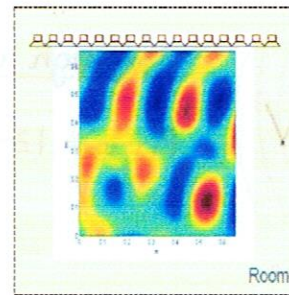
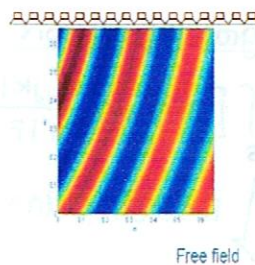


$$Q(\vec{r}_n, \omega) = \underbrace{S(\omega)}_{\text{señal mono}} \cdot \underbrace{\frac{\cos(\theta_n)}{G(\theta_n, \omega)}}_{\text{control de amplitud}} \cdot \underbrace{\sqrt{\frac{jk}{2\pi}}}_{\text{preacentuación de los agudos}} \cdot \underbrace{\sqrt{\frac{\sigma_0}{\sigma_0 + \rho_0}}}_{\text{error por no usar ondas esféricas}} \cdot \underbrace{\frac{e^{-jk|\vec{r}_n|}}{\sqrt{|\vec{r}_n|}}}_{\text{retardo puro (desfase variable con la frecuencia)}}$$

Limitaciones:

- Muchos altavoces: cada uno con su señal y su hardware
- Limitación: no se puede jugar con elevación?
 - ↳ ¿y si hacemos phantom effect en elevación? → No, IID no existe en vertical
 - ↳ mejor solución: combinar con HRTF
- Tamaño de array limitado (truncation)
- Efectos de la sala

Ecos estropean el efecto



Soluciones:

- compensación pasiva: paredes absorbentes
- compensación activa: filtrado inverso
- plane wave decomposition
- multichannel inversion

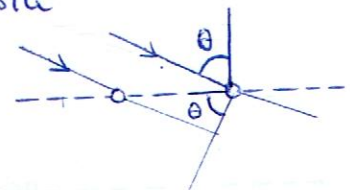
→ spatial aliasing

separación entre altavoces

solo se puede sintetizar bien hasta

$$\Delta x \cos \theta_{\max} \leq \frac{\lambda_a}{2}$$

$$f_a = \frac{c}{2\Delta x \cos \theta_{\max}}$$



Soluciones:

- Ignorar el problema → en el ITEAM $\Delta x = 18\text{cm}$, $f_a = 1\text{kHz}$ la fase errónea solo se nota si corres lateralmente y mirando hacia el array
- usar mas altavoces agudos
- OPSI (Optimized Phantom Source Image)
 - ↳ usar WFS para $f < f_a$
 - ↳ usar efecto phantom para $f > f_a$
- usar DML (Distributed Mode Loudspeaker)
 - ↳ membrana que se mueve entera con MAP (multi actuator panels)

WFS in GTAC

- array de 32 (4x8)
- array de 96 (12x8) → de los mas grandes de Europa
- MAP(DML) → 4 paneles x 6 excitadores

Applications of WFS

- Telepresencia
- Realidad virtual
- Reproducción de música con fidelidad espacial
- Planetarium, IMAX, ... (short term)
- movie theaters (medium term)
- home cinema (long term)

Research lines for WFS

- room compensation
- Discrete Time Modelling (FDTD)
- Authoring & Real Time Tools
- High-power applications
- Improvements in DMLs

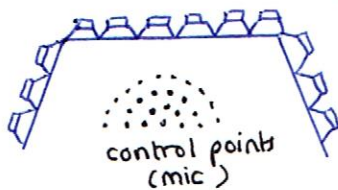
Compensación de la sala

- Plane-wave decomposition →



emitir los ecos "al revés"

- MIMO inverse filters



- measure impulse response between each loudspeaker and each control point.
- obtain bank of inverse filters
- Simulate sound field

corrección multipunto → se han obtenido buenas correcciones en un gran área

Simulación FDTD

- Micrófonos virtuales para obtener RIR
- Superordenadores

High-power applications

- Diameter of loudspeakers
- Directivity

Experiments & Simulations:

- Field analysis
- Precise aliasing Frequencies

Authoring tools for WFS

- standalone application
- VST plugin that works with 'Cubase'

Software Block Diagram

SESIÓN 3. SOUND SOURCE SEPARATION (SSS)

