

サブバンド畳み込みによる高速符号回復: 二値分類の初等的拡張

伊藤 楓馬, 都竹 千尋, 高橋 桂太, 藤井 俊彰

あらまし 画像の符号情報を効率的に圧縮するために、ブロック単位の離散コサイン変換(DCT)の符号検索問題、すなわち、DCT係数の符号をその振幅から再構成する問題を扱う。このため、二値分類機械学習に基づく高速な符号検索手法を提案する。まず、振幅と符号の3次元表現を導入し、同じ周波数帯域に属する振幅/符号を2次元スライスにパックする(サブバンドブロックと呼ぶ)。次に、各符号を2値ラベルとみなし、2値分類によって3次元振幅から符号を検索する。畳み込みニューラルネットワークを用いた2値分類アルゴリズムを実装し、3次元振幅の特徴を効率的に抽出するのに有利である。実験結果は、本手法が圧倒的に低い計算コストで正確な符号検索を達成していることを示している。

キーワード: 画像符号化, 離散コサイン変換, 符号情報, 位相回復, 二値分類, 深層ニューラルネットワーク

1. まえがき

データ中の符号情報を圧縮することは困難な問題であり、幅広い研究分野で基本的な役割を果たしている。画像符号化では、[1]で報告されているように、画像エンコーダによって生成された全ビットの約20%の多数のビットが符号に割り当てられる。このような背景から、多くの先行研究[2-5]は、画像符号化の文脈で符号圧縮法を開発し、特に離散コサイン変換(DCT)[6]係数の符号圧縮に焦点を当てている。符号のビットを減らすために、我々の以前の研究[7, 8]では、図1に要約するように、JPEGのような標準化された画像エンコーダとデコーダをわずかに修正した。エンコーダでは、画像のDCT係数は真の符号と呼ばれる符号と振幅に分離され、後者はそもそも符号化される。エンコーダは、デコーダとアルゴリズム手順を共有する符号検索法を用いて、符号を局所的に再構成する。真の符号を符号化する代わりに、真の符号と検索された符号の間の残差(XOR)を符号化する。デコーダでは、残差と検索された符号の間のXORが真の符号を完全に再構成する、符号化処理の逆パスに従って画像を復号する。

符号がある程度正しく検索された場合、残差は多くのゼロを持つが、少数のゼロを持つ;残差はエントロピー符号化法を用いて圧縮することができる、例えば[9-12]。したがって、検索された標識の精度は、標識圧縮を成功させるために極めて重要である。

図2aにまとめたように、様々な符号検索手法が提案されている。我々の知る限り、これまでの手法は、例外なく、すべて反復最適化に基づいている。例えば、[7, 8]では、画像と符号付きDCT係数が交互に検索され、検索された係数の符号が抽出される。[14]では、符号は変数とみなされ、反復的に最適化される。これまでの手法では、符号の大部分を正確に検索できることが実証されている。例えば、[8]で報告されているように、全符号のおよそ75%を検索することができる。しかし、これらの方法は、主に各反復のDCTと逆DCT(IDCT)を計算するため、過剰な計算コストがかかる。例えば、[8]で報告されているように、 256×256 の符号を検索するのに約10秒かかるため、符号検索はリアルタイムの符号化・復号化シナリオでは実用的でない。

符号検索の速度を高速化するために、DCT-IDCTを用いない高速な手法を提案し、図2bにまとめた。従来の方法とは異なり、本手法は符号を直接的に検索し、符号検索は二値分類問題の変形と見なす。

Received December 20, 2024; Revised March 13, 2025; Accepted April 17, 2025

† Graduate School of Engineering, Nagoya University

(Furo-cho, Chikusa-ku, Nagoya 464-8603, Japan)

【機械翻訳コンテンツの著作権について】

当サイトに掲載されている原著論文の著作権は映像情報メディア学会に帰属します。原著論文および翻訳論文の無断使用は禁止します。
引用の際には、必ず原著論文の書誌情報をご記載ください。詳細は映像情報メディア学会著作権規定をご覧ください。

Fast Sign Retrieval via Sub-band Convolution: An Elementary Extension of Binary Classification

Fuma Ito[†], Chihiro Tsutake (member)[†], Keita Takahashi (member)[†],
Toshiaki Fujii (member)[†]

Abstract To efficiently compress the sign information of images, we address a sign retrieval problem for the block-wise discrete cosine transformation (DCT): reconstruction of the signs of DCT coefficients from their amplitudes. To this end, we propose a fast sign retrieval method on the basis of binary classification machine learning. We first introduce 3D representations of the amplitudes and signs, where we pack amplitudes/signs belonging to the same frequency band into a 2D slice, referred to as the sub-band block. We then retrieve the signs from the 3D amplitudes via binary classification, where each sign is regarded as a binary label. We implement a binary classification algorithm using convolutional neural networks, which are advantageous for efficiently extracting features in the 3D amplitudes. Experimental results demonstrate that our method achieves accurate sign retrieval with an overwhelmingly low computation cost.

Key words: Image coding, discrete cosine transform, sign information, phase retrieval, binary classification, deep neural network.

1. Introduction

Compressing sign information in data is a challenging problem, and it plays a fundamental role in a wide range of research fields. In image coding, a large number of bits, approximately 20% of the total bits, generated by an image encoder are allocated for the signs, as reported in [1]. Against this background, many previous works [2–5] have developed sign compression methods within the context of image coding, specifically focusing on compressing signs of discrete cosine transform (DCT) [6] coefficients.

To reduce the bits for signs, in our earlier work [7,8], we slightly modified a standardized image encoder and decoder, such as JPEG, as summarized in Fig. 1. In the encoder, DCT coefficients of an image are separated into the signs, referred to as the *true signs*, and amplitudes; the latter is encoded in the first place. The encoder locally reconstructs the signs using a *sign retrieval* method whose algorithmic procedure is shared with the decoder. Instead of encoding the true signs, we encode residuals (XORs) between the true signs and retrieved ones. In the decoder, an image is decoded by following the reverse path of the encoding process, where the XORs between the residuals and retrieved signs perfectly reconstruct the true signs. If the signs

are retrieved correctly to some extent, the residuals have many zeros but few ones; the residuals can be compressed using entropy coding methods, e.g., [9–12]. Therefore, the accuracy of the retrieved signs is crucial for successful sign compression.

Various sign retrieval methods have been proposed, as summarized in Fig. 2a. To the best of our knowledge, all of the previous methods, without exception, are based on iterative optimization. For example, in [7,8], an image and signed DCT coefficients are alternatively retrieved, and the signs of the retrieved coefficients are extracted. In [14], the signs are regarded as variables and are optimized in an iterative manner. Previous methods have demonstrated the ability to accurately retrieve a large portion of the signs; for example, as reported in [8], roughly 75% of the total signs can be retrieved. However, these methods come at an excessive computation cost, primarily due to computing the DCT and the inverse DCT (IDCT) for each iteration. For example, as reported in [8], it takes approximately 10 seconds to retrieve 256×256 signs, making sign retrieval impractical for real-time encoding and decoding scenarios.

To accelerate the speed of sign retrieval, we propose a fast method **without the DCT-IDCT iteration**, which is summarized in Fig. 2b. Unlike previous methods, our method retrieves the signs in a direct manner, where sign retrieval is regarded as a variant of a binary classification problem. We first introduce

Received December 20, 2024; Revised March 13, 2025; Accepted April 17, 2025

[†] Graduate School of Engineering, Nagoya University
(Furo-cho, Chikusa-ku, Nagoya 464-8603, Japan)

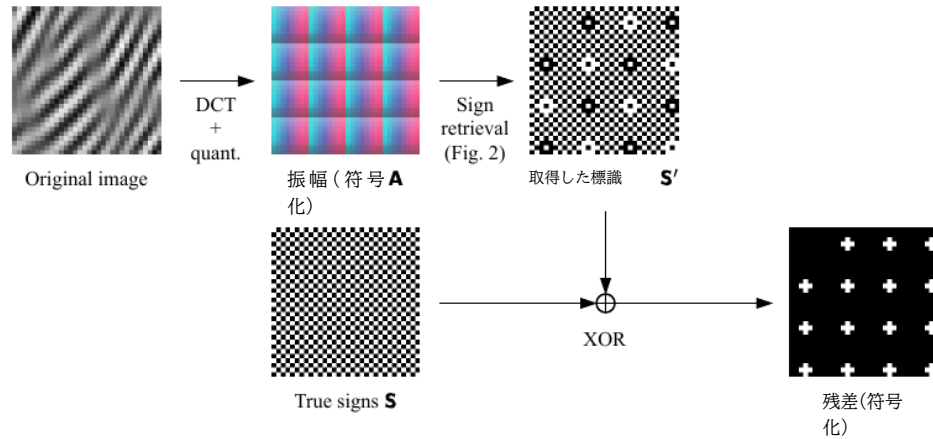
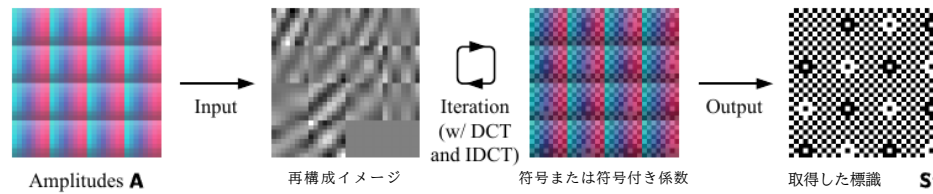
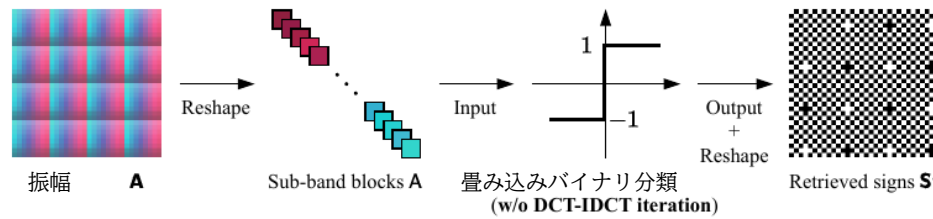


