



(22) Data de Depósito: 19/10/2006
(43) Data da Publicação: 06/12/2011
(RPI 2135)



(51) Int.Cl.:
H04N 7/30

(87) Publicação Internacional: WO 2007/094100de
23/08/2007

Diagrama de um sistema de codificação de vídeo, conforme a Figura 1. O sistema inclui um controlador de codificação (110) que gerencia o processo. A entrada de imagem (115) é processada por um preditor (101) que utiliza uma memória de referência (107) para gerar uma imagem preditiva (102). A diferença entre a entrada e a preditiva é enviada para uma unidade de decisão de modo (109), que também recebe um sinal de erro de previsão (102). Esta unidade seleciona entre uma matriz de quantização (104) e uma matriz de inversão (105) com base em um parâmetro de quantização (117) e um coeficiente de transformação (112). O resultado é enviado para um processador de codificação (111), que gera o armazenamento temporário de saída (120) e os dados de código (115).



Relatório Descritivo da Patente de Invenção para **"MÉTODO E APARELHO E PROGRAMA DE CODIFICAÇÃO/DECODIFICAÇÃO DE VÍDEO"**.

CAMPO DA TÉCNICA

5 A presente invenção refere-se a um método e aparelho de codificação/decodificação de vídeo que utiliza matrizes de quantização.

ANTECEDENTES DA TÉCNICA

10 É proposto um sistema para quantizar um coeficiente de DCT fazendo uma alocação de bit a cada posição de frequência, utilizando uma característica de frequência de coeficientes de DCT providos sujeitando um vídeo a uma transformação ortogonal, por exemplo, uma transformada de co-seno discreta (DCT) (W. H. Chen e C.H.Smith, "Adaptive Coding of Monochrome and Color Images", IEEE Trans. On Comm. Vol. 25, Número 11, Nov. 1977). De acordo com este sistema convencional, muitos bits são alo-

15 cados para um domínio de frequência de baixo nível para obter as informações de coeficiente, enquanto que poucos bits são alocados para um domínio de frequência de passagem alta, por meio de que o coeficiente de DCT é quantizado com eficiência. No entanto, este sistema convencional precisa para preparar uma tabela de alocação de acordo com a grossura de quantização. Portanto, não é sempre um sistema efetivo em termos de quantização robusta.

20

A ITU-TT.81 e a ISO/IEC10918-1 (referidas como JPEG: Joint Photographic Experts Group daqui em diante) induziram em ITU-T e ISO/IEC quantizar igualmente os coeficientes de transformação ao longo da faixa de

25 frequência inteira com a mesma escala de quantização. No entanto, os humanos são comparativamente insensíveis à região de alta frequência de acordo com a propriedade visual humana. Por esta razão, o seguinte sistema é proposto. Isto é, em JPEG, a ponderação é feita a cada domínio de frequência para mudar a escala de quantização, de modo que muitos bits são

30 atribuídos para um domínio de baixa frequência visualmente sensível e a taxa de bits é diminuída no domínio de alta frequência, resultando em aperfeiçoar uma qualidade de imagem subjetiva. Este sistema executa uma

quantização a cada bloco de quantização de conversão. Uma tabela utilizada para esta quantização é referida como uma matriz de quantização.

Ainda em anos recentes o método de codificação de vídeo que aperfeiçoa grandemente a eficiência de codificação é induzido como ITU-TRec.H.264 e ISO/IEC14496-10 (referido como H.264) em combinação com ITU-T e ISO/IEC. Os sistemas de codificação convencionais tais como ISO/IECMPEG-1, 2, 4, e ITU-T.H.261, H.263 quantizam os coeficientes de DCT após a transformada ortogonal para reduzir o número de bits codificados dos coeficientes de transformada. Em um perfil principal de H.264, como a relação entre um parâmetro de quantização e uma escala de quantização está projetada de tal modo que estes tornam-se um intervalo igual em uma longa escala, a matriz de quantização não é introduzida. No entanto, em um alto perfil de H.264, a matriz de quantização deve ser novamente introduzida para aperfeiçoar uma qualidade de imagem subjetiva para a imagem de alta resolução (referir a Jiuhuai Lu, "Proposal of quantization weighting for H.264/MPEG-4 AVC Professional Profiles", JVT of ISO/IEC Mpeg & ITU-T VCEG, JVT-Ko29, março de 2004).

No alto perfil de H.264, um total de oito tipos de diferentes matrizes de quantização pode ser estabelecido em correspondência com dois blocos transformados/quantizados (um bloco de 4x4 pixels e um bloco de 8x8 pixels) para cada modo de codificação (predição intraquadro ou predição interquadros) e para cada sinal (um sinal de luminescência ou um sinal de diferença de cor).

Como a matriz de quantização é empregada para ponderar o pixel de acordo com cada posição de componente de frequência no momento da quantização, a mesma matriz de quantização é necessária no momento da dequantização, também. Em um codificador do alto perfil de H.264, as matrizes de quantização utilizadas são codificadas e multiplexadas e então transmitidas para um decodificador. Concretamente, um valor de diferença é calculado em ordem de escaneamento zigzag ou escaneamento de campo de um componente CC da matriz de quantização, e os dados de diferença obtidos são sujeitos a uma codificação de comprimento variável e mul-

tiplexados como dados de código.

Por outro lado, um decodificador do alto perfil de H.264 de codifica os dados de código de acordo com uma lógica similar ao codificador para reconstruí-los como uma matriz de quantização a ser utilizada no momento da dequantização. A matriz de quantização é finalmente sujeita a uma codificação de comprimento variável. Neste caso, o número de bits codificados da matriz de quantização requer 8 bits no mínimo e não menos do que 1500 bits no máximo na sintaxe.

Um método para transmitir uma matriz de quantização de alto perfil de H.264 pode aumentar um suplemento para codificar a matriz de quantização e assim diminuir grandemente a eficiência de codificação, em uma aplicação utilizada a uma baixa taxa de bits tal como um telefone celular ou um dispositivo móvel.

Um método para ajustar um valor de uma matriz de quantização pela transmissão de uma matriz de quantização de base primeiramente para atualizar a matriz de quantização com um pequeno suplemento e então transmitir um coeficiente k que indica um grau de mudança da matriz de quantização para um decodificador está proposto (referir a JP-A 2003-189308 (KOKAI)).

A JP-A 2003-189308 (KOKAI): "Aparelho de codificação de vídeo, método de codificação, aparelho de decodificação e método de decodificação, e método de transmissão de cadeia de códigos de vídeo" objetiva atualizar uma matriz de quantização a cada tipo de imagem com um pequeno número de bits codificados, e torna possível atualizar a matriz de quantização de base a aproximadamente 8 bits no máximo. No entanto, como este é um sistema para enviar somente um grau de mudança da matriz de quantização de base, a amplitude da matriz de quantização pode ser mudada mas é impossível mudar a sua característica. Ainda, esta precisa transmitir a matriz de quantização de base, e assim o número de bits codificados pode aumentar grandemente devido à situação de codificação.

DESCRIÇÃO DA INVENÇÃO

Quando uma matriz de quantização é codificada por um método prescrito pelo alto perfil de H.264, e transmitida para um decodificador, o número de bits codificados para codificar a matriz de quantização aumenta.

- 5 Quando a matriz de quantização é transmitida a cada quadro novamente, o número de bits codificados para codificar a matriz de quantização cresce adicionalmente. Ainda, quando um grau de mudança da matriz de quantização é transmitido, os graus de liberdade para mudar a matriz de quantização são grandemente limitados. Existe um problema que estes resultados tornam difícil utilizar a matriz de quantização efetivamente.

- Um aspecto da presente invenção provê um método de codificação de vídeo que compreende: gerar uma matriz de quantização utilizando uma função referente à geração da matriz de quantização e um parâmetro relativo à função; quantizar um coeficiente de transformada referente a um
- 15 sinal de imagem de entrada utilizando a matriz de quantização para gerar um coeficiente de transformada quantizado; e codificar o parâmetro e o coeficiente de transformada quantizado para gerar um sinal de código.

BREVE DESCRIÇÃO DOS DESENHOS

- Figura 1 é um diagrama de blocos que ilustra uma estrutura de um aparelho de codificação de vídeo de acordo com uma primeira modalidade.

Figura 2 é um diagrama de blocos que ilustra uma estrutura de uma matriz de quantização de acordo com a primeira modalidade.

- Figura 3 é um fluxograma do aparelho de codificação de vídeo de acordo com a primeira modalidade.

Figura 4A é um diagrama esquemático de uma ordem de predição/forma de bloco relativo à primeira modalidade.

Figura 4B é um diagrama que ilustra uma forma de bloco de 16x16 pixels.

- Figura 4C é um diagrama que ilustra uma forma de bloco de 4x4 pixels.

Figura 4D é um diagrama que ilustra uma forma de bloco de 8x8

pixels.

Figura 5A é um diagrama que ilustra uma matriz de quantização que corresponde a um bloco de 4x4 pixels relativa à primeira modalidade.

Figura 5B é um diagrama que ilustra uma matriz de quantização
5 que corresponde a um bloco de 8x8 pixels.

Figura 6A é um diagrama para explicar um método de geração de matriz de quantização relativo à primeira modalidade.

Figura 6B é um diagrama para explicar outro método de geração de matriz de quantização.

Figura 6C é um diagrama para explicar outro método de geração de matriz de quantização.
10

Figura 7 é um diagrama esquemático de uma estrutura de sintaxe de acordo com a primeira modalidade.

Figura 8 é um diagrama de uma estrutura de dados de uma sintaxe de ajuste de parâmetro de seqüência de acordo com a primeira modalidade.
15

Figura 9 é um diagrama de uma estrutura de dados de uma sintaxe de ajuste de parâmetro de imagem de acordo com a primeira modalidade.

Figura 10 é um diagrama de uma estrutura de dados de uma sintaxe de ajuste de parâmetro de imagem de acordo com a primeira modalidade.
20

Figura 11 é um diagrama de uma estrutura de dados de uma sintaxe mental suplementar de acordo com a primeira modalidade.

Figura 12 é um fluxograma de codificação de múltiplos percursos de acordo com a segunda modalidade.
25

Figura 13 é um diagrama que ilustra uma estrutura de sintaxe em uma sintaxe de cabeçalho de fatia.

Figura 14 é um diagrama que ilustra uma sintaxe de cabeçalho de fatia.
30

Figura 15 é um diagrama que ilustra uma sintaxe de cabeçalho de fatia.

Figura 16 é um diagrama que ilustra um exemplo de um CurrSliceType.

Figura 17 é um diagrama que ilustra uma sintaxe de cabeçalho de fatia.

5 Figura 18 é um diagrama de blocos que ilustra uma estrutura de um aparelho de decodificação de vídeo de acordo com a terceira modalidade da presente invenção.

Figura 19 é um fluxograma do aparelho de decodificação de vídeo de acordo com a terceira modalidade da presente invenção.

10 MELHOR MODO PARA EXECUTAR A INVENÇÃO

Serão agora descritas as modalidades da presente invenção em detalhes em conjunto com os desenhos.

PRIMEIRA MODALIDADE: CODIFICAÇÃO

De acordo com a primeira modalidade mostrada na Figura 1, um
15 sinal de vídeo é dividido em uma pluralidade de blocos de pixels e inserido
em um aparelho de codificação de vídeo 100 como um sinal de imagem de
entrada 116. O aparelho de codificação de vídeo 100 tem, como modos exe-
cutados por um preditor 101, uma pluralidade de modos de predição diferen-
tes em tamanho de bloco ou em método de geração de sinal preditivo. Na
20 presente modalidade é assumido que a codificação é feita da esquerda su-
perior do quadro para a sua direita inferior como mostrado na Figura 4A.

O sinal de imagem de entrada 116 inserido no aparelho de codificação de vídeo 100 é dividido em uma pluralidade de blocos cada um contendo 16x16 pixels como mostrado na Figura 4B. Uma parte do sinal de imagem de entrada 116 é inserida no preditor 101 e codificada por um codificador 111 através de uma unidade de determinação de modo 102, um transformador 103 e um quantizador 104. Este sinal de imagem codificado é armazenado em um armazenamento temporário de saída 120 e então é emitido como dados codificados 115 no tempo de saída controlado por um controlador de codificação 110.

Um bloco de 16x16 pixels mostrado na Figura 4B é referido como um macrobloco e tem um tamanho de bloco de processo básico para o

processo de codificação seguinte. O aparelho de codificação de vídeo 100 lê o sinal de imagem de entrada 116 em unidades de bloco e codifica-o. O macrobloco pode estar em unidades de bloco de 32x32 pixels ou em unidades de bloco de 8x8 pixels.

- 5 O preditor 101 gera um sinal de imagem preditiva 118 com todos os modos selecionáveis no macrobloco utilizando uma imagem de referência codificada armazenada em uma memória de imagem de referência 107. O preditor 101 gera todos os sinais de imagem preditiva para todos os modos de codificação nos quais um bloco de pixels de objeto pode ser codificado.
- 10 No entanto, quando a próxima predição não pode ser feita sem gerar uma imagem decodificada local no macrobloco como a predição intraquadro de H.264 petição de 4x4 pixels (figura 4 C) ou predição de 8x8 pixels (Figura 4D), o preditor 101 pode executar uma transformação ortogonal e quantização, e dequantização e transformação inversa.
- 15 O sinal de imagem preditiva 118 gerado com o preditor 101 é inserido em uma unidade de determinação de modo 102 juntamente com o sinal de imagem de entrada 116. A unidade de determinação de modo 102 insere o sinal de imagem preditiva 118 em um transformador inverso 106, gera um sinal de erro preditivo 119 subtraindo o sinal de imagem preditiva
- 20 118 do sinal de imagem de entrada 116, e insere-o no transformador 103. Ao mesmo tempo, a unidade de determinação de modo 102 determina um modo com base em informações de modo preditas com o preditor 101 e o sinal de erro preditivo 119. Explicando mais concretamente, o modo é determinado utilizando um custo k conhecido pela equação (1) seguinte nesta modalidade.
- 25

$$K = SAD + \lambda \times OH \quad (1)$$

- onde OH indica mais informações, SAD é a soma absoluta de sinais de erro preditivo, e λ é uma constante. A constante λ é determinada com base em um valor de uma largura de quantização ou um parâmetro de
- 30 quantização. Deste modo, o modo é determinado com base no custo K. O modo no qual o custo K indica o menor valor é selecionado como um modo ótimo.

Nesta modalidade, a soma absoluta das informações de modo e do sinal de erro preditivo é utilizada.

Como outra modalidade, o modo pode ser determinado utilizando somente as informações de modo ou somente a soma absoluta do sinal de erro preditivo. Alternativamente, uma transformação de Hadamard pode ser sujeita a estes parâmetros para obter e utilizar um valor aproximado. Ainda o custo pode ser calculado utilizando uma atividade de um sinal de imagem de entrada, e uma função de custo pode ser calculada utilizando uma largura de quantização e um parâmetro de quantização.

