

(12) **FASCÍCULO DE PATENTE DE INVENÇÃO**

(22) Data de pedido: 2002.08.08	(73) Titular(es): SHARP KABUSHIKI KAISHA 22-22 NAGAIKE-CHO ABENO-KU OSAKA 545-8522 JP
(30) Prioridade(s): 2001.08.09 US 311436 P 2001.11.30 US 319018 P 2002.05.02 US 139036	(72) Inventor(es): LOUIS JOSEPH KEROFISKY US
(43) Data de publicação do pedido: 2005.10.19	(74) Mandatário: ANTÓNIO INFANTE DA CÂMARA TRIGUEIROS DE ARAGÃO RUA DO PATROCÍNIO, Nº 94 1399-019 LISBOA PT
(45) Data e BPI da concessão: 2012.03.14 101/2012	

(54) Epígrafe: **QUANTIFICAÇÃO E NORMALIZAÇÃO DE TRANSFORMADA INTEIRA CONJUNTAS UTILIZANDO UMA REPRESENTAÇÃO MANTISSA-EXPOENTE DE UM PARÂMETRO DE QUANTIFICAÇÃO**

(57) Resumo:

PROPORCIONA-SE UM MÉTODO PARA A QUANTIFICAÇÃO DE UM COEFICIENTE. O MÉTODO COMPREENDE: FORNECER UM COEFICIENTE K; FORNECER UM PARÂMETRO DE QUANTIFICAÇÃO (QP); FORMAR UM VALOR DE QUANTIFICAÇÃO (L) A PARTIR DO COEFICIENTE K UTILIZANDO UMA PARTE DE MANTISSA (AM(QP)) E UMA PARTE EXPONENCIAL (XAE(QP)). TIPICAMENTE, O VALOR DE X É 2. EM ALGUNS ASPECTOS DO MÉTODO, A FORMAÇÃO DE UM VALOR DE QUANTIFICAÇÃO (L) A PARTIR DO COEFICIENTE K INCLUI $L=K \cdot A(QP)=K \cdot AM(QP) \cdot (2^{AC(QP)})$. NOUTROS ASPECTOS O MÉTODO COMPREENDE AINDA: A NORMALIZAÇÃO DO VALOR DE QUANTIZAÇÃO EM 2^N COMO SE SEGUE: $LN=L/2^N=K \cdot AM(QP)/2^{(N-AE(QP))}$. EM ALGUNS ASPECTOS, A FORMAÇÃO DO VALOR DE QUANTIFICAÇÃO INCLUI A FORMAÇÃO DE UM CONJUNTO DE FACTORES DE QUANTIFICAÇÃO RECURSIVOS COM O PERÍODO P, EM QUE $A(QP+P)=AC(QP)/X$. A FORMAÇÃO DE FACTORES DE QUANTIFICAÇÃO RECURSIVOS INCLUI A FORMAÇÃO DE FACTORES DE MANTISSA RECURSIVOS EM QUE $AM(QP)=AM(QP \text{ MOD } P)$ E A FORMAÇÃO DE FACTORES DE QUANTIFICAÇÃO RECURSIVOS INCLUI A FORMAÇÃO DE FACTORES EXPONENCIAIS RECURSIVOS, EM QUE $AE(QP)=AE(QP \text{ MOD } P)-QP/P$.

RESUMO

"QUANTIFICAÇÃO E NORMALIZAÇÃO DE TRANSFORMADA INTEIRA CONJUNTAS UTILIZANDO UMA REPRESENTAÇÃO MANTISSA-EXPOENTE DE UM PARÂMETRO DE QUANTIFICAÇÃO"

Proporciona-se um método para a quantificação de um coeficiente. O método compreende: fornecer um coeficiente K ; fornecer um parâmetro de quantificação (QP); formar um valor de quantificação (L) a partir do coeficiente K utilizando uma parte de mantissa ($A_m(QP)$) e uma parte exponencial ($x^{A_e(QP)}$). Tipicamente, o valor de x é 2. Em alguns aspectos do método, a formação de um valor de quantificação (L) a partir do coeficiente K inclui $L = K \cdot A(QP) = K \cdot A_m(QP) \cdot (2^{A_e(QP)})$. Noutros aspectos o método compreende ainda: a normalização do valor de quantização em 2^N como se segue: $L_n = L / 2^N = K \cdot A_m(QP) / 2^{(N - A_e(QP))}$. Em alguns aspectos, a formação do valor de quantificação inclui a formação de um conjunto de factores de quantificação recursivos com o período P , em que $A(QP+P) = A(QP)/x$. A formação de factores de quantificação recursivos inclui a formação de factores de mantissa recursivos em que $A_m(QP) = A_m(QP \bmod P)$ e a formação de factores de quantificação recursivos inclui a formação de factores exponenciais recursivos, em que $A_e(QP) = A_e(QP \bmod P) - QP/P$.

DESCRIÇÃO

"QUANTIFICAÇÃO E NORMALIZAÇÃO DE TRANSFORMADA INTEIRA CONJUNTAS UTILIZANDO UMA REPRESENTAÇÃO MANTISSA-EXPOENTE DE UM PARÂMETRO DE QUANTIFICAÇÃO"

PEDIDOS RELACIONADOS

Este pedido reivindica o benefício de um Pedido Provisório dos Estados Unidos intitulado REDUCED BIT-DEPTH QUANTIZATION, inventado por Louis Kerofsky, N° de Série 06/311436, apresentado em 9 de Agosto de 2001, registo legal N° SLA1110P e um Pedido Provisório dos Estados Unidos intitulado METHODS AND SYSTEMS FOR VIDEO CODING WITH JOINT QUANTIZATION AND NORMALIZATION PROCEDURES, inventado por Louis Kerofsky, N° de Série 06/319018, apresentado em 30 de Novembro de 2001, registo legal N° SLA1110P, e um Pedido Provisório dos Estados Unidos intitulado METHOD FOR REDUCED BIT-DEPTH QUANTIZATION, inventado por Louis Kerofsky, N° de Série 10/139036, apresentado em 2 de Maio de 2002, registo legal N° SLA1110.

ANTECEDENTES DA INVENÇÃO

1. Campo da Invenção

Esta invenção refere-se, de um modo geral, a técnicas de compressão de vídeo e, mais particularmente, a um método para

reduzir o tamanho de bits necessário para o cálculo de transformações de codificação de vídeo.

2. Descrição da Técnica Relacionada

Um formato de informação de vídeo proporciona informação visual adequada para activar um ecrã de televisão ou para armazenamento numa fita de vídeo. De um modo geral, os dados de vídeo estão organizados hierarquicamente. Uma sequência de vídeo é dividida em grupos de tramas e cada grupo pode ser composto por uma série de tramas individuais. Cada trama é aproximadamente equivalente a uma imagem estática, sendo as imagens estáticas actualizadas com uma frequência suficiente para simular uma apresentação de movimento contínuo. Uma trama é, ainda, dividida em fatias, ou secções horizontais, que ajudam a concepção de resistência a erros do sistema. Cada fatia é codificada independentemente para que os erros não se propaguem através das fatias. Uma fatia consiste em macroblocos. Nas normas H.26P e Grupo de Peritos em Imagens em movimento (MPEG)-X, um macrobloco é constituído por 16 x 16 pixéis luma e um conjunto correspondente de pixéis croma, dependendo do formato de vídeo. Um macrobloco tem sempre um número inteiro de blocos, sendo a matriz de 8 x 8 pixéis a menor unidade de codificação.

A compressão de vídeo é um componente crítico para qualquer aplicação que requeira transmissão ou armazenamento de dados de vídeo. As técnicas de compressão compensam o movimento através da reutilização de informação armazenada em diferentes áreas da trama (redundância temporal). A compressão também ocorre através da transformação de dados no domínio espacial para o domínio da

frequência. A compressão de vídeo digital híbrida, explorando redundância temporal por compensação de movimento e redundância espacial por transformação, tal como Transformada Discreta de Cosseno (DCT), foi adaptada como a base nas normas internacionais H.26P e MPEG-X.

Como indicado na Patente US 6317767 (Wang), DCT e transformada discreta de cosseno inversa (IDCT) são operações amplamente utilizadas no processamento de sinais de dados de imagem. Ambas são utilizadas, por exemplo, nas normas internacionais para compressão de vídeo de imagens em movimento, apresentadas pelo MPEG. A DCT tem determinadas propriedades que produzem modelos de codificação simplificados e eficientes. Quando aplicada a uma matriz de dados de pixéis, a DCT é um método de decomposição de um bloco de dados, transformando-o numa soma ponderada de frequências espaciais ou coeficientes de DCT. Inversamente, a IDCT é utilizada para voltar a transformar uma matriz de coeficientes de DCT em dados de pixéis.

Os codecs de vídeo digital (DV) são um exemplo de um dispositivo utilizando um método de compressão de dados com base em DCT. Na etapa de divisão em blocos, a trama de imagem é dividida em blocos N por N de informação de pixel, incluindo, por exemplo, dados de brilho e cor para cada pixel. Um tamanho de bloco comum é de oito pixéis na horizontal por oito pixéis na vertical. Os blocos de pixéis são, depois, "reorganizados" para agrupar vários blocos de diferentes partes da imagem. A reorganização melhora a uniformidade da qualidade da imagem.

Diferentes campos são registados em incidentes de tempo diferentes. Para cada bloco de dados de pixéis, um detector de movimento olha para a diferença entre dois campos de uma trama.