図1:[7]の符号化処理。



(a) 従来方法 [7, 8, 13, 14]



(b) 提案

図2: 符号検索の方法

まず、振幅と符号の3次元表現を導入し、同じ周波数帯域に属する振幅/符号を2次元スライスにパックする(サブバンドブロックと呼ぶ)[2]。次に、各符号を2値ラベルとみなし、2値分類によって3次元振幅から符号を検索する。畳み込みニューラルネットワーク(CNN)を用いた2値分類アルゴリズムを実装し、3次元振幅の特徴を効率的に抽出するのに有利である。実験結果は、我々の手法が圧倒的に低い計算コストで正確な符号検索を達成することを示している:従来の手法にかかった実行時間のわずか0.93%である。

限界: 本論文では、JPEG[15]で行われたような一定サイズのブロックを利用するが、これは高効率ビデオ符号化(HEVC)[16]や汎用ビデオ符号化(VVC)[17]のような最先端の画像符号化規格における可変サイズのものとは互換性がない。

しかし、我々の有望な結果(セクション4で議論)は、画像符号化の分野における更なる進歩に貢献していると考えている。

2. 関連研究

2.1 標識の検索

図1に示す符号検索に基づく方法[7]は、符号圧縮の代表的な研究である。符号検索は位相検索の特殊な例として知られている:すなわち、符号 ± 1 を決定することは、ガウス平面上でそれらの位相 θ_0 と π を検索することと同じである。そのため、符号検索は通常、位相検索アプローチに基づいて対処されてきた。津竹ら[7]は、位相検索に関する先行研究[18, 19]に動機づけられた ℓ_1 -norm最小化法を提案した。Linら[13]は、全変動最小化に基づく同様の方法を提案している。鈴木ら[8]はCNNベースの最小二乗法を提案した。

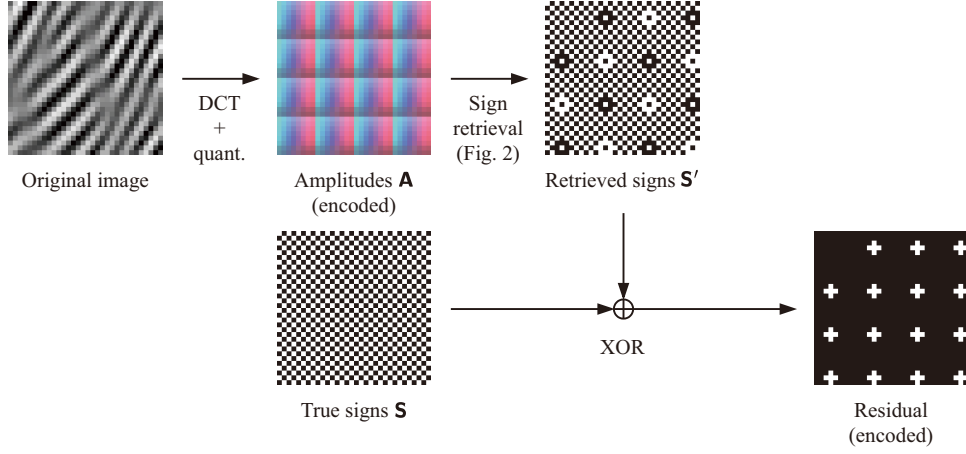


Fig. 1: Encoding process in [7].

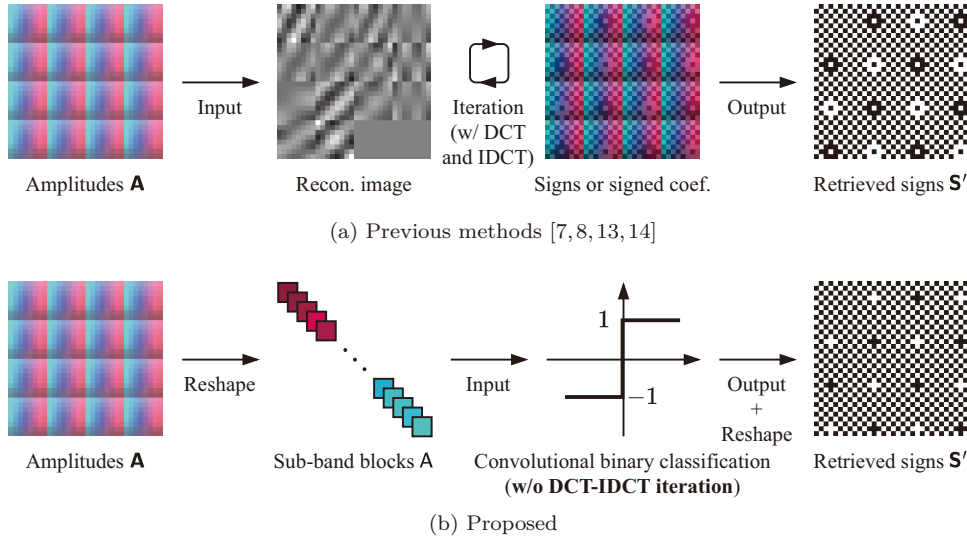


Fig. 2: Sign retrieval methods.

3D representations of the amplitudes and signs, where we pack amplitudes/signs belonging to the same frequency band into a 2D slice, referred to as the *sub-band block* [2]. We then retrieve the signs from the 3D amplitudes via binary classification, where each sign is regarded as a binary label. We implement a binary classification algorithm using convolutional neural networks (CNNs), which are advantageous for efficiently extracting features in the 3D amplitudes. Experimental results demonstrate that our method achieves accurate sign retrieval with an overwhelmingly low computation cost: only **0.93%** of the execution time taken for previous methods.

Limitations: Throughout this paper, we utilize constant-size blocks, as was done in JPEG [15], which is not compatible with the variable-size ones in state-of-the-art image coding standards such as high-efficiency video coding (HEVC) [16] and versatile video cod-

ing (VVC) [17]. However, we believe that our promising results (discussed in Section 4) will contribute to further advancements in the field of image coding.

2. Related work

2.1 Sign retrieval

The sign-retrieval-based method [7], illustrated in Fig. 1, is a seminal work in sign compression. Sign retrieval is known as a special instance of phase retrieval: namely, determining the signs ± 1 is identical to retrieving their phases, 0 and π , on the Gaussian plane. Therefore, sign retrieval has typically been addressed on the basis of phase retrieval approaches. Tsutake et al. [7] proposed an ℓ_1 -norm minimization method, motivated by the prior works on phase retrieval [18, 19]. Lin et al. [13] proposed a similar method based on total variation minimization. Suzuki et al. [8] proposed a CNN-based least squares method. Sidiropoulos et al. [14] pro-

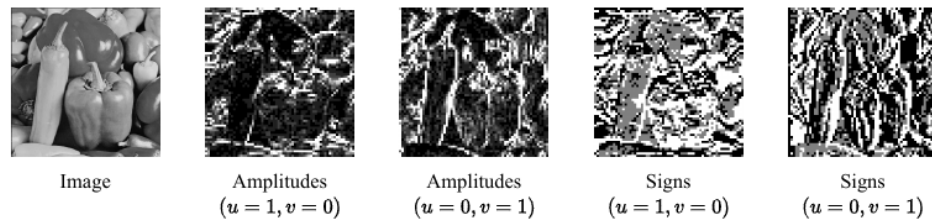


図3: 画像から得られたサブバンドブロック、8倍拡大。振幅が大きいほど輝度が明るくなる。正と負の符号はそれぞれ白と黒で表されている。振幅の黒い領域と符号の灰色の領域は、有意でない係数に対応する。

Sidiropoulosら[14]は二次計画法を提案した。これらの方法は、例外なく、図2aに示すように、繰り返し処理によって符号を取得する。例えば、[7, 8]では、Gerchberg-Saxton法[20]やFienup法[21]と同様に、画像と符号付きDCT係数が交互に計算されている。Frank-Wolfe法[22]は[14]で利用された別の例で、勾配法の変形であり、降下方向を決定するためにDCTを計算する。反復的な性質のため、DCTとIDCTは各反復で実施されるべきであり、大きな計算コストにつながる。

