



TV: TeleVisión – Plan 2010

AVC

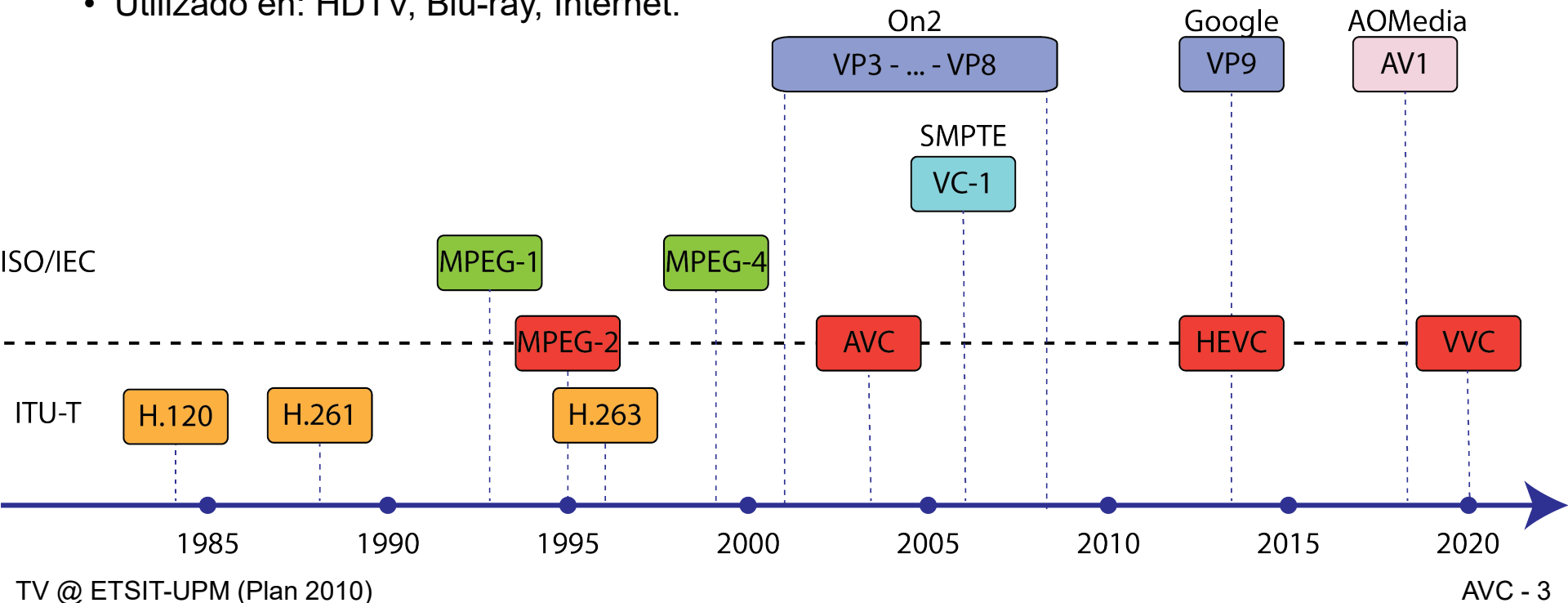
Contenido

1. Introducción
 1. ¿Qué es AVC/H.264?
 2. Aplicaciones
 3. Diferencias con estándares previos
2. Codificador/descodificador
3. Predicción
 1. Intra
 2. Inter
4. Imágenes SP y SI
5. Transformación y cuantificación
6. Filtro de reconstrucción
7. Codificación

1.1. ¿Qué es AVC/H.264?

MPEG-4 – parte 10 / H.264 / AVC (2003):

- JCT-VC: *Joint Collaborative Team on Video Coding*
 - MPEG (ISO + IEC): *Moving Picture Experts Group*
 - VCEG (ITU-T): *Video Coding Experts Group*
- Niveles de compresión muy superiores a sus predecesores.
- Utilizado en: HDTV, Blu-ray, Internet.



1.1. ¿Qué es AVC/H.264?

- AVC → *Advanced Video Coding*
- Es la una de las partes del estándar MPEG-4
- Su objetivo era mejorar la calidad/compresión de los codificadores previos:
 - Incorpora estrategias que habían sido excluidas de codificadores anteriores debido a su complejidad y su alto coste de implementación.
 - Utiliza las mismas etapas clave de codificadores anteriores: predicción, transformación y cuantificación.
 - Para evitar restricciones y dar más libertad a los algoritmos no contempla mantener la compatibilidad con codificadores anteriores.

1.1. ¿Qué es AVC/H.264?

- **MPEG-4:**

- Parte 1 → Sistema (1ª edición en 1999).
- Parte 2 → Visual (1999).
- Parte 3 → Audio (1999).
- Parte 4 → Pruebas de conformidad (2000).
- Parte 5 → Software de referencia (2000).
- ...
- **Parte 10 → AVC (2003) →** Ha seguido en desarrollo (última versión: 14.0 – 22/08/2021)
- ...
- Parte 30 → Timed text and other visual overlays (2014).

1.2. Aplicaciones

- AVC es un estándar que ofrece una gran flexibilidad, tanto en opciones de compresión como de transmisión.
- Es utilizado en numerosas aplicaciones de transporte y almacenamiento:
 - DVDs de alta definición (HD-DVD y Blu-Ray).
 - Televisión de alta definición en Europa (HDTV).
 - Numerosos productos de Apple: itunes, iPod, MacOS...
 - Televisión para móviles.
 - Vídeo por Internet.
 - Videoconferencias.
 - Etc.

1.3. Diferencias con estándares previos

- Predicción:
 - Es mucho más flexible que en estándares anteriores.
 - Modo intra con predicción (bloques de 16×16 o 4×4).
 - Modo inter con macrobloques de múltiples tamaños: desde 16×16 hasta 4×4 .
 - Las imágenes tipo B pueden ser usadas como predicción para otras imágenes.
 - Nuevos tipos de imágenes (SP y SI).
 - Precisión de $1/4$ de píxel.
- Transformación y cuantificación:
 - Se aplica sobre bloques de 8×8 o 4×4 .
 - Nuevos modos para recorrer las matrices de coeficientes.
- Codificación:
 - Ofrece la posibilidad de utilizar codificación de longitud variable (VLC) o codificación aritmética (CABAC).

1.3. Diferencias con estándares previos

- Su capacidad de compresión es muy superior a la de estándares previos:
 - Mejor calidad para una tasa binaria dada.
 - Menor tasa binaria para una calidad dada.
- El precio a pagar es un mayor coste computacional.



MPEG-2



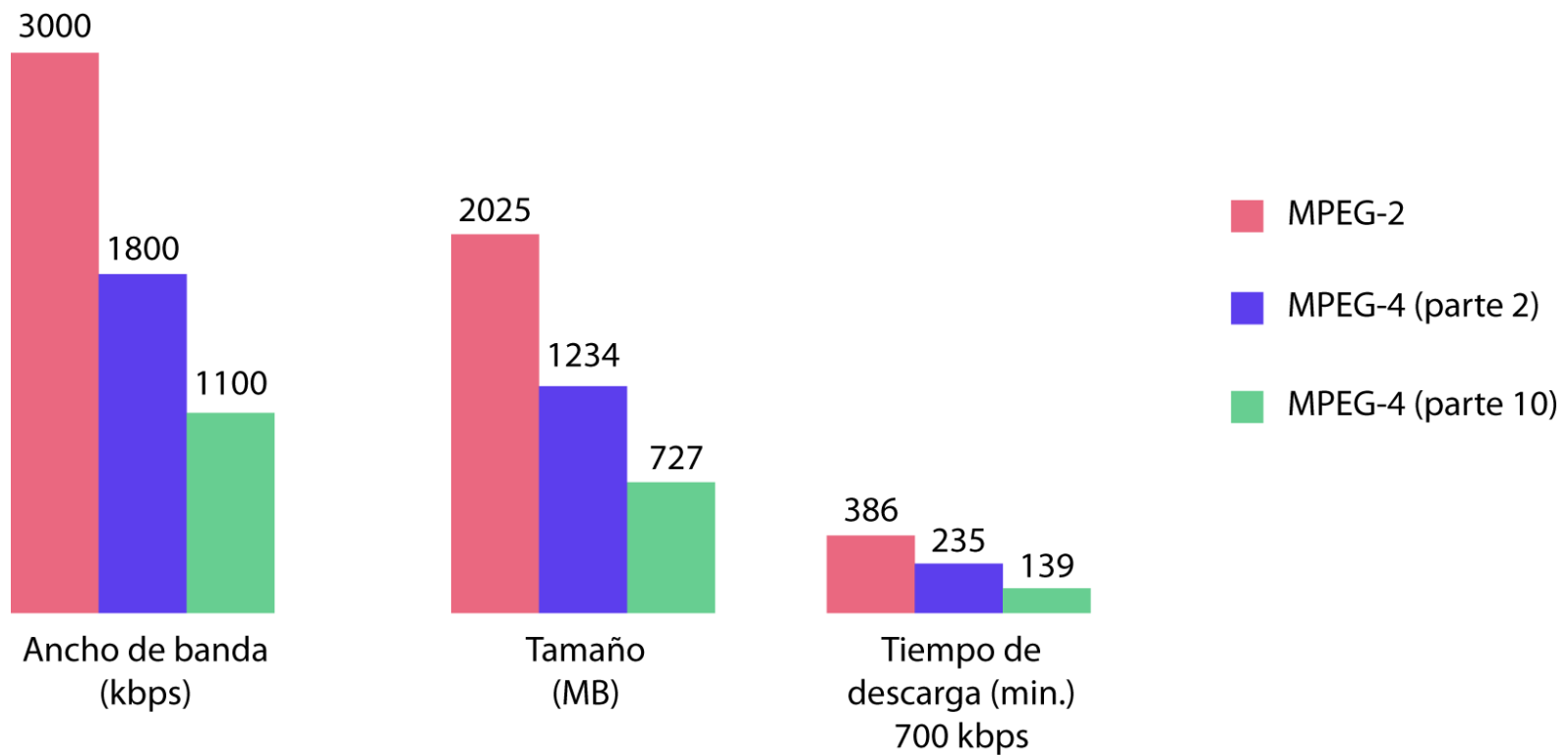
MPEG-4 (parte 2)



MPEG-4 (parte 10)