A informação de movimento é enviada para a fase de processamento seguinte. Na etapa seguinte, a informação de pixéis é transformada utilizando uma DCT. Uma DCT 8-8, por exemplo, recebe oito entradas e devolve oito saídas na direcção vertical e horizontal. Os coeficientes de DCT resultantes são, depois, ponderados pela multiplicação de cada bloco de coeficientes de DCT por constantes de ponderação.

Os coeficientes de DCT ponderados são quantificados na fase seguinte. A quantificação arredonda cada coeficiente de DCT num determinado intervalo de valores para ser o mesmo número. A quantificação tende a definir os componentes de frequência mais elevada da matriz de frequência como zero, o que resulta num armazenamento de muito menos dados. Uma vez que o olho humano é mais sensível a frequências mais baixas, no entanto, muito pouca qualidade de imagem perceptível é perdida por esta fase.

A fase de quantificação inclui a conversão da matriz bidimensional de coeficientes quantificados para um fluxo linear unidimensional de dados através da leitura dos valores da matriz em ziguezague e dividindo o fluxo linear unidimensional de coeficientes quantificados em segmentos, em que cada segmento consiste em uma cadeia de coeficientes nulos seguida por um coeficiente quantificado diferente de zero. Em seguida, realiza-se uma codificação de comprimento variável (VLC) através da transformação de cada segmento, consistindo no número de coeficientes nulos e amplitude do coeficiente diferente de zero no segmento, numa palavra de código de comprimento variável. Finalmente, um processo de enquadramento de tramas reúne cada 30 blocos de coeficientes quantificados codificados por comprimento variável em cinco blocos de sincronização de comprimento fixo.

A descodificação é, essencialmente, o inverso do processo de codificação descrito acima. O fluxo digital é, em primeiro lugar, submetido a um desenquadramento de tramas. A descodificação de comprimento variável (VLD), em seguida, desempacota os dados de modo a poderem ser restaurados para os coeficientes individuais. Após quantificar inversamente os coeficientes, a ponderação inversa e uma transformada discreta de cosseno inversa (IDCT) são aplicadas ao resultado. Os pesos de ponderação inversos são os inversos multiplicativos dos pesos de ponderação que foram aplicados no processo de codificação. A saída da função de ponderação inversa é, depois, processada pela IDCT.

Muito trabalho tem sido feito no que se refere ao estudo de meios para reduzir a complexidade no cálculo da DCT e IDCT. Os algoritmos que calculam IDCT bidimensionais são denominados algoritmos de "tipo I". Os algoritmos de tipo I são fáceis de implementar numa máquina paralela, isto é, um computador formado por uma pluralidade de processadores que funcionam simultaneamente em paralelo. Por exemplo, quando se utilizam N processadores paralelos para executar uma multiplicação matricial em matrizes $N \times N$, podem realizar-se simultaneamente multiplicações de N colunas. Além disso, uma máquina paralela pode ser concebida de modo a conter hardware ou instruções de software especiais para a realização de transposição rápida de matrizes.

Uma desvantagem dos algoritmos tipo I é que são necessárias mais multiplicações. A sequência de cálculo dos algoritmos de tipo I envolve duas multiplicações de matriz separadas por uma transposição de matriz, que, se $N=4$, por exemplo, requer 64 adições e 48 multiplicações para um número total de

112 instruções. Os especialistas na técnica sabem bem que a execução de multiplicações pelos processadores é muito demorada e que o desempenho do sistema é, muitas vezes, otimizado pela redução do número de multiplicações realizadas.

Uma IDCT bidimensional também pode ser obtida por conversão da transposta da matriz de entrada para um vector unidimensional utilizando uma função L. Em seguida, obtém-se o produto tensorial de uma matriz constante. O produto tensorial é, depois, multiplicado pelo vector L unidimensional. O resultado é convertido, novamente, para uma matriz $N \times N$ utilizando a função M. Assumindo, novamente, que $N=4$, o número total de instruções utilizadas por esta sequência de cálculo é de 92 instruções (68 adições e 24 multiplicações). Os algoritmos que executam IDCT bidimensionais utilizando esta sequência de cálculo são denominados algoritmos de "tipo II". Nos algoritmos de tipo II, as duas matrizes constantes são agrupadas e funcionam como uma operação. A vantagem de algoritmos de tipo II é que estes requerem, tipicamente, menos instruções (92 *versus* 112) e, em particular, menos multiplicações dispendiosas (24 *versus* 48). Os algoritmos de tipo II, no entanto, são muito difíceis de implementar eficientemente numa máquina paralela. Os algoritmos de tipo II tendem a reorganizar os dados com muita frequência e a reorganização de dados numa máquina paralela é muito demorada.

Existem inúmeros algoritmos de tipo I e tipo II para a implementação de IDCT, no entanto, a desquantificação foi tratada como uma etapa independente dependendo dos cálculos de DCT e IDCT. Os esforços para proporcionar definições de DCT e IDCT de bits exactos levaram ao desenvolvimento de transformadas inteiras eficientes. Estas transformadas inteiras aumentam,

tipicamente, a gama dinâmica dos cálculos. Como resultado, a implementação destes algoritmos requer processamento e armazenamento de dados que consistem em mais de 16 bits.

Seria vantajoso se os coeficientes quantificados de fase intermédia pudessem ser limitados a um tamanho máximo em processos de transformação.

Seria vantajoso se um processo de quantificação que fosse útil para processadores de 16-bits pudesse ser desenvolvido.

Seria vantajoso se uma implementação de descodificador, desquantificação e transformação inversa pudessem ser implementadas eficientemente com um processador de 16 bits. Da mesma forma, seria vantajoso se a multiplicação pudesse ser realizada com não mais do que 16 bits e, se fosse necessário um acesso à memória, também este não superior a 16 bits.

SUMÁRIO DA INVENÇÃO

A presente invenção é um processo melhorado para compressão de vídeo. Os algoritmos de codificação de vídeo típicos predizem uma trama a partir de tramas anteriormente codificadas. O erro é submetido a uma transformada e os valores resultantes são quantificados. O quantificador controla o grau de compressão. O quantificador controla a quantidade de informação utilizada para representar o vídeo e a qualidade da reconstrução.

O problema é a interacção da transformada e quantificação na codificação de vídeo. No passado, a transformada e quantificador eram concebidos independentemente. A transformada, tipicamente,

a transformada discreta de cosseno, é normalizada. O resultado da transformada é quantificado convencionalmente utilizando quantificação escalar ou vectorial. Em tecnologias anteriores, MPEG-1, MPEG-2, MPEG-4, H.261, H.263, a definição da transformada inversa não tem sido de bits exactos. Isso permite que o implementador tenha alguma liberdade para seleccionar um algoritmo de transformada adequado para a sua plataforma. Um inconveniente desta abordagem é a probabilidade de a incompatibilidade codificador/descodificador danificar o ciclo de predição. Para resolver este problema de incompatibilidade, partes da imagem são periodicamente codificadas sem predição. A tecnologia actual, por exemplo H.26L, tem-se centrado na utilização de transformadas inteiras que permitem uma definição de bits exactos. As transformadas inteiras não podem ser normalizadas. A transformada é concebida de modo a poder utilizar-se um deslocamento final para normalizar os resultados do cálculo, em vez de divisões intermédias. A quantificação também requer divisão. H.26L proporciona um exemplo de como essas transformadas inteiras são utilizadas em conjunto com quantificação.

No actual Modelo de Teste de Longo Prazo (TML) H.26L, a normalização é combinada com quantificação e implementada através de multiplicações de números inteiros e deslocamentos a seguir a transformada directa e quantificação e a seguir a desquantificação e transformada inversa. O TML H.26L utiliza dois conjuntos de números inteiros $A(QP)$ e $B(QP)$ indexados pelo parâmetro de quantificação (QP), ver Quadro 1. Estes valores são limitados pela relação mostrado abaixo na Equação 1.

Quadro 1 parâmetros de quantificação TML

QP	A _{TML} (QP)	B _{TML} (QP)
0	620	3881
1	553	4351
2	492	4890
3	439	5481
4	391	6154
5	348	6914
6	310	7761
7	276	8718
8	246	9781
9	219	10987
10	195	12335
11	174	13828
12	155	15523
13	138	17435
14	123	19561
15	110	21873
16	98	24552
17	87	27656
18	78	30847
19	69	34870
20	62	38807
21	55	43747
22	49	49103

(continuação)

QP	A _{TML} (QP)	B _{TML} (QP)
23	44	54683
24	39	61694
25	35	68745
26	31	77615
27	27	89113
28	24	100253
29	22	109366
30	19	126635
31	17	141533

Equação 1 Relação Normalização/Quantificação Conjunta

$$A(QP) \cdot B(QP) \cdot 676^2 \approx 2^{40}.$$

A normalização e quantificação são realizadas simultaneamente utilizando estes números inteiros e divisões por potências de 2. A codificação por transformada em H.26L utiliza um tamanho de bloco de 4x4 e uma matriz T de transformada de números inteiros, Equação 2. Para um bloco X de 4x4, os coeficientes K de transformada são calculados como na Equação 3. A partir dos coeficientes de transformada, os níveis de quantificação, L, são calculados por multiplicação de números inteiros. No decodificador, os níveis são utilizados para calcular um novo conjunto de coeficientes, K'. Utilizam-se transformadas de matriz de números inteiros adicionais seguidas por um deslocamento para calcular os valores X' reconstruídos. O codificador tem a liberdade de calcular e arredondar a transformada directa. Tanto o codificador como o decodificador

têm que calcular exactamente a mesma resposta para os cálculos inversos.