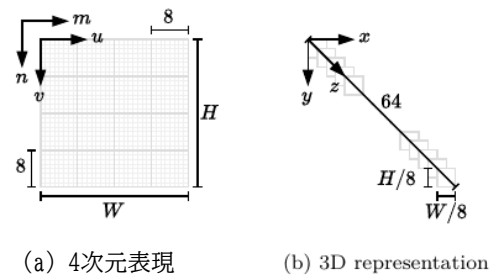


図4: 座標系。

2. 2 符号検索を行わない符号圧縮 Ponomarenkoら[3]は、符号圧縮の前に振幅を符号化する、符号検索に基づく方法[7]と同様の予測に基づく方法を提案した。彼らの方法は、目的関数を最小化する符号を探索する振幅制約付き組合せ問題を解くことによって、符号をブロックごとに予測する。デコーダで真の符号を再構成するために、真の符号と予測された符号の間の残差が符号化される。3]で報告されているように、残差のビット量は真の符号のビット量の $\{60-85\}$ である。その後、[23-25]のような、精度と計算効率を向上させるための研究が開発されている。

Clareら[26]は、符号データ隠蔽として知られる符号圧縮法を提案し、これは最先端の画像符号化規格であるHEVC[16]とVVC[17]の基本コンポーネントである。この方法により、各ブロックのDCT係数の符号を1つずつ符号化することもスキップできる。デコーダはパリティチェックを実行することで、欠落した符号を補正することができる。26]で報告されているように、符号データ隠蔽はHM 4.0アンカー*と比較して、約0.6%のBD率[27](評価指標)を達成している。

符号データの隠蔽は、ブロックごとに複数の符号の符号化をスキップできる[28-30]など、その後の研究で改良されている。

Tu and Tran [2]は、サブバンド符号化[31, 32]に着想を得た符号圧縮法を提案した。彼らは、同じ周波数帯域に属する振幅/符号を含む、いわゆるサブバンドブロックを構築し、真の符号はサブバンド領域で符号化・復号化される。図3は画像から得られたサブバンドブロックの例であり、 u と v はそれぞれ水平方向と垂直方向に沿った周波数インデックスを表す。このように、各要素はその空間的、サブバンド的な近傍と強い相関があることがわかる。このサブバンド表現の顕著な統計的特徴は、我々の手法にインポートされる。

3. 提案手法

3.1 表記法

スカラー変数とスカラー値関数は、通常の手体で示される。一定の値を大文字で表すと、以下のようになる。3次元テンソル、3次元テンソル値関数、およびそれらの要素はサンセリフ体で表記される。4次元テンソルとその要素は太字のサンセリフ体で示されている。集合はギリシャ文字で示される。

図1、図2に示すように、量子化されたDCT係数の振幅を4次元テンソル A_s 、真の符号を s_s 、検索された符号を s_s^0 とそれぞれ表す。図4aは A_s 、 s_s 、と s_s^0 の座標系を示す。

* <https://hevc.hhi.fraunhofer.de>

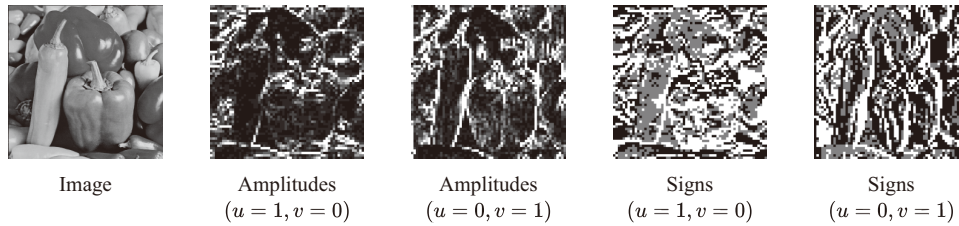


Fig. 3: Sub-band blocks obtained from image, $\times 8$ enlarged. Larger amplitudes are represented by brighter luminance. Positive and negative signs are represented by white and black, respectively. Black regions in amplitudes and gray ones in signs correspond to non-significant coefficients.

posed a quadratic programming method. All of these methods, without exception, retrieve the signs via iterative processes, as illustrated in Fig. 2a. For example, an image and signed DCT coefficients are alternatively computed in [7, 8], similar to the Gerchberg-Saxton method [20] and Fienup method [21]. The Frank-Wolfe method [22], utilized in [14], is another example, which is a variant of a gradient method and computes the DCT to determine the descent direction. Due to the iterative nature, the DCT and IDCT should be conducted in each iteration, leading to significant computation costs.

2.2 Sign compression without sign retrieval

Ponomarenko et al. [3] proposed a prediction-based method similar to the sign retrieval-based method [7], where the amplitudes are encoded before sign compression. Their method predicts the signs block-by-block by solving an amplitude-constrained combinatorial problem, where the signs that minimize an objective function are searched. To reconstruct the true signs at the decoder, residuals between the true signs and predicted ones are encoded. As reported in [3], the bit amount for the residuals is 60–85% of that for the true signs. Subsequent works, such as [23–25], have been developed to enhance the accuracy and computation efficiency.

Clare et al. [26] proposed a sign compression method, known as sign data hiding, which is a fundamental component in state-of-the-art image coding standards, HEVC [16] and VVC [17]. This method allows us to skip encoding a single sign of DCT coefficients for each block. The decoder can compensate for the missing sign by executing a parity check. As reported in [26], sign data hiding achieves a BD-rate [27] (a rate-distortion metric) of approximately -0.6% compared to the HM 4.0 anchor*. Sign data hiding has been improved in subsequent works, such as [28–30], where the

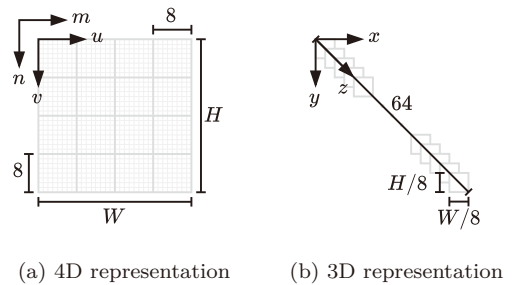


Fig. 4: Coordinate systems.

encoding of multiple signs per block can be skipped.

Tu and Tran [2] proposed a sign compression method inspired by sub-band coding [31, 32]. They constructed the so-called sub-band blocks, which include amplitudes/signs belonging to the same frequency band, and the true signs are encoded and decoded in the sub-band domain. Figure 3 shows examples of sub-band blocks obtained from an image, where u and v represent frequency indices along the horizontal and vertical directions, respectively. As can be seen, each element is strongly correlated with its spatial and sub-band neighbors. This remarkable statistical feature of the sub-band representation is imported into our method.

3. Proposed method

3.1 Notations

A scalar variable and a scalar-valued function are denoted by a regular typeface. A constant value is denoted by a capital case. A 3D tensor, 3D-tensor-valued function, and their elements are denoted by a sans-serif typeface. A 4D tensor and its elements are denoted by a bold sans-serif typeface. A set is denoted by a Greek letter.

As shown in Figs. 1 and 2, the amplitudes of quantized DCT coefficients are denoted by a 4D tensor $\mathbf{A}$, and the true and retrieved signs are denoted by $\mathbf{S}$ and $\mathbf{S}'$, respectively. Figure 4a illustrates the coordinate

* <https://hevc.hhi.fraunhofer.de>

画像の幅と高さをそれぞれ W と H で表す。ブロックサイズは 8×8 に固定し、DCT基底関数の周波数インデックスを $u \in [0, 8]$ 、 $v \in [0, 8]$ とする。ブロックインデックスは $m \in [0, W/8]$ 、 $n \in [0, H/8]$ で示される。 A, S, θ の (m, n) 番目のブロックの (u, v) 番目の要素を $[-]_{u, v, m, n}$ とする。図4bに示すように、 A, S 、と θ の3次元表現を導入し、それぞれ A, S 、と θ とする。その軸は x, y 、と z であり、 $_{\{(x, y, z)\text{-番目の要素は}[-]_{x, y, z}$ で示される。

3.2 サブバンド畳み込みによる符号検索

符号検索を高速化するために、DCT-IDCT反復を用いない高速な手法を提案し、図2bにまとめた。本手法は、二値分類機械学習に基づいて符号を検索する。4次元テンソルをそのまま処理するには大きな計算コストがかかるため、4次元の振幅と符号を3次元のものに以下のように再形成する：

$$A_{x, y, z} = \mathbf{A}_{u(z), v(z), m(x), n(y)} \quad (1)$$

$$S_{x, y, z} = \mathbf{S}_{u(z), v(z), m(x), n(y)}, \quad (2)$$

ここで、 $u(z) = bz/8c$ 、 $v(z) = z \bmod 8$ 、 $m(x) = \lfloor x/8 \rfloor$ 、 $n(y) = \lfloor y/8 \rfloor$ 。 $(u(z), v(z))$ -th番目の周波数帯域に属する振幅/符号は、図3に示すように、サブバンドブロックと呼ばれる z -th番目の2次元スライスに詰められ、3次元テンソルには 64 個のサブバンドブロックがすべて含まれる。DCTカーネルの偶数/奇数対称性を考慮することで、これらのサブバンドを扱うためのより良い方法につながるかもしれないが、それは今後の課題として残す。リシエーピングプロセスの有効性は、4.2節で示す。

3次元振幅 A から、各符号を2値ラベルとみなし、2値分類により符号 S を再構成することを目指す。CNNを用いた2値分類アルゴリズムを実装し、空間方向とサブバンド方向に沿った A の特徴を効率的に抽出する能力を持つ。表1は我々のCNNのネットワーク・アーキテクチャを示しており、 l は層数を表す。Conv- i は、各出力チャンネルに対して $3 \times 3 \times C$ (3×3 :空間カーネルサイズ、 C :入力チャンネル数)のサイズの3Dカーネルを持ち、すべてのチャンネル(周波数サブバンド)が各畳み込みステップで一緒に相互作用することができる。サブバンド領域の空間カーネルサイズ 3×3 は、元の画像領域の 24×24 ピクセルに相当するため、各畳み込みステップで十分に大きな空間近傍が考慮される。

