

**Design of Subjectively Adapted Quantizers
for Two and Three Dimensional Transform Coding
of Image Sequences**

by

**Yacine Rahmouni
B.Eng. (Electrical)**

A thesis submitted to the Faculty of Graduate Studies and Research
in partial fulfillment of the requirements for the degree of
Master of Engineering

Department of Electrical Engineering
McGill University
Montréal, Canada

© Y. Rahmouni, November 1986

Permission has been granted to the National Library of Canada to microfilm this thesis and to lend or sell copies of the film.

The author (copyright owner) has reserved other publication rights, and neither the thesis nor extensive extracts from it may be printed or otherwise reproduced without his/her written permission.

L'autorisation a été accordée à la Bibliothèque nationale du Canada de microfilmer cette thèse et de prêter ou de vendre des exemplaires du film.

L'auteur (titulaire du droit d'auteur) se réserve les autres droits de publication; ni la thèse ni de longs extraits de celle-ci ne doivent être imprimés ou autrement reproduits sans son autorisation écrite.

ISBN 0-315-44378-2

Design of Quantizers for Transform Coding of Image Sequences

Abstract

This thesis addresses the problems of quantization and coding of discrete cosine transform block coefficients. Block quantizers subjectively adapted to the human visual system are investigated, followed by a study of block coders using a combination of run length and Huffman coding techniques. Two and three-dimensional processing are investigated with, however, a strong emphasis on two dimensional processing. Quality of the coded images is given a priority over bandwidth compression. A coder operating at the threshold of impairment perceptibility is sought for natural images and man made objects making the bulk of television programs. Average bit rates below 1.5 bits per pel are achieved with two-dimensional discrete cosine transform and 1.0 bit per pel for three-dimensional processing.

Sommaire

Cette thèse est consacrée à l'étude de quantificateurs et codeurs de blocs pour la transformée en cosinus. Les quantificateurs de blocs sont adaptés spécifiquement aux caractéristiques du système visuel humain. Les codeurs de blocs subséquents font appel aux techniques de codage par plage et de codage de Huffman. Le traitement bidimensionnel et tridimensionnel sont étudiés avec toutefois une nette emphase sur le traitement bidimensionnel. La qualité de l'image a ici préséance sur la réduction du débit. La recherche porte plus spécifiquement sur un codeur opérant au seuil de perception de distorsion pour des images naturelles ou d'objets courant formant le gros des programmes de télévision. Un débit moyen inférieur à 1.5 bit par pixel est atteint en codage bidimensionnel et 1.0 bit par pixel en traitement tridimensionnel.

Acknowledgements

I would like to sincerely thank my thesis supervisor, Dr Eric Dubois, for his guidance and encouragements throughout this project.

I would like to thank my fellow students and staff members, both at McGill University and INRS-Télécommunications for participating in the subjective tests.

I would like to thank INRS-Télécommunications and Bell Northern Research for the use of their facilities.

Table of Contents

<i>Abstract</i>	ii
<i>Sommaire</i>	iii
<i>Acknowledgements</i>	iv
<i>Table of Contents</i>	v
<i>List of Figures</i>	vii
<i>List of Tables</i>	ix
Chapter 1 Introduction	1
Chapter 2 Theoretical Context	6
2.1 Digital Image Communication Systems	7
2.1.1 Digital Images	8
2.1.2 Statistical Characterization of Digital Images	13
2.1.3 Transmission of Digital Images	15
2.1.4 Survey of Some Coding Methods	17
2.2 Transform Coding	18
2.2.1 The Karhunen Loeve transform (KLT)	20
2.2.2 The Discrete Cosine Transform (DCT)	21
2.2.3 Quantization of the Transform Coefficients	23
2.2.4 Source Coding of the Quantized Coefficients	24
2.3 The Human Visual System (HVS)	28
2.3.1 General Description	28
2.3.2 Properties of the Human Visual System	30
2.3.3 The Human Visual System and the DCT	33
Chapter 3 Design of Quantizers for the DCT	36
3.1 Transform Coding System	37
3.1.1 Block Quantizer	37
3.1.2 Frame or Field Processing and Block Size Considerations	39
3.2 Experimental Setting	41
3.2.1 Test Images	41
3.2.2 Some Statistics on the Real Transform Coefficients	42
3.3 Thresholds of Perception for Still Images (2D-DCT)	43
3.4 Thresholds for Time Varying Images(2D-DCT)	50

3.4.1	Extension of the Results Obtained with Still Images	50
3.4.2	Subjective Tests	52
3.4.3	Discussion	55
3.5	Non Uniform Quantization	56
3.6	Extension to 3D-DCT	57
3.6.1	Advantage over Two-Dimensional Processing	57
3.6.2	Experiments	58
Chapter 4	Coding of the DCT Coefficients	60
4.1	A Brief Review of Coding Theory, Huffman and Run-Length Coding	61
4.1.1	Huffman Coding	62
4.1.2	Run-Length Coding	62
4.1.3	Discussion on the Estimation of the Entropy of Quantized Blocks	63
4.2	Some Observations on the Quantized Coefficients	63
4.3	Dual Huffman and Run-Length Coders	68
4.4	Results	68
4.5	Post Filtering	70
4.6	Extension to 3D DCT	72
Chapter 5	Conclusion	74
Appendices		77
Appendix A.	Test Images	78
Appendix B.	Subjective Tests, pictures and contour plots	81
Appendix C.	Subjective Tests, Procedure	87
Appendix D.	Statistics on the transform coefficients	91
Appendix E.	Entropy and Bit Rates for Individual Coefficients	94
References		98

List of Figures

2.1	An Image Communication System	7
2.2	Sampling of a unidimensional signal	9
2.3	Line Interlaced TV Frame	10
2.4	Projection views	10
2.5	Conversion to a Frame Sampling Lattice	11
2.6	Conversion to a Field Sampling Lattice	11
2.7	Digitalization of an Image	12
2.8	Digital Image Communication System	15
2.9	Image Communication by Transform Coding	19
2.10	Block Quantization of the Transform Coefficients	23
2.11	Different Quantizers	24
2.12	Zigzag scanning pattern for the transform coefficients. The upper left corner represents $F(0,0)$ and is the starting point.	27
2.13	The Human Visual System	29
2.14	Contrast Sensitivity Experiment and Results	30
2.15	Mach Bands	31
2.16	Response to sine gratings (From Hall and Hall[1])	32
2.17	A Simple Model for Monochrome Vision (Adapted from [2])	33
2.18	Results of Lohscheller [3]	34
3.1	An example of quantizer, $threshold = 0.75$, $S^{(1)} = 1$, $slope = 1.2$, $saturation = 2.0$ (symmetrical negative side).	38
3.2	View of the monitor and the display window. Here $PH = 0.35m$	42
3.3	Probability density function of $F(2,2,0)$. The sequence used to generate it is LONG-SEQUENCE (described in appendix A) processed with a $16 \times 16 \times 1$ DCT	43
3.4	Examples of one-dimensional distributions, S_0 and S_{15} are fixed.	48
3.5	Contour plot of a block quantizer step size distribution. The values of the parameters are: $S_{0,0,0} = 0.1$, $S_{15,15,0} = 3.0$ and $c = 1.0$	49

3.6	Entropy vs c for the sequences ALBERT, EIA, OSCAR, QUILT, SHORT-SEQUENCE	51
3.7	Subjective rating vs c	54
3.8	Subjective rating vs NSNR	56
4.1	Probability distribution functions for $F(2,2,0)$ quantized, with, and without, the zero value. Different vertical scales used.	65
4.2	Probability distribution functions for $F(7,7,0)$ quantized, with, and without, the zero value. Different vertical scales used.	65
B.1	Level of distortion d_2 corresponding to $c = 2.0$	81
B.2	Level of distortion d_3 corresponding to $c = 1.5$	81
B.3	Level of distortion d_4 corresponding to $c = 1.0$	82
B.4	Level of distortion d_5 corresponding to $c = 0.75$	82
B.5	Level of distortion d_6 corresponding to $c = 0.5$	83
B.6	Level of distortion d_7 corresponding to $c = 0.25$	83

List of Tables

2.1	Subjective Distorsion Scales.....	16
2.2	Typical bit assignment for transform zonal coding in 16×16 pel blocks at a rate of 1.5 bits per pel. (From [2], page 675)	26
3.1	Example 1 of the step size distribution of a block quantizer. Spatially uniform step size of 0.1	44
3.2	Example 2 of the step size distribution of a block quantizer. Intermediate experiment.	45
3.3	Example 3 of the step size distribution of a block quantizer. One of the final experiments.	46
3.4	Table of a block quantizer step size distribution. The values of the parameters are: $S_{0,0,0} = 0.1$, $S_{15,15,0} = 3.0$ and $c = 1.0$	50
3.5	Subjective opinion rating ..	54
4.1	A typical quantized transform coefficient bloc. Taken from ALBERT in the region of his right eye. Parameter $c = 1.0$	64
4.2	Bit rates obtained in <i>bits per pels</i> for different coders. The individual coders are 'integrated' and use a zigzag scanning pattern.	69
4.3	Influence of skipping isolated (± 1) coefficients on the bit rate. Comparison of scanning patterns for the coders with post-filtering.	71
4.4	Comparison of 2D coding and 3D coding for images of approximately same quality.	72
C.1	First session order of presentation	87
C.2	Second session order of presentation	88
D.1	Standard deviation of 2D-DCT coefficients, the test image is LONG-SEQUENCE.	91
D.2	Standard deviation of 3D-DCT coefficients, frame 0, the test image is LONG-SEQUENCE.	92
D.3	Standard deviation of 3D-DCT coefficients, frame 1, the test image is LONG-SEQUENCE ..	92
D.4	Standard deviation of 3D-DCT coefficients, frame 2, the test image is LONG-SEQUENCE.	93
D.5	Standard deviation of 3D-DCT coefficients, frame 3, the test image is LONG-SEQUENCE.	93

E.1	Entropy table for the coefficients quantized with the parameter $c = 1.0$ (2D DCT) and generated from the image LONG-SEQUENCE.	94
E.2	Average bit rate for the coefficients quantized with the parameter $c = 1.0$ (2D DCT) and generated from the image LONG-SEQUENCE.	95
E.3	Average bit rate per coefficients, 3D DCT, frame 0, and generated from the image LONG-SEQUENCE.	95
E.4	Average bit rate per coefficients, 3D DCT, frame 1, and generated from the image LONG-SEQUENCE.	96
E.5	Average bit rate per coefficients, 3D DCT, frame 2, and generated from the image LONG-SEQUENCE.	96
E.6	Average bit rate per coefficients, 3D DCT, frame 3, and generated from the image LONG-SEQUENCE.	97

Chapter 1

Introduction

Digital transmission and storage of images has gained considerable importance in the past decade. Applications are varied: High Definition TV, teleconferencing, transmission of pictures from observation satellites or over telephone lines, storage of pictures in computer databases and many more. The reasons for such an interest lie in the flexibility and ruggedness of digital transmission and processing as well as consistent picture quality.