1.3. Diferencias con estándares previos

Comparación para una película de 90 minutos con calidad DVD



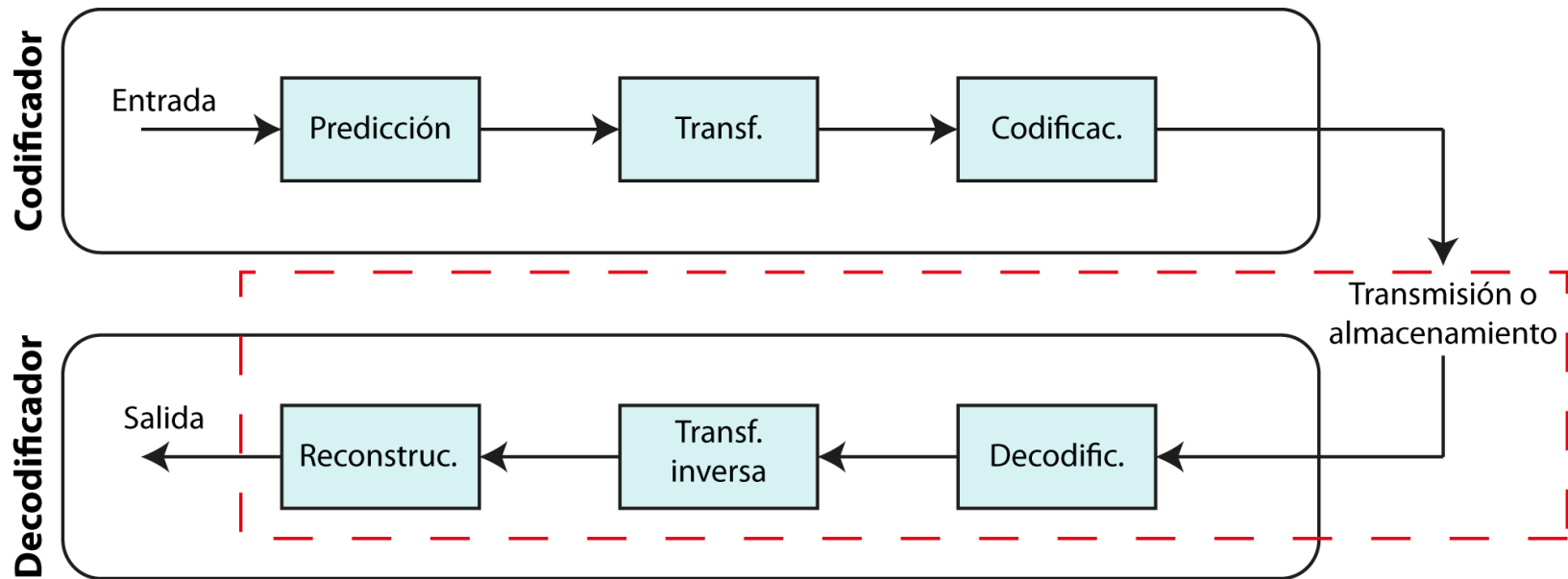
1.3. Diferencias con estándares previos

- **Características principales:**

- La señal de entrada puede ser tanto progresiva como entrelazada.
- Permite 3 esquemas de muestreo de crominancia:
 - 4:2:0
 - 4:2:2
 - 4:4:4
- Permite utilizar 2 esquemas de color:
 - YCrCb
 - RGB
- Es compatible con numerosas resoluciones, velocidades de refresco y velocidades binarias.
 - Desde QCIF (176×144) hasta UHD (8k - 7680×4320).
 - Desde 15 fps hasta 60 fps.
 - Hasta 240 Mbps.

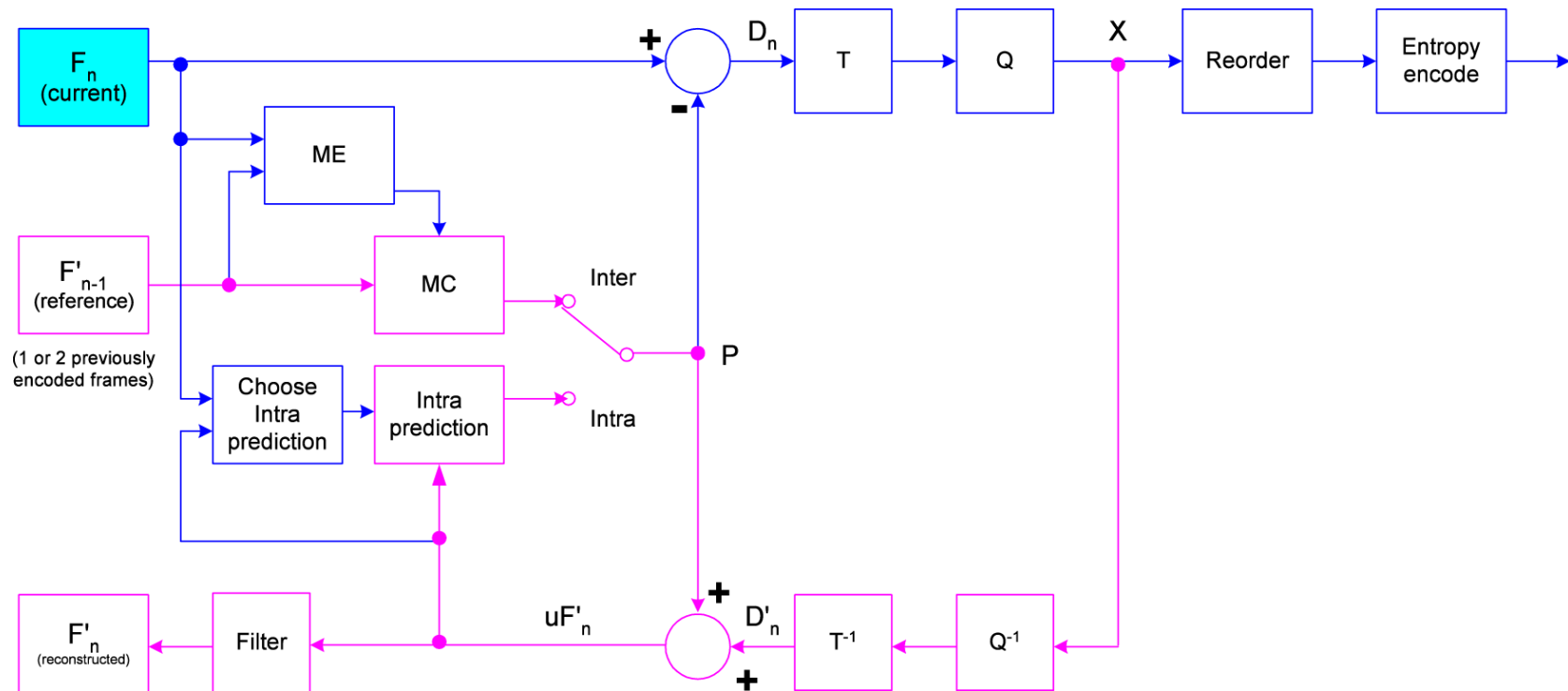
2. Codificador/decodificador

- Al igual que en los estándares previos, no se definen los detalles del codificador, sino que únicamente se define la sintaxis de los datos comprimidos y los métodos para decodificarlos.



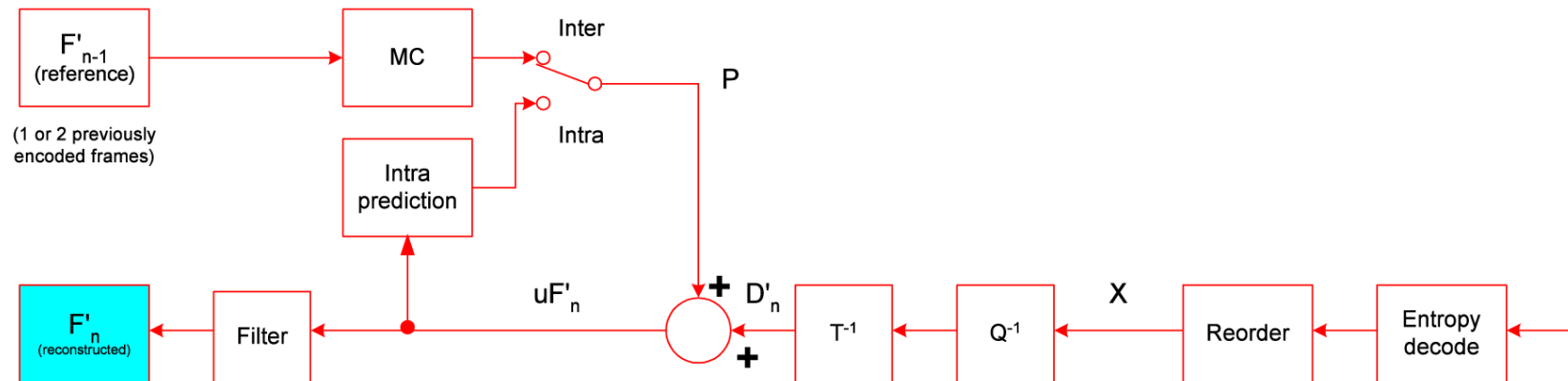
2. Codificador/decodificador

- Codificador:**



2. Codificador/decodificador

- Decodificador:**



3.1. Predicción intra

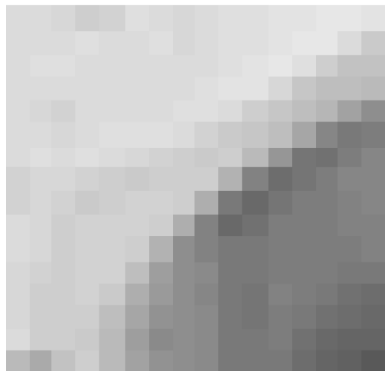
- Se puede utilizar en imágenes I, P y B.
- Un macrobloque codificado en modo intra puede utilizar como predicción otros macrobloques de la misma imagen.
- Los macrobloques de referencia son macrobloques previamente codificados y reconstruidos, pero **sin haber pasado por el filtro de bucle**.
- Sólo se pueden usar **datos de referencia también codificados en modo intra** y pertenecientes a la misma tira → Se reduce la propagación de errores.
- Para el caso de los macrobloques de luminancia hay dos opciones:
 - Predicción a partir de bloques de 4×4 muestras → 9 posibles modos de predicción.
 - Predicción a partir de bloques de 16×16 muestras → 4 posibles modos de predicción.
- Para el caso de los macrobloques de crominancia:
 - 4 posibles modos de predicción que se aplican siempre sobre bloques de 8×8 muestras.

3.1. Predicción intra

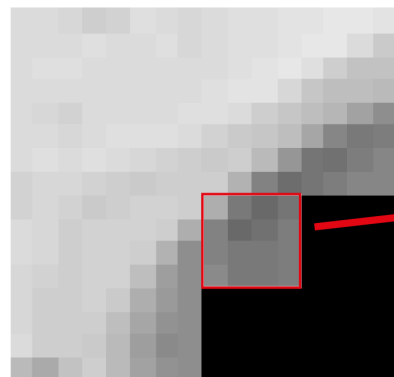
- Luminancia – Modos de predicción para bloques de 4×4 muestras:**



Original MB



Block to be predicted



- Los datos de referencia son muestras vecinas de bloques previamente codificados y reconstruidos.
- Estos datos siempre corresponden a bloques situados por encima y a la izquierda del bloque a predecir.