表1: ネットワーク・アーキテクチャ

Layer	Conv-0	...	Conv- i	...	Conv- I
Ker. size	$3 \times 3 \times 64$	...	$3 \times 3 \times 128$	...	$3 \times 3 \times 128$
In/out ch.	64/128	...	128/128	...	128/63
Act.	ReLU	...	ReLU	...	Sigmoid
Input	3D amp. A	...	Conv $i - 1$	...	Conv $I - 1$

表2: 計算環境。

CPU	Intel Core i9-13900KF
メインメモリ	32 GB
OS	Ubuntu 20.04 LTS
言語とフレームワーク	Python 3.10.11 & PyTorch 1.13.1

我々のCNNは3次元振幅 A をサイズ $W/8 \times H/8 \times 63$ の符号テンソルに写像する。ここで θ は畳み込みカーネルとバイアスを含む学習可能なパラメータの集合を表す。

学習手順は以下の通りである。ここで、 $\Phi = \{A[k], S[k]\} : k \in [0, K]\}$ を学習用の3次元振幅と3次元符号の集合とする。経験的リスクを定義する

$$r(\Phi, \Theta) = \frac{1}{WHK} \sum_{x, y, z} \sum_k l(F_{x, y, z}(A[k], \Theta), S_{x, y, z}[k]), \quad (3)$$

ここで l は損失関数である。経験的リスクを最小化することで、 Θ で示される最適パラメータを得る：

$$\tilde{\Theta} = \underset{\Theta}{\operatorname{argmin}} r(\Phi, \Theta). \quad (4)$$

(4)の解はAdamオプティマイザを用いて得られる。

(3)の損失関数 l は以下のように定義される。ここで

$$b(S_{x, y, z}) = \begin{cases} 1 & \text{if } S_{x, y, z} = 1 \\ 0 & \text{otherwise} \end{cases} \quad (5)$$

は真の符号 $S_{x, y, z}$ のゼロワン表現である。損失関数として2値クロスエントロピー関数を利用する：

$$l(F_{x, y, z}, S_{x, y, z}) = -b(S_{x, y, z}) \log(F_{x, y, z}) - (1 - b(S_{x, y, z})) \log(1 - F_{x, y, z}), \quad (6)$$

ここで (A, Θ) の $F_{x, y, z}$ は表記の便宜上省略した。振幅 (x, y, z) -thが0の場合、対応する損失値を0とする。

テスト段階では、 $F_{x, y, z}(A, \Theta_{302})$ を閾値処理することで、以下のように検索された符号を得る。

【機械翻訳コンテンツの著作権について】

当サイトに掲載されている原著論文の著作権は映像情報メディア学会に帰属します。原著論文および翻訳論文の無断使用は禁止します。
 引用の際には、必ず原著論文の書誌情報をご記載ください。詳細は映像情報メディア学会著作権規定をご覧ください。

system of $\mathbf{A}$, $\mathbf{S}$, and $\mathbf{S}'$. The width and height of an image are denoted by W and H , respectively. The block size is fixed to 8×8 . Frequency indices of the DCT basis function are denoted by $u \in [0, 8)$ and $v \in [0, 8)$. Block indices are denoted by $m \in [0, W/8)$ and $n \in [0, H/8)$. The (u, v) -th element in the (m, n) -th block of $\mathbf{A}$, $\mathbf{S}$, and $\mathbf{S}'$ is denoted by $[\cdot]_{u,v,m,n}$. As shown in Fig. 4b, we introduce 3D representations of $\mathbf{A}$, $\mathbf{S}$, and $\mathbf{S}'$, which are denoted by $\mathbf{A}$, $\mathbf{S}$, and $\mathbf{S}'$, respectively. Their axes are x , y , and z , and the (x, y, z) -th element is denoted by $[\cdot]_{x,y,z}$.

3.2 Sign retrieval via sub-band convolution

To accelerate sign retrieval, we propose a fast method without the DCT-IDCT iteration, which is summarized in Fig. 2b. Our method retrieves the signs based on binary classification machine learning. Because processing 4D tensors as they are requires a large computational cost, we reshape the 4D amplitudes and signs into 3D ones, as follows:

$$\mathbf{A}_{x,y,z} = \mathbf{A}_{u(z),v(z),m(x),n(y)} \quad (1)$$

$$\mathbf{S}_{x,y,z} = \mathbf{S}_{u(z),v(z),m(x),n(y)}, \quad (2)$$

where $u(z) = \lfloor z/8 \rfloor$, $v(z) = z \bmod 8$, $m(x) = x$, and $n(y) = y$. The amplitudes/signs belonging to the $(u(z), v(z))$ -th frequency band are packed in the z -th 2D slice, referred to as the sub-band block, and the 3D tensors contain all the 64 sub-band blocks, as shown in Figure 3. We simply stack these 64 sub-bands along the z dimension; consideration of even/odd symmetry of the DCT kernels might lead to a better method for handling these sub-bands, but we leave it as the future work. The effectiveness of the reshaping process will be demonstrated in Section 4.2.

We aim to reconstruct the signs $\mathbf{S}$ from the 3D amplitudes $\mathbf{A}$ via binary classification, where each sign is regarded as a binary label. We implement a binary classification algorithm using CNNs, which have the capability to efficiently extract features in $\mathbf{A}$ along the spatial and sub-band directions. Table 1 shows the network architecture of our CNN, where I represents the number of layers. The Conv- i has a 3D kernel of the size $3 \times 3 \times C$ (3×3 : spatial kernel size and C : the number of input channels) for each output channel, where all the channels (frequency sub-bands) can interact together in each convolution step. The spatial kernel size 3×3 in the sub-band domain corresponds to 24×24 pixels in the original image domain; thus, a sufficiently large spatial neighbor is considered in each convolution

Table 1: Network architecture.

Layer	Conv-0	...	Conv- i	...	Conv- I
Ker. size	$3 \times 3 \times 64$	...	$3 \times 3 \times 128$	...	$3 \times 3 \times 128$
In/out ch.	64/128	...	128/128	...	128/63
Act.	ReLU	...	ReLU	...	Sigmoid
Input	3D amp. $\mathbf{A}$	...	Conv $i - 1$	...	Conv $I - 1$

Table 2: Computing environment.

CPU	Intel Core i9-13900KF
Main memory	32 GB
OS	Ubuntu 20.04 LTS
Language & framework	Python 3.10.11 & PyTorch 1.13.1

step. Our CNN maps the 3D amplitudes $\mathbf{A}$ into a sign tensor of the size $W/8 \times H/8 \times 63$, where the number of channels, 63, corresponds to the AC components. We denote the inferred signs by $\mathbf{F}(\mathbf{A}, \Theta)$, where Θ represent a set of learnable parameters including convolution kernels and biases.

The training procedure is as follows. Let $\Phi = \{(\mathbf{A}[k], \mathbf{S}[k]) : k \in [0, K)\}$ be a set of 3D amplitudes and 3D signs for training, where K denotes the number of data. We define the empirical risk

$$r(\Phi, \Theta) = \frac{1}{WH} \frac{1}{K} \sum_{x,y,z} \sum_k l(\mathbf{F}_{x,y,z}(\mathbf{A}[k], \Theta), \mathbf{S}_{x,y,z}[k]), \quad (3)$$

where l is a loss function. We obtain the optimal parameter, denoted by $\tilde{\Theta}$, by minimizing the empirical risk:

$$\tilde{\Theta} = \underset{\Theta}{\operatorname{argmin}} r(\Phi, \Theta). \quad (4)$$

The solution to (4) is obtained using the Adam optimizer.

The loss function l in (3) is defined as follows. Let

$$b(\mathbf{S}_{x,y,z}) = \begin{cases} 1 & \text{if } \mathbf{S}_{x,y,z} = 1 \\ 0 & \text{otherwise} \end{cases} \quad (5)$$

be a zero-one representation of the true sign $\mathbf{S}_{x,y,z}$. We utilize the binary cross-entropy function as the loss function:

$$l(\mathbf{F}_{x,y,z}, \mathbf{S}_{x,y,z}) = -b(\mathbf{S}_{x,y,z}) \log(\mathbf{F}_{x,y,z}) - (1 - b(\mathbf{S}_{x,y,z})) \log(1 - \mathbf{F}_{x,y,z}), \quad (6)$$

where $(\mathbf{A}, \Theta)$ of $\mathbf{F}_{x,y,z}$ has been omitted for notation convenience. If the (x, y, z) -th amplitude is zero, we set the corresponding loss value to 0.

At the test phase, we obtain the retrieved signs by thresholding $\mathbf{F}_{x,y,z}(\mathbf{A}, \tilde{\Theta})$ as follows.