Equação 2 matriz de transformada de modelo de teste 8 H.26L

$$T = \begin{pmatrix} 13 & 13 & 13 & 13 \\ 17 & 7 & -7 & -17 \\ 13 & -13 & -13 & 13 \\ 7 & -17 & 17 & -7 \end{pmatrix}$$

Equação 3 TML DCT_LUMA e IDCT_LUMA

$$Y = T \cdot X$$

$$K = Y \cdot T^T$$

$$L = (A_{TML}(QP) \cdot K) / 2^{20}$$

$$K' = B_{TML}(QP) \cdot L$$

$$Y' = T^T \cdot K'$$

$$X' = (Y' \cdot T) / 2^{20}$$

Em que o resultado intermédio Y é o resultado de uma transformada unidimensional e o resultado intermédio Y' é o resultado de uma transformada inversa unidimensional.

A gama dinâmica requerida durante estes cálculos pode ser determinada. A aplicação principal envolve entrada de 9-bits, 8 bits mais sinal, sendo a gama dinâmica requerida por registos intermédios e acessos de memória apresentada no Quadro 2.

Quadro 2 Gama dinâmica de transformada e transformada inversa
TML (bits)

Entrada de 9 bits	Transformada LUMA	Transformada Inversa
Registo	30	27
Memória	21	26

Para manter definições de bits exactos e incorporar quantificação, a gama dinâmica dos resultados intermédios pode ser grande dado que as operações de divisão são adiadas. A presente invenção combina quantificação e normalização para eliminar o crescimento de gama dinâmica de resultados intermédios. Com a presente invenção, as vantagens de definições de bits exactos de transformada inversa e quantificação são mantidas, enquanto se controla a profundidade de bits necessária para estes cálculos. A redução da profundidade de bits exigida reduz a complexidade exigida de uma implementação de hardware e permite uma utilização eficiente de operações de instrução única, múltiplos dados (SIMD), tais como o conjunto de instruções da Intel MMX.

Consequentemente, proporciona-se um método para a quantificação de um coeficiente. O método compreende: fornecer um coeficiente K; fornecer um parâmetro de quantificação (QP); formar um valor de quantificação (L) a partir do coeficiente K utilizando uma parte de mantissa ($A_m(QP)$) e uma parte exponencial ($x^{A_e(QP)}$). Tipicamente, o valor de x é 2.

Em alguns aspectos do método, a formação de um valor de quantificação (L) a partir do coeficiente K inclui:

$$\begin{aligned}
L &= K \cdot A(QP) \\
&= K \cdot A_m(QP) \cdot (2^{A_e(QP)}).
\end{aligned}$$

Em outros aspectos, o método compreende ainda: normalizar o valor de quantificação por 2^N , do seguinte modo:

$$\begin{aligned}
L_n &= L/2^N \\
&= K \cdot A_m(QP)/2^{(N-A_e(QP))}.
\end{aligned}$$

Em alguns aspectos, a formação de um valor de quantificação inclui formar um conjunto de factores de quantificação recursivos com um período P , em que $A(QP+P)=A(QP)/x$. Por conseguinte, a formação de um conjunto de factores de quantificação recursivos inclui formar factores de mantissa recursivos, em que $A_m(QP)=A_m(QP \bmod P)$. Da mesma forma, a formação de um conjunto de factores de quantificação recursivos inclui formar factores exponenciais recursivos, em que $A_e(QP)=A_e(QP \bmod P)-QP/P$.

Mais especificamente, o fornecimento de um coeficiente K inclui o fornecimento de uma matriz de coeficiente $K[i][j]$. Em seguida, a formação de um valor de quantificação (L) a partir da matriz de coeficientes $K[i][j]$ inclui a formação de uma matriz de valor de quantificação ($L[i][j]$) utilizando uma matriz de parte de mantissa ($A_m(QP)[i][j]$) e uma matriz de parte exponencial ($x^{A_e(QP)[i][j]}$).

Da mesma forma, a formação de uma matriz de valor de quantificação ($L[i][j]$) utilizando uma matriz de parte de mantissa ($A_m(QP)[i][j]$) e uma matriz de parte exponencial ($x^{A_e(QP)[i][j]}$) inclui, para cada valor particular de QP , que cada

elemento na matriz de parte exponencial tenha o mesmo valor. Cada elemento na matriz de parte exponencial tem o mesmo valor durante um período (P) de valores de QP, em que $Ae(QP) = Ae(P * (QP/P))$.

Outros pormenores do método acima descrito, incluindo um método para formar um valor de desquantificação (X1) a partir do valor de quantificação utilizando uma parte de mantissa ($Bm(QP)$) e uma parte exponencial ($x^{Be(QP)}$) são proporcionados abaixo.

DESCRIÇÃO RESUMIDA DOS DESENHOS

A Fig. 1 é um fluxograma que ilustra o método da presente invenção para a quantificação de um coeficiente.

As Figs. 2 a 9 mostram formas de realização da presente invenção compreendendo sistemas e métodos para codificação de vídeo.

DESCRIÇÃO PORMENORIZADA DAS FORMAS DE REALIZAÇÃO PREFERIDAS

Os requisitos de gama dinâmica da transformada e quantificação combinadas são reduzidos por factorização dos parâmetros de quantificação $A(QP)$ e $B(QP)$ em termos de uma mantissa e expoente, como mostrado na Equação 4. Com esta estrutura, apenas a precisão devido ao termo de mantissa deve ser preservada durante o cálculo. O termo expoente pode ser incluído no deslocamento de normalização final. Isto é ilustrado no cálculo amostra da Equação 5.

Equação 4 Estrutura de parâmetros de quantificação

$$A_{proposto}(QP) = A_{mantissa}(QP) \cdot 2^{A_{expoente}(QP)}$$

$$B_{proposto}(QP) = B_{mantissa}(QP) \cdot 2^{B_{expoente}(QP)}$$

Equação 5 Transformada LUMA de profundidade-bit reduzida

$$Y = T \cdot X$$

$$K = Y \cdot T^T$$

$$L = (A_{mantissa}(QP) \cdot K) / 2^{2^{17} - A_{expoente}(QP)}$$

$$K' = T^T \cdot L$$

$$Y' = K' \cdot T$$

$$X' = (B_{mantissa}(QP) \cdot Y') / 2^{2^{17} - B_{expoente}(QP)}$$

Para ilustrar a presente invenção, apresenta-se um conjunto de parâmetros de quantificação que reduz o requisito de gama dinâmica de um decodificador H.26L para um acesso de memória de 16 bits. O acesso de memória da transformada inversa é reduzido para 16 bits. Valores para $A_{mantissa}$, $A_{expoente}$, $B_{mantissa}$, $B_{expoente}$, $A_{proposto}$, $B_{proposto}$ são definidos para $QP=0-5$, como mostrado no Quadro 3. Valores adicionais são determinados por recursão, como mostrado na Equação 6. A estrutura destes valores permite gerar novos valores de quantificação para além dos especificados.

Quadro 3 Valores de quantificação 0-5 para TML

QP	A _{mantissa}	A _{expoente}	B _{mantissa}	B _{expoente}	A _{proposto}	B _{proposto}
0	5	7	235	4	640	3760
1	9	6	261	4	576	4176
2	127	2	37	7	508	4736
3	114	2	165	5	456	5280
4	25	4	47	7	400	6016
5	87	2	27	8	348	6912

Equação 6 Relações de recursão

$$\begin{aligned}
 A_{mantissa}(QP + 6) &= A_{mantissa}(QP) \\
 B_{mantissa}(QP + 6) &= B_{mantissa}(QP) \\
 A_{expoente}(QP + 6) &= A_{expoente}(QP) - 1 \\
 B_{expoente}(QP + 6) &= B_{expoente}(QP) + 1
 \end{aligned}$$

Utilizando os parâmetros definidos, os cálculos de transformada podem ser modificados para reduzir a gama dinâmica, como mostrado na Equação 5. Note-se como apenas os valores de mantissa contribuem para o crescimento de gama dinâmica. Os factores de expoente são incorporados na normalização final e não têm impacto na gama dinâmica de resultados intermédios.

Com estes valores e método de cálculo, a gama dinâmica no descodificador é reduzida, pelo que só é necessário um acesso de memória de 16-bits, como visto no Quadro 4.

Quadro 4 Gama dinâmica com quantificação de profundidade de bits reduzida (QP>6)

8-bits	Transformada LUMA	Transformada inversa LUMA
Registro	28	24
Memória	21	16

Vários refinamentos podem ser aplicados ao processo de quantificação/normalização conjunta descrito acima. A técnica geral de factorização dos parâmetros numa mantissa e expoente forma a base destes refinamentos.

A discussão acima assume que todas as funções base da transformada têm uma norma igual e são quantificadas de forma idêntica. Algumas transformadas inteiras têm a propriedade de funções de base diferentes terem normas diferentes. A técnica da presente invenção foi generalizada para suportar transformadas tendo diferentes normas ao substituir os escalares $A(QP)$ e $B(QP)$ acima por matrizes $A(QP)[i][j]$ e $B(QP)[i][j]$. Estes parâmetros são ligados por uma relação de normalização com a forma mostrada abaixo, Equação 7, que é mais geral do que a relação única mostrada na Equação 1.