Um codificador provisório é preparado de acordo com outra modalidade para calcular o custo. Um sinal de erro preditivo é gerado com base no modo de codificação do codificador provisório. O sinal de erro preditivo é realmente codificado para produzir os dados de código. Os dados de imagem decodificada local 113 são produzidos por decodificação local dos dados de código. O modo pode ser determinado utilizando o número de bits codificados dos dados de código e um erro quadrático do sinal de imagem decodificada local e o sinal de vídeo de entrada 116. Uma equação de determinação de modo deste caso está expressa pela equação (2) seguinte.

$$J = D + \lambda \times R \quad (2)$$

onde J indica um custo, D indica uma distorção de codificação que representa um erro quadrático do sinal de vídeo de entrada 116 e do sinal de imagem decodificada local 113, e R representa o número de bits codificados estimados por codificação temporária. Quando este custo J é utilizado, uma escala de circuito aumenta porque a codificação temporária e a decodificação local (dequantização e transformação inversa) são necessárias a cada modo de codificação. No entanto, o número preciso de bits codificados e de distorção de codificação pode ser utilizado, e a alta eficiência de codificação pode ser mantida. O custo pode ser calculado utilizando somente o número de bits codificados ou somente a distorção de codificação. A função de custo pode ser calculada utilizando um valor aproximado a estes parâmetros.

A unidade de determinação de modo 102 está conectada no transformador 103 e no transformador inverso 106. As informações de modo selecionadas com a unidade de determinação de modo 102 e o sinal de erro preditivo 118 são inseridas no transformador 103. O transformador 103
5 transforma o sinal de erro preditivo 118 inserido em coeficientes de transformada e gera os dados de coeficiente de transformada. O sinal de erro preditivo 118 é sujeito a uma transformada ortogonal utilizando uma transformada de co-seno discreta (DCT), por exemplo. Como uma modificação, o coeficiente de transformada pode ser gerado utilizando uma técnica tal como
10 a transformada de ondulação ou a análise de componente independente.

O coeficiente de transformada provido do transformador 103 é enviado para o quantizador 104 e por meio deste quantizado. O parâmetro de quantização necessário para a quantização é ajustado para o controlador de quantização 110. O quantizador 104 quantiza o coeficiente de transformada utilizando a matriz de quantização 114 inserida de um gerador de matriz de quantização 109 e gera um coeficiente de transformada quantizado
15 112.

O coeficiente de transformada quantizado 112 é inserido no processador de codificação 111 juntamente com as informações sobre os métodos de predição tais como as informações de modo e o parâmetro de quantização. O processador de codificação 111 sujeita o coeficiente de transformada quantizado 112 juntamente com as informações de modo de entrada à codificação de entropia (codificação de Huffman ou codificação aritmética). Os dados de código 115 providos pela codificação de entropia do
20 processador de codificação 111 são emitidos do codificador de vídeo 100 para o armazenamento temporário de saída 120 e multiplexados. Os dados de código multiplexados são transmitidos do armazenamento temporário de saída 120.
25

Quando a matriz de quantização 114 a ser utilizada para a quantização é gerada, as informações de instrução que indicam a utilização da matriz de quantização são providas para o gerador de parâmetros de geração 108 pelo controlador de codificação 110. O gerador de parâmetros de
30

geração 108 ajusta o parâmetro de geração de matriz de quantização 117 de acordo com as informações de instrução, e emite-o para o gerador de matriz de quantização 109 e o processador de codificação 111.

5 O parâmetro de geração de matriz de quantização 117 pode ser ajustado por uma unidade de ajuste de parâmetro externa (não-mostrada) controlada pelo controlador de codificação 110. Também, este pode ser atualizado em unidades de bloco de imagem codificada, em unidades de fatia ou em unidades de imagem. o gerador de parâmetros de geração 108 compreende uma função para controlar um tempo de ajuste do parâmetro de
10 geração de matriz de quantização 117.

O gerador de matriz de quantização 109 gera uma matriz de quantização 114 por um método estabelecido para o parâmetro de geração de matriz de quantização 117 e emite-a para o quantizador 104 e o dequantizador 105. Ao mesmo tempo, o parâmetro de geração de matriz de quantização 117 inserido no processador de codificação 111 está sujeito à codificação de entropia juntamente com as informações de modo e o coeficiente de transformada 112 os quais são inseridos do quantizador 104.
15

O dequantizador 105 dequantiza o coeficiente de transformada 112 quantizado com o quantizador 104 de acordo com o parâmetro de quantização ajustado pelo controlador de codificação 110 e a matriz de quantização 114 inserida do gerador de matriz de quantização 109. O coeficiente de transformada dequantizado é enviado para o transformador inverso 106. O transformador inverso 106 sujeita o coeficiente de transformada dequantizado à transformada inversa (por exemplo, uma transformada de co-seno discreta inversa) para decodificar o sinal de erro preditivo.
20
25

O sinal de erro preditivo 116 decodificado com o transformador inverso 106 é somado ao sinal de imagem preditiva 118 para um modo de determinação, o qual é suprido da unidade de determinação de modo 102. O sinal de adição do sinal de erro preditivo e do sinal de imagem preditiva 118 torna-se um sinal decodificado local 113 e é inserido na memória de referência 107. A memória de imagem de referência 107 armazena o sinal decodificado local 113 como uma imagem de reconstrução. A imagem armazenada
30

na memória de imagem de referência 107 deste modo torna-se uma imagem de referência referida quando o preditor 101 gera um sinal de imagem preditiva.

Quando um laço de codificação (um processo a ser executado na ordem do preditor 101 → a unidade de determinação de modo 102 → o transformador 103 → o quantizador 104 → o dequantizador 105 → o transformador inverso 106 → a memória de imagem de referência 107 na Figura 1) é executado para todos os modos selecionáveis para um macrobloco de objeto, um laço é completado. Quando o laço de codificação é completado para o macrobloco, o sinal de imagem de entrada 116 do próximo bloco é inserido e codificado. O gerador de matriz de quantização 108 não precisa gerar uma matriz de quantização a cada macrobloco. A matriz de quantização gerada é mantida a menos que o parâmetro de geração de matriz de quantização 117 ajustado pelo gerador de parâmetro de geração 108 seja atualizado.

O controlador de codificação 110 executa um controle de retorno do número de bits codificados, um seu controle de característica de quantização, um controle de determinação de modo, etc. Também, o controlador de codificação 110 executa um controle de taxa para controlar o número de bits codificados, um controle do preditor 101, e um controle de um parâmetro de entrada externo. Ao mesmo tempo, o controlador de codificação 110 controla o armazenamento temporário de saída 120 para emitir os dados de código para um externo em um tempo apropriado.

O gerador de matriz de quantização 109 mostrado na figura 2 gera a matriz de quantização 114 com base no parâmetro de geração de matriz de quantização de entrada 117. A matriz de quantização é uma matriz como mostrada na Figura 5A ou uma matriz como mostrada na Figura 5B. A matriz de quantização está sujeita à ponderação por um fator de ponderação correspondente a cada ponto de frequência no caso de quantização e dequantização. A Figura 5A mostra uma matriz de quantização que corresponde a um bloco de 4x4 pixels e a Figura 5B mostra uma matriz de quantização que corresponde a um bloco de 8x8 pixels. O gerador de matriz de

quantização 109 compreende uma unidade de decifração de parâmetro gerado 201, uma chave 202 e um ou mais geradores de matriz 203. A unidade de decifração de parâmetro gerado 201 decifra o parâmetro de geração de matriz de quantização 117, e emite as informações de mudança pela chave 5 202 de acordo com cada método de geração de matriz. Estas informações de mudança são ajustadas pelo controlador de geração de matriz de quantização 210 e muda o terminal de saída da chave 202.

A chave 202 é comutada de acordo com as informações de chave providas pela unidade de decifração de parâmetro gerado 201 e ajustadas pelo controlador de geração de matriz de quantização 210. Por exemplo, 10 quando o tipo de geração de matriz do parâmetro de geração de matriz de quantização 117 é um primeiro tipo, a chave 202 conecta o terminal de saída da unidade de decifração de parâmetro gerado 201 com o gerador de matriz 203. Por outro lado, quando o tipo de geração de matriz do parâmetro de 15 geração de matriz de quantização 117 é um enésimo tipo, a chave 202 conecta o terminal de saída da unidade de decifração de parâmetro gerado 201 com o enésimo gerador de matriz 203.

Quando o tipo de geração de matriz do parâmetro de geração de matriz de quantização 117 é um M^{o} tipo ($N < M$) e o M^{o} gerador de matriz 20 203 não está incluído no gerador de matriz de quantização 109, a chave 202 é conectada no gerador de matriz correspondente por um método no qual o terminal de saída da unidade de decifração de parâmetro gerado 201 é determinado com antecedência. Por exemplo, quando um parâmetro de geração de matriz de quantização do tipo que não existe no gerador de matriz de 25 quantização 109 é inserido, a chave 202 sempre conecta o terminal de saída no primeiro gerador de matriz. Quando um tipo de geração de matriz similar é conhecido, esta pode ser conectada no gerador de matriz do L^{o} mais próximo no M^{o} tipo de entrada. Em qualquer caso, o gerado de matriz de quantização 109 conecta o terminal de saída da unidade de decifração de parâmetro gerado 201 do primeiro até o enésimo geradores de matriz 203 de 30 acordo com o parâmetro de geração de matriz de quantização 117 por um método de conexão predeterminado.

Cada gerador de matriz 203 gera a matriz de quantização 114 de acordo com as informações do parâmetro de geração de matriz de quantização correspondente. Concretamente, as informações de parâmetro de geração de matriz de quantização 117 são compostas de informações de parâmetro de um tipo de geração de matriz (T), um grau de mudança (A) de matriz de quantização, um grau de distorção (B) e um item de correção (C). Estes parâmetros são rotulados por diferentes nomes, mas podem ser utilizados em qualquer tipo de modo. Estes parâmetros são definidos como um conjunto de parâmetros expresso pela equação (3) seguinte:

$$QMP = (T, A, B, C) \quad (3)$$

QMP representa as informações de parâmetro de geração de matriz de quantização. O tipo de geração de matriz (T) indica qual gerador de matriz 203 que corresponde a qual tipo deve ser utilizado. Por outro lado, como utilizar o grau de mudança (A), o grau de distorção (B) e o item de correção (C) pode ser livremente definido a cada tipo de geração de matriz. O primeiro tipo de geração de matriz está explicado referindo à Figura 6A.

Uma função de geração de matriz quando o tipo de geração de matriz é 1 está representada pelas seguintes equações (4), (5) e (6):

[fórmula 1]

20 copiar fórmula

$$r = |x + y| \quad (4)$$

$$Q_{4 \times 4}(x, y) = a * r + c \quad (5)$$

$$Q_{8 \times 8}(x, y) = \frac{a}{2} * r + c \quad (6)$$

Ainda, exemplos de conversão de tabela do grau de mudança (A), do grau de distorção (B) e o do item de correção (C) utilizados para o primeiro tipo de matriz estão mostrados pelas seguintes equações (7), (8) e (9):

$$25 \quad a = 0,1 * A \quad (7)$$

$$B = 0 \quad (8)$$

$$c = 16 + C \quad (9)$$

onde o grau de mudança (A) representa um grau de mudança quando a distância do componente de CC para a posição de frequência da matriz de quantização é assumida ser r. Por exemplo, se o grau de mudança (A) for um valor positivo, o valor da matriz aumenta conforme a distância r aumenta. Neste caso, a largura de banda alta pode ser ajustada para um valor grande. Em contraste, se o grau de mudança (A) for um valor negativo, o valor da matriz aumenta com o aumento da distância r. Neste caso, a etapa de quantização pode ser ajustada grosseiramente na largura de banda baixa. No primeiro tipo de geração de matriz, um valor 0 é sempre ajustado sem utilizar o grau de distorção (B). Por outro lado, o item de correção (C) representa um segmento de uma linha reta expresso pelo grau de mudança (A). Como a primeira função de geração de matriz pode ser processada somente por uma operação de multiplicação, de adição, de subtração e de deslocamento é vantajoso porque um custo de hardware pode ser diminuído.

15 A matriz de quantização gerada com base nas equações (7), (8) e (9) no caso de QMP = (1, 40, 0, 0) é expressa pela seguinte equação (10): [fórmula 2]

$$Q_{4 \times 4}(x, y) = \begin{bmatrix} 16 & 20 & 24 & 28 \\ 20 & 24 & 28 & 32 \\ 24 & 28 & 32 & 36 \\ 28 & 32 & 36 & 40 \end{bmatrix} \quad (10)$$

Como a precisão de variável de cada um do grau de mudança (A), do grau de distorção (B) e do item de correção (C) influencia uma escala de hardware, é importante preparar uma tabela que tenha a boa eficiência em uma faixa decidida. Na equação (7), quando o grau de mudança (A) é assumido ser um inteiro não negativo de 6 bits, é possível obter um gradiente de 0 a 6,4. No entanto, um valor negativo não pode ser obtido. Consequentemente, não é possível obter uma faixa de -6,3 a 6,4 bits utilizando uma tabela de tradução que utiliza 7 bits como indicado pela equação (11)

25 seguinte:

$$\alpha = 0,1 \times (A - 63) \quad (11)$$

Se a tabela de tradução do grau de mudança (A), do grau de

distorção (B) e do item de correção (C) que corresponde a um tipo de geração de matriz (T) for provida, e a precisão do grau de mudança (A), do grau de distorção (B) e do item de correção (C) for adquirida a cada tipo de geração de matriz (T), é possível ajustar um parâmetro de geração de matriz de quantização apropriado de acordo com a situação de codificação e o ambiente de utilização. No primeiro tipo de geração de matriz expresso pelas equações (4), (5) e (6), o grau de distorção (B) torna-se sempre 0. Portanto, não é necessário transmitir um parâmetro que corresponda ao grau de distorção (B). Pelo tipo de geração de matriz, o número de parâmetros a serem utilizados pode ser diminuído. Neste caso, os parâmetros não utilizados não são codificados.

Subsequentemente, uma função de geração de matriz de quantização que utiliza uma função quadrática está mostrada como o segundo tipo de geração de matriz. O diagrama esquemático deste tipo de geração de matriz está mostrado na Figura 6C.

[fórmula 3]

$$Q_{4 \times 4}(x, y) = \frac{a}{4} * r^2 + \frac{b}{2} * r + c \quad (12)$$

$$Q_{8 \times 8}(x, y) = \frac{a}{16} * r^2 + \frac{b}{8} * r + c \quad (13)$$

Os parâmetros (A), (B) e (C) relativos às funções a, b e c, respectivamente, representam um grau de mudança, um valor de distorção e de correção da função quadrática. Estas funções são aptas para aumentar grandemente em valor especificamente conforme uma distância aumenta. Quando a matriz de quantização no caso de QMP = (2, 10, 1, 0) é calculada utilizando, por exemplo, as equações (4), (8) e (10), a matriz de quantização da equação (14) seguinte pode ser gerada.

[fórmula 4]

$$Q_{4 \times 4}(x, y) = \begin{bmatrix} 16 & 17 & 18 & 20 \\ 17 & 18 & 20 & 22 \\ 18 & 20 & 22 & 25 \\ 20 & 22 & 25 & 28 \end{bmatrix} \quad (14)$$

Ainda, as seguintes equações (15) e (16) representam exemplos de funções de geração de matriz do terceiro tipo de geração de matriz.