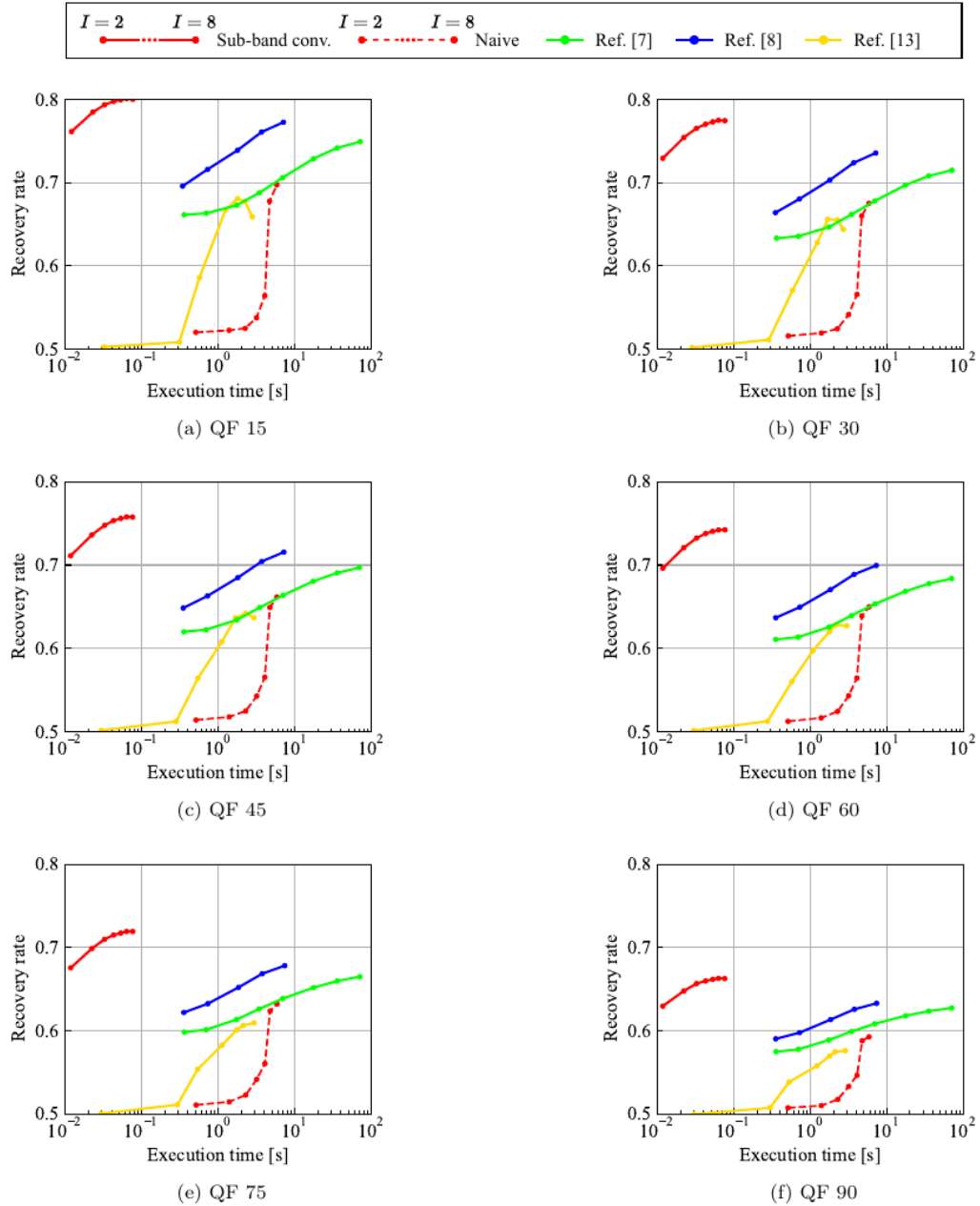


図5: 実行時間に対する回復率。

$$S'_{x,y,z} = \begin{cases} +1 & \text{if } F_{x,y,z} \geq 1/2 \\ -1 & \text{otherwise} \end{cases} \quad (7)$$

$x(m)=m$, $y(n)=n$, $z(u, v)=8u+v$ とする。検索された標識の4次元表現は以下のように計算できる。

$$S'_{u,v,m,n} = S'_{x(m),y(n),z(u,v)} \quad (8)$$

図1に示すように、真の符号 s と検索された符号 s' の間の残差を符号化する。

4. 実験結果

4.1 コンフィギュレーション

For training our CNN, we randomly cropped $K =$

CLIC 2020データセット[33]から $W \times H = 512 \times 512$ のサイズを持つ6,056枚の画像。ビット深度は8ビット/ピクセルで、各ピクセル値は0~1の範囲内で正規化した。画像にはブロック単位のDCTを適用した。DCT係数は品質係数(QF)75を用いて量子化した。(1)と(2)に従って4次元の振幅と符号を再形成した。その結果、 $W/8 \times H/8 \times 64 = 64 \times 64 \times 64$ のサイズの3次元テンソルが得られた。学習率 2×10^{-4} のAdamオプティマイザを利用した。バッチサイズは256、エポック数は15,000であった。

【機械翻訳コンテンツの著作権について】

当サイトに掲載されている原著論文の著作権は映像情報メディア学会に帰属します。原著論文および翻訳論文の無断使用は禁止します。
 引用の際には、必ず原著論文の書誌情報をご記載ください。詳細は映像情報メディア学会著作権規定をご覧ください。

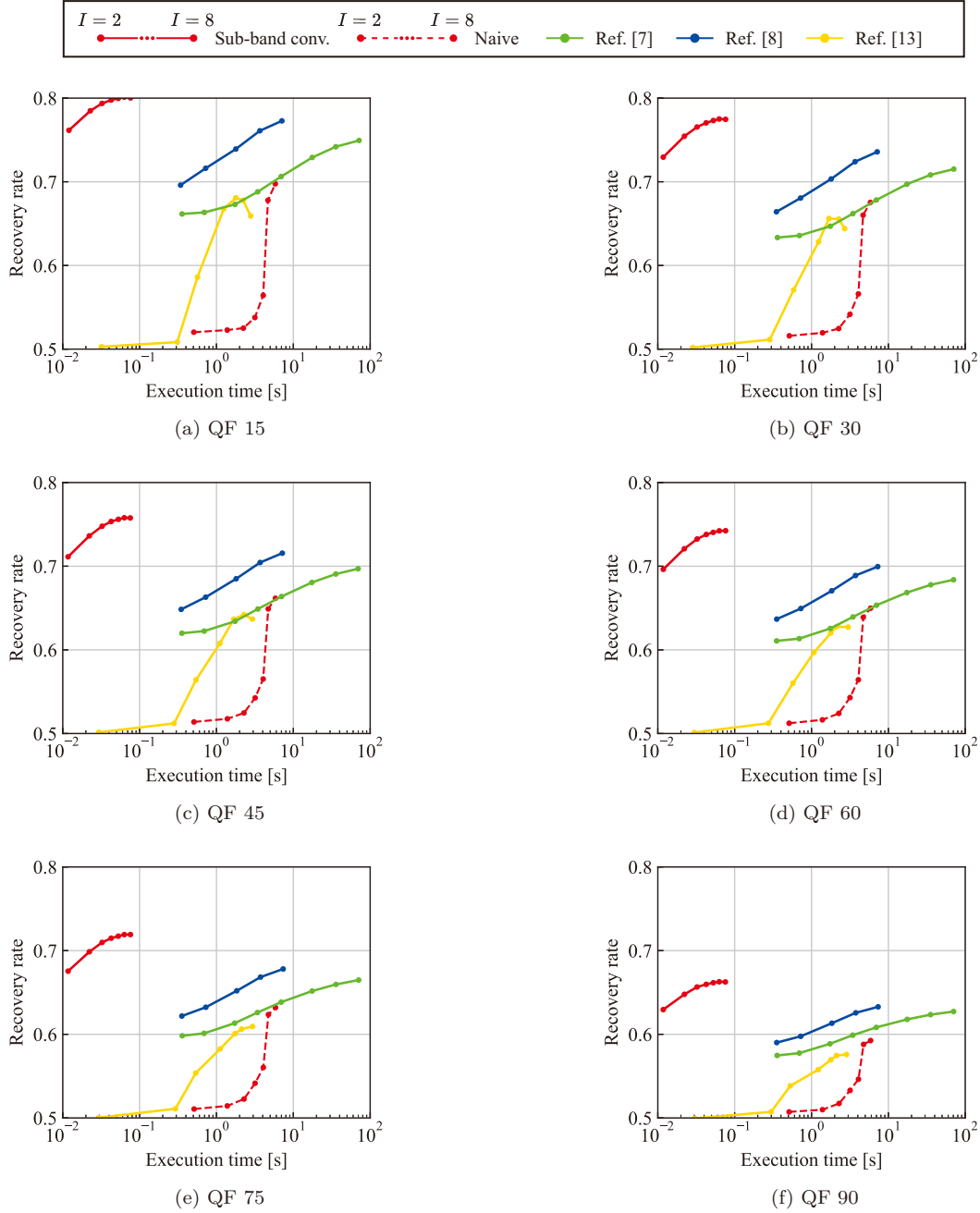


Fig. 5: Recovery rate against execution time.

$$S'_{x,y,z} = \begin{cases} +1 & \text{if } F_{x,y,z} \geq 1/2 \\ -1 & \text{otherwise} \end{cases} \quad (7)$$

Let $x(m) = m$, $y(n) = n$, and $z(u, v) = 8u + v$. A 4D representation of the retrieved signs can be computed as follows.

$$\mathbf{S}'_{u,v,m,n} = S'_{x(m),y(n),z(u,v)} \quad (8)$$

We encode the residuals between the true signs $\mathbf{S}$ and the retrieved ones $\mathbf{S}'$, as shown in Fig. 1.