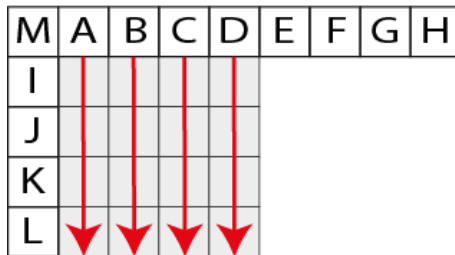
M	A	B	C	D	E	F	G	H
I	a	b	c	d				
J	e	f	g	h				
K	i	j	k	l				
L	m	n	o	p				



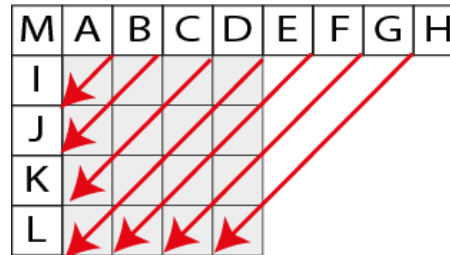
3.1. Predicción intra

- Luminancia – Modos de predicción para bloques de 4x4 muestras:

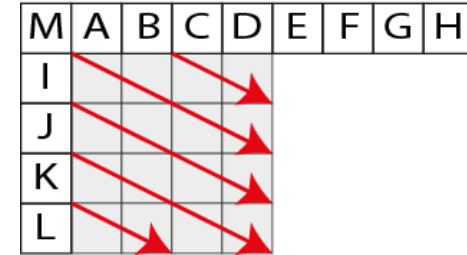
Modo 0 (vertical)



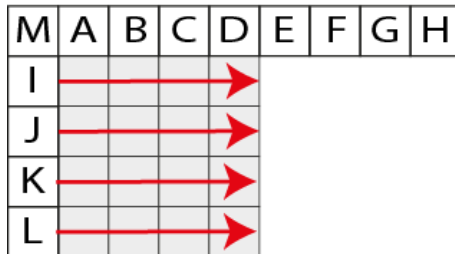
Modo 3 (diagonal down-left)



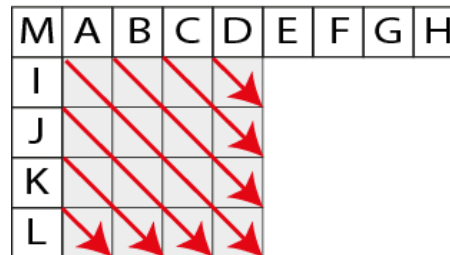
Modo 6 (horizontal-down)



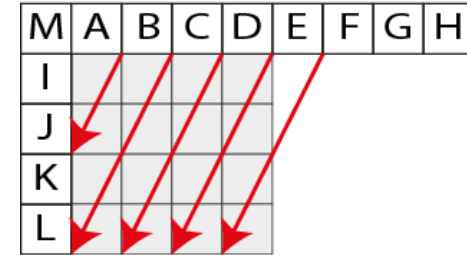
Modo 1 (horizontal)



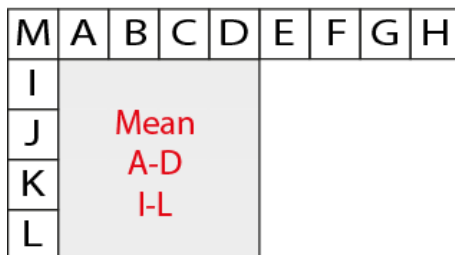
Modo 4 (diagonal down-right)



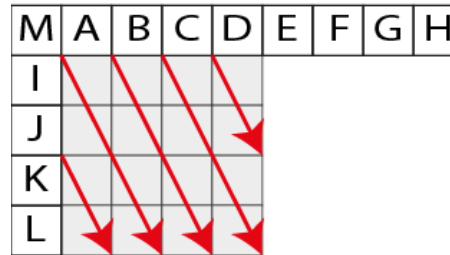
Modo 7 (vertical-left)



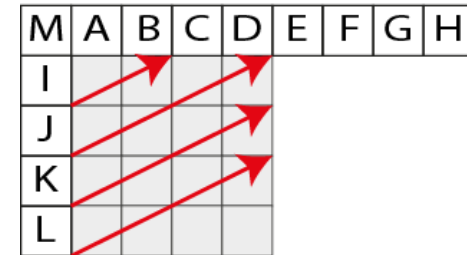
Modo 2 (DC)



Modo 5 (vertical-right)



Modo 8 (horizontal-up)



3.1. Predicción intra

- Luminancia – Modos de predicción para bloques de 4x4 muestras:

- Ejemplo – Modo 4:

Modo 4 (diagonal down-right)

M	A	B	C	D	E	F	G	H
I								
J								
K								
L								

M	A	B	C	D	E	F	G	H
I	a	b	c	d				
J	e	f	g	h				
K	i	j	k	l				
L	m	n	o	p				

Se requiere disponer de las muestras:

- A, B, C, D
- I, J, K, L
- M

$$a = \frac{I}{4} + \frac{M}{2} + \frac{A}{4}$$

$$d = \frac{B}{4} + \frac{C}{2} + \frac{D}{4}$$

3.1. Predicción intra

- Luminancia – Modos de predicción para bloques de 4x4 muestras:

- Ejemplo – Modo 8:

Modo 8 (horizontal-up)

M	A	B	C	D	E	F	G	H
I								
J								
K								
L								

M	A	B	C	D	E	F	G	H
I	a	b	c	d				
J	e	f	g	h				
K	i	j	k	l				
L	m	n	o	p				

Se requiere disponer de las muestras:

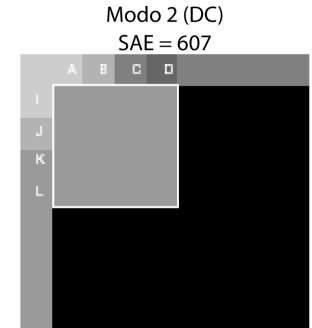
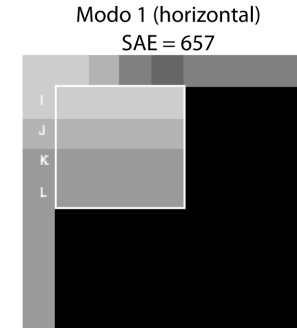
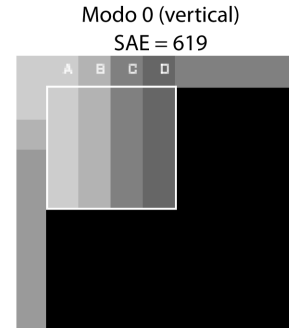
- I, J, K, L

$$a = \frac{I}{2} + \frac{J}{2}$$

$$d = \frac{J}{4} + \frac{K}{2} + \frac{L}{4}$$

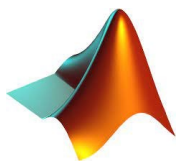
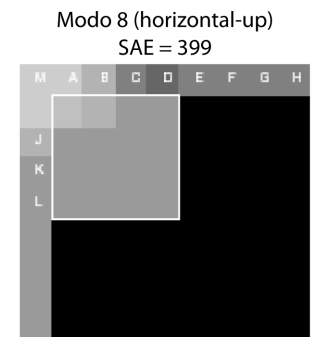
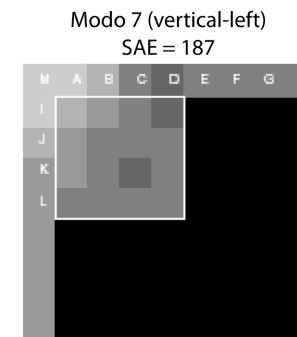
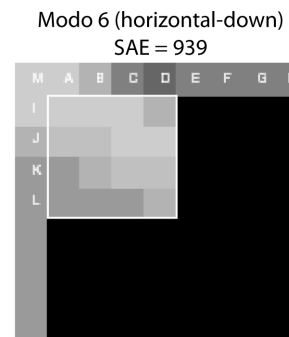
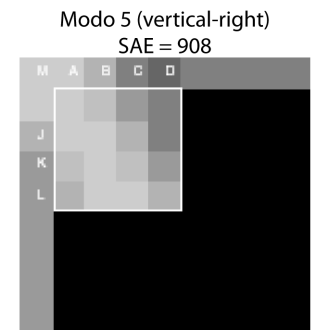
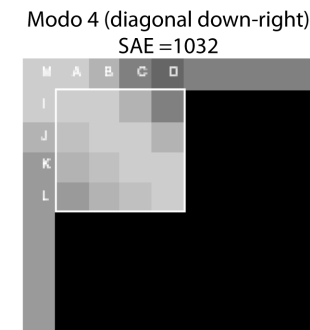
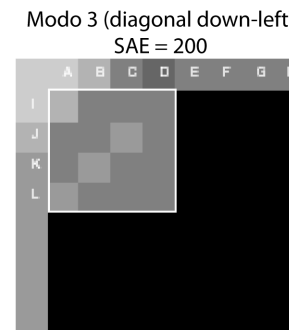
3.1. Predicción intra

- Luminancia – Modos de predicción para bloques de 4x4 muestras:**



- Se selecciona el modo que minimice el SAE (*Sum of Absotute Error*).

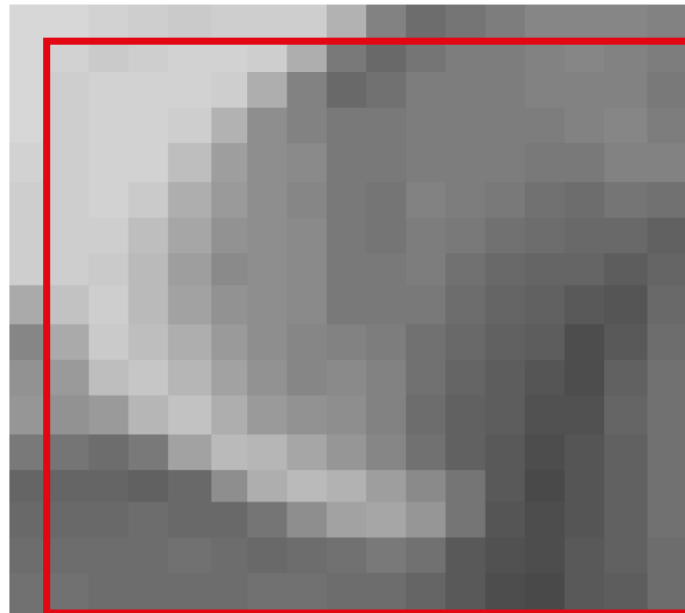
- En el ejemplo mostrado el modo que minimiza dicho error es el 7.



3.1. Predicción intra

- **Luminancia – Modos de predicción para bloques de 16x16 muestras:**

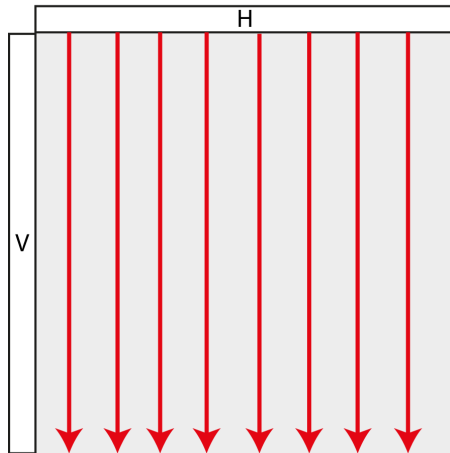
- Se consideran las 16x16 muestras de cada macrobloque.
- La forma de operar es similar a la del caso de bloques 4x4:
 - Se obtienen predicciones a partir de muestras situadas inmediatamente sobre al macrobloque o a su izquierda.
 - Estas muestras han debido ser codificadas y reconstruidas previamente.
- En este caso sólo hay 4 posibles modos de predicción.