Three steps are involved in the process of transmitting images by digital means: First, pictorial data is produced by sampling in space and time and then quantizing in brightness (or color) an analog and continuous scene. An appropriate sampling structure is chosen to minimize the number of pels required and avoid aliasing. For the majority of applications, however, the sampling structure is dictated by the available hardware or compatibility constraints, as for line interlace in television. Secondly, the digital image is coded, transmitted or stored, then received or retrieved, and decoded. Thirdly, interpolation to a continuous and analog form is accomplished by projecting the image on a display and then low pass filtering by the eye and brain of the viewer. The choice of the sampling structure in the first step, should take into account specific properties of the Human Visual System (HVS). In the second step, coding can be further separated into source coding and channel

coding. Source coding removes the redundant information in the source, the digital image, and produces a bit stream. Channel coding converts this bit stream to a format suitable for transmission over the channel, be it an optic fiber, a telephone line, a microwave link etc... It may also add some redundancy for error-control. From now on, coding will mean source coding.

Several methods are presently used and many are being investigated. The simplest of all, Pulse Code Modulation (PCM), consists of fixed length coding. The pictorial data is directly channel coded. The tradeoff for such simplicity is bit rate or equivalently, bandwidth. For example a monochrome, 640×525 pels/frame and 30 frames/second digital TV where the luminance is quantized with 8 bits would require a bit rate of 80.6 Mbps. For color, it would be even more. Transmitting at such high rates may not be economically viable depending on the applications. Luckily, the pictorial data is not totally random since the color of adjacent pels is highly correlated in space and time. In the framework of information and coding theory, the bit rate can be reduced down to the entropy, allowing a saving of about 50% quite easily. However, the order of magnitude is still the same and going below the entropy would result in the alteration of the pictorial data. But does such an alteration affect what the viewer sees? It depends on the type and degree of alteration. The Human Visual System is a complex mechanism that is not very well understood. Psychovisual and neurophysiological experiments have led to some models of the HVS that closely describe certain phenomena but no general model yet exists. Studies have shown that huge bandwidth compression can be achieved by properly taking advantage of the properties of the HVS and transmitting only what is actually perceived. A compression ratio of 100:1 with respect to PCM is envisioned, the problem is in sorting out the relevant information.

Image compression techniques vary greatly in complexity, quality, and efficiency. Simple ones such as DPCM are readily available for real time processing and achieve

a compression ratio of up to 2:1 easily. More elaborate techniques achieve better compression ratios but require complex and expensive equipment. The quality criterion plays an important role in choosing a technique for a particular use. While in broadcast television distortion can not be tolerated, teleconferencing requirements are less stringent.

The work presented here focuses on a particular transform coding technique, namely the Discrete Cosine Transform (DCT) applied to monochrome imagery. The pictorial data, consisting of successive frames of $H \times V$ pels is divided into $N \times M \times K$ blocks of pels. For a still picture, K is equal to one. The DCT is applied to the blocks yielding $N \times M \times K$ real, less correlated coefficients, with most of the 'energy' concentrated into the lower 'frequency' coefficients.[†] The coefficients are then quantized and coded for transmission. Two properties of the coefficients are exploitable for bandwidth reduction. Firstly, due to the energy compaction, many coefficients are often zero or have a small amplitude and may be efficiently coded. Secondly, coefficients have different visual significance depending on their position within the transformed block and need not be quantized with the same quantizer. There are basically two issues in DCT coding: the proper quantization of the block coefficients and their subsequent coding. The problem of quantization can be further separated in two. The selection of the best type of quantizer to use, and the assignment of a particular quantizer to each coefficient. They will be referred to as quantizer type and quantizer distribution. Both are closely related to properties of the HVS.

The goals in this thesis are to determine a quantizer distribution for the DCT that is well adapted to the HVS, and find an efficient coding technique.

Previous work in this area is extensive. Notably Ericson [4], Griswold [5], Man-

[†] Rigorously, the terms in quotes are valid only for the Fourier transform where energy and frequency are defined, here they stand for the equivalent terms for the DCT.

nos and Sakrison [6] have used specific models of the HVS to generate a weighting function. The coefficients are multiplied by the function and then quantized using a unique quantizer. The weighting function is thus equivalent to the quantizers distribution. The results are dependent on the accuracy and completeness of the model. Lohscheller [3] evaluated the visual effect of each coefficient and then generated a weighting function according to visual thresholds. The approach taken in this thesis is original in the sense that no model of the HVS is assumed. Visual experiments and subjective tests are used extensively to study the quantizer distributions that are best suited to the DCT. The quantizer distributions are then modeled using parametrical functions. Both three-dimensional and two-dimensional DCT are investigated with a strong emphasis on the latter. Uniform quantizers were mostly used. A particular type of non uniform quantizer was also investigated.

Concerning the subsequent block coding, the most recent contributions are by Chen and Pratt [7] who used a combination of Huffman and run length coding, and by Saghri and Tescher [8] who used the concept of chain coding to code clusters of zero valued coefficients within a block. Variations on the method of Chen and Pratt are investigated in this thesis.

Chapter 2 presents the concepts of digital images and an introduction to image communication systems. A short survey of some coding method is then followed by a more detailed presentation of transform coding. Since the human viewer is the last element in the chain, a review of the relevant properties of the human visual system concludes the chapter.

Chapter 3 presents the work on the design of block quantizers which are subjectively adapted to the human visual system. A function that automatically generates a block quantizer from a given set of parameters was obtained. This function allows the quality of reproduced images to be easily varied. A subjective test with a group

of viewers was used to formally assess the performance of this block quantization method.

In chapter 4, efficient block coding techniques using Huffman and run length coding are investigated for specific use on the DCT quantized blocks. Average bit rates of 1.5 bit per pels for two-dimensional coding and 1.0 bits per pels for three-dimensional processing were attained for images rated at around the threshold of perceptibility.

Chapter 2

Theoretical Context

The purpose of this chapter is to present an overall view of the fields related to image transmission systems, bandwidth reduction techniques and vision. This will help establish a framework upon which the more specific topics related to quantization and coding for Transform coding can be discussed.

The first section deals with digital communication systems in general. A notation and mathematical framework to characterize images is presented. Requirements for transmission of time varying images over a single channel and some techniques for bandwidth reduction are presented.

The second section focuses on transform coding in general and the discrete cosine transform in particular. It introduces the problems arising from quantizing and coding the transform coefficients.

Finally, section three reveals the present state of knowledge of the mechanism of vision and emphasizes the specific properties of the Human Visual System that are directly applicable to bandwidth reduction with the DCT.

2.1 Digital Image Communication Systems

The term 'image' has been used quite vaguely up to now. Before proceeding any further it is necessary to better define what 'image' means. In the context of this study, an image represents the distribution of light received on a limited portion of a plane known as the receptor. This light may come from light emitting objects like light bulbs, the sun and LEDs ... or from the reflection on objects and the surroundings.

Let $C(x, y, t, \lambda)$ represent the light received at the receptor in $\text{Watt/m}^2/\text{Hz}$. It is a spectrum of electromagnetic waves and $C(x, y, t, \lambda)d\lambda$ represent the power per unit surface contained in frequency interval $(\lambda, \lambda + d\lambda)$ at position (x, y) and time t . In the case of the receptor being the eye, $C(x, y, t, \lambda)$ is converted into a perception of color and intensity. In the case of a camera tube, it is converted into electrical signals. Restraining to monochrome imagery, the intensity of light, or luminance, may be defined as

$$L(x, y, t) = \int_0^\infty V(\lambda) C(x, y, t, \lambda) d\lambda \quad (2.1)$$

where $V(\lambda)$ is the luminous sensitivity of the receptor to frequency λ . For a still image, there is no motion and the luminance signal is independent of time ($L(x, y)$).

In general, the functions of an image communication system are to code the signal $L(x, y, t)$, transmit it over a channel, and decode it. This is summarized in Fig. 2.1.

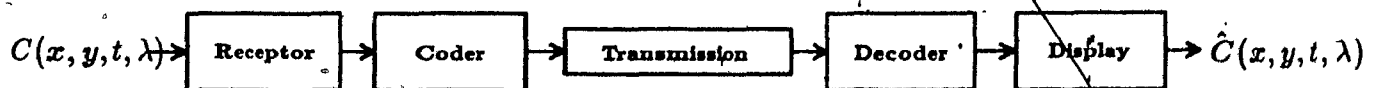


Fig. 2.1 An Image Communication System

The major problem is to reduce the spatial and temporal continuous analog intensity signal $I(x, y, t)$ into a one dimensional signal, namely, a signal dependent on time only. This requires sampling the intensity signal in at least two dimensions. In broadcast television, the signal is sampled in the vertical and temporal dimensions (y and t) in a line interlaced fashion. A discrete image implies that sampling is performed in three dimensions. A digital image is a discrete image in which the luminance signal is also quantized.

2.1.1 Digital Images

As seen in the previous section, a digital image is a quantized sampled image. The present section explains the sampling and quantizing processes.

Sampling

A detailed review of sampling of time varying imagery can be found in [9]. Basically, sampling is the process of converting a continuous representation of the analog intensity signal into a discrete representation. The resulting intensity signal is then defined as a set (Λ) of discrete points called a sampling lattice. An element of the lattice is called a pel (or pixel), meaning 'picture element'. The lattice can be formally represented by

$$\Lambda \doteq \{n_1 \mathbf{v}_1 + n_2 \mathbf{v}_2 + n_3 \mathbf{v}_3 \mid n_1, n_2, n_3 \in \mathbb{Z}\} \quad (2.2)$$

where $\mathbf{v}_1, \mathbf{v}_2, \mathbf{v}_3$ are the basis vectors of the lattice and $\mathbb{Z}$ the set of integer numbers.

A lattice is a discrete additive abelian group. The basis completely characterizes the sampling lattice. For practical applications, only a portion of the lattice is used, n_1, n_2, n_3 have a limited range. In order to better understand this operation, the equivalent one-dimensional operation is pulse modulation used in speech application, telemetry, and in general for multiplexing several time varying signals over the

same channel. The analog signal is sampled at intervals of T seconds, Fig. 2.2, and its sampling lattice may be expressed as:

$$\Lambda = \{n_1 T \mid n_1 \in \mathbb{Z}\} \quad (2.3)$$

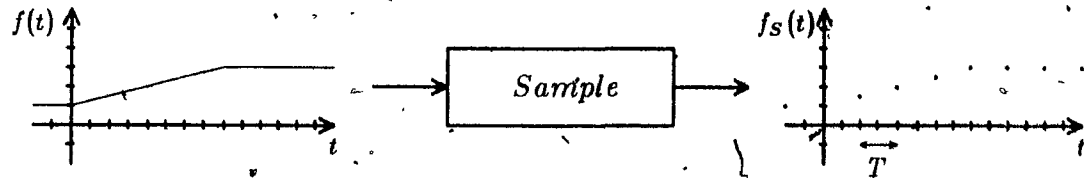


Fig. 2.2 Sampling of a unidimensional signal

In order to be able to exactly reconstruct the signal by interpolation, the sampling interval must be smaller than the inverse of twice the bandwidth of the signal (Nyquist criterion). This is to avoid aliasing.

These limitations also hold for the sampling of images. The basis vectors of the lattice should satisfy certain constraints in order to be able to reconstruct a continuous image from the sampled one without distortion. The reconstruction process is a three dimensional low pass filtering operation. When choosing a sampling structure, one would like to minimize the number of pels while maintaining an adequate visual bandwidth. This means that the pels should be close enough to avoid aliasing and the input signal should be low pass filtered prior to sampling. However, for many applications, compatibility with existing analog systems dictates the choice of the sampling structure. For example, television uses a line interlace scanning with defined vertical and temporal sampling rates. A frame at full vertical resolution consists of two fields sampled at twice the frame sampling rate. Field lines are also vertically displaced so as not to overlap.[†] The compatible sample lattice should use

[†] 525 lines/frame, 60 fields/second in some countries and 625 lines/frame, 50 fields/second in others.

the same vertical and temporal scanning rates, Fig. 2.3. The horizontal sampling rate is chosen by the designer.