4. Experimental results

4.1 Configuration

For training our CNN, we randomly cropped $K =$

6,056 images with a size of $W \times H = 512 \times 512$ from the CLIC 2020 dataset [33]. The bit-depth was 8 bits/pixel, and each pixel value was normalized within the range 0–1. We applied block-wise DCT to the images. We quantized the DCT coefficients using the quality factor (QF) of 75. We reshaped the 4D amplitudes and signs in accordance with (1) and (2); as a result, we obtained 3D tensors with a size of $W/8 \times H/8 \times 64 = 64 \times 64 \times 64$. We utilized the Adam optimizer with the learning rate 2×10^{-4} . The batch size and the number of epochs were 256 and 15,000, respectively. We varied the number of convolution layers I (Table 1) from 2 to 8. For convenience, we refer to the proposed method

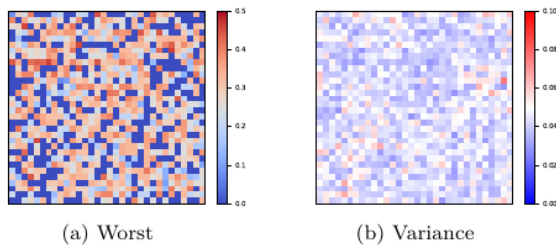


図6:各ブロックの回収率のヒートマップ。

畳み込み層の数(表1)を l_2 から l_8 まで変化させた。便宜上、上記の構成で学習した提案手法をサブバンド畳み込み手法と呼ぶ。

(1)と(2)のリシェーピングプロセスの有効性を評価するために、ナイーブ法と呼ばれる以下の代替アプローチを採用した。図4aに描かれたすべての振幅(または符号)の集まりを、寸法 $W \times H = 512 \times 512$ の平面として扱った。振幅平面から符号平面を再構成し、サブバンド畳み込み法でリシェーピング処理を無効にした。ネットワークアーキテクチャを平面に適応させるため、Conv-0の入力チャンネルとConv-1の出力チャンネルを l_1 に変更した。公平に比較するため、学習プロセスはサブバンド畳み込み法で採用したものと同一構成に従う。

評価のために、CLIC 2021データセット[33]から60画像をサンプリングしたが、これらは学習データセットには含まれていない。画像のDCT係数はQF15、30、45、60、75、90で量子化した。符号は我々の手法(サブバンド畳み込みとナイーブ)と以前の手法[7, 8, 13]を用いて検索した。後者は反復ごとのDCTとIDCTを必要とする反復手法である(セクション2.1で議論)。これまでの手法では、最大反復回数を l_1 から200に変更した。全ての実験は表2の計算機環境で行い、全ての手法は同じCPUで実行した。検索された符号の精度を定量化するために、以下の復元率を定義する:

$$\frac{\#(\text{correctly retrieved signs of AC coefficients})}{\#(\text{signs of AC coefficients})},$$

ここで、DC係数と有意でない係数の符号は除外される。また、符号検索のためのすべての手法の実行時間[s]を測定した。

4.2 結果

図5は、60枚のテスト画像中の全ブロックを平均した、実行時間に対する回復率を示している。本手法の7点は $l_i \in [2, 8]$ の値を変えた場合の結果を表し、従来手法の点は反復回数を変えた場合の結果を表す。

すべてのQFにおいて、サブバンド畳み込み法は、最短の実行時間を達成しながら、回復率の点で従来の方法を完全に上回った。具体的には、最大回復率の実行時間は平均 6.74×10^{-2} 秒、すなわち最先端手法[8]では7.24秒の0.93%であった。サブバンド畳み込み法はナイーブアプローチも上回ったことは注目に値するが、これはリシェーピングプロセスの有効性を示している。

備考1:各ブロックの回復率を報告し、その偏りを調査する。テスト画像から 256×256 領域を切り出し、そのDCT係数の振幅(QF 75)をCNN($l=8$)に入力した。各ブロックの回復率の平均を計算し、60画像に対する最悪値と分散を求めた。図6にその結果を示す。境界には、低い回復率と高い分散値のクラスターが存在することが観察される。

備考2:すべての手法において、QFが低下するにつれて回復率が増加する傾向がある。小さなQFの場合、有意なDCT係数は低周波数帯域に集中し、符号検索が容易になる。したがって、どの手法もQFが小さい場合に高い回収率を達成した。一方、QFが大きい場合、DCT係数はすべての周波数帯域に分布するため、符号検索の難易度が高くなり、すべての手法で低い回復率が得られた。

5. むすび

本研究では、ブロックワイズ離散コサイン変換(DCT)の符号検索問題、すなわち、DCT係数の振幅から符号を再構成することに取り組んだ。符号検索を高速化するために、DCTIDCTを反復しない高速な符号検索手法を提案し、二値分類機械学習に基づいて符号を検索する。まず、振幅と符号の3次元表現を導入し、次に3次元振幅の特徴を抽出するのに有利な畳み込みニューラルネットワークによって本手法を実装した。今後の課題としては、可変サイズブロックを利用した現在の最先端画像符号化規格への本手法の拡張を行う予定である。

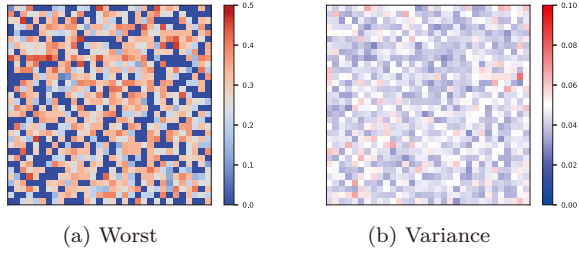


Fig. 6: Heat maps of recovery rates for each block.

trained on the above configuration as the *sub-band convolution method*.

To evaluate the effectiveness of the reshaping processes in (1) and (2), we employed the following alternative approach, referred to as the *naive method*. We treated the collection of all amplitudes (or signs) depicted in Fig. 4a as a plane with dimensions $W \times H = 512 \times 512$. We reconstructed the sign planes from the amplitude planes, where we nullified the reshaping process in the sub-band convolution method. To adapt the network architecture for planes, we changed the input channel of Conv-0 and the output channel of Conv- I to 1. For fair comparison, the training process follows the same configuration as that employed in the sub-band convolution method.

For evaluation, we sampled 60 images from the CLIC 2021 dataset [33], which were not included in the training dataset. The DCT coefficients of the images were quantized with QFs of 15, 30, 45, 60, 75, and 90. The signs were retrieved using our methods (sub-band convolution and naive) and previous methods [7, 8, 13]; the latter were iterative methods requiring per-iteration DCT and IDCT (discussed in Section 2.1). For the previous methods, we changed the maximum number of iterations from 1 to 200. All the experiments were conducted on the computing environment in Table 2; all the methods were executed using the same CPU.

To quantify the accuracy of the retrieved signs, we define the following recovery rate:

$$\frac{\#(\text{correctly retrieved signs of AC coefficients})}{\#(\text{signs of AC coefficients})},$$

where the signs of DC coefficients and non-significant coefficients are excluded. We also measured the execution time [s] of all the methods for sign retrieval.

4.2 Results

Figure 5 shows recovery rates against execution times, which were averaged over all blocks in the 60 test images. The seven points in our method represent re-

sults for different values of $I \in [2, 8]$, while the points in the previous methods represent results for different iteration counts. For all the QFs, the sub-band convolution method perfectly outperformed the previous methods in terms of the recovery rate while achieving the shortest execution time. Specifically, the execution time for the maximum recovery rate was 6.74×10^{-2} seconds on average, i.e., **0.93%** of 7.24 seconds for the state-of-the-art method [8]. It is worth noting that the sub-band convolution method also outperformed the naive approach, which demonstrates the effectiveness of the reshaping process.

Remark 1: We report recovery rates for each block to investigate their biases. We cropped 256×256 regions from the 60 test images and fed amplitudes of their DCT coefficients (QF 75) into our CNN ($I = 8$). We computed the average of recovery rates for each block and obtained the worst value and the variance over the 60 images. Figure 6 illustrates the results. We observe that there are clusters of low recovery rates and high variance values at the boundaries.

Remark 2: For all the methods, the recovery rates tend to increase as QF decreases. For small QFs, significant DCT coefficients are concentrated at the low frequency bands, which simplifies sign retrieval. Therefore, all the methods achieved high recovery rates at small QFs. In contrast, for large QFs, DCT coefficients are distributed across all the frequency bands, which increases the difficulty of sign retrieval; all the methods thus obtained low recovery rates.