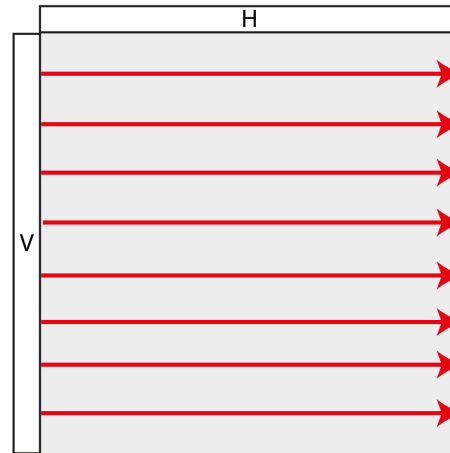
3.1. Predicción intra

• Luminancia – Modos de predicción para bloques de 16x16 muestras:

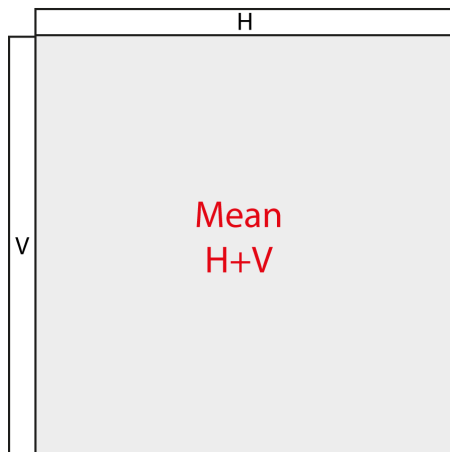
Modo 0 (vertical)



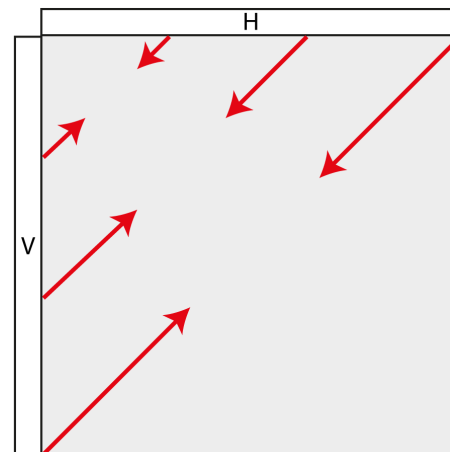
Modo 1 (horizontal)



Modo 2 (DC)



Modo 3 (plane)



○ Modo 0 (vertical) → Extrapolación desde las muestras horizontales.

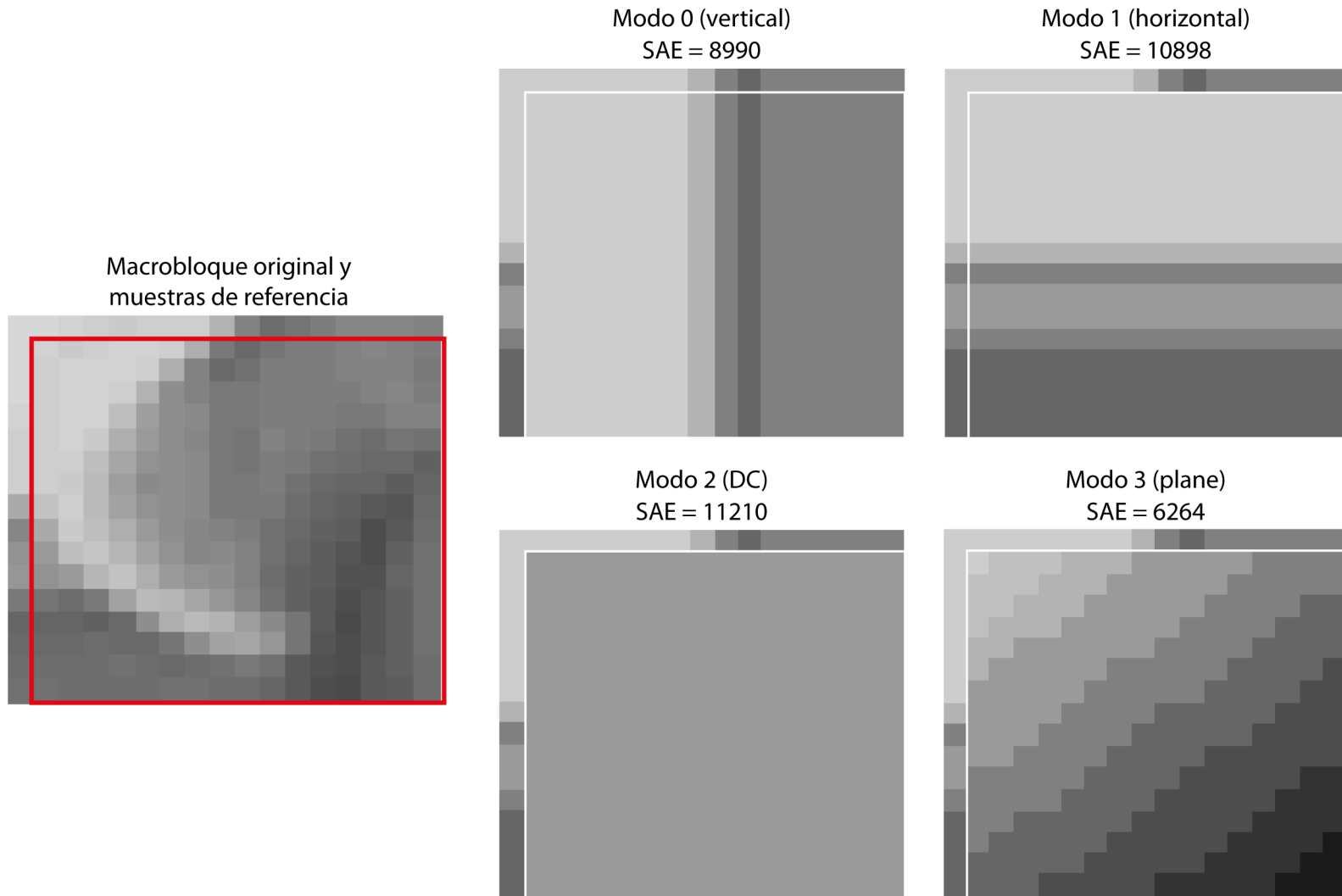
○ Modo 1 (horizontal) → Extrapolación desde las muestras verticales.

○ Modo 2 (DC) → Media de los valores de todas las muestras horizontales y verticales.

○ Modo 3 (plane) → Se aplica una función lineal que da valores a las muestras predichas a partir de muestras de referencia tanto horizontales como verticales (en zonas en las que la luminancia varía suavemente es el que mejor funciona).

3.1. Predicción intra

- Luminancia – Modos de predicción para bloques de 16x16 muestras:



3.1. Predicción intra

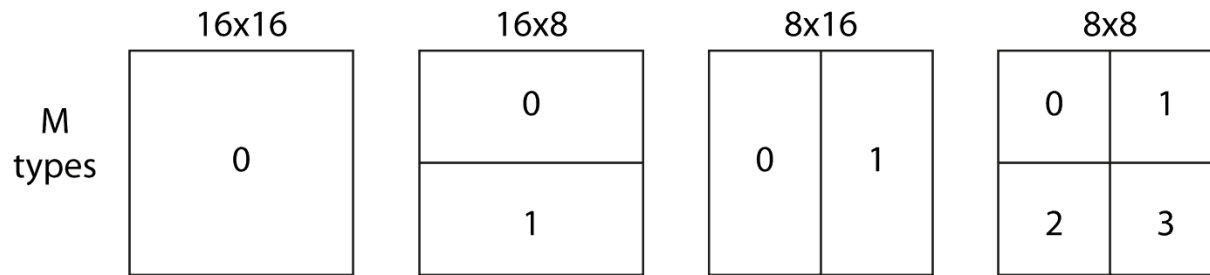
- **Crominancia – Modos de predicción para bloques de 8x8 muestras:**
 - Se aplica siempre sobre bloques de 8x8 muestras.
 - Al igual que en el caso de la luminancia, se utilizan como referencia las muestras (codificadas y reconstruidas) situadas encima y a la izquierda del bloque a predecir.
 - Se utilizan 4 modos similares a los usados para los bloques de 16x16 muestras de luminancia. La única diferencia es el orden de los modos:
 - Modo 0 (DC).
 - Modo 1 (horizontal).
 - Modo 2 (vertical).
 - Modo 3 (plano).
 - Siempre que se aplique predicción intra a un macrobloque, deberá aplicarse tanto a las muestras de luminancia como a las muestras de crominancia.

3.2. Predicción inter

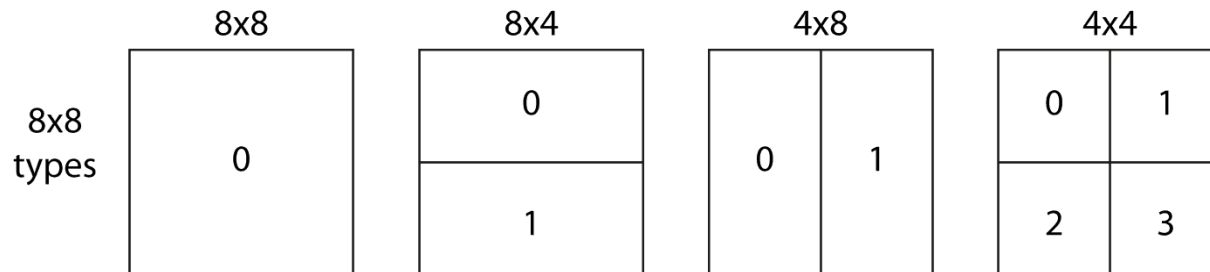
- Se aplica sobre imágenes tipo P y tipo B.
 - A diferencia de codificadores previos, AVC permite usar imágenes B como referencia para las predicciones.
- Se permite el uso de macrobloques de múltiples tamaños (desde 16x16 hasta 4x4):
 - Bloques pequeños en zonas de más movimiento.
 - Bloques grandes en zonas con poco movimiento.
 - Hay que tener en cuenta que si se usan bloques más pequeños se incrementa el número de vectores de movimiento que transmitir.
- Se ofrece la posibilidad de realizar la estimación de movimiento con precisión de $\frac{1}{4}$ de píxel:
 - Mayor coste computacional.
 - Mayor número de bits para representar a los vectores.
 - Menor error de predicción.
- La codificación de los vectores de movimiento se realiza de forma diferencial.