Equação 7 Quantificação/normalização conjunta de matrizes

$$A(QP)[i][j] \cdot B(QP)[i][j] = N[i][j]$$

De acordo com o método descrito anteriormente, cada elemento de cada matriz é factorizado numa mantissa e num termo expoente como ilustrado nas equações abaixo, Equação 8.

Equação 8 Factorização de parâmetros da matriz

$$\begin{aligned} A(QP)[i][j] &= A_{\text{mantissa}}(QP)[i][j] \cdot 2^{A_{\text{expoente}}(QP)[i][j]} \\ B(QP)[i][j] &= B_{\text{mantissa}}(QP)[i][j] \cdot 2^{B_{\text{expoente}}(QP)[i][j]} \end{aligned}$$

Um grande número de parâmetros são necessários para descrever estes parâmetros de quantificação e desquantificação. Várias relações estruturais podem ser utilizadas para reduzir o número de parâmetros livres. O crescimento de quantificador é concebido de modo a que os valores de A sejam reduzidos para metade, após cada período P, e, ao mesmo tempo, os valores de B sejam duplicados, mantendo a relação de normalização. Além disso, os valores de $A_{\text{expoente}}(QP)[i][j]$ e $B_{\text{expoente}}(QP)[i][j]$ são independentes de i, j e (QP) no intervalo [0, P-1]. Esta estrutura é resumida por meio de equações estruturais, Equação 9. Com esta estrutura, existem apenas dois parâmetros $A_{\text{expoente}}[0]$ e $B_{\text{expoente}}[0]$.

Equação 9 Estrutura de termos de expoente

$$\begin{aligned} A_{\text{expoente}}(QP)[i][j] &= A_{\text{expoente}}[0] - QP/P \\ B_{\text{expoente}}(QP)[i][j] &= B_{\text{expoente}}[0] + QP/P \end{aligned}$$

Também se define uma estrutura para os valores de mantissa. Para cada par de índices (i, j), os valores de mantissa são periódicos com o período de P. Isto é resumido pela equação estrutural, Equação 10. Com esta estrutura, existem P matrizes independentes para A_{mantissa} e P matrizes independentes para B_{mantissa} , reduzindo os requisitos de memória e adicionando estrutura aos cálculos.

Equação 10 Estrutura de termos de mantissa

$$A_{mantissa}(QP)[i][j] = A_{mantissa}(QP \% P)[i][j] \quad |$$

$$B_{mantissa}(QP)[i][j] = B_{mantissa}(QP \% P)[i][j] \quad |$$

A transformada inversa pode incluir uma divisão de números inteiros que exija arredondamento. Em casos de interesse, a divisão é por uma potência de 2. O erro de arredondamento é reduzido ao fazer com que os factores de desquantificação sejam múltiplos da mesma potência de 2, o que não dá resto após a divisão.

A desquantificação utilizando os valores de mantissa $B_{mantissa}(QP)$ dá valores desquantificados que são normalizados de forma diferente dependendo de qp . Isto deve ser compensado de acordo com a transformada inversa. Uma forma deste cálculo é mostrada na Equação 11.

Equação 11 Normalização de transformada inversa I

$$K[i][j] = B_{mantissa}(QP \% P)[i][j] \cdot Nível[i][j] \quad |$$

$$X = (T^{-1} \cdot K \cdot T) / 2^{(N-QP/P)}$$

Na Equação 11, o $Nível[i][j]$ é a versão quantificada dos coeficientes de transformada e é denominado "valor de quantificação". $K[i][j]$ é a versão escalada dos coeficientes de transformada e é denominado "valor de desquantificação".

Para eliminar a necessidade de a transformada inversa compensar esta diferença de normalização, a operação de desquantificação é definida de modo a que todos os valores desquantificados tenham a mesma normalização. A forma deste cálculo é mostrada na Equação 12.

Equação 12 Normalização de transformada inversa II

$$\begin{aligned} K[i][j] &= B_{mantissa}(QP \% P)[i][j] \cdot 2^{QP/P} \cdot Nível[i][j] \\ X &= (T^{-1} \cdot K \cdot T) / 2^N \end{aligned} \quad \Bigg|$$

A potência de 2 pode ser calculada utilizando uma operação de deslocamento para a esquerda e o valor de desquantificação $K[i][j]$ na Equação 12 será, então, dado da seguinte forma.

$$K[i][j] = [B_{mantissa} \cdot Nível[i][j]] \ll (QP/P)$$

Segue-se um exemplo que ilustra a utilização de matrizes de quantificação na presente invenção. As transformadas, directa e inversa, definidas na Equação 13 necessitam de uma matriz de quantificação, em vez de um único valor de quantificação escalar. Parâmetros de quantificação e desquantificação de amostra são dados. As Equações 14 e 16, juntamente com cálculos relacionados, ilustram a utilização da presente invenção. Este exemplo utiliza um período $P=6$. Na Equação 14, $A_{mantissa}$ é representado por Q e QP é representado por m . Na Equação 16, $B_{mantissa}$ é representado por R e QP é representado por m .

Equação 13 transformadas

$$T_{directa} = \begin{pmatrix} 1 & 1 & 1 & 1 \\ 2 & 1 & -1 & -2 \\ 1 & -1 & -1 & 1 \\ 1 & -2 & 2 & -1 \end{pmatrix}$$

$$T_{inversa} = \begin{pmatrix} 2 & 2 & 2 & 1 \\ 2 & 1 & -2 & -2 \\ 2 & -2 & -2 & 2 \\ 2 & -1 & 2 & -1 \end{pmatrix}$$

Equação 14 Parâmetros de quantificação

$$Q(m)[i][j] = M_{m,0} \text{ para } (i, j) = \{(0,0), (0,2), (2,0), (2,2)\}$$

$$Q(m)[i][j] = M_{m,1} \text{ para } (i, j) = \{(1,1), (1,3), (3,1), (3,3)\}$$

$$Q(m)[i][j] = M_{m,2} \text{ de outro modo}$$

$$M = \begin{bmatrix} 21844 & 8388 & 13108 \\ 18724 & 7625 & 11650 \\ 16384 & 6989 & 10486 \\ 14564 & 5992 & 9532 \\ 13107 & 5243 & 8066 \\ 11916 & 4660 & 7490 \end{bmatrix}$$

Equação 16 Parâmetros de Desquantificação

$$R(m)[i][j] = S_{m,0} \text{ para } (i, j) = \{(0,0), (0,2), (2,0), (2,2)\}$$

$$R(m)[i][j] = S_{m,1} \text{ para } (i, j) = \{(1,1), (1,3), (3,1), (3,3)\}$$

$$R(m)[i][j] = S_{m,2} \text{ de outro modo}$$

$$S = \begin{bmatrix} 6 & 10 & 8 \\ 7 & 11 & 9 \\ 8 & 12 & 10 \\ 9 & 14 & 11 \\ 10 & 16 & 13 \\ 11 & 18 & 14 \end{bmatrix}$$

A descrição da transformação directa e quantificação directa, Equação 18, é dada a seguir assumindo que a entrada é em X, parâmetro de quantificação QP.

Equação 17 transformada directa

$$K = T_{directa} \cdot X \cdot T_{directa}^T$$

Equação 18 quantificação directa

$$periodo = QP / 6$$

$$fase = QP - 6 \cdot periodo$$

$$Nível[i][j] = (Q(fase)[i][j] \cdot K[i][j]) / 2^{(17+periodo)}$$

A descrição de desquantificação, transformada inversa e normalização para este exemplo é dada abaixo, Equação 19 e 20.

Equação 19 Desquantificação

$$periodo = QP / 6$$

$$fase = QP - 6 \cdot periodo$$

$$K[i][j] = R(fase)[i][j] \cdot Level[i][j] \cdot 2^{periodo}$$

Equação 20 IDCT e normalização

$$X' = T_{inversa} \cdot K \cdot T_{inversa}^T$$

$$X'[i][j] = X[i][j]/2^j$$

A Fig. 1 é um fluxograma que ilustra o método da presente invenção para a quantificação de um coeficiente. Embora este método seja representado como uma sequência de passos numerados para maior clareza, nenhuma ordem deve ser inferida a partir da numeração, salvo indicação em contrário. Deve compreender-se que alguns destes passos podem ser suprimidos, executados em paralelo ou executados sem ser necessário manter uma ordem de sequência estrita. O método inicia-se no Passo 100. No Passo 102, fornece-se um coeficiente K. No Passo 104, fornece-se um parâmetro de quantificação (QP). No Passo 106, forma-se um valor de quantificação (L) a partir do coeficiente K utilizando uma parte (Am(QP)) de mantissa e uma parte ($x^{Ae(QP)}$) exponencial. Tipicamente, a parte ($x^{Ae(QP)}$) exponencial inclui x com o valor 2.

Em alguns aspectos do método, a formação de um valor de quantificação (L) a partir do coeficiente K utilizando uma parte (Am (QP)) de mantissa e uma parte ($x^{Ae(QP)}$) exponencial, no Passo 106, inclui:

$$\begin{aligned} L &= K * A(QP) \\ &= K * Am(QP) * (2^{Ae(QP)}). \end{aligned}$$

Alguns aspectos do método incluem um outro passo. No Passo 108, normaliza-se o valor de quantificação por 2^N , como se segue:

$$\begin{aligned} L_n &= L/2^N \\ &= K * Am(QP) / 2^{(N - Ae(QP))}. \end{aligned}$$

Em outros aspectos, a formação de um valor de quantificação no Passo 106 inclui a formação de um conjunto de factores de quantificação recursivos com um período P , em que $A(QP+P)=A(QP)/x$. Da mesma forma, a formação de um conjunto de factores de quantificação recursivos inclui a formação de factores de mantissa recursivos, em que $A_m(QP)=A_m(QP \bmod P)$. Em seguida, a formação de um conjunto de factores de quantificação recursivos inclui a formação de factores exponenciais recursivos, em que $A_e(QP)=A_e(QP \bmod P)-QP/P$.