5. Conclusion

In this work, we addressed a sign retrieval problem for the block-wise discrete cosine transformation (DCT): reconstruction of the signs of DCT coefficients from their amplitudes. To accelerate sign retrieval, we proposed a fast sign retrieval method without the DCT-IDCT iteration, where the signs are retrieved based on binary classification machine learning. We first introduced 3D representations of the amplitudes and signs and then implemented our method by convolutional neural networks, which are advantageous for extracting features in the 3D amplitudes. Our future work will include the extension of our method to the current state-of-the-art image coding standards, which utilize variable-size blocks. We also need to consider other transformations than the DCT, e.g., the discrete sine transform and low-frequency non-separable secondary

この研究方向は画像符号化の効率を高める可能性を秘めていると我々は考えている。

References

- 1) J. Sole, R. Joshi, N. Nguyen, T. Ji, M. Karczewicz, G. Clare, F. Henry, and A. Duenas, "Transform coefficient coding in HEVC," *IEEE Transactions on Circuits and Systems for Video Technology*, vol. 22, no. 12, pp. 1765–1777, 2012.
- 2) C. Tu and T. D. Tran, "Context-based entropy coding of block transform coefficients for image compression," *IEEE Transactions on Image Processing*, vol. 11, no. 11, pp. 1271–1283, 2002.
- 3) N. N. Ponomarenko, A. V. Bazhyna, and K. O. Egiazarian, "Prediction of signs of DCT coefficients in block-based lossy image compression," *SPIE 6497, Image Processing: Algorithms and Systems V*, 2007.
- 4) J. Koyama, A. Yamori, K. Kazui, S. Shimada, and A. Nakagawa, "Coefficient sign bit compression in video coding," *Picture Coding Symposium*, 2012.
- 5) O. Miroshnichenko, M. Ponomarenko, V. Lukin, and K. Egiazarian, "Compression of signs of DCT coefficients for additional lossless compression of JPEG images," *International Symposium on Electronic Imaging: Image Processing: Algorithms and Systems XVI*, 2018.
- 6) N. Ahmed, T. Natarajan, and K. R. Rao, "Discrete cosine transform," *IEEE Transactions on Computers*, vol. C-23, no. 1, pp. 90–93, 1974.
- 7) C. Tsutake, K. Takahashi, and T. Fujii, "An efficient compression method for sign information of DCT coefficients via sign retrieval," *IEEE International Conference on Image Processing*, 2021.
- 8) K. Suzuki, C. Tsutake, K. Takahashi, and T. Fujii, "Compressing sign information in DCT-based image coding via deep sign retrieval," *ITE Transactions on Media Technology and Applications*, vol. 12, no. 1, pp. 110–122, 2024.
- 9) D. A. Huffman, "A method for the construction of minimum-redundancy codes," *Resonance*, vol. 11, pp. 91–99, 1952.
- 10) R. Rice and J. Plaunt, "Adaptive variable-length coding for efficient compression of spacecraft television data," *IEEE Trans. Communication Technology*, vol. 19, no. 6, pp. 889–897, 1971.
- 11) G. N. N. Martin, "Range encoding: an algorithm for removing redundancy from a digitized message," *Video & Data Recording Conference*, 1979.
- 12) F. Golchin and K. Paliwal, "A context-based adaptive predictor for use in lossless image coding," *IEEE Region 10 International Conference*, 1997.
- 13) R. Lin, S. Liu, J. Jiang, S. Li, C. Li, and C.-C. J. Kuo, "Recovering sign bits of DCT coefficients in digital images as an optimization problem," *Journal of Visual Communication and Image Representation*, vol. 98, no. 103992, 2024.
- 14) N. D. Sidiropoulos, P. A. Karakasis, and A. Konar, "Minimizing low-rank models of high-order tensors: hardness, span, tight relaxation, and applications," *IEEE Transactions on Signal Processing*, vol. 72, pp. 129–142, 2023.
- 15) G. K. Wallace, "The JPEG still picture compression standard," *Communications of the ACM*, vol. 34, no. 4, pp. 30–44, 1991.
- 16) G. J. Sullivan, J.-R. Ohm, W.-J. Han, and T. Wiegand, "Overview of the high efficiency video coding (HEVC) standard," *IEEE Transactions on Circuits and Systems for Video Technology*, vol. 22, no. 12, pp. 1649–1668, 2012.
- 17) B. Bross, Y.-K. Wang, Y. Ye, S. Liu, J. Chen, G. J. Sullivan, and J.-R. Ohm, "Overview of the versatile video coding (VVC) standard and its applications," *IEEE Transactions on Circuits and Systems for Video Technology*, vol. 31, no. 10, pp. 3736–3764, 2021.
- 18) S. Bahmani and J. Romberg, "A flexible convex relaxation for phase retrieval," *Electronic Journal of Statistics*, vol. 11, no. 2, pp. 5254–5281, 2017.
- 19) T. Goldstein and C. Studer, "PhaseMax: convex phase retrieval via basis pursuit," *IEEE Transactions on Information Theory*, vol. 64, no. 4, pp. 2675–2689, 2018.
- 20) R. W. Gerchberg and W. O. Saxton, "A practical algorithm for the determination of phase from image and diffraction plane pictures," *Optik*, vol. 35, pp. 237–246, 1972.
- 21) J. R. Fienup, "Phase retrieval algorithms: a comparison," *Applied Optics*, vol. 21, no. 15, pp. 2758–2769, 1982.

- 22) M. Frank and P. Wolfe, "An algorithm for quadratic programming," *Naval Research Logistics Quarterly*, vol. 3, no. 1–2, pp. 95–110, 1956.
- 23) F. Henry and G. Clare, "Residual coefficient sign prediction," *JVET-D0031*, 2016.
- 24) A. Nakagawa, J. Koyama, K. Kazui, and J. Katto, "Sign information predictive coding of transform coefficients in video compression," *IEICE Transactions on Information and Systems*, vol. J100-D, no. 9, pp. 819–830, 2017.
- 25) A. Filippov, V. Ruftskiy, A. Karabutov, and J. Chen, "Residual sign prediction in transform domain for next-generation video coding," *APSIPA Transactions on Signal and Information Processing*, vol. 8, 2019.
- 26) G. Clare, F. Henry, and J. Jung, "Sign data hiding," *JCTVC-G271*, 2011.
- 27) G. Bjontegaard, "Calculation of average PSNR differences between RD curves," *VCEG-M33*, 2001.
- 28) J. Wang, X. Yu, D. He, F. Henry, and G. Clare, "Multiple sign bits hiding for high efficiency video coding," *IEEE International Conference on Visual Communications and Image Processing*, 2012.
- 29) X. Zhang, O. C. Au, C. Pang, W. Dai, and Y. Guo, "Additional sign bit hiding of transform coefficients in HEVC," *IEEE International Conference on Multimedia and Expo Workshops*, 2013.
- 30) R. Song, Y. Yuan, Y. Li, and Y. Wang, "Extra sign bit hiding algorithm based on recovery of transform coefficients," *Circuits, Systems, and Signal Processing*, vol. 37, no. 9, pp. 4128–4135, 2018.
- 31) J. M. Shapiro, "Embedded image coding using zerotrees of wavelet coefficients," *IEEE Transactions on Signal Processing*, vol. 41, no. 12, pp. 3445–3462, 1993.
- 32) A. Said and W. A. Pearlman, "A new, fast, and efficient image codec based on set partitioning in hierarchical trees," *IEEE Transactions on Circuits and Systems for Video Technology*, vol. 6, no. 3, pp. 243–250, 1996.
- 33) Challenge on Learned Image Compression, <https://compression.cc>.
- 34) M. Koo, M. Salehifar, J. Lim, and S.-H. Kim, "Low frequency non-separable transform (LFNST)," *Picture Coding Symposium*, 2019.



Fuma Ito received a B.E. degree in information and communication engineering from Nagoya University, Japan, in 2023. His primary research interest is image coding and deep learning.



Chihiro Tsutake received B.E., M.E., and Ph.D. degrees from the University of Fukui, Japan, in 2015, 2017, and 2020. Since 2020, he has been with the Graduate School of Engineering, Nagoya University, as an assistant professor. His research interests include multidimensional signal processing and optics.



Keita Takahashi received B.E., M.S., and Ph.D. degrees in information and communication engineering from the University of Tokyo, Japan, in 2001, 2003, and 2006. He was a project assistant professor at the University of Tokyo from 2006 to 2011 and an assistant professor at the University of Electro-Communications from 2011 to 2013. Since 2013, he has been with the Graduate School of Engineering, Nagoya University, as an associate professor. His research interests include image processing, computational photography, and 3D displays. He is a member of the IEEE Computer Society and Signal Processing Society, the Information Processing Society of Japan, and the Institute of Image Information and Television Engineers of Japan.

【機械翻訳コンテンツの著作権について】

当サイトに掲載されている原著論文の著作権は映像情報メディア学会に帰属します。原著論文および翻訳論文の無断使用は禁止します。
引用の際には、必ず原著論文の書誌情報をご記載ください。詳細は映像情報メディア学会著作権規定をご覧ください。

transform [34]. We believe that this research direction has the potential to enhance the efficiency of image coding.