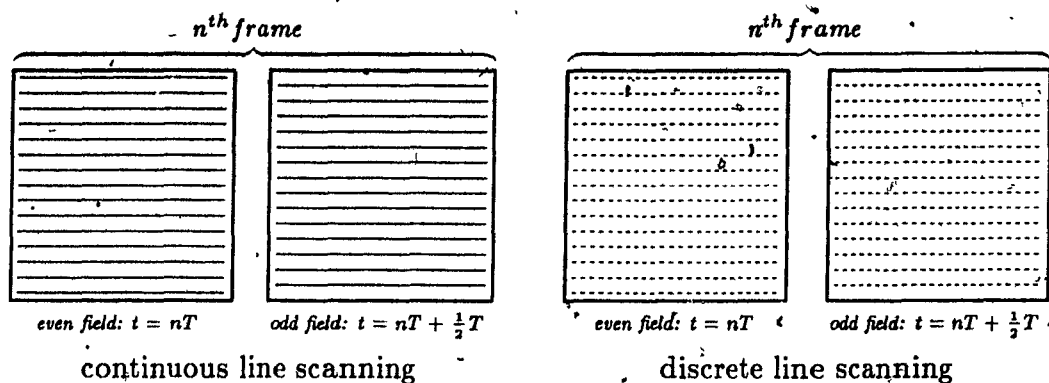


Fig. 2.3 Line Interlaced TV Frame

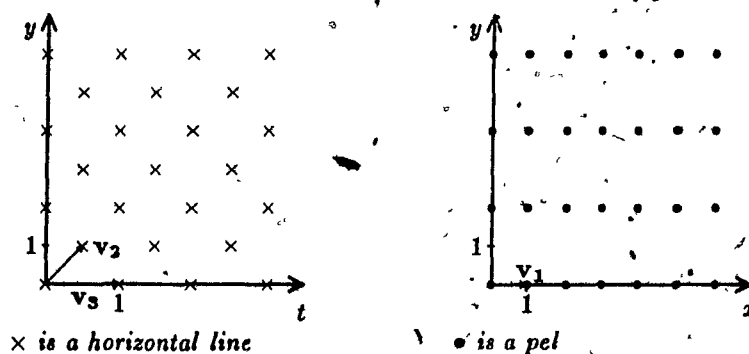


Fig. 2.4 Projection views

In the orthonormal reference system, the basis vectors for this sampling lattice are $(1,0,0)$ $(0,1,\frac{1}{2})$ and $(0,0,1)$, Fig. 2.4. These vectors stand for the horizontal sampling rate, temporal even field sampling rate, and the vertically displaced odd field, respectively. All pels are a linear combination of these vectors. Since this sampling lattice does not lend itself to easy mathematical tractability, it is often preferred to use a simpler lattice with the orthonormal basis vectors $(1,0,0)$ $(0,1,0)$

and $(0,0,1)$ for mathematical processing. Two solutions are used to convert the image sampled in the line interlaced fashion to an orthonormal sampling. The first solution consists of merging the odd field and the even field to get a full frame. This will be called 'frame sampling' in this thesis, Fig. 2.5. In the second solution known as 'field sampling', Fig. 2.6, the odd fields and even fields are vertically aligned.

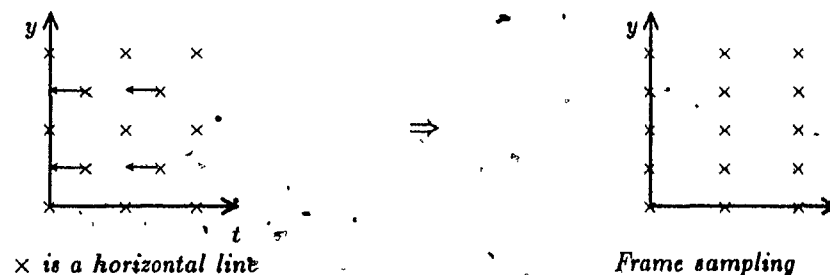


Fig. 2.5 Conversion to a Frame Sampling Lattice

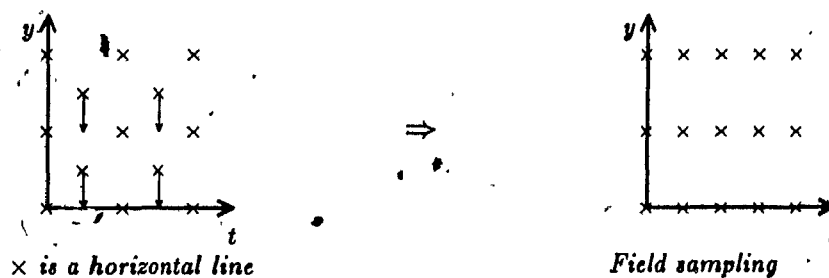


Fig. 2.6 Conversion to a Field Sampling Lattice

In both cases, digital processing of the images is much easier. At display time, the reverse operation is performed to return to line interlace scanning. The price to pay is that, depending on the coding and processing done, the quality of the reconstructed image may be affected.

Quantization

This section is about *scalar* quantization of the original image so that further block quantization can be performed by digital processing. Quantization is the conversion of an analog signal into a discrete representation. Quantization is a lossy operation since the analog signal cannot be converted back exactly from its discrete representation. Quantization thus introduces a distortion called *quantization noise*. (However this quantization noise may be imperceptible to the viewer.) A quantizer is fully described by its input-output relationship. For any real input value $x_n \leq x < x_{n+1}$ the output value $\hat{x}_n$ is chosen so that the quantization noise is minimized. The set $\{\hat{x}_n\}$ can be mapped into the set of integers for convenience.

If the step size $x_{n+1} - x_n$, ($n \geq 1$) is constant, the quantizer is called uniform. Otherwise it is a non-uniform quantizer. Prior to quantization, the luminance signal goes through an exponential non linearity, $z = y^{\frac{1}{\gamma}}$, known as *gamma correction* in order to compress its dynamic range. This is usually done by the camera itself. The resulting signal is then uniformly quantized using 64 (6 bits) or 256 (8 bits) grey levels. The number of grey levels is a compromise between quality, storage and transmission requirements. This method is subjectively near optimum [10].

To summarize, Fig. 2.7, a digital image is realized by sampling the luminance signal and performing both gamma correction and uniform quantization. Note that gamma correction may be done prior to sampling.



Fig. 2.7 Digitalization of an Image

From now on, if not otherwise stated, 'image' will actually mean digital image sampled with an orthonormal lattice and will be denoted by $f(i, j, k)$, where f is the intensity of the signal at the pel spatially located in (i, j) , ($i = 0, \dots, N - 1$), ($j = 0, \dots, M - 1$), at time k .

2.1.2 Statistical Characterization of Digital Images

It is convenient to regard a real image as a sample of a stochastic process. For continuous images, the function $f(x, y, t)$ is assumed to be a member of a continuous three-dimensional stochastic process with space variables (x, y) and time variable (t) . For digital images, $f(i, j, k)$ is defined at discrete points and has an integer value.

The mean, variance, autocovariance and autocorrelation are widely used characterizations of random processes and are defined below for discrete images.

Mean

$$\mu_f(i, j, k) = E[f(i, j, k)] \quad (2.4)$$

Variance

$$\sigma_f^2(i, j, k) = E[(f(i, j, k))^2] - (E[f(i, j, k)])^2 \quad (2.5)$$

Autocovariance

$$K_f(i_1, j_1, k_1; i_2, j_2, k_2) = E[f(i_1, j_1, k_1) \cdot f(i_2, j_2, k_2)] - E[f(i_1, j_1, k_1)] \cdot E[f(i_2, j_2, k_2)] \quad (2.6)$$

Autocorrelation

$$R_f(i_1, j_1, k_1; i_2, j_2, k_2) = E[f(i_1, j_1, k_1) \cdot f(i_2, j_2, k_2)] \quad (2.7)$$

An image is said to be stationary in the strict sense if all its moments are unaffected by shifts in the space and time origin. It is stationary in the wide sense

(less stringent) if its mean is constant and its autocorrelation is dependent only on the difference in the image coordinates, $(i_1 - i_2), (j_1 - j_2), (k_1 - k_2)$ and not on their individual values.

For convenience, images are often modeled as samples of a 'first order Markov process. A random process is first order Markov if it is stationary and the correlation between points in the image frames is proportional to the distance between them. While this approximation is close enough to model certain natural scenes, it is not well suited for man made objects which exhibit patterns and high horizontal and vertical correlation. For example, a typical movie frame has a background that is out of focus but its objects or actors of interest are in focus and usually exhibit details, sharp edges, patterns etc... Any global mathematical model of images should be used with great caution.

It is subjectively clear that some images convey more 'information' than others. For example, a view of a chess board with the players in the middle of a game conveys more information than a view of sand dunes in the Sahara. The latter one offers uniformity, few shades and few details while the former one contains many more details to be analyzed (for example who is winning the game, who are the players etc...). It would be helpful to be able to measure the amount of information, but as usual, the hardest part is to define rigorously what 'information' is. A generalized measure of information is the 'entropy' $H(S)$ defined as follows:

$$H(S) = - \sum_{i=1}^N P(s_i) \log_2 P(s_i) \quad (2.8)$$

where S is a source with elements s_i ($i = 1, \dots, N$) having the probability of occurrence $P(s_i)$. This defines how many bits are necessary to code source elements on the average, and thus gives a quantitative value of the information contained in the source. It should be noted that this definition corresponds to the zero-order entropy. It applies if the source is independent and identically distributed (iid) and

is a sample of a stationary random process. Information cannot be coded below entropy without introducing distortion. For a digital image source, s_i stands for the integer grey scale value of the luminance $f(i, j, k)$. In general, images are not iid processes and thus strictly speaking zero order entropy is not appropriate in this case. Higher order entropy definition should be used that takes into account the pel statistics. However, for comparison purposes, this assumption may be made in order to determine the minimum number of bits per pel required for transmitting the digital image.

2.1.3 Transmission of Digital Images

Given a digital image, consisting of a sequence of frames of $N \times M$ pels linearly quantized with b bits, the problem is to transmit it as efficiently as possible with the least perceptible distortion. Figure 2.8 shows a typical communication system.

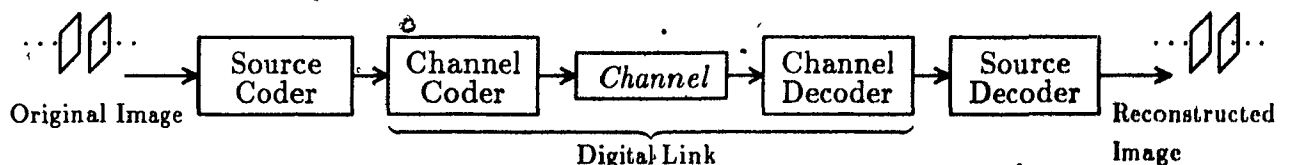


Fig. 2.8 Digital Image Communication System

The source coder converts the digital image into a stream of bits. An efficient source coder will minimize the average number of bits per pels at its output. This bit stream is transmitted over a digital link that is assumed to be error free. At the receiver, the source decoder yields a reconstructed digital image $\hat{f}$ which can then be displayed on a digital monitor. If the source coder does not introduce any distortion, $f = \hat{f}$, otherwise an error is introduced.