Em alguns aspectos, a formação de um valor de quantificação inclui a formação de um conjunto de factores de quantificação recursivos com um período P , em que $A(QP+P)=A(QP)/2$. Em outros aspectos, a formação de um conjunto de factores de quantificação recursivos inclui a formação de factores de mantissa recursivos, em que $P=6$. Da mesma forma, a formação de um conjunto de factores de quantificação recursivos inclui a formação de factores exponenciais recursivos, em que $P=6$.

Em alguns aspectos do método, o fornecimento de um coeficiente no Passo 102 inclui o fornecimento de uma matriz de coeficiente $K[i][j]$. Em seguida, a formação de um valor de quantificação (L) a partir da matriz de coeficiente $K[i][j]$ utilizando uma parte ($A_m(QP)$) de mantissa e uma parte ($x^{A_e(QP)}$) exponencial, no Passo 106, inclui a formação de uma matriz de valor de quantificação ($L[i][j]$) utilizando uma matriz de parte de mantissa ($A_m(QP)[i][j]$) e uma matriz de parte exponencial ($x^{A_e(QP)[i][j]}$). Da mesma forma, a formação de uma matriz de valor de quantificação ($L[i][j]$), utilizando uma matriz de parte de mantissa ($A_m(QP)[i][j]$) e uma matriz de parte exponencial ($x^{A_e(QP)[i][j]}$) inclui, para cada valor particular de QP , que cada

elemento na matriz de parte exponencial tenha o mesmo valor. Tipicamente, cada elemento na matriz de parte exponencial é o mesmo valor para um período (P) de valores QP, em que $Ae(QP) = Ae(P * (QP/P))$.

Alguns aspectos do método incluem um outro passo. No Passo 110, forma-se um valor de desquantificação (X1) a partir do valor de quantificação, utilizando uma parte (Bm(QP)) de mantissa e uma parte ($x^{Be(QP)}$) exponencial. Novamente, a parte ($x^{Be(QP)}$) exponencial inclui, tipicamente, x com o valor 2.

Em alguns aspectos do método, a formação de um valor de desquantificação (X1) a partir do valor de quantificação, utilizando uma parte (Bm(QP)) de mantissa e uma parte ($2^{Be(QP)}$) exponencial, inclui:

$$\begin{aligned} X1 &= L * B(QP) \\ &= L * Bm(QP) * (2^{Be(QP)}). \end{aligned}$$

Outros aspectos do método incluem outro passo, Passo 112, de desnormalização do valor de quantificação por 2^N , da seguinte forma:

$$\begin{aligned} X1d &= X1 / 2^N \\ &= X1 * Bm(QP) / 2^N. \end{aligned}$$

Em alguns aspectos, a formação de um valor de desquantificação, no Passo 110, inclui a formação de um conjunto de factores de desquantificação recursivos com um período P, em que $B(QP+P) = x * B(QP)$. Em seguida, a formação de um conjunto de factores de desquantificação recursivos inclui a formação de

factores de mantissa recursivos, em que $Bm(QP) = Bm(QP \bmod P)$. Além disso, a formação de um conjunto de factores de desquantificação recursivos inclui a formação de factores exponenciais recursivos, em que $Be(QP) = Be(QP \bmod P) + QP/P$.

Em alguns aspectos, a formação de um conjunto de factores de quantificação recursivos com um período P inclui o valor de x igual a 2 e a formação de factores recursivos de mantissa inclui o valor de P igual a 6. Em seguida, a formação de um conjunto de factores de desquantificação recursivos inclui a formação de factores exponenciais recursivos, em que $Be(QP) = Be(QP \bmod P) + QP/P$.

Em alguns aspectos do método, a formação de um valor de desquantificação ($X1$), a partir do valor de quantificação, utilizando uma parte ($Bm(QP)$) de mantissa e uma parte ($x^{Be(QP)}$) exponencial, no Passo 110, inclui a formação de uma matriz de valor de desquantificação ($X1[i][j]$), utilizando uma matriz de parte de mantissa ($Bm(QP)[i][j]$) e uma matriz de parte exponencial ($x^{Be(QP)[i][j]}$). Da mesma forma, a formação de uma matriz de valor de desquantificação ($X1[i][j]$) utilizando uma matriz de parte de mantissa ($Bm(QP)[i][j]$) e uma matriz de parte exponencial ($x^{Be(QP)[i][j]}$) inclui, para cada valor particular de QP , que cada elemento na matriz de parte exponencial tenha o mesmo valor. Em alguns aspectos, cada elemento na matriz de parte exponencial tem o mesmo valor para um período (P) dos valores de QP , em que $Be(QP) = Be(P * (QP/P))$.

Outro aspecto da invenção inclui um método para a desquantificação de um coeficiente. No entanto, o processo é

essencialmente igual ao dos Passos 110 e 112 acima e não é repetido por uma questão de brevidade.

Foi apresentado um método para a quantificação de um coeficiente. Um exemplo é dado ilustrando um processo de desquantificação e normalização combinadas aplicado à norma de codificação de vídeo H.26L com o objectivo de reduzir a profundidade de bit necessária no decodificador para 16 bits. Os conceitos da presente invenção também podem ser utilizados para satisfazer outros objectivos de concepção de acordo com a H.26L. Em geral, a presente invenção é aplicada à combinação de cálculos de normalização e quantificação.

Formas de realização da presente invenção podem ser implementadas como hardware, firmware, software e outros tipos de implementação. Algumas formas de realização podem ser implementadas em dispositivos informáticos de utilização geral ou em dispositivos informáticos concebidos especificamente para a implementação destas formas de realização. Algumas formas de realização podem ser armazenadas em memória como um meio de armazenamento da forma de realização ou com a finalidade de executar a forma de realização num dispositivo informático.

Algumas formas de realização da presente invenção compreendem sistemas e métodos para codificação de vídeo, como mostrado na Figura 2. Nestas formas de realização, os dados 130 de imagem são subtraídos de 132 com os dados representando tramas 145 de vídeo anteriores, resultando numa imagem 133 diferencial, que é enviada para um módulo 134 de transformada. O módulo 134 de transformada pode utilizar DCT ou outros métodos de transformada para transformar a imagem. Geralmente, o resultado do processo de transformada será coeficientes (K), que

são, depois, enviados para um módulo 136 de quantificação para quantificação.

O módulo 136 de quantificação pode ter outras entradas, tais como entradas 131 de utilizador para estabelecer parâmetros de quantificação (QP) e para outras entradas. O módulo 136 de quantificação pode utilizar os coeficientes de transformação e os parâmetros de quantificação para determinar níveis de quantificação (L) na imagem de vídeo. O módulo 136 de quantificação pode utilizar métodos que empregam uma parte de mantissa e uma parte exponencial, no entanto, também se podem empregar outros métodos de quantificação nos módulos 136 de quantificação de formas de realização da presente invenção. Estes níveis 135 de quantificação e parâmetros 137 de quantificação são enviados para um módulo 138 de codificação, bem como para um módulo 140 de desquantificação (DQ).

A saída para o módulo 138 de codificação é codificada e transmitida para fora do codificador para descodificação imediata ou armazenamento. O módulo 138 de codificação pode utilizar codificação de comprimento variável (VLC) nos seus processos de codificação. O módulo 138 de codificação pode utilizar codificação aritmética no seu processo de codificação. A saída do módulo 138 de codificação são dados 139 codificados que podem ser transmitidos para o descodificador ou armazenados no dispositivo de armazenamento.

A saída do módulo 136 de quantificação também é recebida no módulo 140 de desquantificação para iniciar a reconstrução da imagem. Isto é feito para manter uma contabilização exacta de tramas anteriores. O módulo 140 de desquantificação executa um processo com, essencialmente, o efeito inverso do módulo 136 de

quantificação. Os níveis de quantificação ou valores (L) são desquantificados produzindo coeficientes de transformada. Os módulos 140 de desquantificação podem utilizar métodos que empregam uma parte de mantissa e uma parte exponencial, como aqui descrito.

Os coeficientes de transformada provenientes do módulo 140 de desquantificação são enviados para um módulo 142 de transformação inversa (IT), onde são submetidos a uma transformada inversa para uma imagem 141 diferencial. Esta imagem 141 diferencial é, depois, combinada com dados de tramas 145 de imagem anteriores para formar uma trama 149 de vídeo, que pode ser introduzida numa memória 146 de tramas para referência relativamente a uma sucessão de tramas.

A trama 149 de vídeo também pode servir como entrada num módulo 147 de estimativa de movimento, que também recebe dados 130 de imagem. Estas entradas podem ser utilizadas para prever semelhanças de imagem e ajudar a comprimir dados de imagem. A saída do módulo 147 de estimativa de movimento é enviada para o módulo 148 de compensação de movimento e combinada com dados de saída do módulo 138 de codificação, que são enviados para uma posterior descodificação e eventual visualização de imagens.