3.2. Predicción inter

- Estructura en árbol de macrobloques:
- Cada macrobloque de 16x16 muestras puede ser dividido de 4 formas distintas:

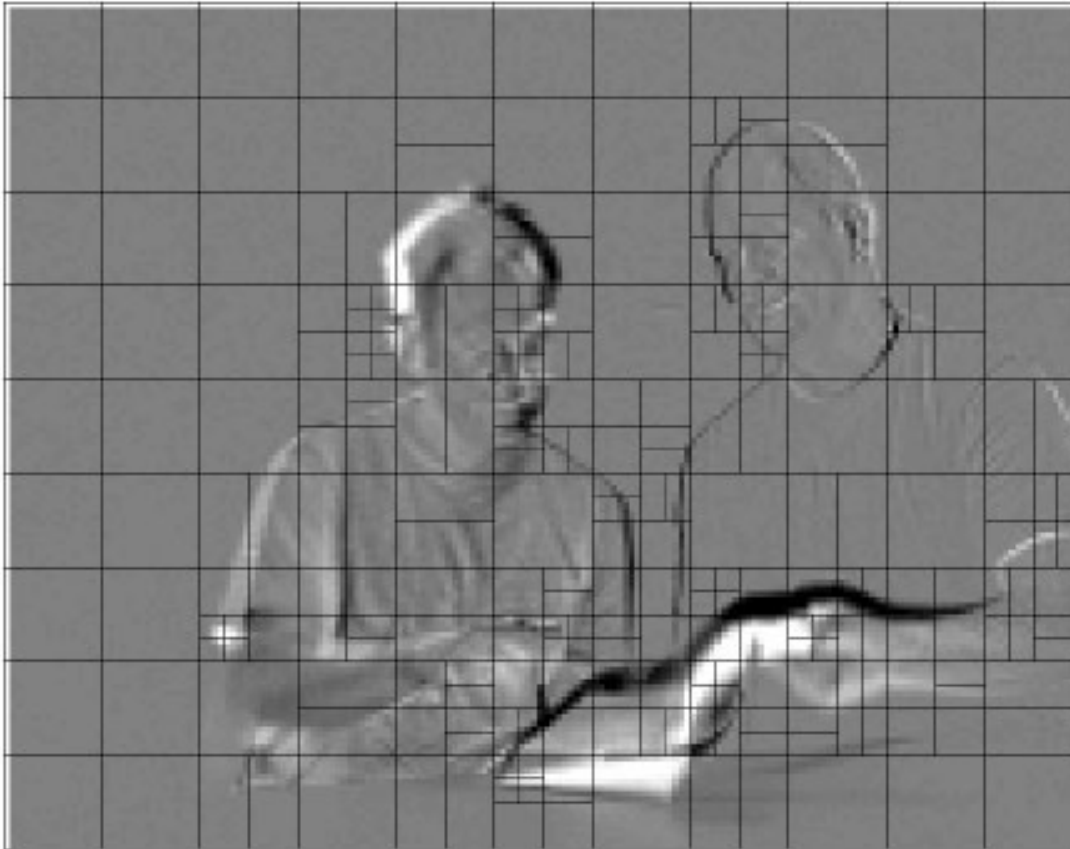


- El modelo de 8x8 puede volver a ser dividido de 4 posibles modos:



3.2. Predicción inter

- Estructura en árbol de macrobloques:

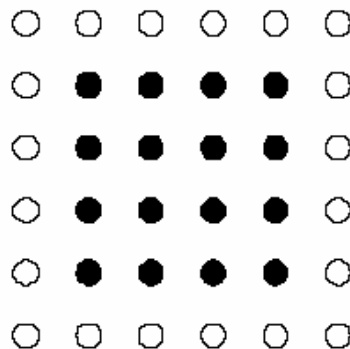


- Se analiza el error de predicción sin aplicar compensación de movimiento:
 - Se asignan MBs grandes a las regiones con menos movimiento.
 - Se asignan MBs pequeños a las regiones con más variaciones.
- La partición mostrada es la decisión que proporciona el mejor compromiso entre error de predicción y número de vectores de movimiento.

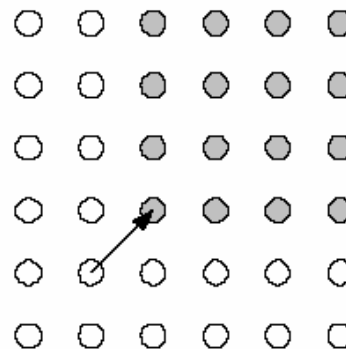
3.2. Predicción inter

- **Vectores de movimiento – precisión a nivel de subpíxel:**

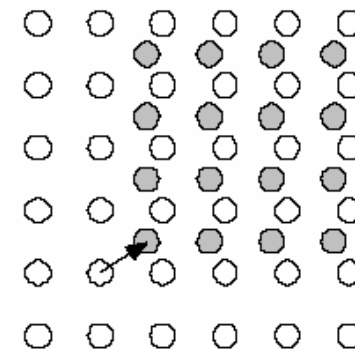
- En una primera etapa se generan muestras de referencia con precisión de $\frac{1}{2}$ píxel:
 - A partir de muestras de referencia con precisión de píxel $\rightarrow$ Se utiliza un filtro FIR.
- En una segunda etapa se generan muestras de referencia con precisión de $\frac{1}{4}$ píxel:
 - A partir de muestras de referencia con precisión de píxel.
 - A partir de las muestras con precisión de $\frac{1}{2}$ píxel obtenidas en la primera etapa.
- Se reduce el error de predicción, pero aumenta el coste computacional y el número de bits necesarios para representar los vectores (2 bits más por componente).



(a) 4x4 block in current frame



(b) Reference block: vector (1, -1)



(c) Reference block: vector (0.75, -0.5)

3.2. Predicción inter

- **Vectores de movimiento – precisión a nivel de subpíxel:**



- Precisión de medio píxel:

- Se aplica un filtro FIR de componentes:

$$(1, -5, 20, 20, -5, 1) / 32$$

- Filtrado horizontal y vertical:

$$b = \frac{E - 5F + 20G + 20H - 5I + J}{32}$$

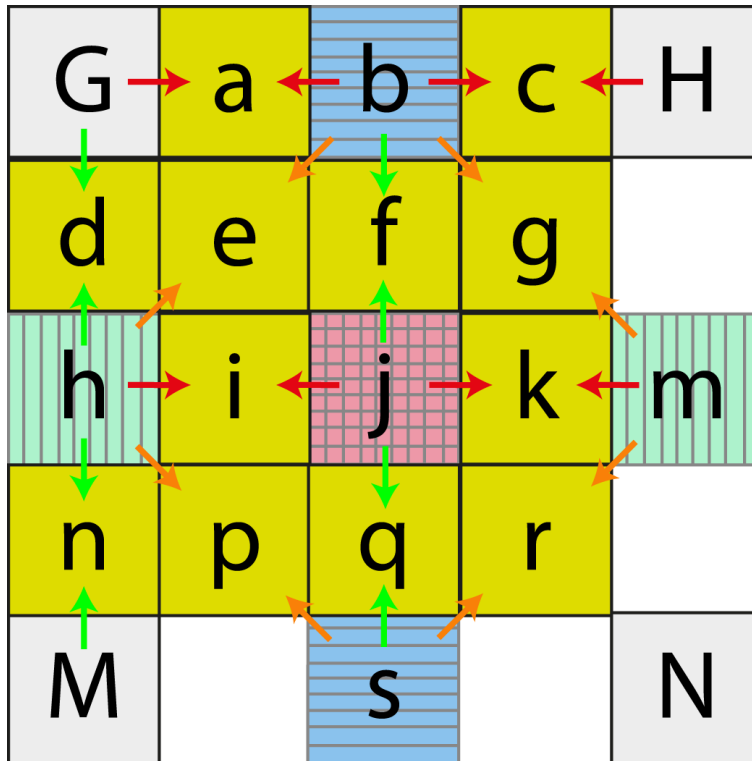
$$h = \frac{A - 5C + 20G + 20M - 5R + T}{32}$$

- Filtrado diagonal:

$$j = \frac{cc - 5dd + 20h + 20m - 5ee + ff}{32}$$

3.2. Predicción inter

- Vectores de movimiento – precisión a nivel de subpíxel:

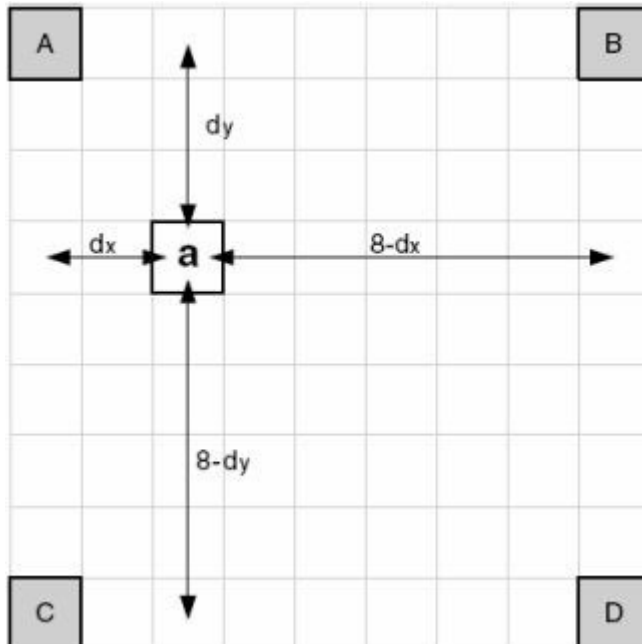


- Precisión de $\frac{1}{4}$ píxel:

- Media de las 2 muestras más próximas.
- En el caso de las medias en “diagonal” se utilizan únicamente las muestras con precisión de $\frac{1}{2}$ píxel obtenidas mediante un único filtrado (en el ejemplo, la muestra ‘j’ no se utiliza).

3.2. Predicción inter

- **Vectores de movimiento – precisión a nivel de subpíxel:**



- Crominancia:

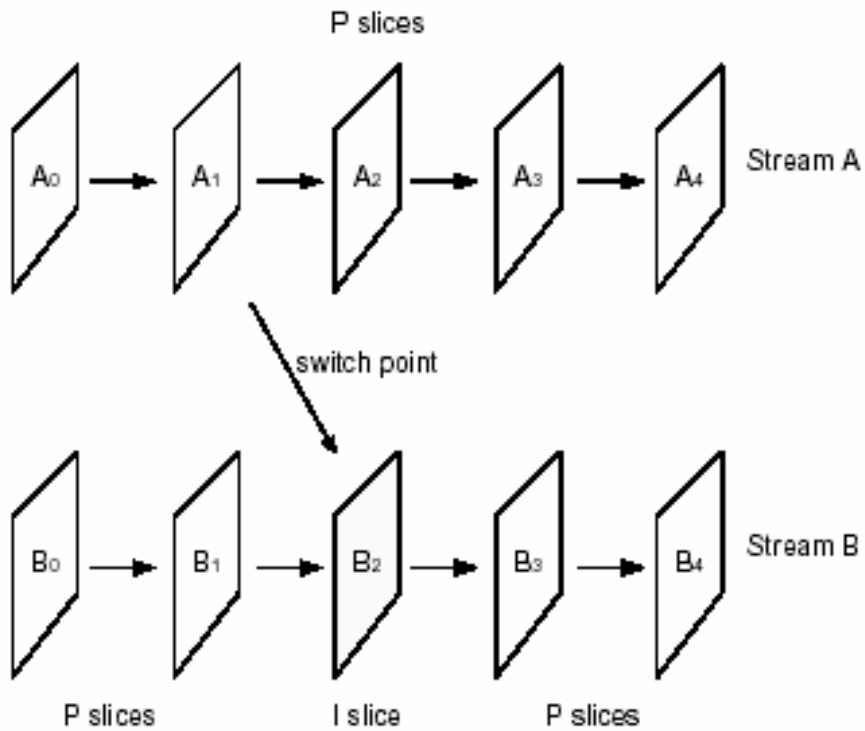
- Se aplica una interpolación bilineal a partir de muestras de crominancia de precisión entera.
- Si se usa muestreo 4:2:0 o 4:1:1, las muestras de crominancia tendrán menor resolución espacial → La precisión de las muestras de crominancia interpoladas podrá ser de hasta $1/8$.