There exist methods to estimate the impact of such errors. Objective measures give an exact value of the error while subjective measures rely on the personal opinion of observers as to the quality of the reconstructed image or its degradation with respect to the original. The most common objective measure is the normalized mean square error $NMSE$ and its derivative, the normalized signal-to-noise ratio $NSNR$ given by:

$$NMSE = \frac{\sum_{i=1}^N \sum_{j=1}^M \sum_{k=1}^K [f(i, j, k) - \hat{f}(i, j, k)]^2}{\sum_{i=1}^N \sum_{j=1}^M \sum_{k=1}^K [f(i, j, k)]^2} \quad (2.9)$$

$$NSNR = 10 \log_{10}(NMSE) \quad (2.10)$$

Although this measure is mathematically tractable, it does not always correlate well with subjective measures. Attempts have been made to introduce objective measures taking into account a model of the visual system [6] [11] [12].

The most appropriate measure of image distortion is the subjective measure that involves the opinion of the viewer. The viewer is asked to express his opinion regarding the quality of the reconstructed image according to a 'distortion scale'. The CCIR has defined a few standard tests oriented toward particular needs and set guidelines as to the manner these tests should be conducted (recommendations 450 and 500). The two most often used distortion scales are the five-grade quality and impairment scales in Table 2.1.

Impairment Scale	
5	Imperceptible
4	Perceptible, but not annoying
3	Slightly annoying
2	Annoying
1	Very annoying

Quality Scale	
5	Excellent
4	Good
3	Fair
2	Poor
1	Bad

Table 2.1 Subjective Distorsion Scales

The quality scale is used to judge the overall quality of an image without any reference. The impairment scale is used to compare the quality of a processed (and presumably impaired) image with respect to the original, for example, to assess the quality of a coder.

2.1.4 Survey of Some Coding Methods

Putting aside PCM (pulse code modulation) which is not really a compression technique, four main classes of coding techniques exist, namely, predictive coding, transform coding, interpolative coding, and finally, miscellaneous techniques. A brief description of each is given below; a more comprehensive review can be found in [13] [14] [15].

Predictive coding methods are spatiotemporal methods. The image is scanned and for each pel a value corresponding to a prediction error is transmitted. The prediction error is formed by the difference between the luminance of the pel and a prediction signal based on the luminance of previously transmitted pels. This is in contrast with PCM where the luminance value is sent for each pel. The error has a smaller range and variance than the luminance signal. Its quantized value may be coded for transmission using fixed or variable length codewords. The prediction may be one, two or three dimensional, with increasing complexity and storage (buffer) requirements. Adaptive quantization and/or predictors can be used to enhance performances and adapt to the image being transmitted. Predictive coding can be implemented quite easily for real time use.

In transform coding, the image is parsed into blocks and each block is mapped to a transform block by a linear transformation. Several transforms have nice properties that make them attractive for picture coding. They can decorrelate the block to a certain extent and pack energy into fewer coefficients. The coefficients

are thus a more compact representation of the image if adequately quantized and coded. A more thorough review will be presented in the next section.

In interpolative coding, a subset of pels are transmitted and the remaining ones are interpolated. For example fixed subsampling may be either horizontal, spatial, temporal or a combination of them. In more complex schemes, the set of transmitted elements may depend on the image itself in an adaptive fashion. Temporal subsampling with an algorithm that interpolate the skipped frame using motion compensation techniques show promise.

In hybrid coding, a combination of the above coding technique is used to further reduce the bit rate. For example spatial DPCM and temporal subsampling with motion compensation.

2.2 Transform Coding

The transform coding process is indirect. A unitary linear transformation is performed on the image data to produce a set of transform coefficients which are then quantized and coded for transmission. In practical systems, the image is divided into smaller blocks and the transform is performed on the blocks rather than on the whole image. The receiver decodes the block and performs the inverse transformation. Fig. 2.9 portrays a communication system based on transform

coding.

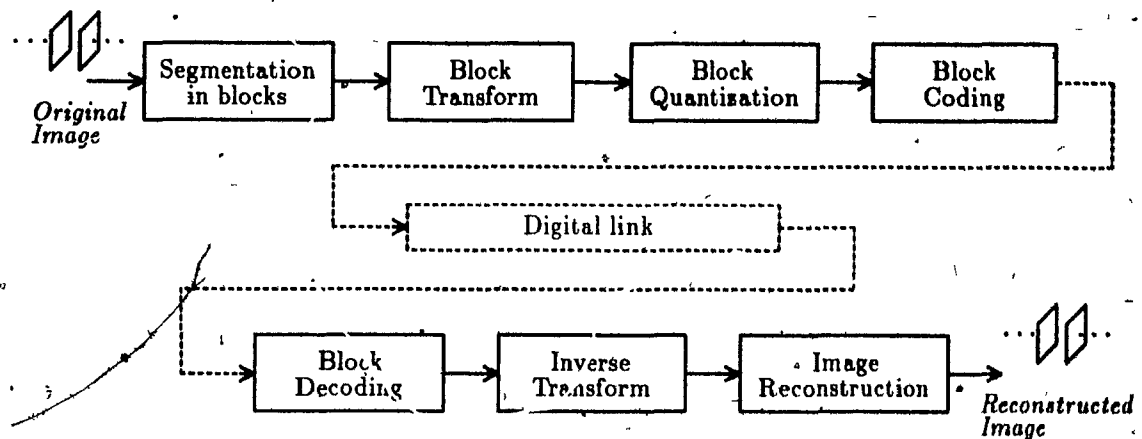


Fig. 2.9 Image Communication by Transform Coding

The concept of transform coding for images emerged in the late sixties. The Fourier transform was the first to be investigated. Its transform coefficients represent the spatial spectral distribution. The basic concept behind the Fourier transform is that for natural images, most of the spectral energy is concentrated at low frequencies, and a small portion at the higher frequencies. These latter terms can be coded with only a few number of code symbols or entirely discarded with only negligible image distortion. Hence bandwidth reduction is accomplished.

Other useful transforms are the Karhunen-Loeve, Walsh, Hadamard, Slant, Cosine and Sine transforms. The Karhunen-Loeve transform is optimal when some assumptions regarding image statistics are made and is often used as a reference. However, it is impractical due to the dependency on image content and computation requirements. The choice of a particular transform depends on the specific needs. If processing speed is an important factor, the Hadamard transform is adequate since

it uses only integer arithmetic. However, other transforms may yield better quality images while achieving the same bit rate.

The Cosine transform has emerged as a favorite because it theoretically approaches the bandwidth reduction of the Karhunen-Loeve transform while having fast computational algorithms.

2.2.1 The Karhunen Loeve transform (KLT)

Although the Karhunen-Loeve transform is seldom used in practice, it is a good example to gain some insight to the theory behind transform coding. The material contained in this subsection comes largely from [16]. For transform coding, the key to data compression is signal representation, which concerns the representation of a given class (or classes) of signal in an efficient manner. In a general framework, if a discrete signal is comprised of N sampled values (a block), then it can be thought of as being a vector X in an N -dimensional space. For a more efficient representation of the signal, an orthogonal transform T_o is applied to X yielding the vector Y . The objective is to select a subset of M components of Y , where M is substantially less than N . The remaining $(N - M)$ components can be discarded without introducing objectionable error when the signal is reconstructed using the retained M components. This achieves data compression. The most often used error criterion for judging or comparing orthogonal transforms is the mean square error. Simply stated, an orthogonal transform is a basis change (in particular, a rotation) for the representation of the vector signal X in the N -dimensional space. It can be shown that the optimal basis with respect to the mean square error criterion consists of the eigenvectors of the autocovariance matrix of X . In that basis, the autocovariance matrix of Y is diagonal, meaning that the transform vector components y_i are uncorrelated. The M components that should be retained correspond simply to the

highest M eigenvalues. An alternative to discarding coefficients would be to keep them all but use an optimal bit allocation as described in Segall [17].

The Karhunen-Loeve transform is the transform corresponding to this 'optimal' basis. Two objections can be raised to its use for image coding. The first one is that images are not a stationary process. This implies that the autocovariance matrix of the image blocks changes and must be constantly updated. This adds an overhead since the new matrix should be transmitted along with the coded blocks. If it is updated too often, the overhead offsets any gain. The second objection concerns the adequacy of the mean square error criterion used to determine the optimality of the Karhunen-Loeve transform. Depending on the level and kind of impairments, the mean square error may not correlate well with subjective error measures.

2.2.2 The Discrete Cosine Transform (DCT)

The Discrete Cosine Transform (DCT) is a deterministic transform in the sense that its basis vectors are fixed. It has gained much interest in recent years due to its excellent performance regarding compression, quality of the reconstructed images and to a lesser extent, the availability of fast algorithms. The DCT is an outgrowth of the Discrete Fourier Transform (DFT). Their definitions for a $N \times M \times K$ block is given below.

DCT

$$F(u, v, w) = \frac{1}{\sqrt{N}} \frac{1}{\sqrt{M}} \frac{1}{\sqrt{K}} \sum_{n=0}^{N-1} \sum_{m=0}^{M-1} \sum_{k=0}^{K-1} C(u)C(v)C(w) \cdot f(n, m, k) \cdot \cos\left(\frac{\pi(2n+1)u}{2N}\right) \cos\left(\frac{\pi(2m+1)v}{2M}\right) \cos\left(\frac{\pi(2k+1)w}{2K}\right) \quad (2.11)$$

Inverse DCT

$$f(n, m, k) = \frac{1}{\sqrt{N}} \frac{1}{\sqrt{M}} \frac{1}{\sqrt{K}} \sum_{u=0}^{N-1} \sum_{v=0}^{M-1} \sum_{w=0}^{K-1} C(u)C(v)C(w) \cdot F(u, v, w) \cdot \cos\left(\frac{\pi(2u+1)n}{2N}\right) \cos\left(\frac{\pi(2v+1)m}{2M}\right) \cos\left(\frac{\pi(2w+1)k}{2K}\right) \quad (2.12)$$

with $C(0) = 1$ and $C(l) = \sqrt{2}$ for $l \neq 0$.

DFT

$$F(u, v, w) = \frac{1}{\sqrt{N}} \frac{1}{\sqrt{M}} \frac{1}{\sqrt{K}} \sum_{n=0}^{N-1} \sum_{m=0}^{M-1} \sum_{k=0}^{K-1} f(n, m, k) \cdot \exp\left(\frac{-2\pi i u n}{N}\right) \exp\left(\frac{-2\pi i v m}{M}\right) \exp\left(\frac{-2\pi i w k}{K}\right) \quad (2.13)$$

Inverse DFT

$$f(i, j, k) = \frac{1}{\sqrt{N}} \frac{1}{\sqrt{M}} \frac{1}{\sqrt{K}} \sum_{u=0}^{N-1} \sum_{v=0}^{M-1} \sum_{w=0}^{K-1} F(u, v, w) \cdot \exp\left(\frac{+2\pi i u n}{N}\right) \exp\left(\frac{+2\pi i v m}{M}\right) \exp\left(\frac{+2\pi i w k}{K}\right) \quad (2.14)$$

with $i = \sqrt{-1}$.

The problem with the DFT is that the periodic signal consisting of the shifting of the block in three dimensions is not 'continuous' in general. The DFT produces non-negligible high frequency terms which produce the so called Gibbs phenomenon and causes visible block effects (the block boundaries are not smooth).