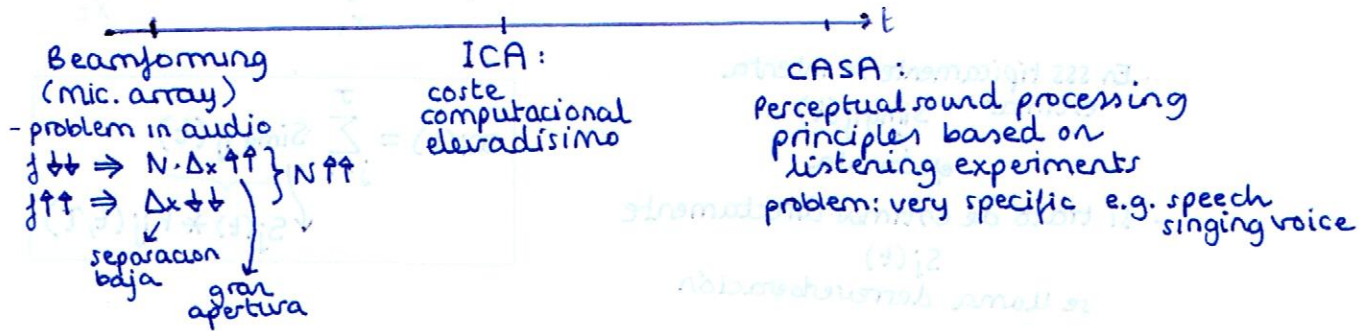
INTRODUCTION

- Recovering each source signal from a given mixture
- Difficult → active research

Applications:

- remixing
- speech enhancement (e.g. hearing devices)
- speech recognition
 - ↳ singer recognition and melody extraction
- Wave-Field-Synthesis

Evolution:



Sparse methods:

Idea: baja probabilidad de que fuentes distintas sean simultáneas en tiempo o frecuencia o ambos

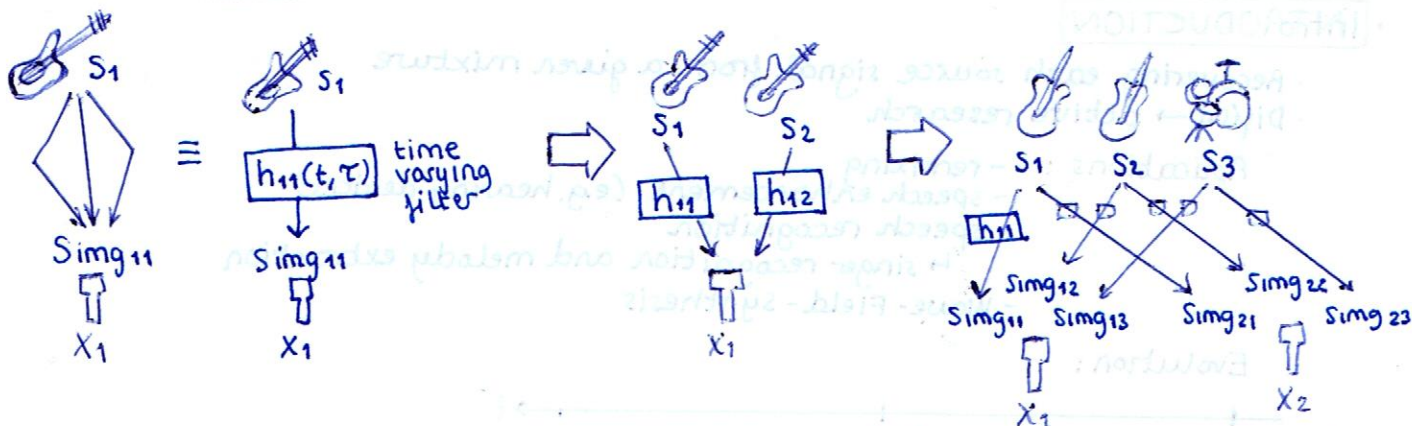
↳ time/frequency masking

MIXTURES

→ Audio Sources

- **VOZ**: secuencia de fonemas
 - ↳ cuerdas vocales (pitch fundamental + armónicos)
 - ↳ ruidos de susurro (también tiene formantes)
 - ↳ transitorios (labios)
- **Music**: musical instruments + singers
 - ↳ tonos/notas = harmonic partials → acordes = mezcla notas
 - ↳ transient signals (ej: tambor)
- **Sonidos naturales**: muy distintos
 - ej: pito de coche
 - ej: lluvia
- **Studio recordings**: multitrack + downmix recording
- **Live recordings**: Puede intentarse la separación con micrófonos
 - ej: 2 micrófonos X-Y → apuntando ↙ ↘ cardioides en 8
 - ej: 2 micrófonos spaced → sencillos y rápidos
- **Synthetic mixtures**: tabla de mezclas