References

- 1) J. Sole, R. Joshi, N. Nguyen, T. Ji, M. Karczewicz, G. Clare, F. Henry, and A. Duenas, "Transform coefficient coding in HEVC," *IEEE Transactions on Circuits and Systems for Video Technology*, vol. 22, no. 12, pp. 1765–1777, 2012.
- 2) C. Tu and T. D. Tran, "Context-based entropy coding of block transform coefficients for image compression," *IEEE Transactions on Image Processing*, vol. 11, no. 11, pp. 1271–1283, 2002.
- 3) N. N. Ponomarenko, A. V. Bazhyna, and K. O. Egiazarian, "Prediction of signs of DCT coefficients in block-based lossy image compression," *SPIE 6497, Image Processing: Algorithms and Systems V*, 2007.
- 4) J. Koyama, A. Yamori, K. Kazui, S. Shimada, and A. Nakagawa, "Coefficient sign bit compression in video coding," *Picture Coding Symposium*, 2012.
- 5) O. Miroshnichenko, M. Ponomarenko, V. Lukin, and K. Egiazarians, "Compression of signs of DCT coefficients for additional lossless compression of JPEG images," *International Symposium on Electronic Imaging: Image Processing: Algorithms and Systems XVI*, 2018.
- 6) N. Ahmed, T. Natarajan, and K. R. Rao, "Discrete cosine transform," *IEEE Transactions on Computers*, vol. C-23, no. 1, pp. 90–93, 1974.
- 7) C. Tsutake, K. Takahashi, and T. Fujii, "An efficient compression method for sign information of DCT coefficients via sign retrieval," *IEEE International Conference on Image Processing*, 2021.
- 8) K. Suzuki, C. Tsutake, K. Takahashi, and T. Fujii, "Compressing sign information in DCT-based image coding via deep sign retrieval," *ITE Transactions on Media Technology and Applications*, vol. 12, no. 1, pp. 110–122, 2024.
- 9) D. A. Huffman, "A method for the construction of minimum-redundancy codes," *Resonance*, vol. 11, pp. 91–99, 1952.
- 10) R. Rice and J. Plaunt, "Adaptive variable-length coding for efficient compression of spacecraft television data," *IEEE Trans. Communication Technology*, vol. 19, no. 6, pp. 889–897, 1971.
- 11) G. N. N. Martin, "Range encoding: an algorithm for removing redundancy from a digitized message," *Video & Data Recording Conference*, 1979.
- 12) F. Golchin and K. Paliwal, "A context-based adaptive predictor for use in lossless image coding," *IEEE Region 10 International Conference*, 1997.
- 13) R. Lin, S. Liu, J. Jiang, S. Li, C. Li, and C.-C. J. Kuo, "Recovering sign bits of DCT coefficients in digital images as an optimization problem," *Journal of Visual Communication and Image Representation*, vol. 98, no. 103992, 2024.
- 14) N. D. Sidiropoulos, P. A. Karakasis, and A. Konar, "Minimizing low-rank models of high-order tensors: hardness, span, tight relaxation, and applications," *IEEE Transactions on Signal Processing*, vol. 72, pp. 129–142, 2023.
- 15) G. K. Wallace, "The JPEG still picture compression standard," *Communications of the ACM*, vol. 34, no. 4, pp. 30–44, 1991.
- 16) G. J. Sullivan, J.-R. Ohm, W.-J. Han, and T. Wiegand, "Overview of the high efficiency video coding (HEVC) standard," *IEEE Transactions on Circuits and Systems for Video Technology*, vol. 22, no. 12, pp. 1649–1668, 2012.
- 17) B. Bross, Y.-K. Wang, Y. Ye, S. Liu, J. Chen, G. J. Sullivan, and J.-R. Ohm, "Overview of the versatile video coding (VVC) standard and its applications," *IEEE Transactions on Circuits and Systems for Video Technology*, vol. 31, no. 10, pp. 3736–3764, 2021.
- 18) S. Bahmani and J. Romberg, "A flexible convex relaxation for phase retrieval," *Electronic Journal of Statistics*, vol. 11, no. 2, pp. 5254–5281, 2017.
- 19) T. Goldstein and C. Studer, "PhaseMax: convex phase retrieval via basis pursuit," *IEEE Transactions on Information Theory*, vol. 64, no. 4, pp. 2675–2689, 2018.
- 20) R. W. Gerchberg and W. O. Saxton, "A practical algorithm for the determination of phase from image and diffraction plane pictures," *Optik*, vol. 35, pp. 237–246, 1972.
- 21) J. R. Fienup, "Phase retrieval algorithms: a comparison," *Applied Optics*, vol. 21, no. 15, pp. 2758–2769, 1982.
- 22) M. Frank and P. Wolfe, "An algorithm for quadratic programming," *Naval Research Logistics Quarterly*, vol. 3, no. 1–2, pp. 95–110, 1956.
- 23) F. Henry and G. Clare, "Residual coefficient sign prediction," *JVET-D0031*, 2016.
- 24) A. Nakagawa, J. Koyama, K. Kazui, and J. Katto, "Sign information predictive coding of transform coefficients in video compression," *IEICE Transactions on Information and Systems*, vol. J100-D, no. 9, pp. 819–830, 2017.
- 25) A. Filippov, V. Ruftskiy, A. Karabutov, and J. Chen, "Residual sign prediction in transform domain for next-generation video coding," *APSIPA Transactions on Signal and Information Processing*, vol. 8, 2019.
- 26) G. Clare, F. Henry, and J. Jung, "Sign data hiding," *JCTVC-G271*, 2011.
- 27) G. Bjontegaard, "Calculation of average PSNR differences between RD curves," *VCEG-M33*, 2001.
- 28) J. Wang, X. Yu, D. He, F. Henry, and G. Clare, "Multiple sign bits hiding for high efficiency video coding," *IEEE International Conference on Visual Communications and Image Processing*, 2012.
- 29) X. Zhang, O. C. Au, C. Pang, W. Dai, and Y. Guo, "Additional sign bit hiding of transform coefficients in HEVC," *IEEE International Conference on Multimedia and Expo Workshops*, 2013.
- 30) R. Song, Y. Yuan, Y. Li, and Y. Wang, "Extra sign bit hiding algorithm based on recovery of transform coefficients," *Circuits, Systems, and Signal Processing*, vol. 37, no. 9, pp. 4128–4135, 2018.
- 31) J. M. Shapiro, "Embedded image coding using zerotrees of wavelet coefficients," *IEEE Transactions on Signal Processing*, vol. 41, no. 12, pp. 3445–3462, 1993.
- 32) A. Said and W. A. Pearlman, "A new, fast, and efficient image codec based on set partitioning in hierarchical trees," *IEEE Transactions on Circuits and Systems for Video Technology*, vol. 6, no. 3, pp. 243–250, 1996.
- 33) Challenge on Learned Image Compression, <https://compression.cc>.
- 34) M. Koo, M. Salehifar, J. Lim, and S.-H. Kim, "Low frequency non-separable transform (LFNST)," *Picture Coding Symposium*, 2019.



Fuma Ito received a B.E. degree in information and communication engineering from Nagoya University, Japan, in 2023. His primary research interest is image coding and deep learning.



Chihiro Tsutake received B.E., M.E., and Ph.D. degrees from the University of Fukui, Japan, in 2015, 2017, and 2020. Since 2020, he has been with the Graduate School of Engineering, Nagoya University, as an assistant professor. His research interests include multidimensional signal processing and optics.



Keita Takahashi received B.E., M.S., and Ph.D. degrees in information and communication engineering from the University of Tokyo, Japan, in 2001, 2003, and 2006. He was a project assistant professor at the University of Tokyo from 2006 to 2011 and an assistant professor at the University of Electro-Communications from 2011 to 2013. Since 2013, he has been with the Graduate School of Engineering, Nagoya University, as an associate professor. His research interests include image processing, computational photography, and 3D displays. He is a member of the IEEE Computer Society and Signal Processing Society, the Information Processing Society of Japan, and the Institute of Image Information and Television Engineers of Japan.

こちらは英語原著論文 (J-STAGE掲載) の機械翻訳版です。
次ページが原著論文で、翻訳版と交互に展開されます。機械翻訳のため、誤字や誤訳、翻訳が未反映の部分が含まれている可能性があります。
引用の際には、必ず原著論文の書誌情報をご記載ください。



Toshiaki Fujii received B.E., M.E., and Dr.E. degrees in electrical engineering from the University of Tokyo, Japan, in 1990, 1992, and 1995. In 1995, he joined the Graduate School of Engineering, Nagoya University, where he is currently a professor. From 2008 to 2010, he was with the Graduate School of Science and Engineering, Tokyo Institute of Technology. From 2019 to 2021, he also served as a senior science and technology policy fellow of the Cabinet Office, Government of Japan. His current research interests include multidimensional signal processing, multi-camera systems, multi-view video coding and transmission, free-viewpoint video, and their applications. He is a member of the IEEE Signal Processing Society, the ISO/IEC JTC1/SC29/WG4, WG1 (MPEG-I Visual, JPEG) standardization committee of Japan, the Information Processing Society of Japan, and the Institute of Image Information and Television Engineers of Japan.

【機械翻訳コンテンツの著作権について】

当サイトに掲載されている原著論文の著作権は映像情報メディア学会に帰属します。原著論文および翻訳論文の無断使用は禁止します。
引用の際には、必ず原著論文の書誌情報をご記載ください。詳細は映像情報メディア学会著作権規定をご覧ください。



Toshiaki Fujii received B.E., M.E., and Dr.E. degrees in electrical engineering from the University of Tokyo, Japan, in 1990, 1992, and 1995. In 1995, he joined the Graduate School of Engineering, Nagoya University, where he is currently a professor. From 2008 to 2010, he was with the Graduate School of Science and Engineering, Tokyo Institute of Technology. From 2019 to 2021, he also served as a senior science and technology policy fellow of the Cabinet Office, Government of Japan. His current research interests include multidimensional signal processing, multi-camera systems, multi-view video coding and transmission, free-viewpoint video, and their applications. He is a member of the IEEE Signal Processing Society, the ISO/IEC JTC1/SC29/WG4, WG1 (MPEG-I Visual, JPEG) standardization committee of Japan, the Information Processing Society of Japan, and the Institute of Image Information and Television Engineers of Japan.