The idea behind the DCT is to build a 'superblock' by taking the symmetry of the $N \times M \times K$ block with respect to the planes $x = -\frac{1}{2}$ $y = -\frac{1}{2}$ $t = -\frac{1}{2}$ and apply the DFT to that block of $2N \times 2M \times 2K$ pels. If $K = 1$ then only $2N \times 2M$ is necessary. This removes the artificially high frequency content and the blocking effect. Using the fact that luminance values are real in nature and the symmetrical nature of the 'superblock' yields the formula of the DCT from the one of the DFT. Fast algorithms exist to compute the DCT, these are either based on the fast DFT or are specific to the DCT [18] [19] [20]. The formula given above corresponds to the *even* DCT since its block sizes are even. If the 'superblock' is constructed by taking the symmetry with respect to the planes $x = 0$ $y = 0$ and $t = 0$, the block dimensions become $2N - 1$, $2M - 1$ and $2K - 1$ and the DCT is said to be *odd*. It is seldom used because it does not lend itself to fast computational algorithms.

In general for all transforms, energy compaction approaches the KLT asymptotically with increasing block size. However, the DCT is best in the sense that it performs very well for small block sizes.

2.2.3 Quantization of the Transform Coefficients

The step following the mathematical transform is the quantization of the generally real valued transform coefficients. The purpose of this subsection is to formally define a few terms and describe some often used block quantizers.

Block quantization is the operation that consists in mapping each real valued transform coefficient to a code symbol belonging to a finite set, usually an integer. The output of a block quantizer is a block of code symbols that will be called the quantized transform coefficients, Fig. 2.10. An $N \times M \times K$ block quantizer consists of NMK individual quantizers Q_{uvw} ($u = 0, \dots, N-1$ $v = 0, \dots, M-1$ $w = 0, \dots, K-1$).

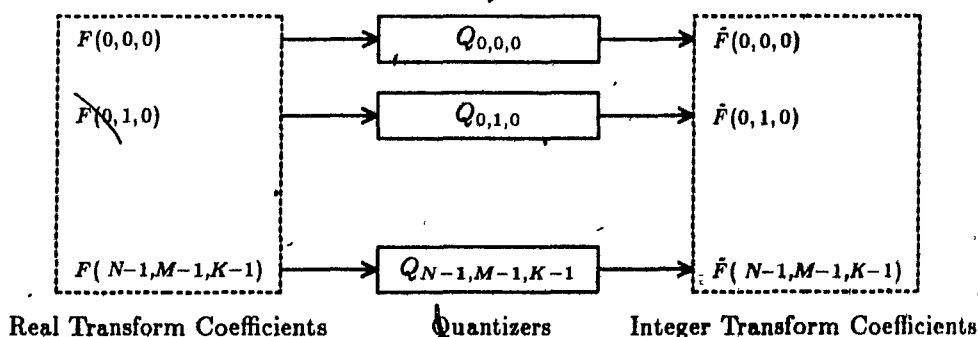


Fig. 2.10 Block Quantization of the Transform Coefficients

As previously seen for the quantization of the luminance signal, a quantizer is totally defined by its input-output relationship. A block quantizer will be said to be spatially uniform[†] if Q_{uvw} is the same for all coefficients. In practice, this implies that only one quantizer is needed hence reducing the system complexity. Uniform and threshold quantizers are most often used. The uniform quantizer has

[†] Not to confuse with an individual uniform quantizer which has a constant step size

a uniform step size throughout its input range, Fig. 2.11. The threshold quantizer is a uniform quantizer with a different input step size corresponding to a null output. Other types of quantizers will just be called non-uniform.

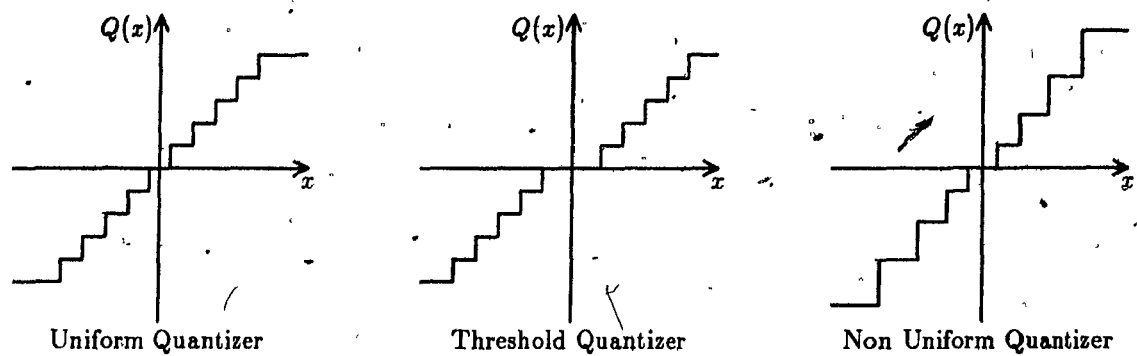


Fig. 2.11 Different Quantizers

The threshold quantizer is primarily used to select only those coefficients having an amplitude in absolute value greater than a certain threshold. A high value for the threshold will augment the probability of having zero valued outputs that are usually coded efficiently for transmission. The threshold may be controlled in a feedback loop by some parameters dependent on image constant or system state [7]. The choice and design of the quantizer depend upon the specific transform used, the coding strategy and most important, subjective experiments. It should be noted that the quantization process is a lossy operation and may introduce visible distortion.

2.2.4 Source Coding of the Quantized Coefficients

Source coding is the process of producing a bit stream from the block of quantized coefficients. The design of a coder is intimately dependent on the quantization scheme and the specific transform used. An efficient coder will minimize the number of bits required to code a block. Two issues affect the complexity of a transform

coder, namely, block quantization and block (or vector) coding. For block quantization, the designer has to choose the type of quantizer to use, such as MMSE, uniform, threshold quantizer, and the quantizer distribution for the coefficients. Block coding is often simplified to scalar coding of each coefficient using fixed or variable length codewords. Some examples are given below.

Zonal Coding

In this method, MMSE[†] (or Max[21]) quantizers are used. The outputs of the quantizers are coded with fixed length codewords. Let b_{uvw} be the number of bits allocated to the coefficient $F(u, v, w)$, and b the number of bits used to code the luminance value of a pel $f(i, j, k)$. The compression ratio achieved is the given by:

$$C = \frac{\sum_{u=0}^{N-1} \sum_{v=0}^{M-1} \sum_{w=0}^{K-1} b_{uvw}}{NMKb} \quad (2.15)$$

The choice of the number of bits for the MMSE quantizers is left to the designer and is a tradeoff between visual quality (more bits) and compression (less bits). However the total number of bits in a block is fixed by the channel bandwidth. The problem is thus to optimally allocate bits to each coefficient. Examples of algorithms are presented in [2], [22], [23] and [17]. An example with a $16 \times 16 \times 1$ block is given below in Fig. 2.2. Zero bit means that the coefficient is discarded.

This method is simple and yields a constant bit rate. A major disadvantage is the fact that the bit allocation is fixed. The MMSE quantizers are designed with statistics from a set of images and are optimum in the mean square error sense for that set of images. The objection is that if an image differs too much from the design set, its quality may degrade quickly. In particular, if a set of coefficients is allocated no bits, but is important for the given block, the distortion for that block will be high and visible and may attract the attention of the viewer. Also, as seen before, the mean square error is not a good distortion measure.

[†] Minimum Mean Square Error

8	8	8	7	7	7	5	5	4	4	4	4	4	4	4	4
8	8	7	6	5	5	3	3	3	3	3	2	2	2	2	2
8	7	6	4	4	4	3	3	2	2	2	2	2	2	2	2
7	6	4	3	2	2	2	2	1	1	1	1	0	0	0	0
7	5	4	2	2	2	2	1	1	1	0	0	0	0	0	0
7	5	4	2	2	2	1	1	0	0	0	0	0	0	0	0
5	3	3	2	2	1	1	0	0	0	0	0	0	0	0	0
5	3	3	2	1	1	0	0	0	0	0	0	0	0	0	0
4	3	2	1	1	0	0	0	0	0	0	0	0	0	0	0
4	3	2	1	1	0	0	0	0	0	0	0	0	0	0	0
4	3	2	1	0	0	0	0	0	0	0	0	0	0	0	0
4	3	2	1	0	0	0	0	0	0	0	0	0	0	0	0
4	2	2	0	0	0	0	0	0	0	0	0	0	0	0	0
4	2	2	0	0	0	0	0	0	0	0	0	0	0	0	0
4	2	2	0	0	0	0	0	0	0	0	0	0	0	0	0
4	2	2	0	0	0	0	0	0	0	0	0	0	0	0	0

Table 2.2 Typical bit assignment for transform zonal coding in 16×16 pel blocks at a rate of 1.5 bits per pel. (From [2], page 675)

Threshold Sampling Methods

These methods are meant to provide some adaptivity in selecting the coefficients to be transmitted. Only 'relevant' coefficients containing a significant amount of energy are chosen for transmission. They can be selected by using threshold quantizers where only non-zero valued coefficients are transmitted along with their location. A particular case consists of transmitting only L highest valued coefficients.

The Method of Chen and Pratt

The following is the method used by Chen and Pratt [7] in their 'scene adaptive coder'. The quantization is performed as follows. First, the coefficients with amplitude greater than a threshold value t are scaled by a scaling parameter s . Second, the quantization is performed by doing a simple floating point to integer roundoff

conversion of the scaled coefficients. The parameters t and s depend upon the status of an output buffer and are used to control the 'coarseness' of the quantization. A good image quality resulting from a fine quantization tends to yield more output bits hence filling up the buffer, and vice versa. This regulates the output bit rate going to the channel coder. The coder uses a combination of Huffman coding and run length coding. The coefficients are scanned sequentially following a zigzag pattern, Fig. 2.12. Only the coefficient greater than the threshold t are quantized. The output is coded with a Huffman coder. The coefficients less than t , the zero valued ones, are coded with a run length coder followed by a Huffman coder. Once a zero is encountered in the scanning process, the number of succeeding zeros is counted up to the occurrence of the next non zero value. The length of the sequence of zeros is then coded using a separate Huffman coder. Both amplitude and run length coder are then integrated into a unique coder. A similar coding method is used in this thesis (chapter 5).

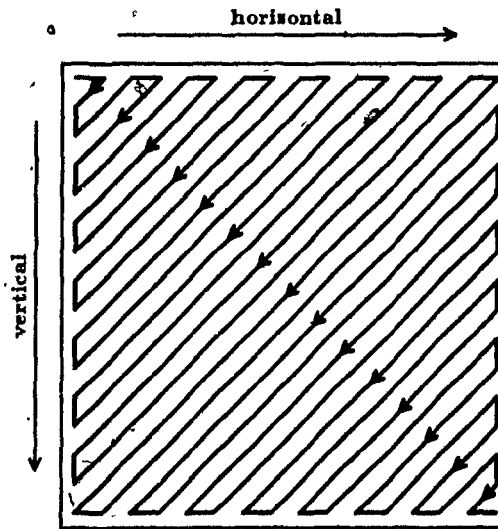


Fig. 2.12 Zigzag scanning pattern for the transform coefficients. The upper left corner represents $F(0,0)$ and is the starting point.

The use of a variable length codeword coder and adaptive processing makes

automatic and dynamical bit allocation possible to better use the channel bandwidth available. Saghri and Tesher [8] applied the concept of chain coding to code the clusters of zero valued coefficients.