The mixing process



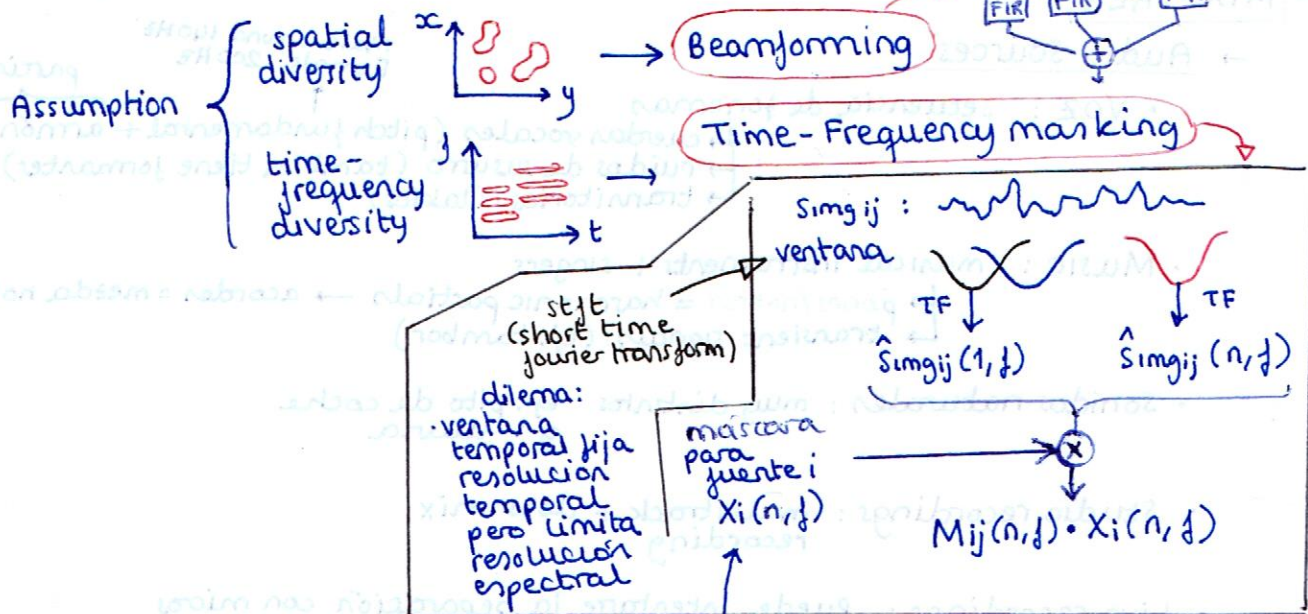
- En SSS típicamente se intenta estimar $S_{ij}(t)$ efecto sala
- si trato de estimar directamente $S_j(t)$ se llama derreverberación

$$x_i(t) = \sum_{j=1}^J S_{ij}(t) = \sum_{j=1}^J S_j(t) * h_{ij}(t, \tau)$$

BLIND SOURCE SEPARATION

separar haciendo el menor número de suposiciones posible

FILTERING TECHNIQUES



problema: estimar una máscara

→ MULTI-CHANNEL SEPARATION → You can exploit spatial clues

III - 2

• ICA (Independent Component Analysis)

Assumption: the sources are statistically independent

↳ minimize mutual information between estimated source signals

↳ find optimal demixing matrix (MAP)

Mirar estadísticos de alto orden (>2)
(más allá de la varianza)

ejemplo:

- una gaussiana da cero
- suma de 2 gaussianas no da cero, pero si las separo cada una da cero

• ADRes (Azimuth Discrimination and Resynthesis)

Idea stereo: $\left\{ \begin{array}{l} \text{Left } l(t) = \text{guitarra} + \text{voz} \\ \text{Right } r(t) = \text{guitarra} + \text{voz} \end{array} \right.$

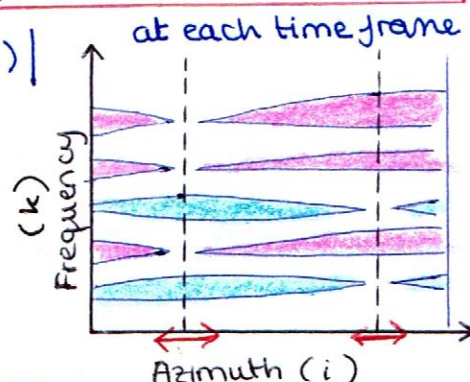
• $l(t) - g_j r(t) \Rightarrow$ ~~causes the cancellation of the j-th source in the left channel~~

• $r(t) - g_j l(t)$ causes the cancellation of the j-th source in the right channel

We know how to cancel a source, but how can we recover it?