[fórmula 5]

$$Q_{4 \times 4}(x, y) = a * r + b(\sin(\frac{\pi}{16}r)) + c \quad (15)$$

$$Q_{8 \times 8}(x, y) = \frac{a}{2} * r + \frac{b}{2}(\sin(\frac{\pi}{32}r)) + c \quad (16)$$

- O item de distorção mostrado na Figura 6B é somado ao primeiro tipo de matriz. A amplitude de distorção (B) representa a magnitude da amplitude de uma função de seno. Quando b é um valor positivo, o efeito que uma linha reta fica empenada no lado inferior emerge. Por outro lado, quando b é um valor negativo, um efeito que a linha reta fica empenada no lado superior emerge. É necessário mudar a fase correspondente por um bloco de 4x4 pixels ou um bloco de 8x8 pixels. Várias distorções podem ser geradas pela mudança de fase. Quando a matriz de quantização no caso de QMP = (3, 32, 7, -6) é calculada utilizando as equações (4) e (25), a matriz de quantização da equação (17) seguinte pode ser gerada.

[fórmula 6]

$$Q_{4 \times 4}(x, y) = \begin{bmatrix} 10 & 14 & 19 & 23 \\ 14 & 19 & 23 & 27 \\ 19 & 23 & 27 & 31 \\ 23 & 27 & 31 & 35 \end{bmatrix} \quad (17)$$

- Apesar de uma função de seno ser utilizada nesta modalidade, uma função de co-seno ou outras funções podem ser utilizadas, e uma fase ou um período pode ser mudado. A amplitude de distorção (B) pode utilizar várias funções tais como uma função sigmóide, uma função Gaussiana, uma função logarítmica e uma função N-dimensional. Ainda, quando as variáveis do grau de mudança (A) que incluem a amplitude de distorção (B) e o item de correção (C) são um valor inteiro, uma tabela de tradução pode ser preparada com antecedência para evitar o processo de computação da alta carga de processamento tal como as funções de seno.

A função utilizada para o tipo de geração de matriz está sujeita a

um cálculo de número real. Conseqüentemente, quando um cálculo de função de seno é feito a cada codificação, o processo de cálculo aumenta. Ainda, um hardware para executar o cálculo de função de seno deve ser preparado. Assim, uma tabela de tradução de acordo com a precisão de um parâmetro a ser utilizado pode ser provida.

Como o cálculo de ponto flutuante aumenta em custo em comparação com o cálculo de inteiros, os parâmetros de geração de matriz de quantização são definidos por valores inteiros respectivamente, e um valor correspondente é extraído de uma tabela de tradução individual que corresponde a um tipo de geração de matriz.

Quando um cálculo de precisão de número real é possível, a distância pode ser computada pela equação (18) seguinte.

[fórmula 7]

$$r = \sqrt{x^2 + y^2} \quad (18)$$

Ainda, é possível mudar os valores nas direções vertical e lateral da matriz de quantização ponderando-a de acordo com a distância. Colocando uma grande importância sobre, por exemplo, uma direção vertical, uma função de distância como indicada pela equação (19) seguinte é utilizada.

[fórmula 8]

$$r = |2x + y| \quad (19)$$

Quando uma matriz de quantização no caso de QMP = (2, 1, 2, 8) é gerada pela equação acima, uma matriz de quantização expressa pela equação (20) seguinte é provida.

[fórmula 9]

$$Q_{4 \times 4}(x, y) = \begin{bmatrix} 8 & 10 & 14 & 20 \\ 9 & 12 & 17 & 24 \\ 10 & 14 & 20 & 28 \\ 12 & 17 & 24 & 33 \end{bmatrix} \quad (20)$$

As matrizes de quantização 204, geradas com o primeiro até o enésimo geradores de matriz 203 são emitidas do gerador de matriz de

quantização 109 seletivamente. O controlador de geração de matriz de quantização 210 controla a chave 202 para comutar o terminal de saída da chave 202 de acordo com cada tipo de geração de matriz decifrado com a unidade de decifração de parâmetro de geração 201. Ainda, o controlador de
 5 geração de matriz de quantização 210 verifica se a matriz de quantização que corresponde ao parâmetro de geração de matriz de quantização é apropriadamente gerado.

As configurações do aparelho de codificação de vídeo 100 e do gerador de matriz de quantização 109 de acordo com a modalidade estão
 10 aqui acima explicadas. Um exemplo para executar um método de codificação de vídeo com o aparelho de codificação de vídeo 100 e o gerador de matriz de quantização 109 será descrito referindo a um fluxograma da Figura 3.

Primeiramente, um sinal de imagem de um quadro é lido de uma
 15 memória externa (não-mostrada), e inserido no aparelho de codificação de vídeo 100 como o sinal de imagem de entrada 116 (etapa S001). O sinal de imagem de entrada 116 é dividido em macroblocos cada um composto de 16x16 pixels. Um parâmetro de geração de matriz de quantização 117 é ajustado para o aparelho de codificação de vídeo 100 (S002). Isto é, o contro-
 20 lador de codificação 110 envia as informações que indicam utilizar uma matriz de quantização para o quadro corrente para o gerador de parâmetro 108. Quando recebendo esta informação, o gerador de parâmetro 108 envia o parâmetro de geração de matriz de quantização para o gerador de matriz de quantização 109. O gerador de matriz de quantização 109 gera uma matriz
 25 de quantização de acordo com um tipo do parâmetro de geração de matriz de quantização.

Quando o sinal de imagem de entrada 116 é inserido no aparelho de codificação de vídeo 100, a codificação é iniciada em unidades de um bloco (etapa S003). Quando um macrobloco do sinal de imagem de entrada
 30 116 é inserido no preditor 101, a unidade de determinação de modo 102 inicializa um índice que indica um modo de codificação e um custo (etapa S004). Um sinal de imagem preditiva 118 de um modo de predição selecio-

nável em unidades de bloco é gerado pelo preditor 101 utilizando o sinal de imagem de entrada 116 (etapa S005). Uma diferença entre este sinal de imagem preditiva 118 e o sinal de imagem de entrada 116 é calculada por meio de que um sinal de erro preditivo 119 é gerado. Um custo é calculado da soma de valor absoluto SAD deste sinal de erro preditivo 119 e do número de bits codificados OH do modo de predição (etapa S006). De outro modo, uma decodificação local é feita para gerar um sinal decodificado local 113, e o custo é calculado do número de bits codificados D do sinal de erro que indica um valor diferencial entre o sinal decodificado local 113 e o sinal de imagem de entrada 116, e o número de bits codificados R de um sinal codificado obtido codificando temporalmente o sinal de imagem de entrada.

A unidade de determinação de modo 102 determina se o custo calculado é menor do que o menor custo min_cost (etapa S007). Quando este é menor (a determinação é SIM), o menor custo é atualizado pelo custo calculado, e um modo de codificação que corresponde ao custo calculado é mantido como um índice de best_mode (etapa S008). Ao mesmo tempo uma imagem preditiva é armazenada (etapa S009). Quando o custo calculado é maior do que o menor custo min_cost (a determinação é NÃO) o índice que indica um número de modo é incrementado, e é determinado se o índice após o incremento é o último modo (etapa S010).

Quando o índice é maior do que MAX que indica o número do último modo (a determinação é SIM), as informações de modo de codificação de best_mode e o sinal de erro preditivo 109 são enviados para o transformador 103 e para o quantizador 104 para serem transformadas e quantizadas (etapa S011). O coeficiente de transformada quantizado 112 é inserido no processador de codificação 111 e codificado por entropia juntamente com as informações preditivas com o processador de codificação 111 (etapa S012). Por outro lado, quando o índice é menor do que MAX indicando o número do último modo (a determinação é NÃO), o sinal de imagem preditiva 118 de um modo de codificação indicado pelo próximo índice é gerado (etapa S005).

Quando a codificação é feita em `best_mode`, o coeficiente de transformada quantizado 112 é inserido no dequantizador 105 e na transformada inversa 106 para ser dequantizado e transformado inverso (etapa S013), por meio de que o sinal de erro preditivo é decodificado. Este sinal de erro preditivo decodificado é adicionado ao sinal de imagem preditiva de `best_mode` provida da unidade de determinação de modo 102 para gerar um sinal decodificado local 113. Este sinal decodificado local 113 é armazenado na memória de imagem de referência 107 como uma imagem de referência (etapa S014).

Se a codificação de um quadro terminou é determinado (etapa S015). Quando o processo está completo (a determinação é SIM), um sinal de imagem de entrada do próximo quadro é lido, e então o processo retorna para a etapa S002 para codificação. Por outro lado, quando o processo de codificação de um quadro não está completo (a determinação é NÃO) o processo retorna para a etapa S003, e então o próximo bloco de pixels é inserido e o processo de codificação é continuado.

O acima é o resumo do aparelho de codificação de vídeo 100 e do método de codificação de vídeo na modalidade da presente invenção.

Na modalidade acima, o gerador de matriz de quantização 108 gera e utiliza uma matriz de quantização para codificar um quadro. No entanto, uma pluralidade de matrizes de quantização pode ser gerada para um quadro ajustando uma pluralidade de parâmetros de geração de matriz de quantização. Neste caso, como uma pluralidade de matrizes de quantização gerada em diferentes tipos de geração de matriz com o primeiro até o *enésimo* geradores de matriz 203 pode ser comutada em um quadro, uma quantização flexível torna-se possível. Concretamente, o primeiro gerador de matriz gera uma matriz de quantização que tem um peso uniforme, e o segundo gerador de matriz gera uma matriz de quantização que tem um grande valor em uma largura de banda alta. O controle de quantização é permitido em uma menor faixa pela mudança destas duas matrizes a cada bloco a ser codificado. Como o número de bits codificados transmitidos para gera a matriz de quantização é de diversos bits, a alta eficiência de codificação pode ser

mantida.

Na modalidade da presente invenção, está explicada uma técnica de geração de matriz de quantização de um tamanho de bloco de 4x4 pixels e um tamanho de bloco 8x8 pixels para a geração de uma matriz de quantização referente a um componente de luminância. No entanto, a geração da matriz de quantização é possível por um esquema similar para um componente de diferença de cor. Então de modo a evitar o aumento de suplemento para a multiplexação do parâmetro de geração de matriz de quantização do componente de diferença de cor com sintaxe, a mesma matriz de quantização que o componente de luminância pode ser utilizada, e a matriz de quantização com um deslocamento que corresponde a cada posição de frequência pode ser feita e utilizada.

Na presente modalidade da presente invenção, está explicado um método de geração de matriz de quantização que utiliza uma função trigonométrica (uma função de seno) no enésimo gerador de matriz 203. No entanto, a função a ser utilizada pode ser uma função sigmóide e uma função Gaussiana. É possível fazer uma matriz de quantização mais complicada de acordo com um tipo de função. Ainda, quando o tipo de geração de matriz (T) correspondente entre os parâmetros de geração de matriz de quantização QMP providos do controlador de geração de matriz de quantização 210 não puder ser utilizado no aparelho de codificação de vídeo, é possível fazer uma matriz de quantização substituindo o tipo de geração de matriz bem semelhante ao tipo de geração de matriz (T). Concretamente, o segundo tipo de geração de matriz é uma função que um grau de distorção que utiliza uma função de seno é adicionado ao primeiro tipo de geração de matriz, e similar à tendência da matriz de quantização gerada. Portanto, quando o terceiro gerador de matriz não puder ser utilizado no aparelho de codificação quando $T = 3$ é inserido, o primeiro gerador de matriz é utilizado.

Na modalidade da presente invenção, quatro parâmetros do tipo de geração de matriz (T), o grau de mudança (A) de matriz de quantização, o grau de distorção (B) e o item de correção (C) são utilizados. No entanto, parâmetros além destes parâmetros podem ser utilizados, e o número de

parâmetros decidido pelo tipo de geração de matriz (T) pode ser utilizado. Ainda, uma tabela de tradução de parâmetros decididos pelo tipo de geração de matriz (T) com antecedência pode ser provida. O número de bits codificados para codificar os parâmetros de geração de matriz de quantização diminui conforme o número de parâmetros de geração de matriz de quantização a serem transmitidos diminui e a precisão diminui. No entanto, como ao mesmo tempo o grau de liberdade da matriz de quantização diminui, o número de parâmetros de geração de matriz de quantização e a sua precisão precisam somente ser selecionados em consideração do equilíbrio entre um perfil e uma escala de hardware a ser aplicada.

Na modalidade da presente invenção, um quadro a ser processado é dividido em blocos retangulares de 16x16 pixels de tamanho, e então os blocos são codificados de uma esquerda superior de uma tela para uma sua direita inferior, sequencialmente. No entanto, a seqüência de processamento pode ser outra seqüência. Por exemplo, os blocos podem ser codificados da direita inferior para a esquerda superior, ou em uma forma de rolagem do meio da tela. Ainda, os blocos podem ser codificados da direita superior para a esquerda inferior, ou da parte periférica para a sua parte central.

Na modalidade da presente invenção, o quadro está dividido em macroblocos de um tamanho de bloco de 16x16 pixels, e um bloco de 8x8 pixels ou um bloco de 4x4 pixels é utilizado como uma unidade de processamento para uma predição intraquadro. No entanto, o bloco a ser processado não precisa ser feito em uma forma de bloco uniforme, e pode ser feito em um tamanho de bloco de pixels tal como um tamanho de bloco de 16x8 pixels, um tamanho de bloco 8x16 pixels, um tamanho de bloco de 8x4 pixels, um tamanho de bloco de 4x8 pixels. Por exemplo, um bloco de 8x4 pixels e um bloco de 2x2 pixels estão disponíveis sob uma estrutura similar. Ainda, não é necessário tomar um tamanho de bloco uniforme em um macrobloco, e diferentes tamanhos de bloco podem ser selecionados. Por exemplo, um bloco de 8x8 pixels e um bloco de 4x4 pixels podem coexistir no macrobloco. Neste caso, apesar do número de bits codificados para codificar

os blocos divididos aumenta com o aumento do número de blocos divididos, uma predição de precisão mais alta é possível, resultando na redução de um erro de predição. Conseqüentemente, um tamanho de bloco precisa somente ser selecionado em consideração de equilíbrio entre o número de bits codificados de coeficientes de transformada e a imagem decodificada local.

Na modalidade da presente invenção, o transformador 103, o quantizador 104, o dequantizador 105 e o transformador inverso 106 estão providos. No entanto, o sinal de erro preditivo não precisa sempre ser sujeito à transformação, quantização, transformação inversa e dequantização, e o sinal de erro preditivo pode ser codificado com o processador de codificação 109 como está, e a quantização e a quantização inversa podem ser omitidas. Similarmente, a transformação e a transformação inversa não precisa ser feita.