4. Imágenes SP y SI

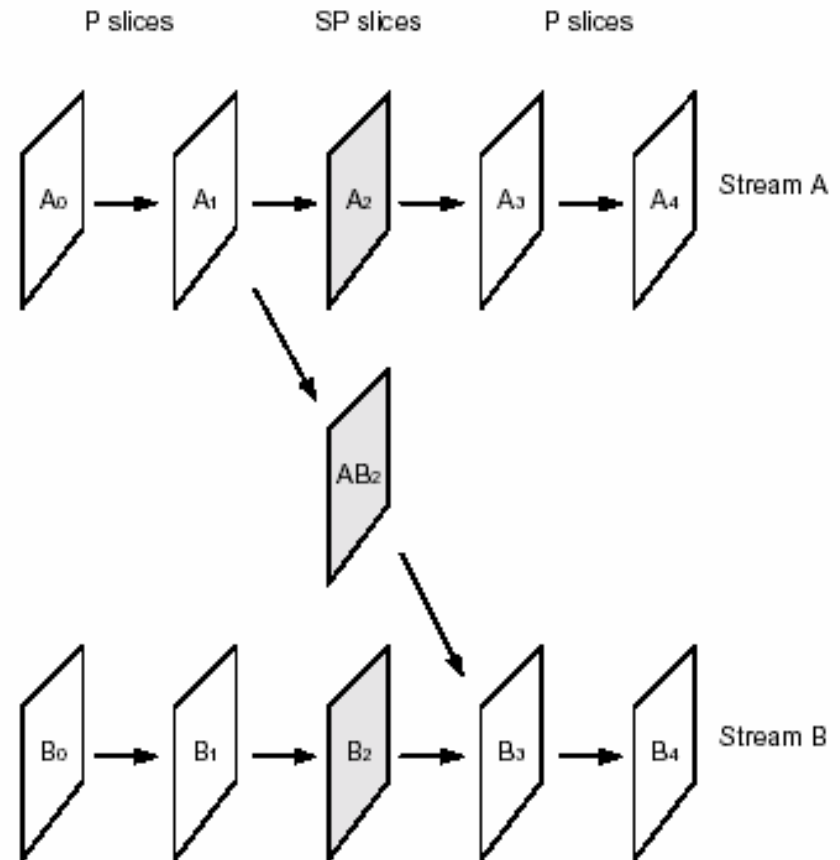
- **SP** → Switching P → Predicción con estimación y compensación de movimiento.
- **SI** → Switching I → Predicción intra 4x4.
- Permiten al decodificador la conmutación eficiente entre distintos flujos de vídeo.
 - Ejemplo: Se tiene un vídeo codificado con distintas tasas binarias y se transmite por Internet. El decodificador podrá ir saltando de un flujo a otro para adaptarse a la ocupación de la red.
- Hasta ahora esto sólo podía hacerse mediante las imágenes de tipo I:
 - ✓ Eran las únicas que podían descodificarse sin utilizar información de otras imágenes.
 - ✗ Incrementan muy notablemente la tasa binaria → Posibles retardos y pérdidas en el proceso de conmutación entre flujos.

4. Imágenes SP y SI

Cambio de flujo utilizando imágenes tipo I

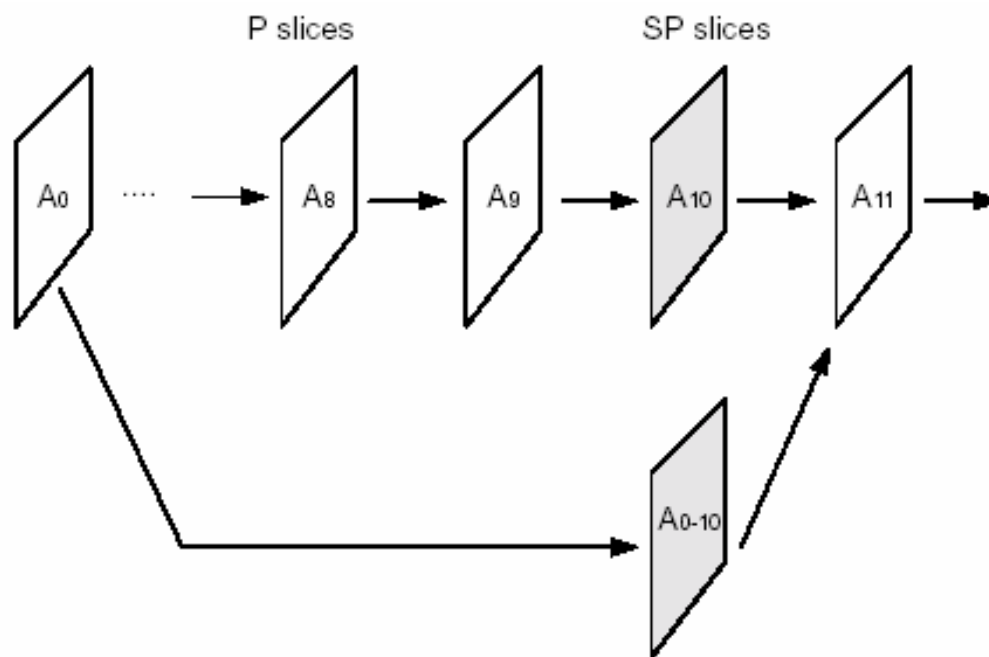


Cambio de flujo utilizando imágenes tipo SP o SI



4. Imágenes SP y SI

Avance rápido y acceso aleatorio



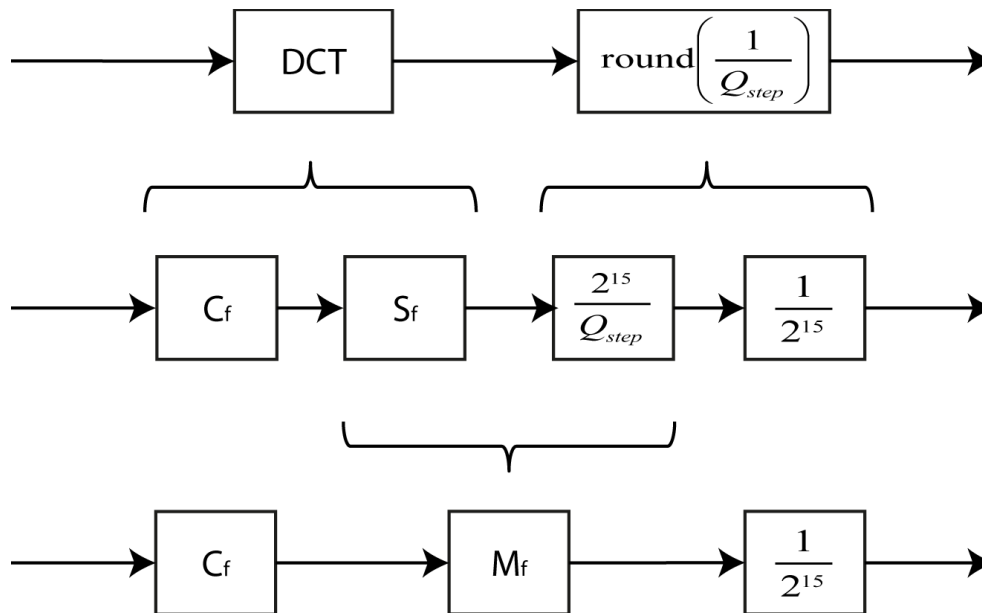
5. Transformación y cuantificación

- En los estándares anteriores se utiliza la transformada DCT bidimensional aplicada sobre bloques de 8x8 muestras.
- AVC utiliza **3 tipos de transformadas**, dependiendo del tipo de datos que se esté codificando:
 - Transformada similar a la DCT (transformación principal) → Se aplica sobre bloques de 4x4 muestras (todas las muestras).
 - Transformada Hadamard 4x4 → Se aplica sobre matrices de 4x4 muestras DC de luminancia.
 - Transformada Hadamard 2x2 → Se aplica sobre matrices de 2x2 muestras DC de crominancia.
- En el caso de que los macrobloques tengan tamaños distintos a 16x16 muestras de luminancia (4x8, 8x4, 8x8, 16x8, etc.) se aplican pequeñas variaciones sobre estas 3 transformadas.
- En los perfiles altos también se permite el uso de transformadas de 8x8 muestras.

5. Transformación y cuantificación

- **Transformación principal:**

- Se utiliza una versión escalada de la DCT → El objetivo es utilizar valores enteros para:
 - Reducir la complejidad de los cálculos.
 - Evitar pérdidas por errores de redondeo.
- Se consigue mediante una reestructuración de los procesos de transformación y cuantificación.



$$M_f \approx \frac{S_f 2^{15}}{Q_{step}}$$

5. Transformación y cuantificación

- Transformación principal:

- Supongamos una matriz de 4x4 muestras, $\mathbf{X}$. Su DCT bidimensional se puede calcular como:

$$\mathbf{Y} = \mathbf{A} \cdot \mathbf{X} \cdot \mathbf{A}^T \quad \mathbf{A} = \begin{pmatrix} a & a & a & a \\ b & c & -c & -b \\ a & -a & -a & a \\ c & -b & b & -c \end{pmatrix} \quad \begin{aligned} a &= \frac{1}{2} \\ b &= \sqrt{\frac{1}{2}} \cos\left(\frac{\pi}{8}\right) \\ c &= \sqrt{\frac{1}{2}} \cos\left(\frac{3\pi}{8}\right) \end{aligned}$$

- Este cálculo requiere redondear números irracionales → Alto coste.
- AVC propone aplicar un factor de escala (2.5) y redondear al entero más próximo. De este modo la matriz $\mathbf{A}$ se convierte en la matriz $\mathbf{C}_f$:

$$\mathbf{C}_f = \begin{pmatrix} 1 & 1 & 1 & 1 \\ 2 & 1 & -1 & -2 \\ 1 & -1 & -1 & 1 \\ 1 & -2 & 2 & -1 \end{pmatrix}$$

5. Transformación y cuantificación

- Transformación principal:

- Las filas de la matriz de transformación tienen distintas normas → hay que escalarlas para que sean normales y mantener la propiedad de ortonormalidad:

$$\mathbf{A} = \mathbf{C}_f * \mathbf{R}_f \quad \mathbf{R}_f = \begin{pmatrix} \frac{1}{2} & \frac{1}{2} & \frac{1}{2} & \frac{1}{2} \\ \frac{1}{\sqrt{10}} & \frac{1}{\sqrt{10}} & \frac{1}{\sqrt{10}} & \frac{1}{\sqrt{10}} \\ \frac{1}{2} & \frac{1}{2} & \frac{1}{2} & \frac{1}{2} \\ \frac{1}{\sqrt{10}} & \frac{1}{\sqrt{10}} & \frac{1}{\sqrt{10}} & \frac{1}{\sqrt{10}} \end{pmatrix}$$

- Por lo tanto:

$$\begin{aligned} \mathbf{Y} &= (\mathbf{C}_f * \mathbf{R}_f) \cdot \mathbf{X} \cdot (\mathbf{C}_f^T * \mathbf{R}_f^T) \\ \mathbf{Y} &= (\mathbf{C}_f \cdot \mathbf{X} \cdot \mathbf{C}_f^T) * (\mathbf{R}_f * \mathbf{R}_f^T) \\ \mathbf{Y} &= (\mathbf{C}_f \cdot \mathbf{X} \cdot \mathbf{C}_f^T) * \mathbf{S}_f \end{aligned} \quad \mathbf{S}_f = \begin{pmatrix} \frac{1}{4} & \frac{1}{2\sqrt{10}} & \frac{1}{4} & \frac{1}{2\sqrt{10}} \\ \frac{1}{2\sqrt{10}} & \frac{1}{10} & \frac{1}{2\sqrt{10}} & \frac{1}{10} \\ \frac{1}{4} & \frac{1}{2\sqrt{10}} & \frac{1}{4} & \frac{1}{2\sqrt{10}} \\ \frac{1}{2\sqrt{10}} & \frac{1}{10} & \frac{1}{2\sqrt{10}} & \frac{1}{10} \end{pmatrix}$$

5. Transformación y cuantificación

- Transformación principal:

- Escalado + cuantificación:

$$\mathbf{Z}_{i,j} = \mathbf{Y}_{i,j} * \text{round}\left(\frac{\mathbf{S}_f}{Q_{step}}\right)$$

- El estándar define un total de 52 posibles valores para Q_{step} .

5. Transformación y cuantificación

- **Transformación DC:**

- Además de la transformación vista anteriormente, se les aplica otra adicional.
- Esta transformación adicional ayuda a explotar la gran correlación existente entre los DC de los distintos bloques.

- A cada bloque se le aplica la siguiente transformación:

- Bloque de 4x4 DC de luminancia →

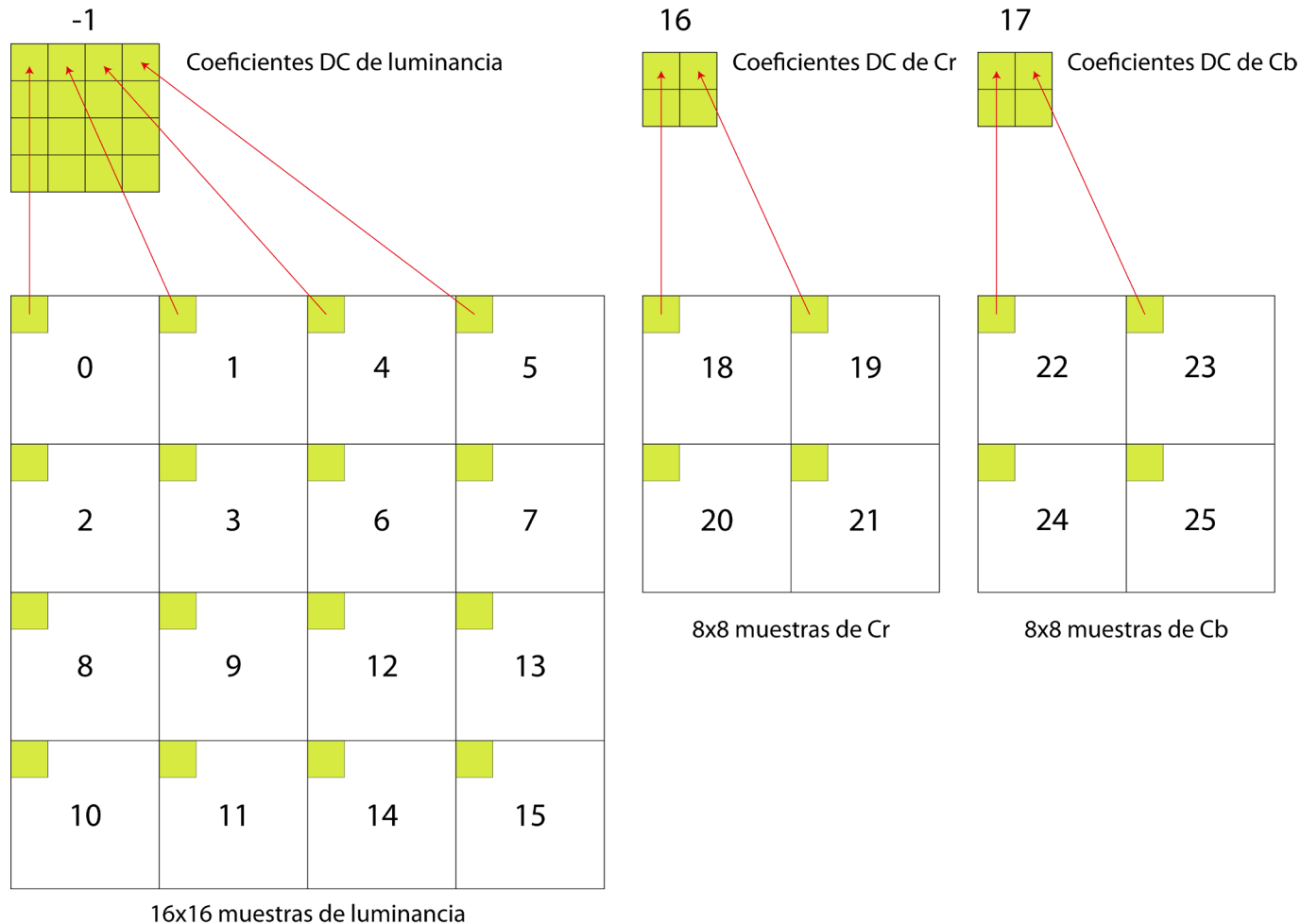
$$\mathbf{Y}_D = \begin{pmatrix} 1 & 1 & 1 & 1 \\ 1 & 1 & -1 & -1 \\ 1 & -1 & -1 & 1 \\ 1 & -1 & 1 & -1 \end{pmatrix} \cdot \mathbf{W}_D \cdot \begin{pmatrix} 1 & 1 & 1 & 1 \\ 1 & 1 & -1 & -1 \\ 1 & -1 & -1 & 1 \\ 1 & -1 & 1 & -1 \end{pmatrix} \cdot \frac{1}{2}$$

- Bloque de 2x2 DC de crominancia →

$$\mathbf{Y}_D = \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix} \cdot \mathbf{W}_D \cdot \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix}$$

5. Transformación y cuantificación

- Orden de transmisión de los datos:



5. Transformación y cuantificación

- ¿Por qué utilizar bloques de 4x4?:
 - La predicción en AVC es mucho mejor que en algoritmos previos → Las imágenes de error tienen menor correlación espacial → La correlación existente se aprovecha igual de bien con matrices de tamaño 4x4 que con matrices de tamaño 8x8.
 - Se reduce el efecto de bloques, ya que las variaciones debidas a cuantificaciones más groseras se notan menos cuanto más pequeños son los bloques.
 - El coste computacional y la complejidad es menor que en codificadores previos:
 - Codificadores anteriores → Aritmética de 32 bits.
 - AVC → Aritmética de 16 bits.

6. Filtro de reconstrucción

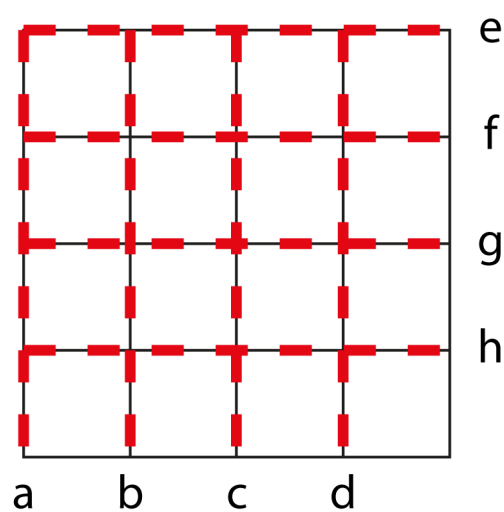
- El objetivo del filtro es reducir la distorsión introducida en la cuantificación.
- Se aplica sobre bloques de 4x4 muestras previamente codificados y descodificados, antes de ser almacenados para su posible utilización en la predicción de las siguientes imágenes codificadas.
 - Se aplica tanto en el codificador como en el descodificador.
- A diferencia que en estándares previos:
 - Se aplica antes de realizar la estimación y compensación de movimiento.
 - Es de inclusión obligatoria.
- Aporta dos beneficios:
 - Al reducir el efecto de bordes entre bloques mejora la apariencia de las imágenes descodificadas, especialmente cuando se aplican altos grados de compresión (cuantificación más grosera).
 - Al realizarse la estimación y compensación sobre macrobloques con bloques filtrados el resultado de la predicción es de mayor calidad.

6. Filtro de reconstrucción

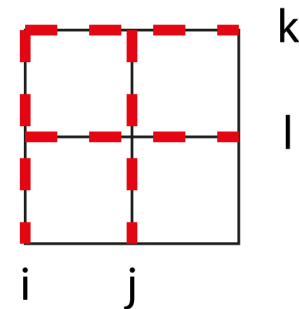
- **Orden de filtrado:**

- Se aplica sobre las fronteras horizontales y verticales de bloques de tamaño 4x4 dentro de cada macrobloque.

1. Filtrado vertical de las fronteras entre bloques de luminancia → Fronteras **a**, **b**, **c** y **d**.
2. Filtrado horizontal de las fronteras entre bloques de luminancia → Fronteras **e**, **f**, **g** y **h**.
3. Filtrado vertical de las fronteras entre bloques de crominancia → Fronteras **i** y **j**.
4. Filtrado horizontal de las fronteras entre bloques de crominancia → Fronteras **k** y **l**.