① $Az_r(k, i) = |L(k) - \underbrace{\frac{i}{\beta}}_{g(i)} R(k)|$
 $i \in [0, \beta]$

y lo mismo con canal left
 $Az_l(k, i)$



② Convertir nulos en picos

Buscar picos en un azimuth subspace width

↳ ya sabemos a qué frecuencias está cada fuente en un determinado time frame

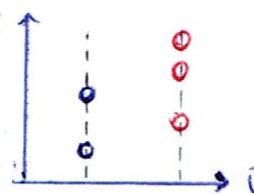
$$Az'(k, i) = \begin{cases} Az(k)_{\max} - Az(k)_{\min} & \text{if } Az(k, i) = Az(k)_{\min} \\ 0 & \text{resto} \end{cases}$$

③ señal separada:

$$Y_R(k) = \sum_i Az_r'(k, i)$$

$$Y_L(k) = \sum_i Az_l'(k, i)$$

↑
módulo
(la fase se cogerá del original $R(k), L(k)$)



(+) Good quality of separation

(+) Works with stereo → casi todo está en estéreo

(+) Very simple

Problems: - sources with same pan position are extracted together
- musical noise in extracted sources

DUET (Degenerate Unmixing Estimation Technique)

Assumption: only one source is present at any time-frequency point (W-disjoint orthogonality)

It is possible to determine the source in each point from the spatial location information

Interchannel Intensity Difference

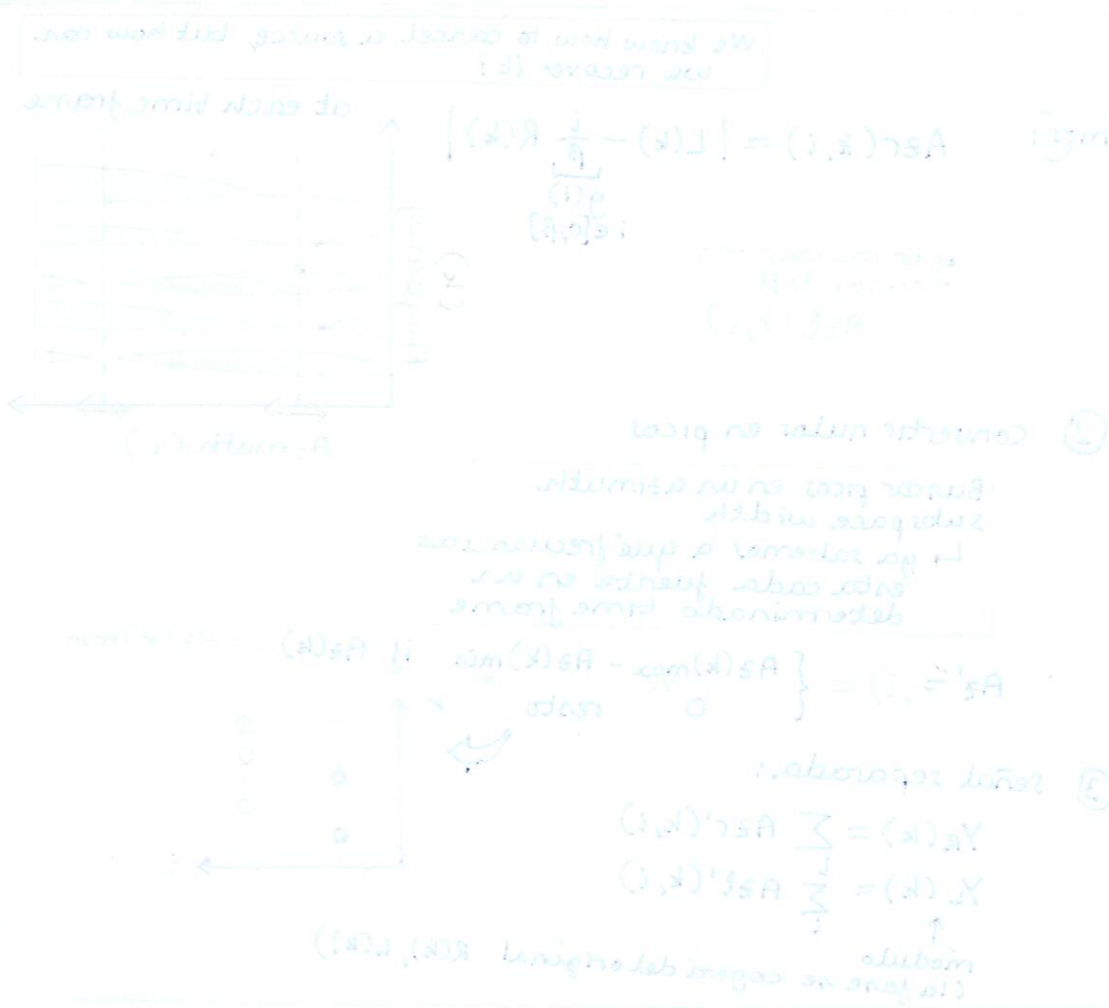
$$IID(n, f) = 20 \log \frac{X_2(n, f)}{X_1(n, f)}$$