2.3 The Human Visual System (HVS)

A good knowledge of the psychovisual properties of the human visual system is helpful to design efficient coders. Specifically, knowing what kind of distortion is introduced by a coder and its subjective effects on the viewer permit one to define thresholds for perception. If the distortion magnitude is below the threshold, it is not visible or bothersome. When its magnitude is above the threshold the distortion becomes visible or annoying. Thresholds are not absolute values but depend on the specific applications and subjective test criteria.

2.3.1 General Description

This subsection is meant to provide a very simplified description of the HVS. A more thorough review can be found in [24]. Since vision is a complex process that involves both physical interaction of the visual sensors and thought, it is not always easy to point out the boundaries between matter and abstraction. Fig. 2.13 attempts to present a representation of the HVS that can be loosely broken into the optical system, the retina which performs low level processing and the brain which performs high level processing. The separations here are not really physical but rather conceptual. For example, the eye comprises both the optical system and the retina. The retina is separated from the brain by the optic nerve which can also

be considered as part of the brain.

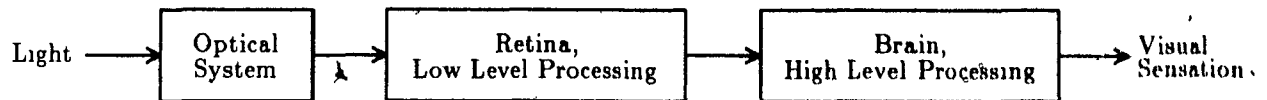


Fig. 2.13 The Human Visual System

The optical system is responsible for focusing light on the retina and controlling the amount of light projected on the retina.

The retina comprises a layer of cells called photoreceptors and a few layers between the photoreceptors and the optical nerve responsible for *low level processing* of the received image. The photoreceptors convert the light energy into electrical signals. There exist two kinds of photoreceptors; the *rods* and the *cones*. The rods (about 130 million on the retina) are sensitive to the overall shape of objects and low level of light. They are responsible for 'night vision' called *scotopic vision*. The cones are responsible for *photopic vision*. There are about 6.5 million cones mostly concentrated in a small region in the focusing area called *the fovea*. They are sensitive to color and details and account for most of the day vision. There are around 120 cones per degree in the fovea resulting in a visual resolution better than $\frac{1}{100}$ of a degree. The low level processing is responsible for 'multiplexing' the huge amount of visual information received for transmission to the brain via the optic nerve. For example, local operators recognize edges and certain shapes.

The brain is responsible for the sensation of visual perception. It is not yet well understood how it works.

When analyzing a scene, the eyes scan quickly the image and fixate parts of the image on the fovea in a rapid succession of *fixations*. The scanning pattern is controlled by the brain with a complex feedback mechanism dependent on the image content, the memory of the viewer etc.

2.3.2 Properties of the Human Visual System

Listed here are some general properties of the human visual system that are relevant to image coding. Since most are experimental results, it is probably helpful to describe both the experimental setting and the outcome.

Contrast Sensitivity

Depending on the brightness of the background, variation in brightness are perceived differently. The experiment consists of displaying a uniform background with luminance L and a circle with luminance $L + \Delta L$ and asking the viewer to adjust ΔL so that the circle be at the threshold of perceptibility, Fig. 2.14. It was found that the ratio $\frac{\Delta L}{L}$ was constant in a large range, Fig. 2.14. This suggests that the dynamic range of the input light is compressed by a logarithmic function. The contrast sensitivity is defined as $\frac{L}{\Delta L}$.

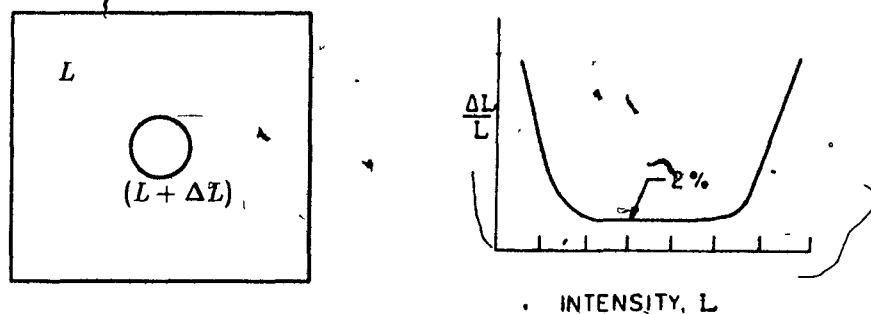


Fig. 2.14 Contrast Sensitivity, Experiment and Results

Edge Enhancement

The perception of edges and details is enhanced. The Mach band phenomenon is a good example. In an image consisting of adjacent rectangular bands of different gray levels, Fig 2.15, the perceived gray level near the edges is different than in the middle of the rectangles. Edge enhancement may be viewed as equivalent to highpass filtering in the spatial domain.

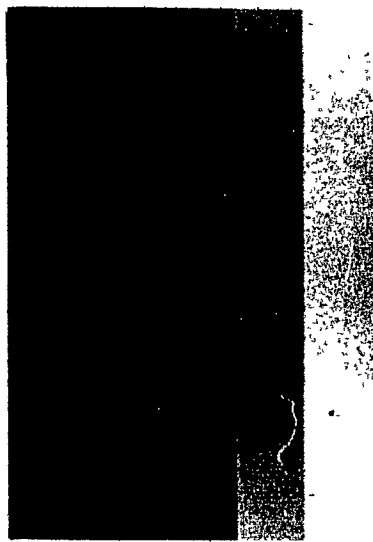


Fig. 2.15 Mach Bands

Frequency Response

A sinusoidal grating of amplitude ΔL and horizontal frequency f is added to a uniform background of luminance L . The viewer is asked to adjust ΔL at the threshold of perceptibility. It was found that $\frac{\Delta L}{L}(f)$ had the characteristics of a lowpass filter, Fig 2.16. Moreover, replacing the sine grating by a square grating, thus adding harmonics, did not have significant effects. In the temporal domain,

the human visual system is a lowpass filter with a cutoff frequency around 70Hz .

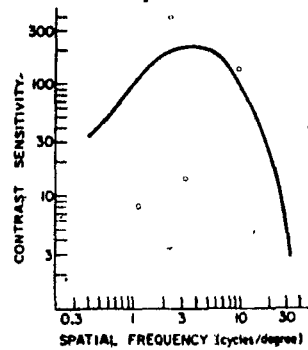


Fig. 1. Spatial frequency response function of human visual system.

Fig. 2.16 Response to sine gratings (From Hall and Hall[1])

Masking Properties

An image or region of an image will be called active if it is highly detailed, with many edges, patterns etc... and/or changes rapidly in time. For example, in a scene of a tree with many branches and leaves in a field on a windy day, the region of the tree may be considered as active. In active regions, the visibility of impairments is drastically reduced. This is a manifestation of masking, which in psychophysics is defined as the reduction in visibility of a stimulus by spatial or temporal non-uniformity in its surroundings [11].

Conclusion

These results combined with other experiments have led to models of the human visual system. A simple one is presented in Fig.2.17 and uses the following observations. The contrast sensitivity experiment implies a logarithmic dynamic range compression. The edge awareness-effect implies a spatial highpass filtering. The sine

grating experiment implies a lowpass filtering. Also, an upper limit on the visual resolution is due to the discrete nature of the photoreceptors. Lukas and Budrikis [11] have developed a model of the HVS that is more accurate to describe properties of the threshold vision. It is aimed at facilitating the estimation of impairment for black-and-white television pictures. This model is acceptable for 'monochrome vision'. For color vision, a similar multichannel model exists. More sophisticated models include banks of parallel bandpass filters acting on non-overlapping frequency ranges.

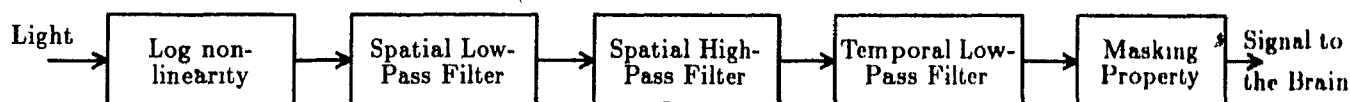


Fig. 2.17 A Simple Model for Monochrome Vision (Adapted from [2])

2.3.3 The Human Visual System and the DCT

As previously seen, the DCT allows coding to be made in a frequency space. The transform coefficients in a block thus represent the amplitude of frequency components rather than the luminance amplitudes. By taking advantage of the fact that the frequency response of the HVS is not flat, the transform coefficients are not perceived equally and need not be quantized using the same quantizer. In other words, the tolerance to the quantization error is not uniform. The experiments conducted with sine gratings give some hints in that higher frequency coefficients can be quantized more coarsely (or less finely) than lower frequency terms. However, no meaningful quantitative information as to the specific quantizers to use can be extracted from these experiments.

For the simplicity of its design, the first DCT coders used uniform block quantizers. The step size was usually dictated by the most sensitive coefficients and was usually too fine for the majority. This yields an unnecessarily high bit rate since imperceptible information is transmitted. Several researchers have tried to use the properties of the HVS to design better coders basing their work on specific models. Ericsson[4] working with $8 \times 8 \times 1$ blocks introduced a weighting function derived from detection thresholds for circular gratings and taking into account blocking effects. The transform coefficients are multiplied by their respective weights prior to quantization. This approach was inspired from the work of Mannos and Sakrison [6]. Lohscheller [3] measured the visibility threshold for each transform coefficient against a uniform background for still pictures, Fig 2.18. The inverse of the visibility threshold was used as a weighting function. This approach, like Mannos and Sakrison's, does not take into account the combined effect when *all* the coefficients are quantized.

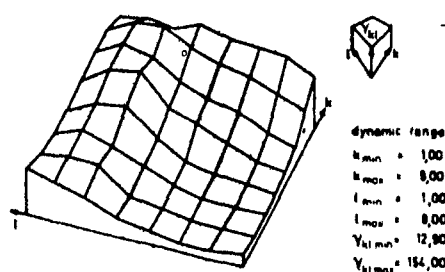


Fig 7. Expected values of the visibility threshold γ_{vis} for the spectral coefficients of a DCT (block size 8×8 pels)

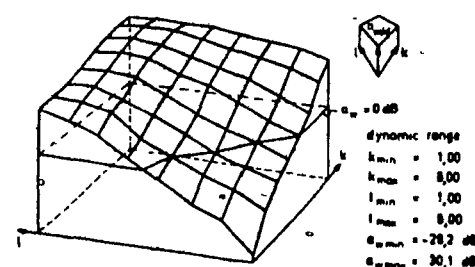


Fig 8 Image oriented visibility measure α_v

Fig. 2.18 Results of Lohscheller [3]

Griswold [5] defined a cosine energy function from a model of the HVS and used it to generate a bit allocation map for the coefficients. Bit allocation implicitly characterizes a quantizer as in zonal coding. Nill [25] incorporated a model of the HVS in a refined version of the mean square error quality assessment that

helped him find a weighting function for the transform coefficients. The trend so far has been to use weighting functions derived from specific models of the HVS, or more accurately, general models obtained from specific psychovisual experiments.

• Lohscheller achieved the most by using thresholds specific to the DCT.

The originality of the work presented in this thesis is to take the reverse approach. Although keeping in mind the properties of the HVS, no specific model is assumed. The weighting distribution is determined empirically through extensive experiments and viewing sessions on a set of images. This takes into account the combined effect of all coefficients. Once a satisfactory distribution is reached, a parametrical function is sought that closely approximates the distribution. The function can be viewed as a specific model of the visual perceptibility of the HVS for the DCT. Moreover, it permits a certain flexibility in the specification of the block quantizer that may be useful in adaptive quantization.