SEGUNDA MODALIDADE: CODIFICAÇÃO

Uma codificação de múltiplos percursos referente à segunda modalidade está explicada referindo a um fluxograma da Figura 12. Nesta modalidade, a descrição detalhada do fluxo de codificação que tem a mesma função que a primeira modalidade da Figura 3, isto é, as etapas S002 - S015, é omitida. Quando a matriz de quantização ótima é ajustada a cada imagem, a matriz de quantização deve ser otimizada. Por esta razão, a codificação de múltiplos percursos é eficiente. De acordo com esta codificação de múltiplos percursos, o parâmetro de geração de matriz de quantização pode ser efetivamente selecionado.

Nesta modalidade, para a codificação de múltiplos percursos, as etapas S101 - S108 são adicionadas antes da etapa S002 da primeira modalidade como mostrado na Figura 12. Em outras palavras, primeiramente, o sinal de imagem de entrada 116 de um quadro é inserido no aparelho de codificação de vídeo 100 (etapa S101), e codificado sendo dividido em macroblocos de 16x16 pixels de tamanho. Então, o controlador de codificação 110 inicializa um índice do parâmetro de geração de matriz de quantização utilizado para o quadro corrente para 0, e inicializa min_costQ que representa o custo mínimo, também (etapa S102). Então, o controlador de geração

de matriz de quantização 210 seleciona um índice do parâmetro de geração de matriz de quantização mostrado em PQM_idx de um conjunto de parâmetros de geração de matriz de quantização, e envia-o para o gerador de matriz de quantização 109. O gerador de matriz de quantização 109 gera a matriz de quantização de acordo com um esquema do parâmetro de geração de matriz de quantização inserido (etapa S103). Um quadro é codificado utilizando a matriz de quantização gerada desta vez (etapa S104). Um custo acumula a cada macrobloco para calcular um custo de codificação de um quadro (etapa S105).

10 É determinado se o custo calculado é menor do que o menor custo min_costQ (etapa 106). Quando o custo calculado é menor do o menor custo (a determinação é SIM), o menor custo é atualizado pelo custo do cálculo. Neste momento, o índice do parâmetro de geração de matriz de quantização é mantido como um índice Best_PQM_idx (etapa S107). Quando o
15 custo calculado é maior do que o menor custo min_costQ (a determinação é NÃO), o PQM_index é incrementado e é determinado se o PQM_idx incrementado é o último (etapa S108). Se a determinação NÃO, o índice do parâmetro de geração de matriz de quantização é atualizado, e uma codificação adicional é continuada. Por outro lado, se a determinação for SIM, o
20 Best_PQM_idx é inserido no gerador de matriz de quantização 109 novamente, e o fluxo de codificação principal, isto é, as etapas S002 - S015 da Figura 3 são executadas. Quando os dados de código codificados em Best_PQM_idx no momento do processo de múltiplos percursos é mantido o fluxo de codificação principal não precisa ser executado, e assim é possível
25 terminar a codificação do quadro atualizando os dados de código.

Na segunda modalidade, quando a codificação é feita em múltiplos percursos, não é necessário sempre codificar o quadro inteiro. Um parâmetro de geração de matriz de quantização disponível pode ser determinado por distribuição de coeficiente de transformada obtida em unidade de
30 bloco. Por exemplo, quando os coeficientes de transformada gerados a uma baixa taxa são quase 0, porque a propriedade dos dados de código não muda mesmo se a matriz de quantização não for utilizada, o processo pode ser

grandemente reduzido.

Será explicado um método de codificação de um parâmetro de geração de matriz de quantização. Como mostrado na Figura 7, a sintaxe é compreendida de três partes principalmente. Uma sintaxe de alto nível (401) é empacotada com informações de sintaxe de camadas de nível mais alto do que um nível de fatia. Uma sintaxe de nível de fatia (402) descreve as informações necessárias a cada fatia. Uma sintaxe de nível de macrobloco (403) descreve um valor de mudança de parâmetro quantização ou informações de modo necessário para cada macrobloco. Estas sintaxes são configuradas por sintaxes adicionalmente detalhadas. Em outras palavras, a sintaxe de alto nível (401) está compreendida de seqüências tais como a sintaxe de ajuste de parâmetro de seqüência (404) e a sintaxe de ajuste de parâmetro de imagem (405), e uma sintaxe de nível de imagem. Uma sintaxe de nível de fatia (402) está compreendida de uma sintaxe de cabeçalho de fatia (406), uma sintaxe de dados de fatia (407), etc. Ainda, a sintaxe de nível de macrobloco (403) está compreendida de uma sintaxe de cabeçalho de macrobloco (408), uma sintaxe de dados de macrobloco (409), etc.

As sintaxes acima são componentes que são absolutamente essenciais para a decodificação. Quando as informações de sintaxe estão faltando, torna-se impossível reconstruir corretamente os dados no momento da decodificação. Por outro lado, existe uma sintaxe suplementar para multiplexar as informações que não são sempre necessárias no momento da decodificação. Esta sintaxe descreve os dados estatístico de uma imagem, os parâmetros de câmera, etc., e está preparada como um papel para filtrar e ajustar os dados no momento da decodificação.

Nesta modalidade, as informações de sintaxe necessárias são a sintaxe de ajuste de parâmetro de seqüência (404) e a sintaxe de ajuste de parâmetro de imagem (405). Cada sintaxe está aqui após descrita.

ex_seq_scaling_matrix_flag mostrado na sintaxe de ajuste de parâmetro de seqüência da Figura 8 é um sinalizador que indica se a matriz de quantização é utilizada. Quando este sinalizador é VERDADEIRO, a matriz de quantização pode ser mudada em unidades de seqüência. Por outro

lado, quando o sinalizador é FALSO, a matriz de quantização não pode ser utilizada na seqüência. Quando `ex_seq_scaling_matrix_flag` é VERDADEIRO, `ex_matrix_type`, `ex_matrix_A`, `ex_matrix_B` e `ex_matrix_C` adicionais são enviados. Estes correspondem ao tipo de geração de matriz (T), ao grau de mudança (A) de matriz de quantização, ao grau de distorção (B) e ao item de correção (C), respectivamente.

`ex_pic_scaling_matrix_flag` mostrado na sintaxe de ajuste de parâmetro de imagem da Figura 9 é um sinalizador que indica se a matriz de quantização é mudada a cada quadro. Quando este sinalizador é VERDADEIRO, a matriz de quantização pode ser mudada em unidades de imagem. Por outro lado, quando o sinalizador é FALSO, a matriz de quantização não pode ser mudada a cada imagem. Quando `ex_pic_scaling_matrix_flag` é VERDADEIRO, `ex_matrix_type`, `ex_matrix_A`, `ex_matrix_B` e `ex_matrix_C` adicionais são transmitidos. Estes correspondem ao tipo de geração de matriz (T), ao grau de mudança (A) para a matriz de quantização, ao grau de distorção (B) e ao item de correção (C), respectivamente.

Um exemplo de que uma pluralidade de parâmetros de geração de matriz de quantização é enviada está mostrado na Figura 10 como outro exemplo da sintaxe de ajuste de parâmetro de imagem. `ex_pic_scaling_matrix_flag` mostrado na sintaxe de ajuste de parâmetro de imagem é um sinalizador que indica se a matriz de quantização é mudada a cada imagem. Quando o sinalizador é VERDADEIRO, a matriz de quantização pode ser mudada em unidades de imagem. Por outro lado, quando o sinalizador é FALSO, a matriz de quantização não pode ser mudada a cada imagem. Quando `ex_pic_scaling_matrix_flag` é VERDADEIRO, ainda, `ex_num_of_matrix_type` é enviado. Este valor representa o número de conjuntos de parâmetros de geração de matriz de quantização. Uma pluralidade de matrizes de quantização pode ser enviada pela combinação de conjuntos. `ex_matrix_type`, `ex_matrix_A`, `ex_matrix_B` e `ex_matrix_C`, os quais são enviados sucessivamente, são enviados por um valor de `ex_num_of_matrix_type`. Como um resultado, uma pluralidade de matrizes de quantização pode ser provida em uma imagem. Ainda, quando a matriz de

quantização deve ser mudada em unidades de bloco, os bits podem ser transmitidos a cada bloco pelo número de matrizes de quantização correspondentes, e trocados. Por exemplo, se `ex_num_of_matrix_type` for 2, uma sintaxe de 1 bit é adicionada à sintaxe de cabeçalho de macrobloco. A matriz de quantização é mudada conforme este valor é VERDADEIRO ou FALSO.

Na presente modalidade, quando uma pluralidade de parâmetros de geração de matriz de quantização são mantidos em um quadro como acima descrito, estes podem ser multiplexados em uma sintaxe suplementar. Um exemplo que uma pluralidade de parâmetros de geração de matriz de quantização são enviados utilizando a sintaxe suplementar está mostrado na Figura 11. `ex_sei_scaling_matrix_flag` mostrado na sintaxe suplementar é um sinalizador que indica se uma pluralidade de matrizes de quantização está mudada. Quando este sinalizador é VERDADEIRO, as matrizes de quantização podem ser mudadas. Por outro lado, quando o sinalizador é FALSO, as matrizes de quantização não podem ser mudadas. Quando `ex_sei_scaling_matrix_flag` é VERDADEIRO, ainda, `ex_num_of_matrix_type` é enviado. Este valor indica o número de conjuntos de parâmetros de geração de matriz de quantização. Uma pluralidade de matrizes de quantização pode ser enviada pela combinação de conjuntos. Quanto ao `ex_matrix_type`, `ex_matrix_A`, `ex_matrix_B` e `ex_matrix_C` os quais são enviados sucessivamente, um único valor de `ex_num_of_matrix_type` é enviado. Como um resultado, uma pluralidade de matrizes de quantização pode ser provida na imagem.

Nesta modalidade, a matriz de quantização pode ser retransmitida pela sintaxe de cabeçalho de fatia na sintaxe de nível de fatia mostrada na Figura 7. Um exemplo de tal caso será explicado utilizando a Figura 13. A Figura 13 mostra a estrutura de sintaxe na sintaxe de cabeçalho de fatia. O `slice_ex_scaling_matrix_flag` mostrado na sintaxe de cabeçalho de fatia da Figura 13 é um sinalizador que indica se uma matriz de quantização pode ser utilizada na fatia. Quando o sinalizador é VERDADEIRO, a matriz de quantização pode ser mudada na fatia. Quando o sinalizador é FALSO, a matriz de quantização não pode ser mudada na fatia. O sli-

ce_ex_matrix_type é transmitido quando o slice_ex_scaling_matrix_flag é VERDADEIRO. Esta sintaxe corresponde a um tipo de geração de matriz (T). Sucessivamente, slice_ex_matrix_A, slice_ex_matrix_B e slice_ex_matrix_C são transmitidos. Estes correspondem a um grau de mudança (A) de uma matriz de quantização (C) um seu grau de distorção (B) e um seu item de correção respectivamente. NumOfMatrix na Figura 13 representa o número de matrizes de quantização disponíveis na fatia. Quando a matriz de quantização é mudada em uma região menor no nível de fatia, esta é mudada em componente de luminância e componente de cor, esta é mudada em tamanho de bloco de quantização, esta é mudada a cada modo de codificação, etc., o número de matrizes de quantização disponíveis pode ser transmitido como um parâmetro de modelagem da matriz de quantização que corresponde ao número. Para propósitos de exemplo, quando existem dois tipos de blocos de quantização de um tamanho de bloco de 4x4 pixels e um tamanho de bloco de 8x8 pixels na fatia, e diferentes matrizes de quantização podem ser utilizadas para os blocos de quantização, o valor de NumOfMatrix é ajustado para 2.

Nesta modalidade da presente invenção, a matriz de quantização pode ser mudada no nível de fatia utilizando a sintaxe de cabeçalho de fatia mostrada na Figura 14. Na Figura 14, três parâmetros de modelagem a serem transmitidos estão preparados comparado com a Figura 13. Quando uma matriz de quantização é gerada com a utilização, por exemplo, da equação (5), o parâmetro não precisa ser transmitido porque o grau de distorção (B) é sempre ajustado para 0. Portanto, o codificador e o decodificador podem gerar a matriz de quantização idêntica mantendo um valor inicial de 0 como um parâmetro interno.

Nesta modalidade, o parâmetro pode ser transmitido utilizando a sintaxe de cabeçalho de fatia expressa na Figura 15. Na Figura 15, PrevSliceExMatrixType, PrevSliceExMatrix_A, PrevSliceExMatrix_B (adicionalmente, PrevSliceExMatrix_C) são adicionados à Figura 13. Explicando mais concretamente, slice_ex_scaling_matrix_flag é um sinalizador que indica se uma matriz de quantização é utilizada ou não na fatia, e quando este sinalizador é

VERDADEIRO, um parâmetro de modelagem é transmitido para um decodificador como mostrado nas Figuras 13 e 14. Por outro lado, quando o sinalizador é FALSO, PrevSliceExMatrixType, PrevSliceExMatrix_A, PrevSliceExMatrix_B (adicionalmente, PrevSliceExMatrix_C) são ajustados. Estes significados são interpretados como segue.

PrevSliceExMatrixType indica um tipo de geração (T) utilizado no momento quando os dados são codificados no mesmo tipo de fatia que aquela de uma antes da fatia corrente em ordem de codificação. Esta variável é atualizada imediatamente antes que esta codificação da fatia seja terminada. O valor inicial é ajustado para 0.

PrevSliceExMatrix_A indica um grau de mudança (A) utilizado no momento quando os dados são codificados no mesmo tipo de fatia que aquela de uma antes da fatia corrente em ordem de codificação. Esta variável é atualizada imediatamente antes que esta codificação da fatia seja terminada. O valor inicial é ajustado para 0.

PrevSliceExMatrix_B indica um grau de distorção (B) utilizado no momento quando os dados são codificados no mesmo tipo de fatia que aquela de uma antes da fatia corrente em ordem de codificação. Esta variável é atualizada imediatamente antes que esta codificação da fatia seja terminada. O valor inicial é ajustado para 0.

PrevSliceExMatrix_C indica um item de correção (C) utilizado no momento quando os dados são codificados no mesmo tipo de fatia que aquela de uma antes da fatia corrente em ordem de codificação. Esta variável é atualizada imediatamente antes que esta codificação da fatia seja terminada. O valor inicial é ajustado para 16.

CurrSliceType indica um tipo de fatia da fatia codificada corrente, e um índice correspondente é atribuído a cada um de, por exemplo, I-Slice, P-Slice e B-slice. Um exemplo de CurrSliceType está mostrado na Figura 16. Um valor é atribuído a cada um dos respectivos tipos de fatia. 0 é atribuído a I-Slice utilizando somente uma predição intra-imagem, por exemplo. Ainda, 1 é atribuído a P-Slice capaz de utilizar uma única predição direcional do quadro codificado anteriormente codificado em ordem de tempo e predição de

intra-imagem. Por outro lado, 2 é atribuído a B-Slice capaz de utilizar uma predição bidirecional, uma predição unidirecional e uma predição intra-imagem.

5 Deste modo, o parâmetro de modelagem da matriz de quantização codificada no mesmo tipo de fatia que aquele da fatia imediatamente antes da fatia corrente é acessado e restaurado. Como um resultado, é possível reduzir o número de bits codificados necessários para a transmissão do parâmetro de modelagem.