Interchannel Phase Difference

$$IPD(n, f) = \angle \left[\frac{X_2(n, f)}{X_1(n, f)} \right]$$

Problems:

- Rarely holds
- musical noise artifacts
- it cannot deal with convolutive mixers where mixing delay is larger than one sample



Problems: - sources with same position are extracted together
- musical noise in extracted sources

→ SINGLE CHANNEL SEPARATION

- Spatial clues cannot be included
- more challenging problem

• Hidden Markov Models (HMM)

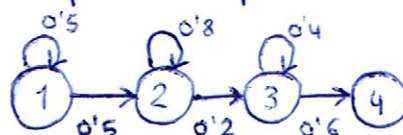
Markov model: states correspond to random processes whose outcomes are observable

HMM: the states are not observable directly but rather have a certain probability distribution of the observation

- In speech:
- observations are significant speech parameters
 - Each state models a particular phoneme or particular chord

e.g.
HMM modelling
a vowel

e.g.
HMM modelling a
vowel in a singing
voice has 10 states



↑ probability of changing state

Three probabilities:

A: state transitions

B: distribution of observations for a given state

π : probability of initial state

steps:

- ①: Train individual HMM beforehand to get A, B, π
- ②: calculate state path (most likely succession of states) (e.g. Viterbi)
- ③: Each state corresponds to a short-term spectrum of the source

(+) Good separation of male and female voice

(+) Acceptable separation of singing voice

(-) Difficult to determine optimum number of states for each case

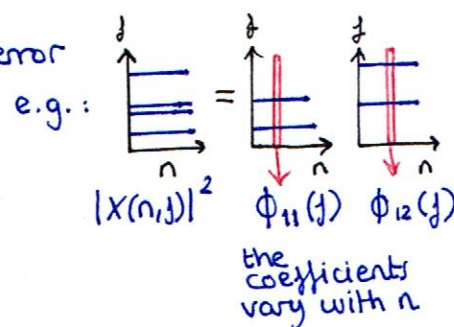
(-) Train beforehand with the solo voice

• Spectral decomposition

spectrum as weighted sum of normalised bases plus error

$$|X(n, f)|^2 = \left(\sum_j \sum_k e_{jk}(n) \cdot \phi_{jk}(f) \right) + \epsilon(n, f)$$

time varying
coefficients



↳ Non-negative matrix factorization

$$\underline{X} = \underline{B} \cdot \underline{G}$$

↑
↑

basis functions
time varying coefficients

• Estimate B and G (non-negative)

• once estimated, each source will correspond to set of basis spectra

Computational Auditory Stream Analysis (CASA)

→ Inspirado en la capacidad de los humanos para distinguir sonidos
 ↳ experimentos de audición → progresos durante 10 años

- ① Descomponer la muestra en componentes tiempo-frecuencia
- ② Agrupar los componentes en sus correspondientes fuentes

in speech: — observations are significant speech parameters
 — Each state models a particular phoneme or particular chunk



↑ probability of changing state

Three probabilities:
 A: state transition
 B: distribution of observation for a given state
 π: probability of initial state

HMM modelling a word in a sentence
 voice has to follow

①: Train individual HMMs beforehand for each A, B, π

②: Calculate state paths (most likely sequence of states) (e.g. Viterbi)
 ③: Evaluate corresponds to a short-term spectrum of the source

(+) Good separation of male and female voice
 (+) Accurate separation of singing voice
 (-) Difficult to determine optimum number of states for each case
 (-) Trains beforehand with the voice

Spectral decomposition

spectrum as weighted sum of normalized basis functions

$$|X(n)|^2 = \left(\sum_k c_k(n) \cdot \phi_k(n) \right)^2 + \epsilon(n)$$

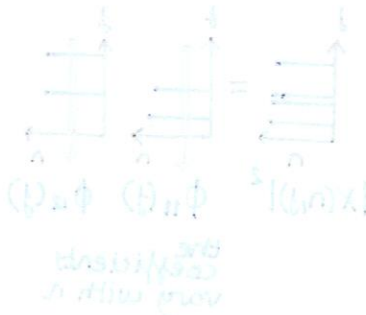
time varying coefficients

↳ Non-negative matrix factorization

$$X = B \cdot C$$

↑ basis function
 ↑ time varying coefficients

• Estimate Band G (non-negative)
 • once estimated, each source will correspond to sep. of basis vectors