O módulo 148 de compensação de movimento utiliza os dados de imagem preditos para reduzir requisitos de dados de trama; a sua saída é subtraída de dados 130 de imagem de entrada.

Algumas formas de realização da presente invenção compreendem sistemas e métodos para descodificação de vídeo, como mostrado na Figura 3. Um descodificador de formas de realização da presente invenção pode receber dados 150

codificados para um módulo 152 de descodificação. Os dados 150 codificados podem compreender dados que foram codificados por um codificador 100, tal como descrito com referência à Figura 2.

O módulo 152 de descodificação pode empregar métodos de descodificação de comprimento variável, se forem utilizados no processo de codificação. Outros métodos de descodificação também podem ser utilizados como ditado pelo tipo de dados 150 codificados. O módulo 152 de descodificação executa, essencialmente, o processo inverso do módulo 138 de codificação. A saída do módulo 152 de descodificação pode compreender parâmetros 156 de quantificação e valores 154 de quantificação. Outra saída pode compreender dados de estimativa de movimento e dados de predição de imagem que podem ser enviados directamente para um módulo 166 de compensação de movimento.

Tipicamente, os parâmetros 156 de quantificação e valores 154 de quantificação são enviados para um módulo 158 de desquantificação, onde os valores de quantificação voltam a ser convertidos para coeficientes de transformada. O módulo 158 de desquantificação pode utilizar métodos que empregam uma parte de mantissa e uma parte exponencial, como aqui descrito. Estes coeficientes são, depois, enviados para um módulo 160 de transformação inversa para conversão, novamente, para os dados 161 de imagem de domínio espacial.

A unidade 166 de compensação de movimento utiliza dados de vector de movimento e a memória 164 de tramas para construir uma imagem 165 de referência.

Os dados 161 de imagem representam uma imagem diferencial que tem que ser combinada com os dados 165 de imagem anteriores

para formar uma trama 163 de vídeo. Esta trama 163 de vídeo é enviada 168 para posterior processamento, exibição ou para outros fins e pode ser armazenada na memória 164 de tramas e utilizada para referência para tramas subsequentes.

Em algumas formas de realização da presente invenção, como ilustrado na Figura 4, os dados 102 de imagem podem ser enviados para um codificador ou parte 104 de codificação, para os diversos processos de transformação, quantificação, codificação e outros processos típicos de codificação de vídeo, como descrito acima para algumas formas de realização da presente invenção. A saída do codificador pode, depois, ser armazenada em quaisquer meios 106 de armazenamento legíveis por computador. Os meios 106 de armazenamento podem funcionar como uma memória tampão de curto prazo ou como um dispositivo de armazenamento de longo prazo.

Quando desejado, os dados de vídeo codificados podem ser extraídos dos meios 106 de armazenamento e decodificados por um decodificador ou parte 108 de decodificação para serem enviados 110 para um monitor ou outro dispositivo.

Em algumas formas de realização da presente invenção, como ilustrado na Figura 5, os dados 112 de imagem podem ser enviados para um codificador ou parte 114 de codificação para os vários processos de transformação, quantificação, codificação e outros processos típicos de codificação de vídeo, como descrito acima para algumas formas de realização da presente invenção. A saída do codificador pode, depois, ser enviada através de uma rede, tal como uma LAN, WAN ou a Internet 116. Um dispositivo de armazenamento, tal como os meios 106 de armazenamento, pode fazer parte de uma rede. Os dados de vídeo codificados podem ser

recebidos e decodificados por um decodificador ou parte 118 de decodificação, que também comunica com a rede 116. O decodificador 118 pode, então, decodificar os dados para consumo 120 local.

Em algumas formas de realização da presente invenção, como ilustrado na Figura 6, um método ou aparelho de quantificação compreende uma parte 172 de mantissa e uma parte 174 exponencial. Os parâmetros 176 de quantificação são inseridos em ambas as partes 172 e 174. Um coeficiente K 170 é inserido na parte 172 de mantissa, onde é modificado utilizando o parâmetro de quantificação e outros valores, como explicado acima. O resultado desta operação é combinado com o resultado produzido na parte exponencial utilizando o parâmetro de quantificação, produzindo, desse modo, um nível ou valor L 178 de quantificação.

Em algumas formas de realização da presente invenção, como ilustrado na Figura 7, um método ou aparelho de quantificação compreende uma parte 182 de mantissa e uma parte 184 de deslocamento. Os parâmetros 186 de quantificação são inseridos em ambas as partes 182 e 184. Um coeficiente K 180 é inserido na parte 182 de mantissa, onde é modificado utilizando o parâmetro de quantificação e outros valores, como explicado acima. O resultado desta operação é, ainda, processado na parte de deslocamento utilizando o parâmetro de quantificação, produzindo, desse modo, um nível ou valor L 188 de quantificação.

Algumas formas de realização da presente invenção, como ilustrado na Figura 8, compreendem um método ou aparelho de desquantificação com uma parte 192 de mantissa e uma parte 194

exponencial. Os parâmetros 196 de quantificação são inseridos em ambas as partes 192 e 194. Um valor L 190 de quantificação é inserido na parte 192 de mantissa, onde é modificado utilizando o parâmetro de quantificação e outros valores, como explicado acima. O resultado desta operação é, ainda, processado na parte exponencial utilizando o parâmetro de quantificação, produzindo, desse modo, um coeficiente X1 198.

Algumas formas de realização da presente invenção, como ilustrado na Figura 9, compreendem um método ou aparelho de desquantificação com uma parte 202 de mantissa e uma parte 204 de deslocamento. Os parâmetros 206 de quantificação são inseridos em ambas as partes 202 e 204. Um valor L 200 de quantificação é inserido na parte 202 de mantissa, onde é modificado utilizando o parâmetro de quantificação e outros valores, como explicado acima. O resultado desta operação é, ainda, processado na parte exponencial utilizando o parâmetro de quantificação, produzindo, desse modo, um coeficiente X1 208.

Algumas formas de realização da presente invenção podem ser armazenadas em meios legíveis por computador, tais como meios magnéticos, meios ópticos e outros meios, bem como combinações de meios. Algumas formas de realização também podem ser transmitidas como sinais através de redes e meios de comunicação. Estas transmissões e acções de armazenamento podem ocorrer como parte da operação de formas de realização da presente invenção ou como uma forma de transmitir a forma de realização para um destino.

Lisboa, 17 de Maio de 2012

REIVINDICAÇÕES

1. Método de descodificação de vídeo para a reconstrução de uma amostra X' de um valor L quantificado, compreendendo:

um passo (190) para a introdução do referido valor L quantificado;

um passo para realizar a transformação inteira inversa do referido valor L para obter um valor Y' , em que todas as funções de base da transformada inteira inversa têm uma norma igual, e

um passo para realizar a desquantificação e normalização combinadas do valor Y' de transformada inversa do valor L quantificado, para obter um valor X' desquantificado e normalizado representando a amostra reconstruída, em que o passo de desquantificação e normalização combinadas inclui:

introduzir um parâmetro QP de quantificação (196),
e

obter o referido o valor X' desquantificado e normalizado representando a amostra reconstruída;

caracterizado por

se obter o referido valor X' desquantificado e normalizado representando a amostra reconstruída por utilização de uma parte $B_m(QP)$ de mantissa que é uma função do referido

parâmetro QP de quantificação e uma parte Be(QP) exponencial que é uma função do referido parâmetro QP de quantificação e representando um deslocamento de normalização, como

$$X' = Y' * Bm(QP) * 2^{Be(QP)} * 2^{-N}, \quad |$$

em que N é um número inteiro, a multiplicação por 2^{-N} representa um deslocamento de normalização final e a referida parte Bm(QP) de mantissa e referida parte Be(QP) exponencial satisfazem as relações de recursão