10 Nesta modalidade da presente invenção, a Figura 17 pode ser utilizada. A Figura 17 mostra uma estrutura que NumOfMatrix é removido da Figura 5. Quando somente uma matriz de quantização está disponível para a fatia codificada, esta sintaxe simplificada mais do que a Figura 15 é utilizada. Esta sintaxe mostra aproximadamente a mesma operação que o caso em que NumOfMatrix1 é 1 na Figura 15.

15 Quando uma pluralidade de matrizes de quantização pode ser mantida na mesma imagem com o decodificador, o parâmetro de geração de matriz de quantização é lido da sintaxe suplementar para gerar uma matriz de quantização correspondente. Por outro lado, quando uma pluralidade de matrizes de quantização não puder ser mantida na mesma imagem, a matriz de quantização gerada pelo parâmetro de geração de matriz de quantização
20 descrito na sintaxe de ajuste de parâmetro de imagem é utilizado sem decodificação da sintaxe suplementar.

 Na modalidade acima discutida, a matriz de quantização é gerada de acordo com um tipo de geração de matriz correspondente. Quando o
25 parâmetro de geração da matriz de quantização é codificado, o número de bits codificados utilizados para enviar a matriz de quantização pode ser reduzido. Ainda, torna-se possível selecionar adaptavelmente as matrizes de quantização na imagem. A codificação capaz de lidar com vários usos tais como a quantização feita em consideração de uma imagem de subjetividade
30 e uma codificação feita em consideração da eficiência de codificação torna-se possível. Em outras palavras, uma codificação preferida de acordo com o conteúdo de bloco de pixels pode ser executada.

Como acima mencionado, quando a codificação é executada em um modo selecionado, um sinal de imagem decodificada precisa somente ser gerado somente para o modo selecionado. Este não precisa ser sempre executado em um laço para a determinação de um modo de predição.

- 5 O aparelho de decodificação de vídeo de acordo com o aparelho de codificação de vídeo será aqui após explicado.

TERCEIRA MODALIDADE: DECODIFICAÇÃO

- De acordo com um aparelho de decodificação de vídeo 300 referente à presente modalidade mostrado na Figura 18, um armazenamento temporário de entrada 309 uma vez salva os dados de código enviados do
- 10 aparelho de codificação de vídeo 100 da Figura 1 através de um meio de transmissão ou um meio de gravação. Os dados de código salvos são lidos do armazenamento temporário de entrada 309, e inseridos em um processador de decodificação 301 sendo separados com base em sintaxe em cada
- 15 quadro. O processador de decodificação 301 decodifica uma cadeia de códigos de cada sintaxe dos dados de código para cada uma de uma sintaxe de alto nível, uma sintaxe de nível de fatia e uma sintaxe de nível de macrobloco de acordo com a estrutura de sintaxe mostrada na Figura 7. Devido a esta decodificação, o coeficiente de transformada quantizado, o parâmetro de
- 20 geração de matriz de quantização, o parâmetro de quantização, as informações de modo de predição, as informações de comutação de predição, etc., são reconstruídos.

- O processador de decodificação 301 produz, da sintaxe decodificada, um sinalizador que indica se uma matriz de quantização é utilizada
- 25 para um quadro correspondente, e insere-o em uma unidade de ajuste de parâmetro de geração 306. Quando este sinalizador é VERDADEIRO, um parâmetro de geração de matriz de quantização 311 é inserido na unidade de ajuste de parâmetro de geração 306 do processador de decodificação 301. A unidade de ajuste de parâmetro de geração 306 tem uma função de
- 30 atualização do parâmetro de geração de matriz de quantização 311, e insere um conjunto dos parâmetros de geração de matriz de quantização 311 em um gerador de matriz de quantização 307 com base na sintaxe decodificada

pelo processador de decodificação 301. O gerador de matriz de quantização 307 gera uma matriz de quantização 318 que corresponde ao parâmetro de geração de matriz de quantização 311 de entrada, e emite-o para um dequantizador 302.

5 O coeficiente de transformada quantizado emitido do processador de decodificação 301 é inserido no dequantizador 302, e dequantizado por meio disto com base nas informações decodificadas utilizando a matriz de quantização 318, o parâmetro de quantização, etc. O coeficiente de transformada dequantizado é inserido em um transformador inverso 303. O transformador inverso 303 sujeita o coeficiente de transformada dequantizado a
10 uma transformada inversa (por exemplo, uma transformada de co-seno discreta inversa) para gerar um sinal de erro 313. A transformação ortogonal inversa é aqui utilizada. No entanto, quando o codificador executa uma transformação de ondulação ou uma análise de componente independente,
15 o transformador inverso 303 pode executar uma transformação de ondulação inversa ou uma análise de componente de independência inversa. O coeficiente sujeito à transformação inversa com o transformador inverso 303 é enviado para um somador 308 como um sinal de erro 313. O somador 308 soma o sinal preditivo 315 emitido do preditor 305 e o sinal de erro 313, e
20 insere um sinal de adição em uma memória de referência 304 como um sinal decodificado 314. A imagem decodificada 314 é enviada do decodificador de vídeo 300 para o exterior e armazenada no armazenamento temporário de saída (não-mostrado). A imagem decodificada armazenada no armazenamento temporário de saída é lida e o tempo gerenciado pelo controlador de
25 decodificação 310.

 Por outro lado, as informações de predição 316 e as informações de modo as quais são decodificadas com o processador de decodificação 301 são inseridas no preditor 305. O sinal de referência 317 já codificado é suprido da memória de referência 304 para o preditor 305. O preditor
30 305 gera o sinal preditivo 315 com base nas informações de modo inseridas, etc. e supre-o para o somador 308.

 O controlador de decodificação 310 controla o armazenamento

temporário de entrada 307, o tempo de saída, o tempo de decodificação, etc.

O aparelho de decodificação de vídeo 300 da terceira modalidade está configurado como acima descrito, e o método de decodificação de vídeo executado com o aparelho de decodificação de vídeo é explicado referindo ao fluxograma de Figura 19.

Os dados de código de um quadro são lidos do armazenamento temporário de entrada 309 (etapa S201), e decodificados de acordo com uma estrutura de sintaxe (etapa S202). É determinado por um sinalizador se a matriz de quantização é utilizada para o quadro de leitura com base na sintaxe decodificada (etapa S203). Quando esta determinação for SIM, um parâmetro de geração de matriz de quantização é ajustado para o gerador de matriz de quantização 307 (etapa S204). O gerador de matriz de quantização 307 gera a matriz de quantização que corresponde ao parâmetro de geração (etapa S205). Para esta geração de matriz de quantização, um gerador de matriz de quantização 307 que tem a mesma configuração que o gerador de matriz de quantização 109 mostrado na Figura 2 o qual é utilizado para o aparelho de codificação de vídeo é empregado, e executa o mesmo processo que o aparelho de codificação de vídeo para gerar uma matriz de quantização. Um gerador de parâmetro de geração 306 para suprir um parâmetro de geração para o gerador de matriz de quantização 109 tem a mesma configuração que o gerador de parâmetro de geração 108 para o aparelho de codificação.

Em outras palavras, no gerador de parâmetro de geração 306, a sintaxe é formada de três partes principalmente, isto é, uma sintaxe de alto nível (401), uma sintaxe de nível de fatia (402) e uma sintaxe de nível de macrobloco (403) como mostrado na Figura 7. Estas sintaxes estão compreendidas de sintaxes adicionalmente detalhadas como o aparelho de codificação.

As sintaxes acima mencionadas são componentes que são absolutamente imperativos no momento da decodificação. Estas informações de sintaxe faltarem, os dados não podem ser corretamente decodificados no momento da decodificação. Por outro lado, existe uma sintaxe suplementar

para multiplexar as informações que não são sempre necessárias no momento da decodificação.

As informações de sintaxe as quais são necessárias nesta modalidade contêm uma sintaxe de ajuste de parâmetro de sequência (404) e uma sintaxe de ajuste de parâmetro de imagem (405). As sintaxes estão compreendidas de uma sintaxe de ajuste de parâmetro de sequência e uma sintaxe de ajuste de parâmetro de imagem, respectivamente, como mostrado nas Figuras 8 e 9 como o aparelho de codificação de vídeo.

Como outro exemplo da sintaxe de ajuste de parâmetro de imagem pode ser utilizada a sintaxe de ajuste de parâmetro de imagem utilizada para enviar uma pluralidade de parâmetros de geração de matriz de quantização mostrada na Figura 10 como descrito no aparelho de codificação de vídeo. No entanto, se a matriz de quantização for mudada em unidades de bloco, os bits precisam somente ser transmitidos pelo número matrizes de quantização correspondentes para cada bloco, e trocados. Quando, por exemplo, `ex_num_of_matrix_type` é 2, uma sintaxe de 1 bit é adicionada na sintaxe de cabeçalho de macrobloco, e a matriz de quantização é mudada conforme este valor é VERDADEIRO ou FALSO.

Quando, nesta modalidade, uma pluralidade de parâmetros de geração de matriz de quantização é mantida em um quadro como acima descrito, os dados multiplexados com a sintaxe suplementar podem ser utilizados. Como descrito no aparelho de codificação de vídeo, é possível utilizar a pluralidade de parâmetros de geração de matriz de quantização utilizando as sintaxes suplementares mostradas na Figura 11.

Nesta modalidade da presente invenção, o recebimento de uma matriz de quantização pode ser feito por meio de uma sintaxe de cabeçalho de fatia na sintaxe de nível de fatia mostrado na Figura 7. Um exemplo de tal caso é explicado utilizando a Figura 13. A Figura 13 mostra uma estrutura de sintaxe em uma sintaxe de cabeçalho de fatia. O `slice_ex_scaling_matrix_flag` mostrado na sintaxe de cabeçalho de fatia da Figura 13 é um sinalizador que indica se uma matriz de quantização é utilizada na fatia. Quando o sinalizador é VERDADEIRO, a matriz de quantiza-

ção pode ser mudada na fatia. Por outro lado, quando o sinalizador é FALSO, é impossível mudar a matriz de quantização na fatia.

Quando o `slice_ex_scaling_matrix_flag` é VERDADEIRO, `slice_ex_matrix_type` é adicionalmente recebido. Esta sintaxe corresponde a um tipo de geração de matriz (T). Sucessivamente, `slice_ex_matrix_A`, `slice_ex_matrix_B` e `slice_ex_matrix_C` são recebidos. Estes correspondem a um grau de mudança (A), um grau de distorção (B) e um item de correção de uma matriz de quantização (C), respectivamente. NumOfMatrix na Figura 13 representa o número de matrizes de quantização disponíveis na fatia. Quando a matriz de quantização é mudada em uma região menor no nível de fatia, esta é mudada em componente de luminância e componente de cor, esta é mudada em tamanho de bloco de quantização, e esta é mudada a cada modo de codificação, etc., o número de matrizes de quantização disponíveis pode ser recebido como um parâmetro de modelagem da matriz de quantização que corresponde ao número. Para propósitos de exemplo, quando existem dois tipos de blocos de quantização de um tamanho de bloco de 4x4 pixels e um tamanho de bloco de 8x8 pixels na fatia, e diferentes matrizes de quantização podem ser utilizadas para os blocos de quantização, o valor de NumOfMatrix é ajustado para 2.

Nesta modalidade da presente invenção, a matriz de quantização pode ser mudada no nível de fatia utilizando a sintaxe de cabeçalho de fatia mostrada na Figura 14. Na Figura 14, três parâmetros de modelagem a serem transmitidos estão preparados comparado com a Figura 13. Quando uma matriz de quantização é gerada com a utilização, por exemplo, da equação (5), o parâmetro não precisa ser transmitido porque o grau de distorção (B) é sempre ajustado para 0. Portanto, o codificador e o decodificador podem gerar a matriz de quantização idêntica mantendo um valor inicial de 0 como um parâmetro interno.

Nesta modalidade da presente invenção, o parâmetro pode ser recebido utilizando a sintaxe de cabeçalho de fatia expressa na Figura 15. Na Figura 15, `PrevSliceExMatrixType`, `PrevSliceExMatrix_A` e `PrevSliceExMatrix_B` (adicionalmente, `PrevSliceExMatrix_C`) são adicionados à Figura

13. Explicando mais concretamente, `slice_ex_scaling_matrix_flag` é um sinalizador que indica se uma matriz de quantização é utilizada ou não na fatia. Quando este sinalizador é VERDADEIRO, um parâmetro de modelagem é recebido como mostrado nas Figuras 13 e 14. Por outro lado, quando o sinalizador é FALSO, `PrevSliceExMatrixType`, `PrevSliceExMatrix_A` e `PrevSliceExMatrix_B` (adicionalmente, `PrevSliceExMatrix_C`) são ajustados. Estes significados são interpretados como segue.

`PrevSliceExMatrixType` indica um tipo de geração (T) utilizado no momento quando os dados são decodificados no mesmo tipo de fatia que aquela de uma antes da fatia corrente em ordem de decodificação. Esta variável é atualizada imediatamente antes que esta decodificação da fatia seja terminada. O valor inicial é ajustado para 0.

`PrevSliceExMatrix_A` indica um grau de mudança (A) utilizado no momento quando a fatia é decodificada no mesmo tipo de fatia que aquela de uma antes da fatia corrente em ordem de decodificação. Esta variável é atualizada imediatamente antes que esta decodificação da fatia seja terminada. O valor inicial é ajustado para 0.

`PrevSliceExMatrix_B` indica um grau de distorção (B) utilizado no momento quando a fatia é decodificada no mesmo tipo de fatia que aquela de uma antes da fatia corrente em ordem de decodificação. Esta variável é atualizada imediatamente antes que esta decodificação da fatia seja terminada. O valor inicial é ajustado para 0.

`PrevSliceExMatrix_C` indica um item de correção (C) utilizado no momento quando os dados são decodificados no mesmo tipo de fatia que aquela de uma antes da fatia corrente em ordem de decodificação. Esta variável é atualizada imediatamente antes que esta decodificação da fatia seja terminada. O valor inicial é ajustado para 16.

`CurrSliceType` indica um tipo de fatia da fatia corrente. Respetivos índices são atribuídos a, por exemplo, I-Slice, P-Slice e B-slice respectivamente. Um exemplo de `CurrSliceType` está mostrado na Figura 16. Respetivos valores são atribuídos a respetivos tipos de fatia. 0 é atribuído a I-Slice utilizando somente uma predição intra-imagem, por exemplo. Ainda, 1

é atribuído a P-Slice capaz de utilizar uma única predição direcional do quadro codificado anteriormente codificado em ordem de tempo e predição de intra-imagem. Enquanto isso, 2 é atribuído a B-Slice capaz de utilizar uma predição bidirecional, uma predição unidirecional e uma predição intra-imagem.

Deste modo, o parâmetro de modelagem da matriz de quantização codificada no mesmo tipo de fatia que aquele da fatia imediatamente antes da fatia corrente é acessado e ajustado novamente. Como um resultado, é possível reduzir o número de bits codificados necessários para receber o parâmetro de modelagem.