Chapter 3 Design of Quantizers for the DCT

The goal of this chapter is to study some quantization schemes for the DCT coefficients that are subjectively adapted to the human visual system. An empirical approach is taken. Through subjective testing experiments with DCT coded images, a model of the HVS's sensitivity to the kind of distortion introduced by the DCT is sought. The term 'model' being quite ambitious, the study focuses rather on finding parametrical functions to generate block quantizers appropriate to DCT coding.

This is carried out in three steps: The first step consists of preliminary experiments on still pictures with $N \times N$ two-dimensional DCT. Uniform quantizers are used. The quantizer distribution is zonal: the block is divided into zones in which coefficients use quantizers with the same step size. Zones and step sizes are adjusted to produce images with distortion at the threshold of perceptibility for experienced viewers.

The second step consists of approximating the previously found zonal distribution with a parametrical function. The quantizer step size for a coefficient is a function of its position within the block and a number of parameters. The function and the parameters should allow an easy specification of a block quantizer and permit a range of distortions around the threshold of perceptibility to be obtained by varying the parameters. Ultimately the number of parameters should be reduced to one.

The third step consists in extending the experiments to time varying images and evaluating the performance of the quantization scheme through formal subjective tests.

Based on the results obtained, non-uniform quantizers are investigated. The study is then extended to $N \times N \times K$ (three-dimensional) DCT.

3.1 Transform Coding System

3.1.1 Block Quantizer

As seen in Sect. 2.2.3, a block quantizer is specified by assigning a particular quantizer $Q_{u,v,w}$ to each of the $N \times M \times K$ coefficients of the block. In this study, a particular type of non-uniform quantizer is used which is described by a set of six parameters: $Q_{u,v,w}(\text{minimum}_{u,v,w}, \text{maximum}_{u,v,w}, \text{starting step}_{u,v,w}, \text{threshold}_{u,v,w}, \text{slope}_{u,v,w}, \text{saturation}_{u,v,w})$. These parameters are illustrated in Fig. 3.1

The parameters are used to 'customize' the quantizer for a particular coefficient and are defined as follows:

- 1 $\text{minimum}_{u,v,w}$ is the minimum input value
- 2 $\text{maximum}_{u,v,w}$ is the maximum input value
- 3 $S_{u,v,w} = S_{u,v,w}^{(1)}$ is the quantization step just after the threshold.
- 4 $\text{threshold}_{u,v,w}$: for an input value x such that $-\text{threshold}_{u,v,w} \leq x \leq \text{threshold}_{u,v,w}$, the output is zero
- 5,6 $\text{slope}_{u,v,w}$ and $\text{saturation}_{u,v,w}$ are defined by:
 $S_{u,v,w}^{(n+1)} = \min(\text{slope}_{u,v,w}^n, \text{saturation}_{u,v,w}) S_{u,v,w}^{(1)}$ where $S_{u,v,w}^{(n+1)}$ is the $(n+1)^{\text{th}}$ quantization step.

Minimum and maximum depend on the measured dynamic range of the coefficient according to statistics generated with a set of test images. The other parameters are assigned by the designer and are used to control the 'shape' of the quantizer. Uniform quantizers will be used mostly, in which case $threshold_{u,v,w} = \frac{S_{u,v,w}}{2}$ and $slope_{u,v,w} = 1$.

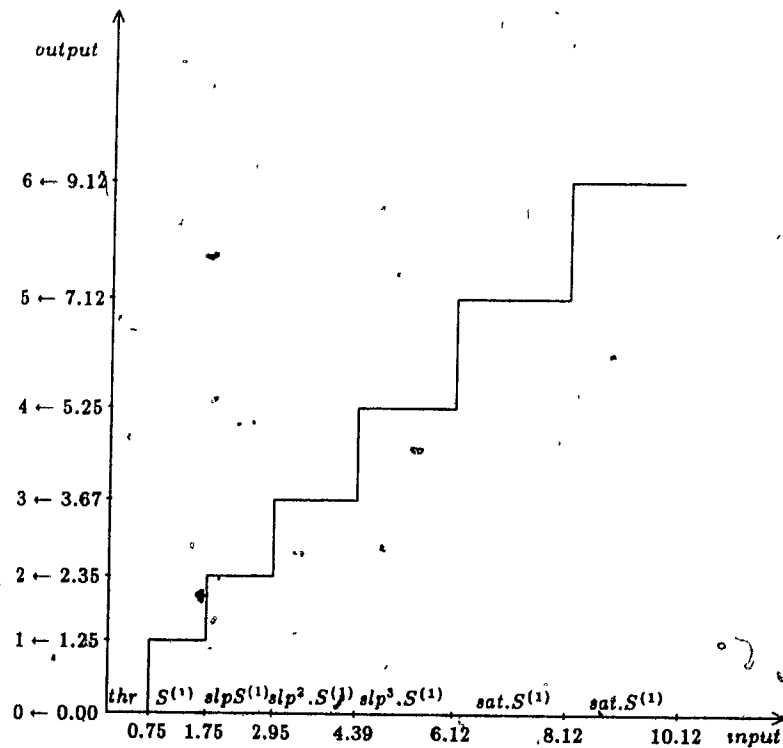


Fig. 3.1 An example of quantizer, $threshold = 0.75$, $S^{(1)} = 1$, $slope = 1.2$, $saturation = 2.0$ (symmetrical negative side)

The representative levels of the quantizer are real values $\hat{x}_n$ defined by

$$\hat{x}_n = \begin{cases} \frac{x_n + x_{n+1}}{2} & n \geq 1 \\ 0 & n = 0 \\ \frac{x_{n-1} + x_n}{2} & n \leq -1 \end{cases}$$

The output of the quantizer is given by $Q(x) = \hat{x}_n$ if $x \in i_n$ where

$$i_n = \begin{cases} [x_n, x_{n+1}] & n \geq 1 \\ [x_{-1}, x_1] & n = 0 \\ [x_{n-1}, x_n] & n \leq -1 \end{cases}$$

To make the coding process easier, these real values are mapped into the set of integers, $\hat{x}_n \mapsto y_n = n$. From now on, a quantized coefficient will actually refer to the integer value.

The definition of the DCT given earlier has been slightly modified so that the dynamic range of the transform coefficients be of the same order of magnitude as the dynamic range of the luminance value of the pels. For example, using this definition, the value of the coefficient $F(0,0,0)$ is the actual average luminance value of the pels in the block. This only represents a scaling of the transform coefficients.

DCT

$$F(u, v, w) = \frac{1}{N} \frac{1}{M} \frac{1}{K} \sum_{n=0}^{N-1} \sum_{m=0}^{M-1} \sum_{k=0}^{K-1} C(u) C(v) C(w) \cdot f(n, m, k) \cdot \cos\left(\frac{\pi(2n+1)u}{2N}\right) \cos\left(\frac{\pi(2m+1)v}{2M}\right) \cos\left(\frac{\pi(2k+1)w}{2K}\right) \quad (3.1)$$

Inverse DCT

$$f(n, m, k) = \sum_{u=0}^{N-1} \sum_{v=0}^{M-1} \sum_{w=0}^{K-1} C(u) C(v) C(w) \cdot F(u, v, w) \cdot \cos\left(\frac{\pi(2u+1)n}{2N}\right) \cos\left(\frac{\pi(2v+1)m}{2M}\right) \cos\left(\frac{\pi(2w+1)k}{2K}\right) \quad (3.2)$$

with $C(0) = 1$ and $C(l) = \sqrt{2}$ for $l \neq 0$.

The evaluation of a block quantizer is carried out in three steps: The first step involves the specification of the quantizer for each coefficient in terms of the relevant parameters. The second step consists of performing the DCT, quantization and inverse DCT for a set of test images. The final step consists of evaluating the processed image based on a number of criteria. Subjective evaluation is the principal criterion used.

3.1.2 Frame or Field Processing and Block Size Considerations

For the DCT coding of line interlaced images, two major decisions must be taken. The first one concerns the block size to be used. The second one is whether to use field or frame processing, or in other words convert the line interlaced sampling

lattice to a frame sampling lattice or a field sampling lattice. The terms *intraframe* processing will apply to bidimensional processing ($N \times M \times 1$) using a frame sampling lattice while *interframe* processing will apply to three-dimensional processing. The terms *intrafield* and *interfield* have the same meaning when using a field sampling lattice. The inconvenience with frame processing is that the line interlace produces images with jagged edges where there is motion due to the superposition of the two displaced fields. This increases the energy contained in vertical high frequency coefficients, thus reducing the energy compaction. The inconvenience with field processing is that the correlation with vertically adjacent pels is reduced since the distance between them is actually doubled. This also makes energy compaction less effective.

In this study, frame processing was chosen for the following reasons. When there is little or no motion and not too many edges and details, the effects of the line interlace are not significant. Moreover, even though the high frequency content of the frame is increased when there is motion, there is no need to accurately reproduce the jagged edges. Their perceptibility is reduced since they are in a region of temporal activity.

Larger block sizes better exploit the energy compaction property of the DCT. As mentioned earlier, due to the non-stationary nature of images, adaptive systems outperform non-adaptive ones. Adaptivity is better achieved with smaller block sizes. Smaller block sizes also allow faster computation which is extremely important for real time processing. A spatial block size of 16×16 was chosen since it offers a good compromise and is used extensively by other researchers. For a three-dimensional DCT, a temporal dimension of four was chosen to yield block of $16 \times 16 \times 4$. This small value also reduces memory requirements for computer simulation and processing time to an acceptable level.

3.2 Experimental Setting

3.2.1 Test Images

A set of six images was used for the experiments, namely, EIA, QUILT, OSCAR, TOYS, YVON and ALBERT. The first frame of these images can be found in Appendix A. EIA is a test pattern that is used in the TV industry. It is meant to present critical areas where distortion, if any, is easy to detect. There is translational and rotational motion. QUILT presents diverse and strong patterns with much high frequency content. It is a difficult image to code with transform methods. There is horizontal motion. OSCAR is a furry puppet on a turntable in which rotational motion is present. TOYS comprises a cube with letters on it, a few toys on a piece of wood. ALBERT and YVON are head and shoulder scenes typical of teleconferencing.

These images are meant to present typical characteristics of television and teleconferencing scenes where the distortion introduced by the quantization of the DCT coefficients can be studied. For the experiments on still images, only the first frame was used. For experiments on time varying images, a sequence of 12 frames was used and displayed in a palindromic fashion.

The HDVS system of the BNR/INRS laboratory was used. The HDVS permits real time recording and reproduction of digital high resolution color images. In the display mode, up to eight image sequences stored on disks are accessible and one may switch to any sequence instantly. Any frame or field can be displayed and values of individual pels are available. Due to hardware limitations, only a window of 256×212 pels/frame can be displayed in real time, Fig 3.2. In monochrome mode, the images are quantized with eight bits (256 grey levels).