$$Bm(QP + 6) = Bm(QP), \quad |$$

e

$$Be(QP + 6) = Be(QP) + 1.$$

Lisboa, 17 de Maio de 2012

FIG.1

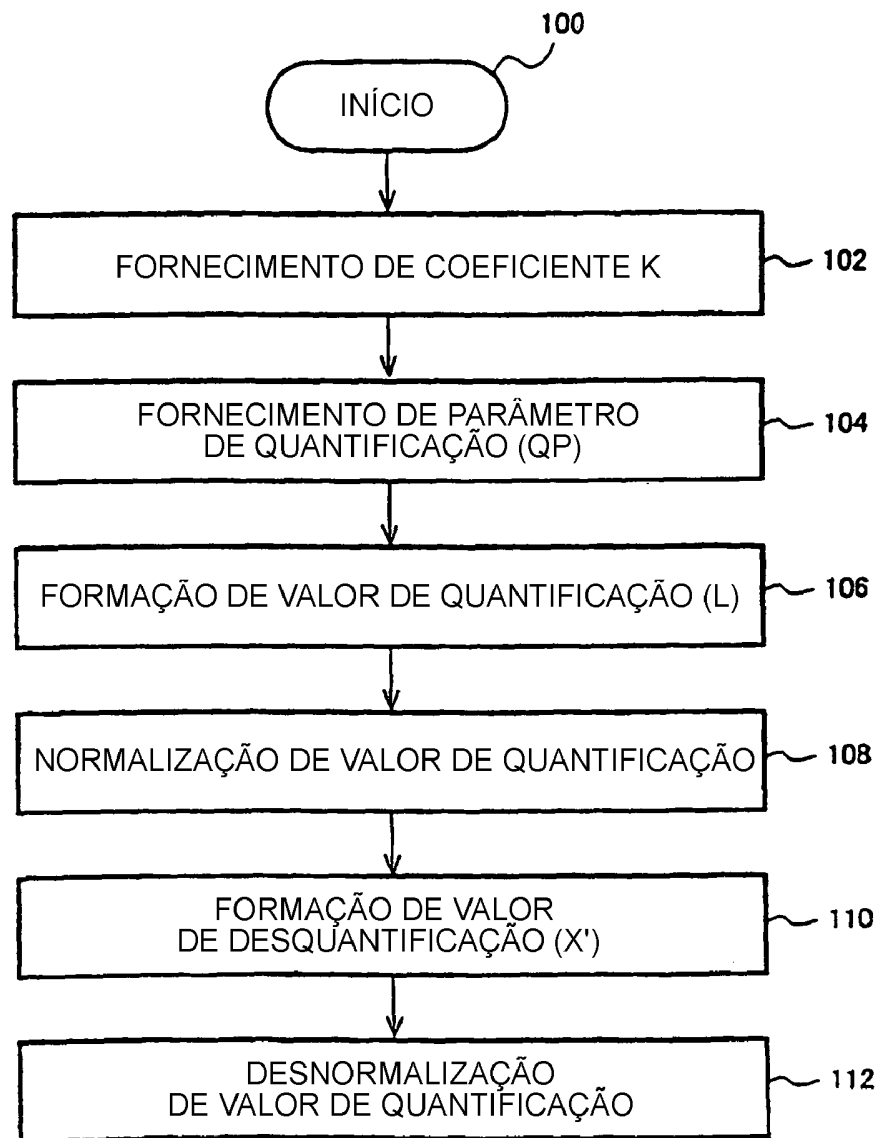
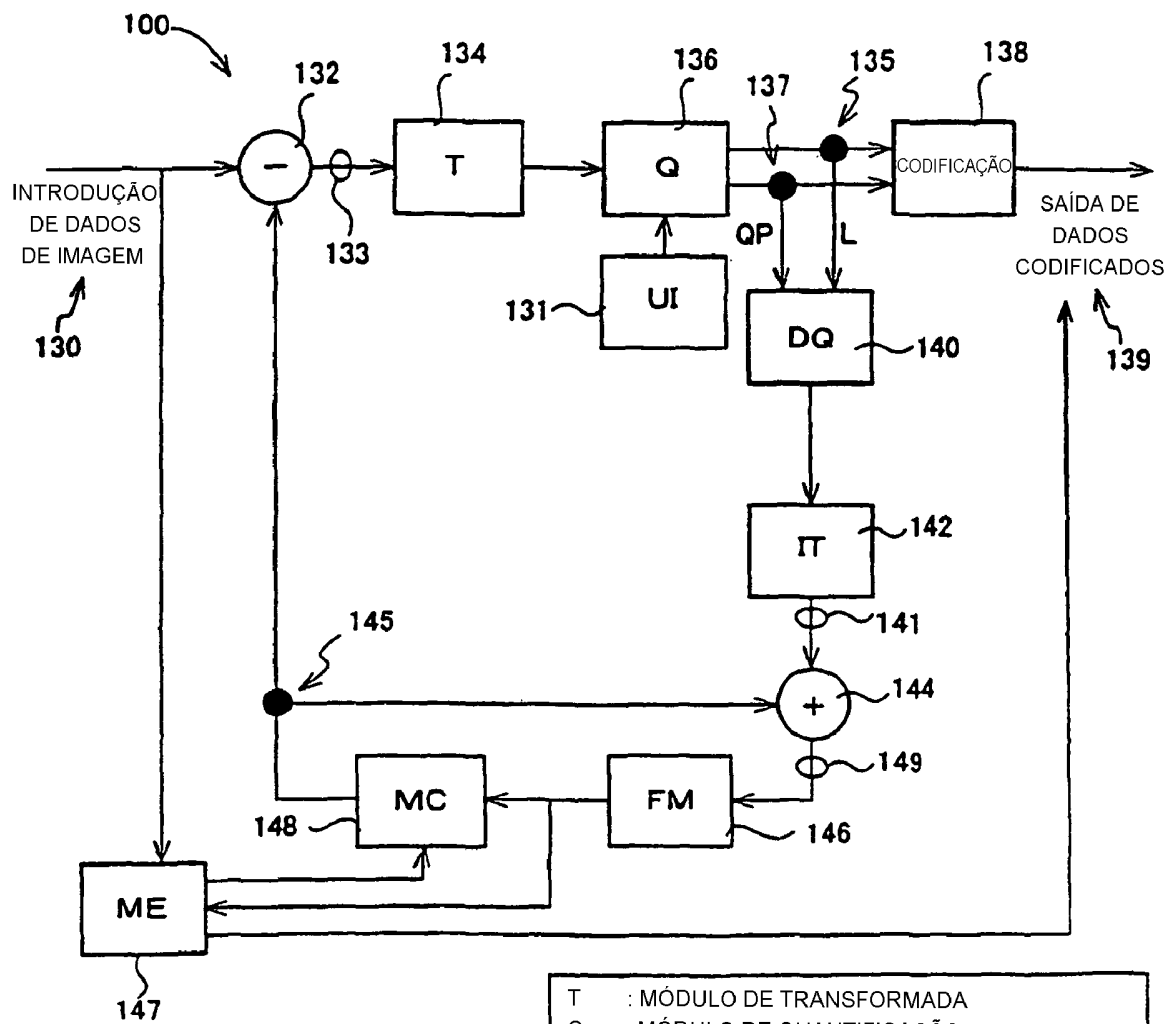


FIG.2



T	: MÓDULO DE TRANSFORMADA
Q	: MÓDULO DE QUANTIFICAÇÃO
DQ	: MÓDULO DE DESQUANTIFICAÇÃO
IT	: MÓDULO DE TRANSFORMADA INVERSA
FM	: MEMÓRIA DE TRAMAS
MC	: MÓDULO DE COMPENSAÇÃO DE MOVIMENTO
ME	: MÓDULO DE ESTIMATIVA DE MOVIMENTO