Nesta modalidade da presente invenção, a Figura 17 pode ser utilizada. A Figura 17 mostra uma estrutura que NumOfMatrix é removido da Figura 5. Quando somente uma matriz de quantização está disponível para a fatia, esta sintaxe simplificada mais do que a Figura 16 é utilizada. Esta sintaxe mostra aproximadamente a mesma operação que o caso em que NumOfMatrix1 é 1 na Figura 15.

Quando a matriz de quantização é gerada como acima descrito, o coeficiente de transformada decodificado 312 é dequantizado pela matriz de quantização (etapa S206), e sujeito a uma transformação inversa com o transformador inverso 303 (etapa S207). Como um resultado, o sinal de erro é reproduzido. Então uma imagem preditiva é gerada pelo preditor 305 com base nas informações de predição 316 (etapa S209). Esta imagem preditiva e o sinal de erro são somados para reproduzir dados de imagem decodificados (etapa S209). Este sinal de imagem de decodificação é armazenado na memória de referência 304, e emitido para um dispositivo externo.

Nesta modalidade como acima discutido, uma matriz de quantização é gerada com base nos dados de código de entrada de acordo com o tipo de geração de matriz correspondente e utilizada em dequantização, por meio de que o número de bits codificados da matriz de quantização pode ser reduzido.

Uma função de cada parte acima descrita pode ser executada por um programa armazenado em um computador.

Nas modalidades acima, a codificação de vídeo é explicada. No entanto, a presente invenção pode ser aplicada para a gravação de imagens estáticas.

De acordo com a presente invenção, uma pluralidade de matrizes de quantização são geradas utilizando um ou mais parâmetros tal como um índice de geração de função para gerar uma matriz de quantização, um grau de mudança que indica um grau de mudança de uma matriz de quantização, um grau de distorção e um item de correção. A quantização e a de-quantização são executadas utilizando a matriz de quantização. O ajuste ótimo do parâmetro de geração de matriz de quantização é codificado e transmitido. Como um resultado, a presente invenção pode executar uma eficiência de codificação mais alta do que o método de transmissão de matriz de quantização convencional.

De acordo com a presente invenção, está provido um método e aparelho de codificação/decodificação de vídeo que tornam possível aperfeiçoar a eficiência de codificação na baixa taxa de bits que pode ser feita uma realidade possível.

APLICABILIDADE INDUSTRIAL

A invenção pode ser aplicada à codificação e decodificação de um filme, uma imagem estática, um áudio, etc. sobre cada campo tais como os dispositivos de vídeo, de áudio, equipamentos móveis, difusão, terminais de informações, ou redes.

REIVINDICAÇÕES

1. Método de codificação de vídeo para quantizar um coeficiente de transformada que utiliza uma matriz de quantização, o método de codificação de vídeo compreendendo:

5 uma etapa para selecionar um tipo de geração de uma matriz de quantização;

 uma etapa para obter uma pluralidade de funções de geração pelo ajuste de um parâmetro para uma pluralidade de funções preparadas com antecedência em correspondência com o tipo de geração;

10 uma etapa de geração da matriz de quantização para gerar uma pluralidade de matrizes de quantização utilizando uma pluralidade de funções de geração;

 uma etapa de quantização para produzir um coeficiente de transformada de quantização pela quantização de um coeficiente de transformada relativo a um sinal de imagem de entrada utilizando a pluralidade de matrizes de quantização; e

15 um a etapa de codificação para produzir um sinal codificado pela multiplexação e codificação do coeficiente de transformada quantizado, informações que indicam o tipo de geração e as informações do parâmetro.

20 2. Método de codificação de vídeo, de acordo com a reivindicação 1, em que o parâmetro inclui pelo menos um de um grau de mudança que representa um grau de mudança da matriz de quantização, um grau de distorção e um item de correção.

25 3. Método de codificação de vídeo, de acordo com a reivindicação 1, onde a etapa de obter a função de geração obtém a função de geração pelo ajuste do parâmetro para a função definida utilizando qualquer uma de uma função seno, uma função de co-seno, uma função N-dimensional, uma função sigmóide e uma função Gaussiana.

30 4. Método de codificação de vídeo, de acordo com a reivindicação 1, em que a etapa de geração de matriz de quantização inclui uma etapa de mudar a precisão de operação na produção da matriz de quantização em correspondência com uma precisão do parâmetro ajustado

para a função de geração.

5. Método de codificação de vídeo, de acordo com a reivindicação 1, em que a etapa de obter a função de geração seleciona uma tabela na qual os valores de cálculo das funções de geração plurais correspondem ao parâmetro, e a etapa de geração de matriz de quantização executa um cálculo para gerar a matriz de quantização plural referindo-se à tabela selecionada.

6. Método de codificação de vídeo, de acordo com a reivindicação 1, em que a etapa de codificação codifica as informações que indicam a matriz de quantização das matrizes de quantização plurais, a qual é utilizada na etapa de quantização como informações de cabeçalho de qualquer uma de uma seqüência codificada, uma imagem codificada ou uma fatia codificada.

7. Método de decodificação de vídeo para dequantizar um coeficiente de transformada utilizando uma matriz de quantização que corresponde a cada posição de freqüência do coeficiente de transformada, o método de decodificação de vídeo compreendendo:

uma etapa de decodificação para adquirir as informações de um parâmetro de uma função de geração para gerar um coeficiente de transformada quantizado, informações que indicam um tipo de geração de uma matriz de quantização e informações de um parâmetro de uma função de geração para gerar uma matriz de quantização;

uma etapa para obter as funções de geração plurais ajustando o parâmetro da função de geração para uma pluralidade de funções preparadas com antecedência em correspondência com o tipo de geração;

uma etapa de geração de matriz de quantização para gerar uma pluralidade de matrizes de quantização utilizando as funções de geração plurais;

uma etapa de dequantização para obter um coeficiente de transformada pela dequantização do coeficiente de transformada quantizado utilizando a pluralidade de matrizes de quantização geradas;

uma etapa de geração de imagem decodificada para gerar uma

imagem decodificada com base no coeficiente de transformada.

5 8. Método de decodificação de vídeo, de acordo com a reivindicação 7, em que o parâmetro inclui pelo menos um de um grau de mudança que representa um grau de mudança de uma matriz de quantização, um grau de distorção e um item de correção.

9. Método de decodificação de vídeo, de acordo com a reivindicação 7, em que a etapa de obter a função de geração obtém a função de geração pelo ajuste do parâmetro para a função definida utilizando qualquer uma de uma função de seno, uma função de co-seno,
10 uma função N-dimensional, uma função sigmóide e uma função Gaussiana.

10. Método de codificação de vídeo, de acordo com a reivindicação 7, em que a etapa de geração de matriz de quantização inclui uma etapa de mudar uma precisão de operação na produção da matriz de quantização em correspondência com uma precisão do parâmetro ajustado
15 para a função de geração.

11. Método de codificação de vídeo, de acordo com a reivindicação 7, em que a etapa de obter a função de geração seleciona uma tabela na qual os valores de cálculo das funções de geração plurais correspondem ao parâmetro, e a etapa de geração de matriz de quantização
20 executa um cálculo para gerar a matriz de quantização plural referindo-se à tabela selecionada

12. Método de codificação de vídeo, de acordo com a reivindicação 7, em que a etapa de geração da matriz de quantização inclui uma etapa de gerar uma matriz de quantização pela substituição de uma
25 função de geração disponível quando uma função de geração que corresponde a um índice de função de geração dentro de um parâmetro de geração de uma matriz de quantização decodificada não está disponível na decodificação.

13. Método de codificação de vídeo, de acordo com a reivindicação 7, em que a etapa de codificação codifica as informações que
30 indicam a matriz de quantização das matrizes de quantização plurais, a qual é utilizada na etapa de quantização como informações de cabeçalho de

qualquer uma de uma seqüência codificada, uma imagem codificada e uma fatia codificada.

14. Aparelho de codificação de vídeo para quantizar um coeficiente de transformada que utiliza uma matriz de quantização, o
5 aparelho de codificação de vídeo compreendendo:

uma unidade de seleção para selecionar um tipo de geração de uma matriz de quantização;

- uma unidade de aquisição de função de geração para adquirir uma pluralidade de funções de geração pelo ajuste de um parâmetro para
10 uma pluralidade de funções preparadas com antecedência em correspondência com o tipo de geração;

uma unidade de geração de matriz de quantização para gerar uma pluralidade de matrizes de quantização utilizando a pluralidade de funções de geração;

- 15 uma unidade de quantização para produzir um coeficiente de transformada de quantização pela quantização de um coeficiente de transformada relativo a um sinal de imagem de entrada utilizando a pluralidade de matrizes de quantização; e

- uma unidade de codificação para produzir um sinal codificado
20 pela multiplexação e codificação do coeficiente de transformada quantizado, informações que indicam o tipo de geração e informações do parâmetro.

15. Aparelho de decodificação de vídeo para dequantizar um coeficiente de transformada que utiliza uma matriz de quantização que corresponde a cada posição de freqüência do coeficiente de transformada, o
25 aparelho de decodificação de vídeo compreendendo:

- uma unidade de decodificação para adquirir as informações de um parâmetro de uma função de geração para gerar um coeficiente de transformada quantizado, informações que indicam um tipo de geração de uma matriz de quantização e informações de um parâmetro de uma função
30 de geração para gerar uma matriz de quantização;

uma unidade de aquisição de função de geração para adquirir funções de geração plurais pelo ajuste do parâmetro da função de geração

para uma pluralidade de funções preparadas com antecedência em correspondência com o tipo de geração;

5 uma unidade de geração de matriz de quantização para gerar uma pluralidade de matrizes de quantização utilizando as funções de geração plurais;

 uma unidade de dequantização para obter um coeficiente de transformada pela dequantização do coeficiente de transformada quantizado utilizando a pluralidade de matrizes de quantização geradas; e

10 uma unidade de geração de imagem decodificada para gerar uma imagem decodificada com base no coeficiente de transformada.