**Filtrado en un macrobloque
de luminancia de 16x16**



**Filtrado en un macrobloque
de crominancia de 8x8**

6. Filtro de reconstrucción

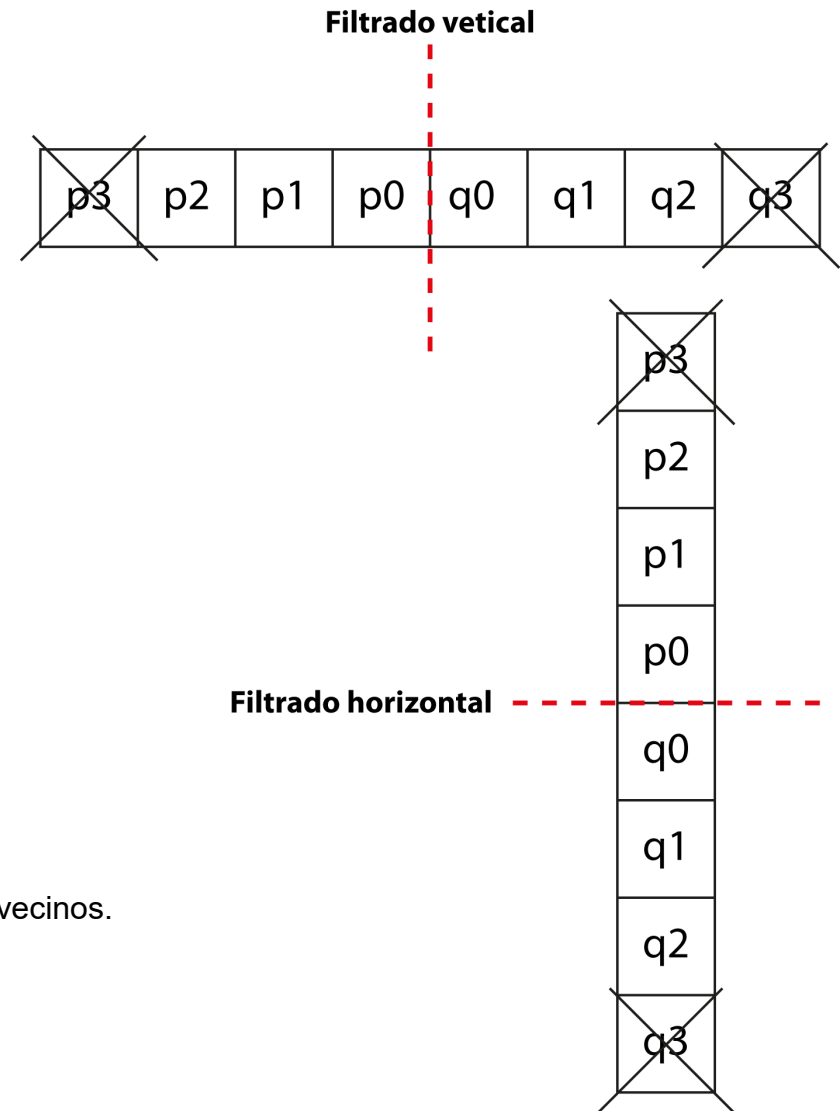
- **Modo de aplicación:**

- Cada operación de filtrado afecta a distinto número de muestras:

- Desde ninguna.
- Hasta un máximo de 3 a cada lado del borde:
 - ❑ En el ejemplo: p2, p1, p0, q0, q1 y q2 darían lugar a P2, P1, P0, Q0, Q1 y Q2 .

- Se utilizan más o menos muestras en función de:

- El cuantificador utilizado.
- Los modos de codificación utilizados en los bloques vecinos.
- Los valores de gradiente a lo largo del borde.



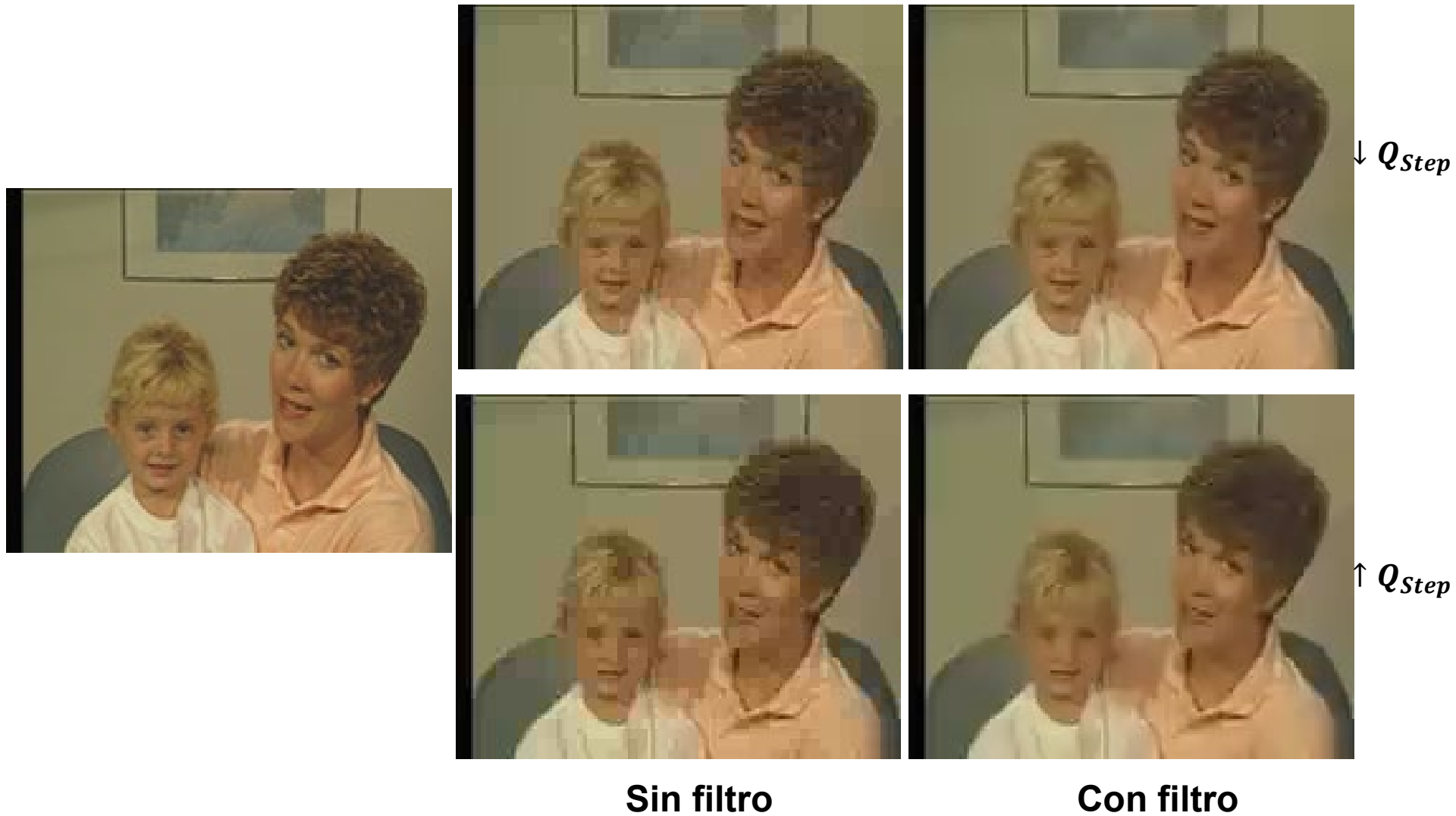
6. Filtro de reconstrucción

• Tipos de filtrado:

• El bloque p o el q han sido codificados en modo intra. • El borde a filtrar es una frontera entre macrobloques	$B_s = 4$ (filtrado más robusto)
• El bloque p o el q han sido codificados en modo intra. • El borde a filtrar no es una frontera entre macrobloques	$B_s = 3$
• Ni el bloque p ni el q han sido codificados en modo intra. • El bloque p o el q contienen coeficientes codificados (trans-cuant).	$B_s = 2$
• Ni el bloque p ni el q han sido codificados en modo intra. • Ni el bloque p ni el q contienen coeficientes codificados. • Los bloques p y q tienen: distintas imágenes de referencia, o distinto número de imágenes de referencia o distintos valores de vectores de movimiento.	$B_s = 1$
• Ni el bloque p ni el q han sido codificados en modo intra. • Ni el bloque p ni el q contienen coeficientes codificados. • Los 2 bloques tienen la misma imagen de referencia e idénticos vectores de movimiento.	$B_s = 0$ (no hay filtrado)

- El filtrado es más robusto (mayor B_s) cuanto mayor es la probabilidad de la aparición de discontinuidades.
- El número de muestras filtradas ($P_3, P_2, \dots, Q_3$) y las muestras originales ($p_3, p_2, \dots, q_3$) utilizadas para obtener cada muestra filtrada depende tanto de B_s como de algunos umbrales que define el estándar y que varían en función de los parámetros de cuantificación.

6. Filtro de reconstrucción



7. Codificación

- AVC permite utilizar distintos algoritmos de codificación:
 - **VLC (*Variable Length Coding*):**
 - Es el que definen los codificadores anteriores.
 - Asigna más bits a los símbolos menos frecuentes y menos bits a los más frecuentes.
 - La asignación se hace en base a una única tabla de referencia.
 - **CAVLC (*Context-based Adaptive Variable Length Coding*):**
 - Similar a VLC pero teniendo en cuenta las condiciones de contexto de los datos.
 - Tiene más tablas de referencia y escoge la más conveniente en función del contenido de cada imagen.
 - **CABAC (*Context Adaptive Binary Arithmetic Coding*):**
 - Está basado en codificación aritmética.
 - Es el más potente (menor *bit-rate*), pero también el más complejo.
 - También utiliza las condiciones de contexto para tomar decisiones.

7. Codificación

- **CABAC:**

- Consigue muy buenos resultados gracias a:
 - Utiliza distintos modelos (tablas) para codificar cada elemento, teniendo en cuenta su contexto.
 - Estima la probabilidad de forma adaptativa, en función del contexto de los datos.
 - Utiliza codificación aritmética.
- Consta de los siguientes pasos:
 - Binarización: Cualquier símbolo resultante de la codificación se binariza.
 - Selección del modelo de contexto: El modelo de contexto es un modelo de probabilidad que será aplicado sobre uno o varios símbolos binarios. Se selecciona de entre un grupo de modelos disponibles, en función de las estadísticas de otros símbolos previamente codificados.
 - Codificación aritmética: Se codifica cada bin (cada símbolo binario) de acuerdo con el modelo de contexto seleccionado.
 - Actualización de la probabilidad: Se actualizan los valores de probabilidad del modelo de contexto en función del cada nuevo valor codificado.

Bibliografía

- I. E. Richardson, “The H.264 Advanced Video Compression Standard”, *John Wiley & Sons*, 2011.
- T. Wiegand, G. J. Sullivan, G. Bjontegaard, and A. Luthra, “Overview of the H.264/AVC Video Coding Standard”, *IEEE Trans. Circuits and Systems for Video Technology*, vol. 13, no. 7, pp. 560-676, 2003.
- H. J. Ochoa-Domínguez, J. Mireles-García y J. D. Cota-Ruiz, “Descripción del nuevo estándar de vídeo H.264 y comparación de su eficiencia de codificación con otros estándares”, *Ingeniería Investigación y Tecnología*, vol. 8, no. 3, pp. 157-180, 2007.