FIG.3

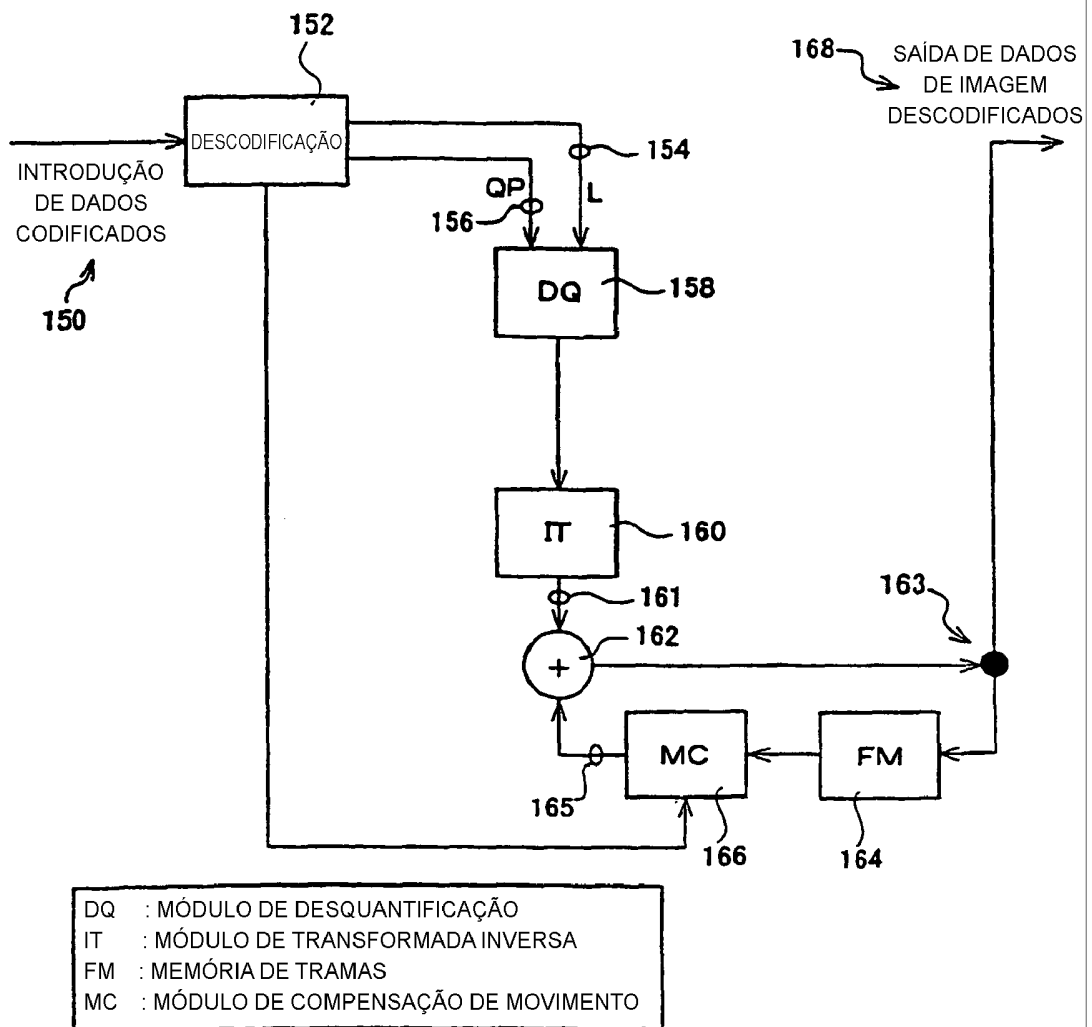


FIG.4

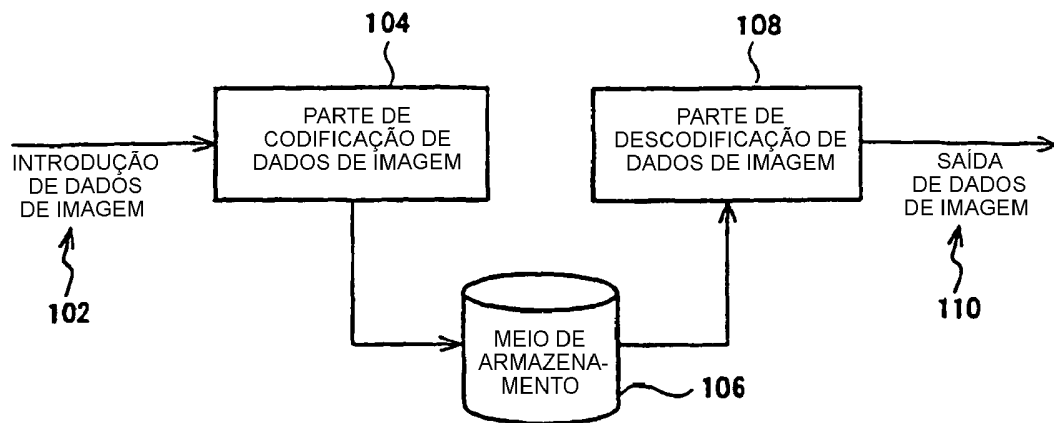


FIG.5

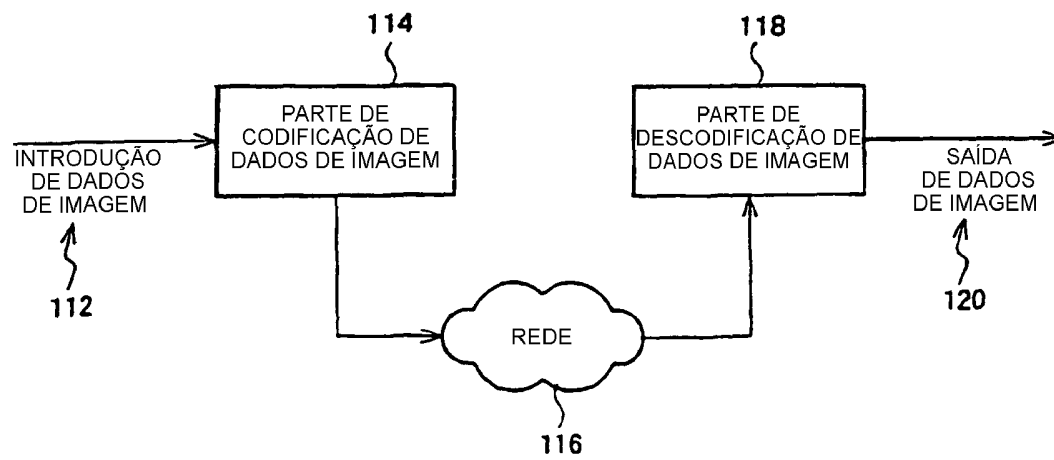


FIG.6

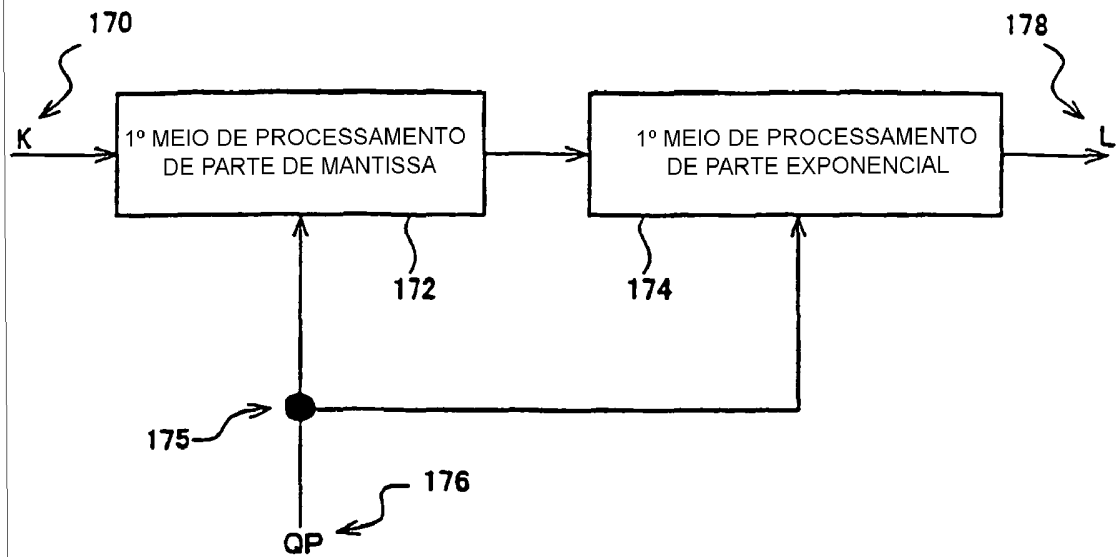


FIG.7

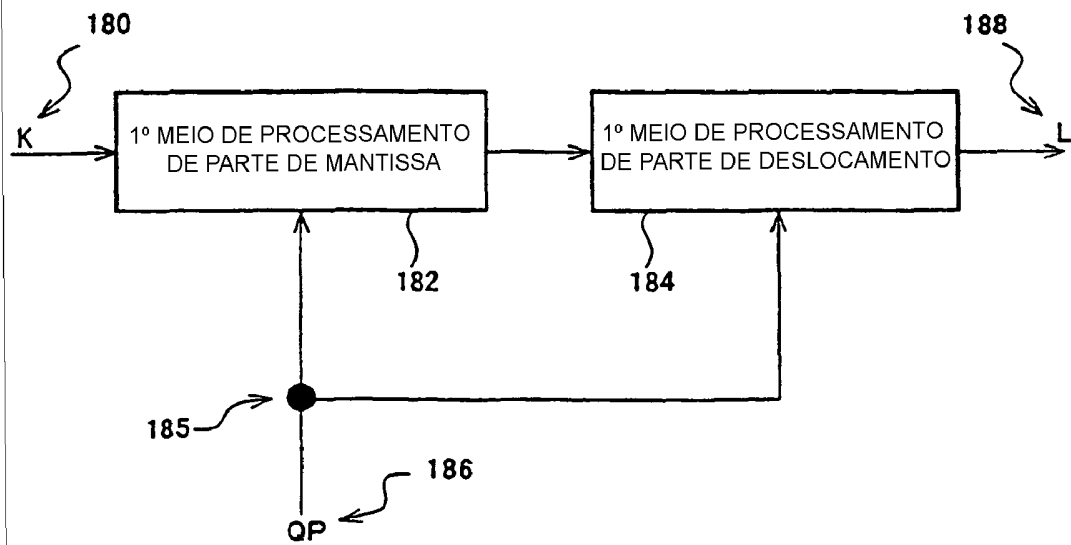


FIG.8

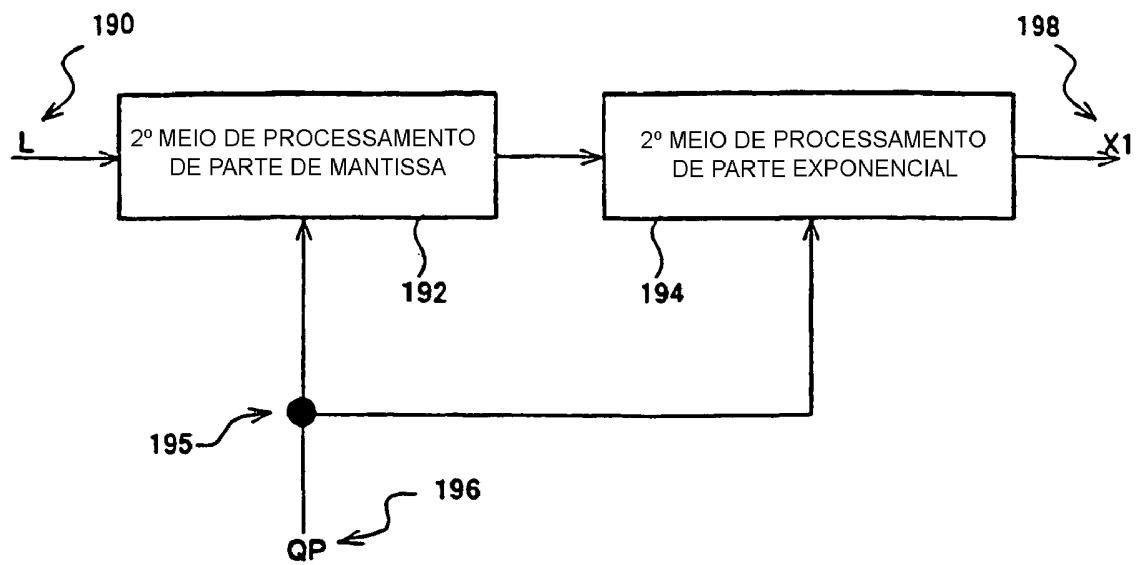


FIG.9