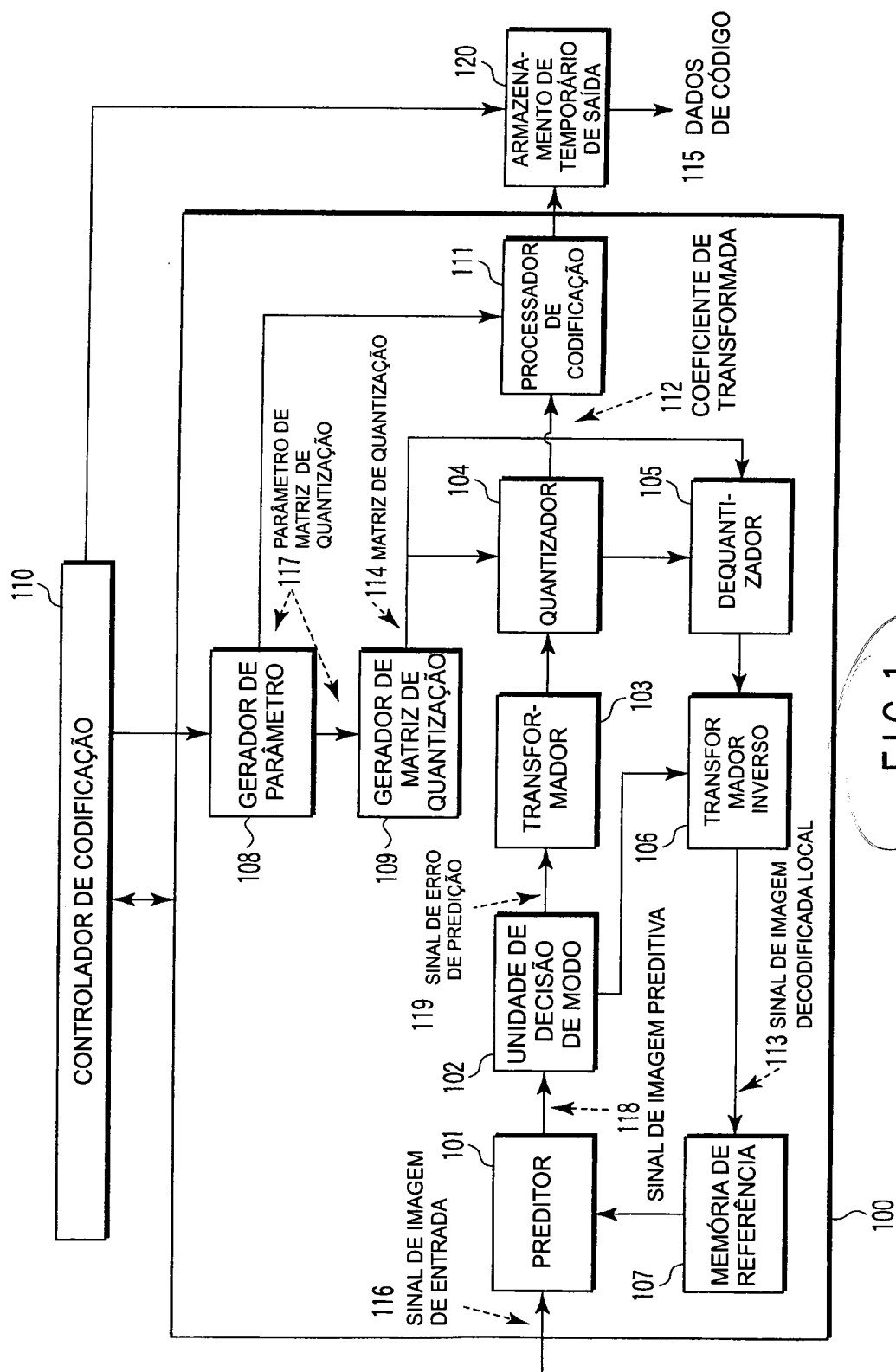


FIG. 1

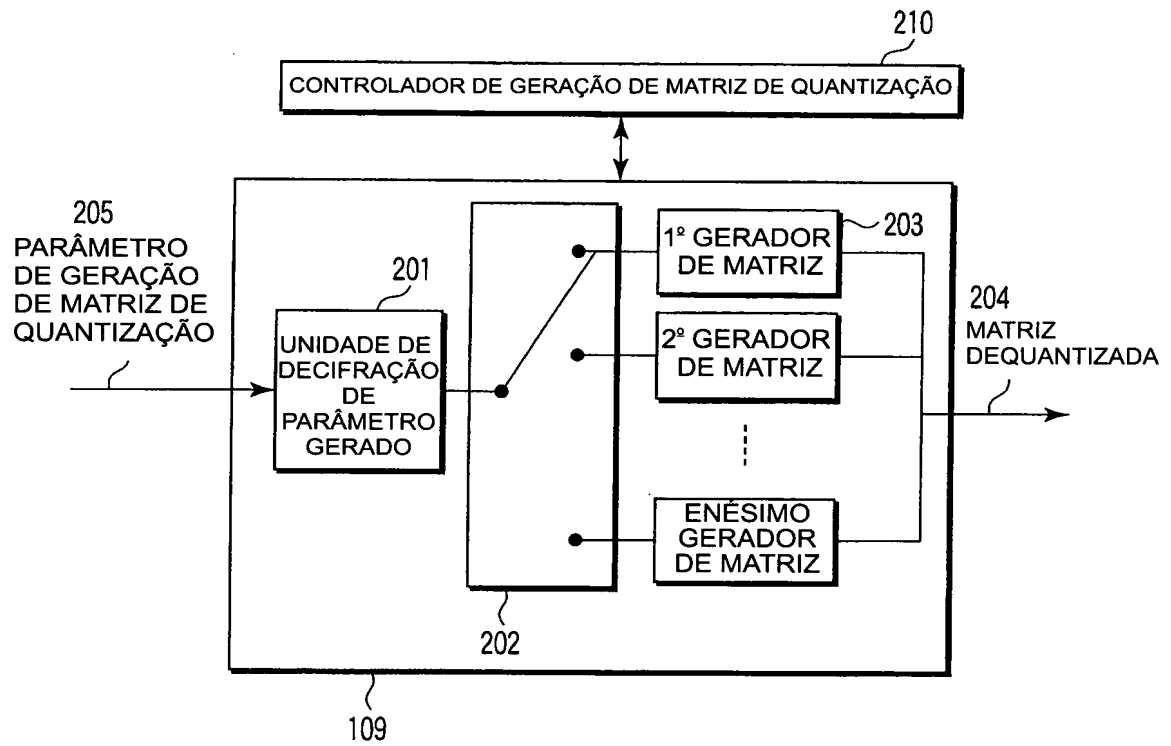


FIG. 2

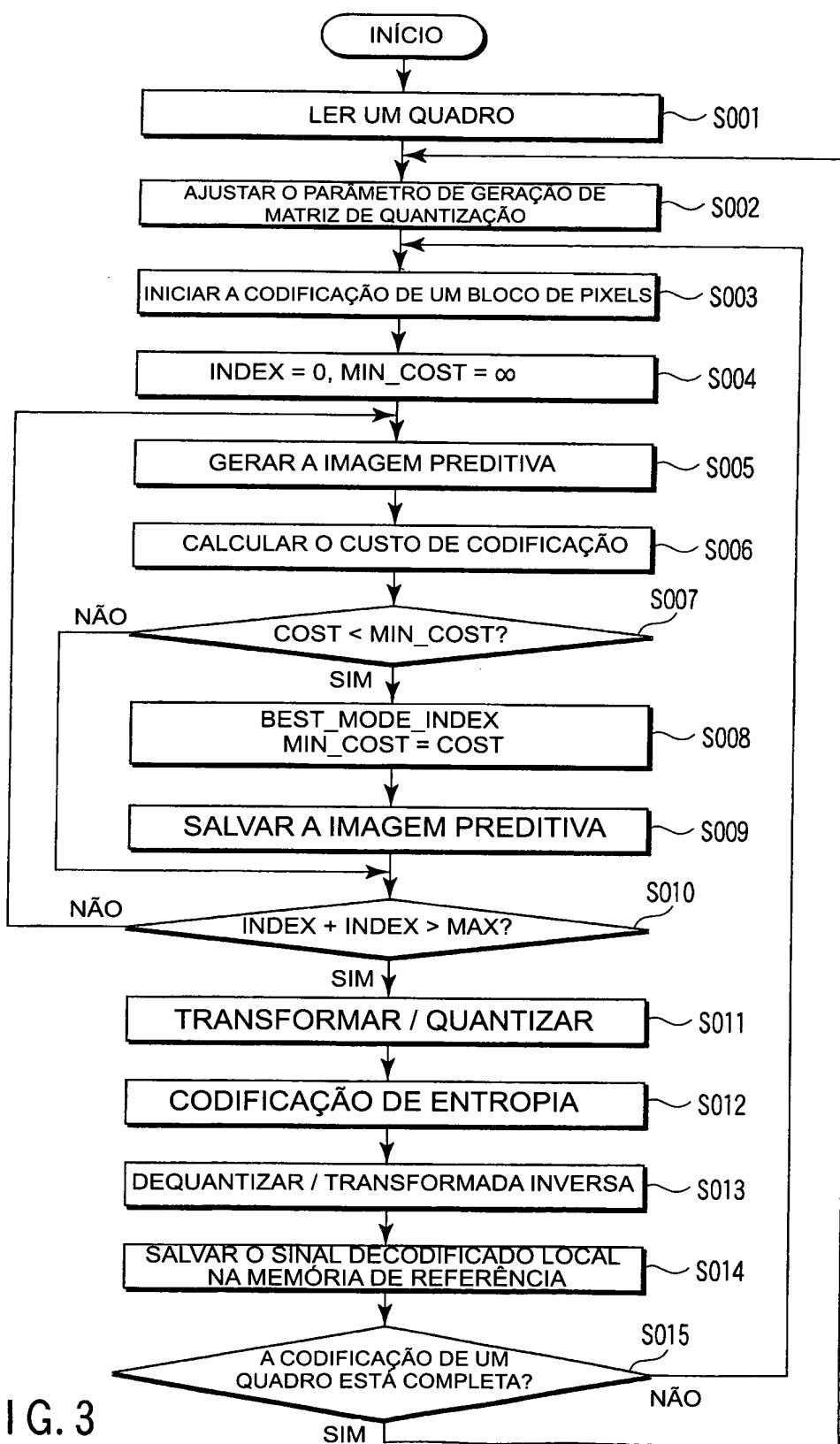


FIG. 3

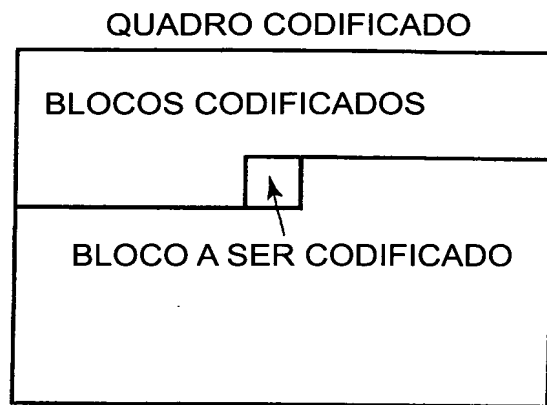


FIG. 4A

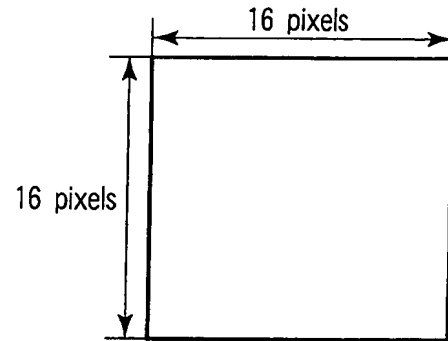


FIG. 4B

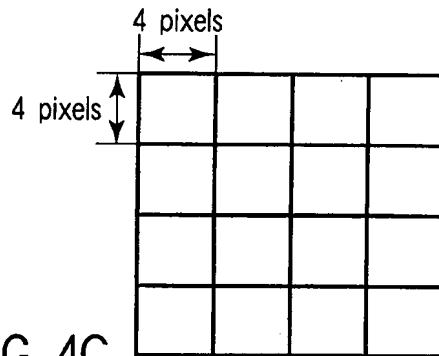


FIG. 4C

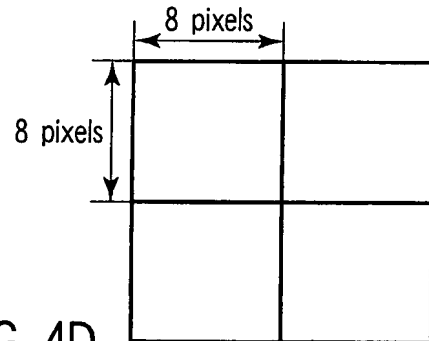


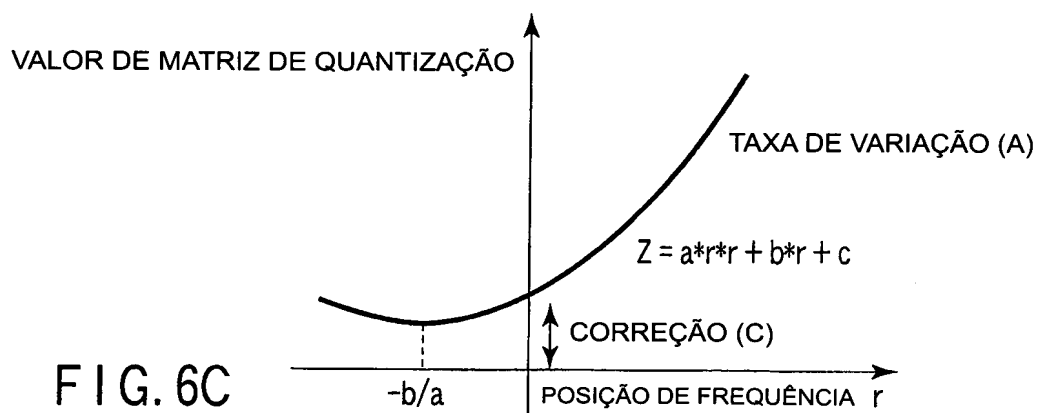
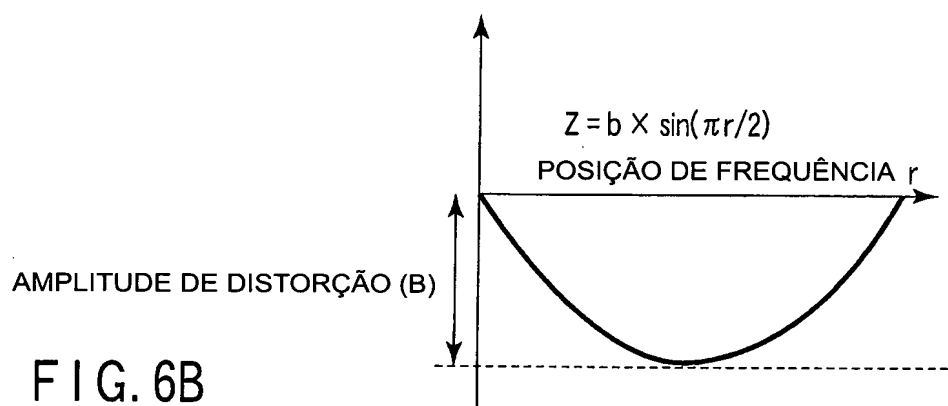
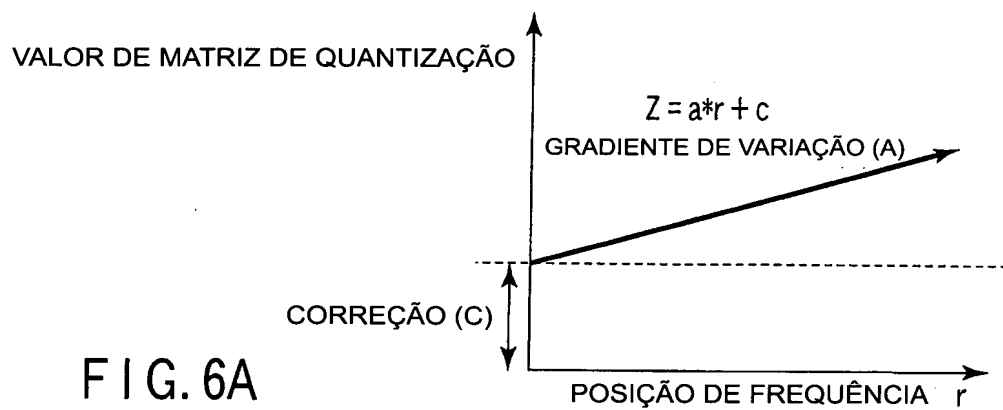
FIG. 4D

FIG. 5A

$$Q_{4 \times 4}(i, j) = \begin{bmatrix} 6 & 12 & 20 & 27 \\ 12 & 20 & 27 & 32 \\ 20 & 27 & 32 & 37 \\ 27 & 32 & 37 & 41 \end{bmatrix}$$

FIG. 5B

$$Q_{8 \times 8}(i, j) = \begin{bmatrix} 9 & 13 & 15 & 17 & 19 & 21 & 22 & 24 \\ 13 & 13 & 17 & 19 & 21 & 22 & 24 & 25 \\ 15 & 17 & 19 & 21 & 22 & 24 & 25 & 27 \\ 17 & 19 & 21 & 22 & 24 & 25 & 27 & 28 \\ 19 & 21 & 22 & 24 & 25 & 27 & 28 & 30 \\ 21 & 22 & 24 & 25 & 27 & 28 & 30 & 32 \\ 22 & 24 & 25 & 27 & 28 & 30 & 32 & 33 \\ 24 & 25 & 27 & 28 & 30 & 32 & 33 & 35 \end{bmatrix}$$



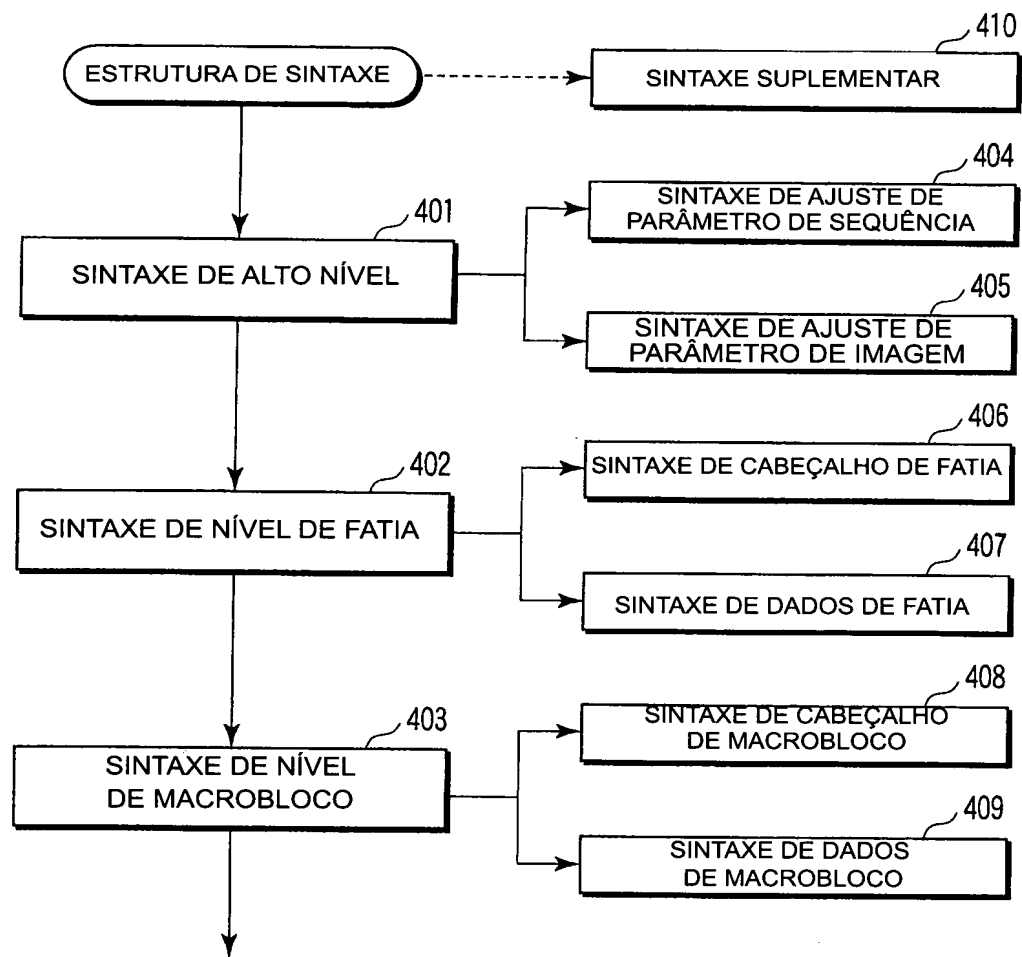


FIG. 7

```
seq_parameter_set(){  
    . . .  
    ex_seq_scaling_matrix_flag  
    if(ex_seq_scaling_matrix_flag){  
        ex_matrix_A  
        ex_matrix_B  
        ex_matrix_C  
    }  
    . . .  
}
```