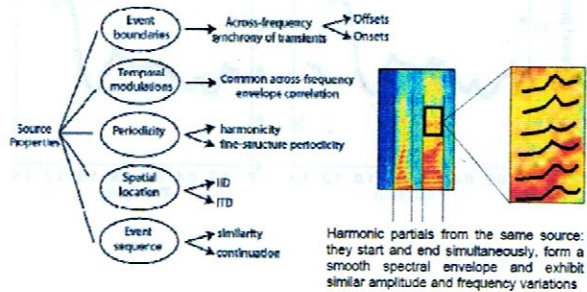


CASA : Resumen del Proceso

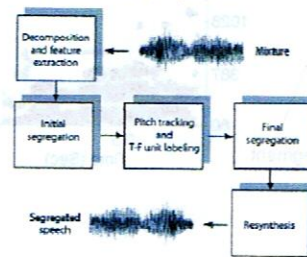
III - 4

10 years of progress

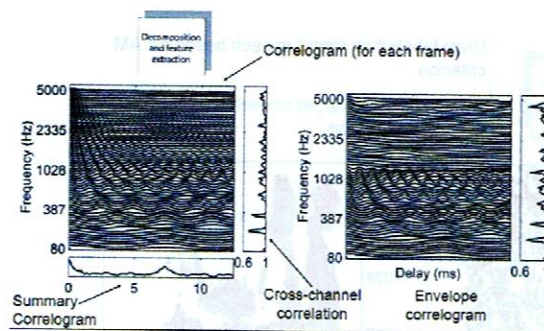
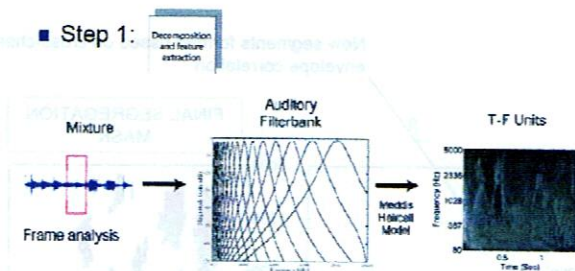
- 100 mixture set used in Cooke (1991)
- Speech + { white noise, tone, telephone, impulses, siren, music, speech, ... }
- Largely data-driven systems below
- Unrealistic corpus: mainly voiced speech + intrusion



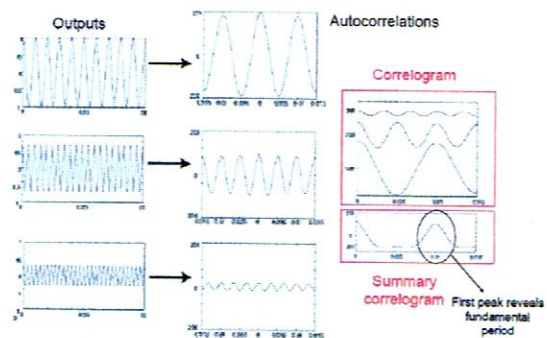
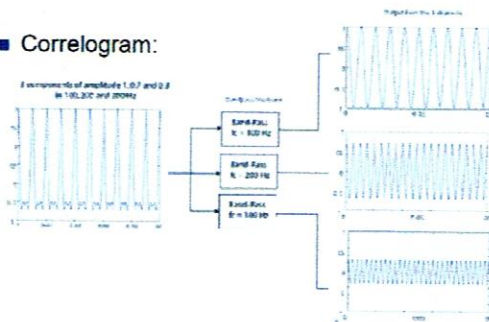
Monaural Speech Segregation:



Step 1:



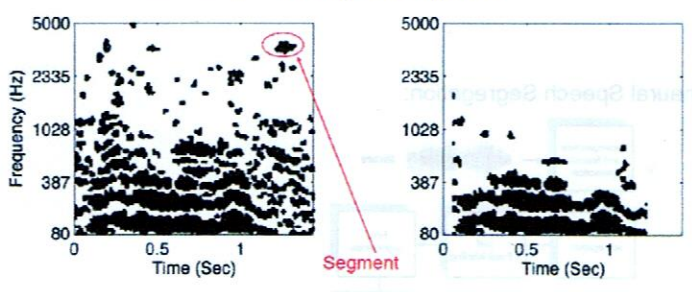
Correlogram:



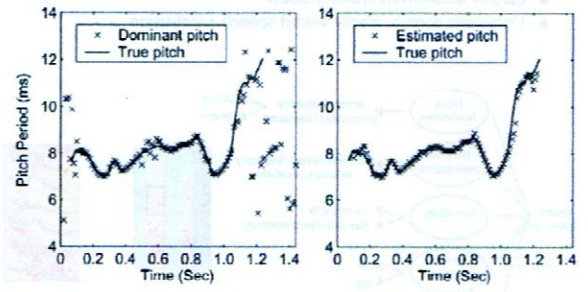
Step 2:

Initial segregation

If a T-F unit has correlagram at zero-lag and cross-channel correlation greater than a certain threshold, the unit is selected → Adjacent units are merged into segments



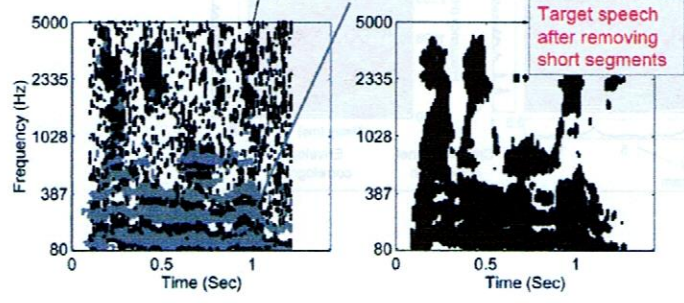
Pitch tracking → pitch is estimated in the plausible pitch range according to 2 constraints: agreement with the dominant pitch and slow pitch change.



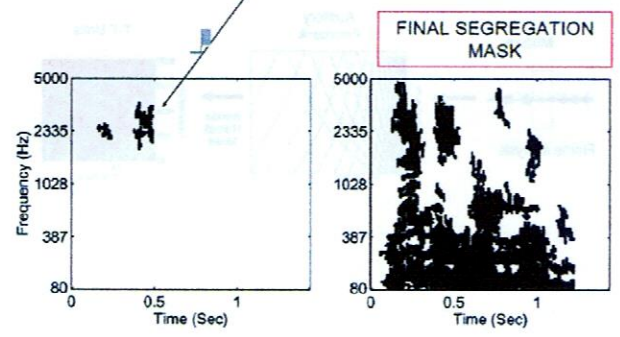
Step 3:

Pitch tracking and T-F unit labelling

Units labeled as target speech based on AM criterion
Units labeled as target speech based on periodicity criterion



New segments formed based on cross-channel envelope correlation



SESIÓN 4 : SPATIAL AUDIO CODING (SAO) SPATIAL AUDIO OBJECT CODING (SAOC)

Spatial Audio Coding (mp3 Surround)

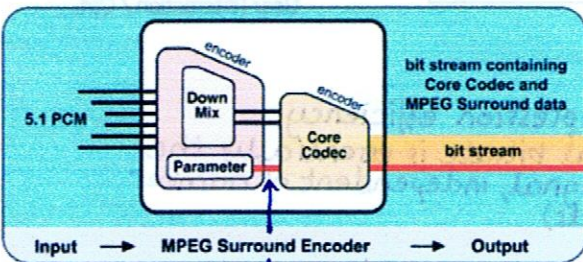
(MPEG Surround)
(MPS)

CODEC:

- Codifica multicanal en un stream stereo + info adicional

El downmix puede ser automático o externo (a gusto del artista)

↑ spatial parameters



↑ info adicional (reduced set of spatial parameters)

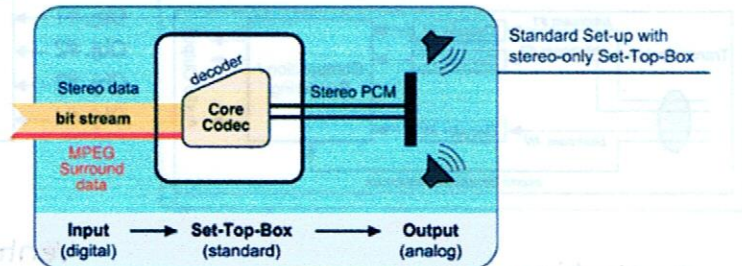
Ventajas y aplicaciones

- Tasa de bits muy cercana a estéreo → bueno para broadcast IPTV
 - Fácil de integrar por los broadcasters
 - Está en estándares de broadcast ej: DAB
- Dispositivos móviles:
 - surround binaural
- Music and video downloads
 - same music file plays on conventional stereo player

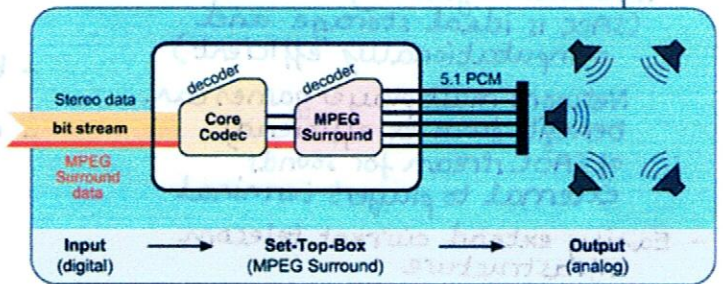
DECODEC

- Legacy decoder:

- ignora info adicional
- backward compatibility

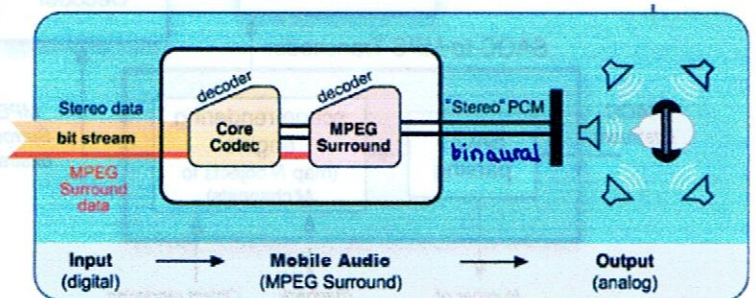


MPEG surround Decoder



Capacidad adicional En Sonido

- A partir del stream estéreo + parámetros genera 2 canales con HRTF para usar cascos (perfecto para dispositivo móvil!)



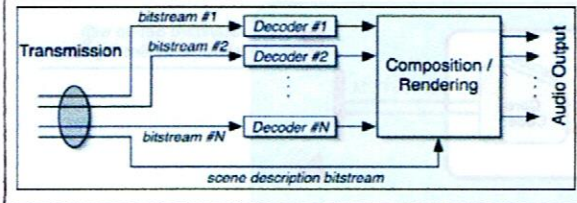
lo hace "directamente" sin pasar por 5.1

Spatial Audio Object Coding (SAOC)

Object based representation:

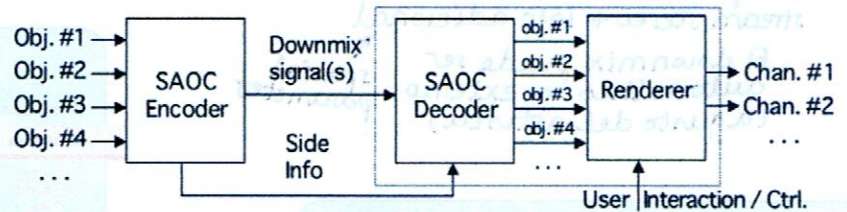
- Individual dry sources + scene description or user input

↓
Context based interactivity



SAOC busca lograr la misma funcionalidad que el Object Based Representation, de forma eficiente

- ex:
- stereo stream
 - side info → ayuda a separar los distintos dry sources



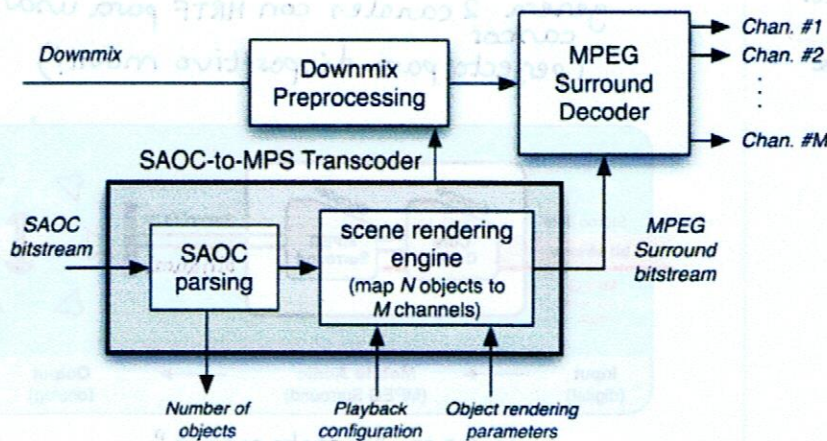
Applications:

- personal interactive remixes
- interactive gaming (SAOC is ideal storage and computationally efficient)
Network multiplayer games can benefit from tx efficiency of SAOC stream for sounds external to player's terminal
- Easily extend current telecom infrastructure

Ventajas:

- high compression efficiency (the total bitrate is essentially the stereo signal, independent of number of objects)
- backward compatibility
- el SAOC decoder + renderer pueden integrarse para evitar el paso intermedio de separación en objetos → reduce complejidad computacional

MPEG SAOC Architecture



Objective:

- Integrated Decoder to avoid complexity
- N objects to M channels
 $N > M$
by simple using an M-channel MPS decoder driven by parameters