FIG. 8

```
pic_parameter_set(){  
    . . .  
    ex_pic_scaling_matrix_flag  
    if(ex_pic_scaling_matrix_flag){  
        ex_matrix_type  
        ex_matrix_A  
        ex_matrix_B  
        ex_matrix_C  
    }  
    . . .  
}
```

FIG. 9

```
pic_parameter_set(){  
    . . .  
    ex_pic_scaling_matrix_flag  
    if(ex_pic_scaling_matrix_flag){  
        ex_num_of_matrix_type  
        for(i=0;i< ex_num_of_matrix_type;i++){  
            ex_matrix_type[i]  
            ex_matrix_A[i]  
            ex_matrix_B[i]  
            ex_matrix_C[i]  
        }  
    }  
    . . .  
}
```

FIG. 10

```
Supplemental_header(){  
    . . .  
    ex_sei_scaling_matrix_flag  
    if(ex_sei_scaling_matrix_flag){  
        ex_num_of_matrix_type  
        for(i=0;i< ex_num_of_matrix_type;i++){  
            ex_matrix_type[i]  
            ex_matrix_A[i]  
            ex_matrix_B[i]  
            ex_matrix_C[i]  
        }  
    }  
    . . .  
}
```

FIG. 11

FLUXO DE PROCESSO DE MÚLTIPLAS PASSADAS

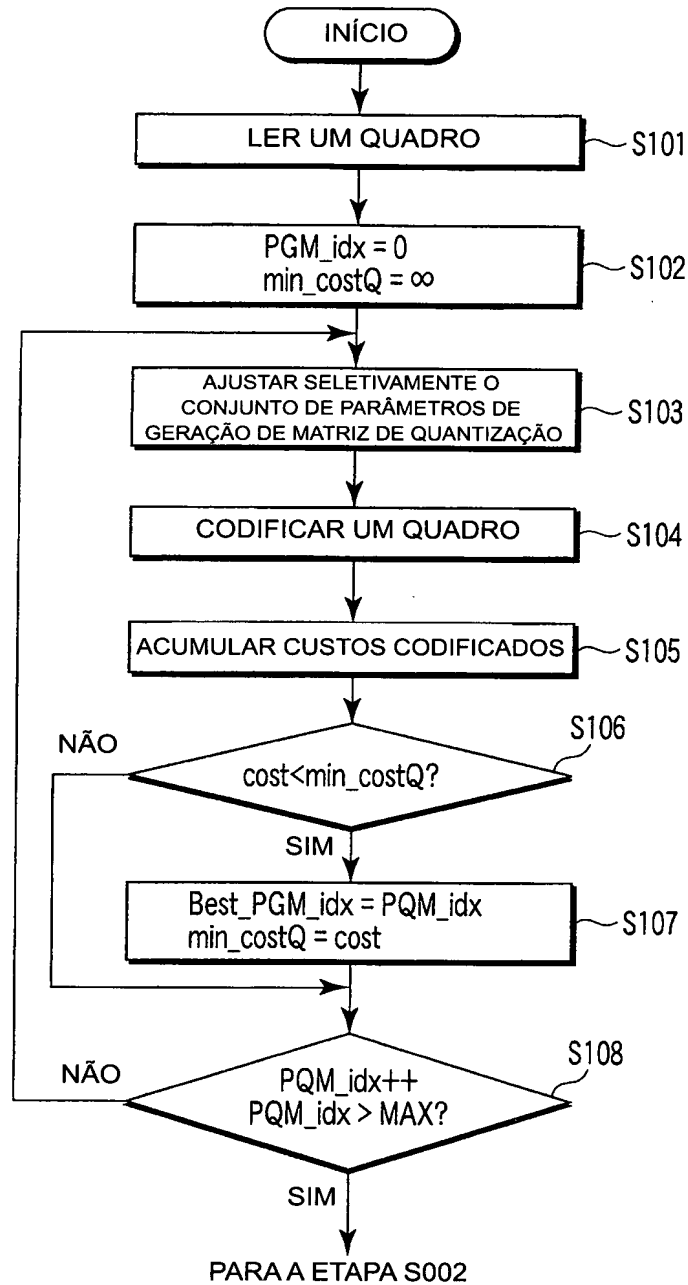


FIG. 12

```
Slice_header(){  
    . . .  
    slice_ex_scaling_matrix_flag  
    if(slice_ex_scaling_matrix_flag){  
        for(j=0;j<NumOfMatrix;j++){  
            slice_ex_matrix_type[j]  
            slice_ex_matrix_A[i]  
            slice_ex_matrix_B[i]  
            slice_ex_matrix_C[i]  
        }  
    }  
    . . .  
}
```

FIG. 13

```
Slice_header(){  
    . . .  
    slice_ex_scaling_matrix_flag  
    if(slice_ex_scaling_matrix_flag){  
        for(j=0;j<NumOfMatrix;j++){  
            slice_ex_matrix_type[j]  
            slice_ex_matrix_A[i]  
            slice_ex_matrix_B[i]  
        }  
    }  
    . . .  
}
```

FIG. 14

```

Slice_header(){
    . . .
    slice_ex_scaling_matrix_flag
    if(slice_ex_scaling_matrix_flag){
        for(j=0;j<NumOfMatrix;j++){
            slice_ex_matrix_type[j]
            slice_ex_matrix_A[i]
            slice_ex_matrix_B[i]
            slice_ex_matrix_C[i]
        }
    }else{
        for(j=0;j<NumOfMatrix;j++){
            PrevSliceExMatrixType[CurrSliceType][j]
            PrevSliceExMatrix_A[CurrSliceType][i]
            PrevSliceExMatrix_B[CurrSliceType][i]
            PrevSliceExMatrix_C[CurrSliceType][i]
        }
    }
    . . .
}

```

FIG. 15

Tipo de fatia	Valor de Tipo de CurrSlice
I-Slice (Fatia para somente predicção intra-imagem)	0
P-Slice (Fatia aplicável à predicção de imagem unidirecional)	1
B-Slice (Fatia aplicável à predicção de imagem bidirecional)	2

FIG. 16

```
Slice_header(){  
    . . .  
    slice_ex_scaling_matrix_flag  
    if(slice_ex_scaling_matrix_flag){  
        slice_ex_matrix_type  
        slice_ex_matrix_A  
        slice_ex_matrix_B  
        slice_ex_matrix_C  
    }else{  
        PrevSliceExMatrixType[CurrSliceType]  
        PrevSliceExMatrix_A[CurrSliceType]  
        PrevSliceExMatrix_B[CurrSliceType]  
        PrevSliceExMatrix_C[CurrSliceType]  
    }  
    . . .  
}
```

FIG. 17

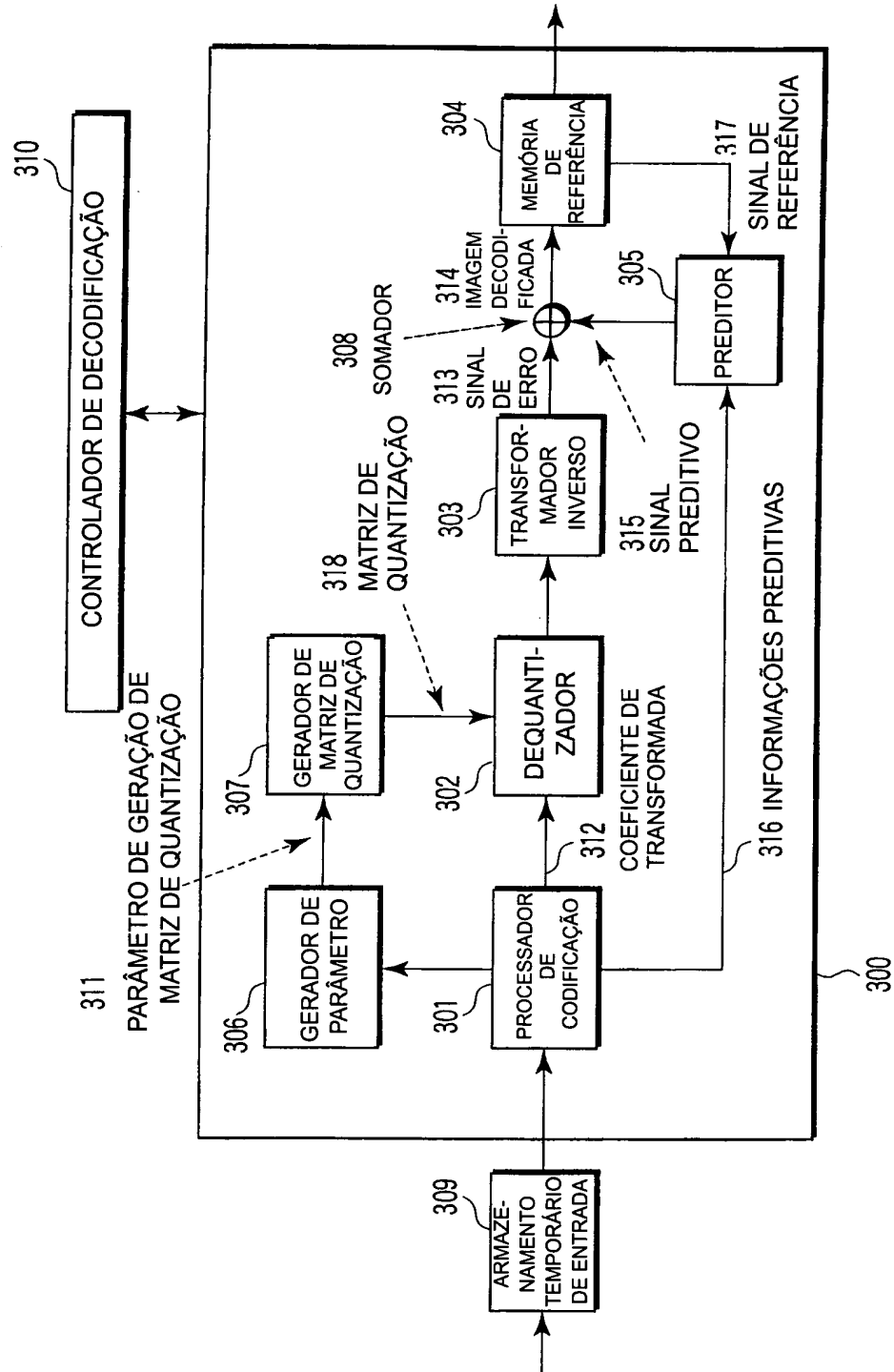


FIG. 18

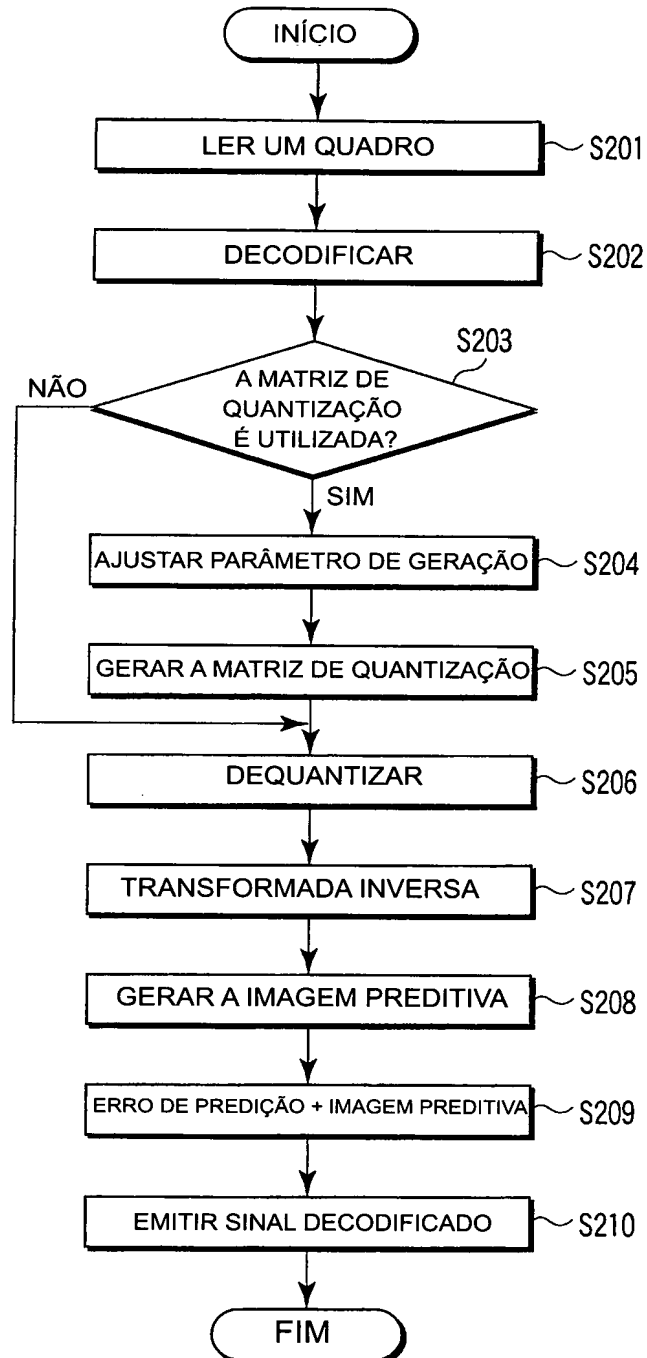


FIG. 19

RESUMO

Patente de Invenção: "MÉTODO E APARELHO E PROGRAMA DE CODIFICAÇÃO/DECODIFICAÇÃO DE VÍDEO".

- A presente invenção refere-se a um método de codificação de vídeo para executar a dequantização de um coeficiente de transformada que utiliza uma matriz de quantização do coeficiente de transformada que corresponde a cada posição de frequência, que inclui uma etapa de geração de matriz de quantização para gerar uma matriz de quantização utilizando uma função de geração utilizada para a geração da matriz de quantização e um parâmetro de geração, uma etapa de quantização para quantizar o coeficiente de transformada, utilizando a matriz de quantização gerada, e uma etapa de codificar o coeficiente de transformada quantizado para gerar um sinal de código.