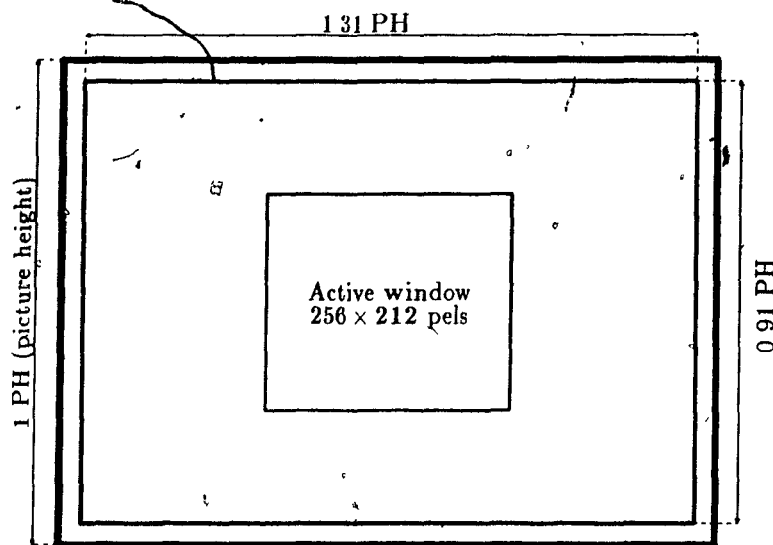


Fig. 3.2 View of the monitor and the display window.
Here $PH = 0.35m$

3.2.2 Some Statistics on the Real Transform Coefficients

The minimum, maximum, average and standard deviations for the real coefficients (before quantization) were computed for a $16 \times 16 \times 1$ block using the test images mentioned earlier and other ones as well. Results are given in Appendix D. The standard deviation gives an idea of the energy contained in a coefficient. The standard deviation table clearly shows the energy compaction property of the DCT.

The probability density function of the coefficients is also of interest. Except for $F(G, 0, 0)$, the probability density function resembles a decaying exponential. An example is given in Fig. 3.3. Depending on the value of *threshold* for the particular coefficient, a fair percentage of the quantized values can have the value zero.

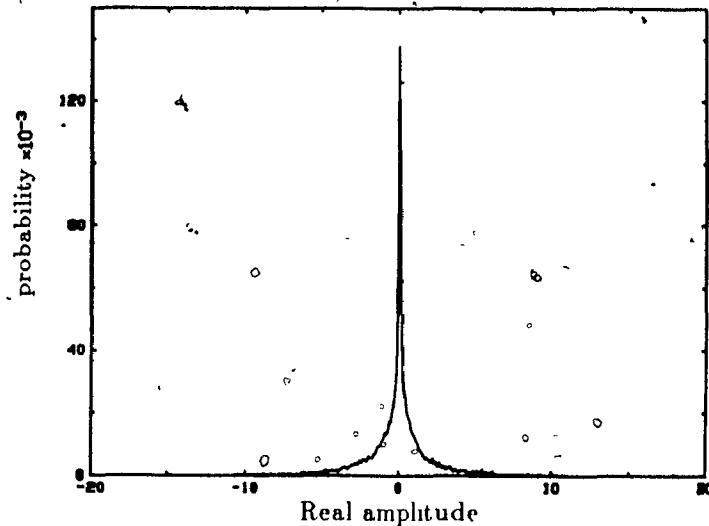


Fig. 3.3 Probability density function of $F(2, 2, 0)$. The sequence used to generate it is LONG-SEQUENCE (described in appendix A) processed with a $16 \times 16 \times 1$ DCT

3.3 Thresholds of Perception for Still Images (2D-DCT)

Experimenting with still images has two advantages. First, it is easier to study the distortion in the coded image. Secondly, there are no temporal masking effects. As a starting point for this study, uniform quantizers will be used. In that case, only one parameter is required to characterize a quantizer, namely, its step size $S_{u,v,w}$.

Assume it is possible to assess the effect of each individual coefficient on the reconstructed image. The value of the quantization step size $St_{u,v,w}$ such that if $S_{u,v,w} > St_{u,v,w}$, distortion becomes visible will be said to be at the threshold of perception for the coefficient $F(u, v, w)$. The block quantizer with such quantization steps is what is sought in this chapter and defines an upper limit on the values of the quantization steps for the transform coefficients.

Needless to say, it is virtually impossible to find such a block quantizer since

$v \backslash h$	0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
0	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1
1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1
2	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1
3	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1
4	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1
5	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1
6	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1
7	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1
8	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1
9	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1
10	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1
11	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1
12	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1
13	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1
14	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1
15	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1

Table 3.1 Example 1 of the step size distribution of a block quantizer. Spatially uniform step size of 0.1 .

it is difficult (if not impossible) to assess the perceptibility of a single coefficient and the threshold will depend on the viewing distance, the viewer and the image itself. Therefore, a conservative method to evaluate the outcome of the experiments is used. The viewing distance is less than the one used for subjective tests later on. A distortion will be considered visible if it is visible on any of the test images.

Typical distortions introduced by quantizing the transform coefficients are:

- **Block effects:** The effect of segmenting the image into smaller blocks is visible. This is usually due to a coarse quantization of low frequency coefficients.
- **Artifacts:** Introduction or modification of patterns and in general small changes in areas with details.
- **Loss of resolution:** Generally introduced by discarding high frequency coefficients.
- **Noise:** Introduced by the additive effect of quantizing the coefficients. It is always present. For images already noisy, it may be not distinguishable from the original noise.

The types of distortion are listed in decreasing order of annoyance level. In general, block effects are quite annoying while noise may be more easily tolerated.

Experiments

First, the maximum step size for a spatially uniform block quantizer was determined. This value is dictated by the low frequency coefficients that need to be finely quantized to avoid block effects and artifacts. Secondly, quantizing all the coefficients with the above mentioned step size, a zone (or a set) of coefficients is more coarsely quantized. The maximum step size introducing distortion is sought. This operation is repeated until the zones cover the whole block. Thirdly, the coefficients are quantized with the step size found for the zone they belong to. If distortion becomes visible, some 'tuning' is applied. The zones are first chosen to be large and as experience and 'feeling' is gained, they are made smaller.

This constitutes a lengthy procedure that comprises over 150 experiments.[†] Some intermediate results are given in Table 3.1, 3.2 and 3.3.

$v \backslash h$	0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
0	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1
1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.2	0.2	0.2	0.2	0.3	0.3	0.3	0.4
2	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.2	0.2	0.2	0.2	0.3	0.3	0.3	0.4	0.4
3	0.1	0.1	0.1	0.1	0.1	0.1	0.2	0.2	0.2	0.2	0.3	0.3	0.3	0.4	0.4	0.4
4	0.1	0.1	0.1	0.1	0.2	0.2	0.2	0.2	0.3	0.3	0.3	0.3	0.4	0.4	0.4	0.4
5	0.1	0.1	0.1	0.2	0.2	0.2	0.2	0.3	0.3	0.3	0.3	0.4	0.4	0.4	0.4	0.4
6	0.1	0.1	0.2	0.2	0.2	0.2	0.3	0.3	0.3	0.3	0.4	0.4	0.4	0.4	0.4	0.4
7	0.1	0.2	0.2	0.2	0.2	0.3	0.3	0.3	0.3	0.4	0.4	0.4	0.4	0.4	0.4	0.4
8	0.1	0.2	0.2	0.3	0.3	0.3	0.3	0.3	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4
9	0.1	0.2	0.3	0.3	0.3	0.3	0.3	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4
10	0.1	0.3	0.3	0.3	0.3	0.3	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4
11	0.1	0.3	0.3	0.3	0.3	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4
12	0.1	0.3	0.3	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4
13	0.1	0.3	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4
14	0.1	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4
15	0.1	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4	0.4

Table 3.2 Example 2 of the step size distribution of a block quantizer. Intermediate experiment.

[†] An experiment: assessing and analyzing the effect of a block quantizer on the test images.

$u \backslash h$	0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
0	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1	0.1
1	0.1	0.1	0.1	0.1	0.1	0.1	0.2	0.2	0.2	0.2	0.2	0.2	0.2	0.2	0.2	0.2
2	0.1	0.1	0.1	0.1	0.1	0.2	0.2	0.2	0.2	0.2	0.2	0.2	0.2	0.2	0.2	0.2
3	0.1	0.1	0.1	0.1	0.2	0.2	0.2	0.3	0.3	0.3	0.4	0.4	0.4	0.4	0.5	0.5
4	0.1	0.1	0.1	0.2	0.2	0.2	0.2	0.3	0.3	0.3	0.4	0.4	0.4	0.4	0.5	0.5
5	0.1	0.1	0.2	0.2	0.2	0.2	0.3	0.3	0.3	0.4	0.4	0.4	0.4	0.5	0.5	0.5
6	0.1	0.2	0.2	0.2	0.2	0.3	0.3	0.3	0.4	0.4	0.4	0.4	0.5	0.5	0.5	0.5
7	0.1	0.2	0.2	0.2	0.3	0.3	0.3	0.4	0.4	0.4	0.4	0.5	0.5	0.5	0.5	0.5
8	0.1	0.2	0.2	0.3	0.3	0.3	0.4	0.4	0.4	0.4	0.5	0.5	0.5	0.5	0.5	0.5
9	0.1	0.2	0.2	0.3	0.3	0.4	0.4	0.4	0.4	0.5	0.5	0.5	0.5	0.5	0.5	0.5
10	0.1	0.2	0.2	0.3	0.4	0.4	0.4	0.4	0.5	0.5	0.5	0.5	0.5	0.5	0.5	0.5
11	0.1	0.2	0.2	0.4	0.4	0.4	0.4	0.5	0.5	0.5	0.5	0.5	0.5	0.5	0.5	0.5
12	0.1	0.2	0.2	0.4	0.4	0.4	0.5	0.5	0.5	0.5	0.5	0.5	0.5	0.5	0.5	0.5
13	0.1	0.2	0.2	0.4	0.4	0.5	0.5	0.5	0.5	0.5	0.5	0.5	0.5	0.5	0.5	0.5
14	0.1	0.2	0.2	0.4	0.5	0.5	0.5	0.5	0.5	0.5	0.5	0.5	0.5	0.5	0.5	0.5
15	0.1	0.2	0.2	0.5	0.5	0.5	0.5	0.5	0.5	0.5	0.5	0.5	0.5	0.5	0.5	0.5

Table 3.3 Example 3 of the step size distribution of a block quantizer. One of the final experiments.

The first distribution in Table. 3.1 corresponds to the starting point, spatially uniform and having a quantization step of 0.1. The second in Table. 3.2 corresponds to a tentative zonal distribution. The third one in Table. 3.3 is the last result after some 'tuning'. The symmetry was kept on purpose for the sake of simplicity. The last distribution does not represent the ultimate threshold distribution. In fact, it is a conservative distribution that will serve as a reference for future experiments.

An increase in the quantization step of a few low frequency coefficients will generally be much more perceptible than for higher frequencies. In general for high frequency coefficients ($u, v, > 7$), the effect is quite hard to pinpoint and, for a non-experienced viewer is not perceptible at a first glance. The crescent shape is due to the fact that the eye is more sensitive to horizontal and vertical patterns than diagonal ones. While experimenting, the author has gained a knowledge of the possible distortions on the test images and his judgement has sometimes become overcritical. For the last experiments, a small level of distortion, assumed gener-

ally imperceptible to a non-experienced viewer, was tolerated. Frequent informal viewing sessions with other viewers were held.

The zonal approach to finding a Q-step distribution at the threshold of perceptibility presents several drawbacks. The distribution is discontinuous at the zone boundaries. The zones have to be selected in a more or less arbitrary fashion. The quantization step assignment is not very precise and is very tedious. The use of a parametrical function to assign values to the quantization steps eliminates these problems. In that case, a Q-step is given by $S_{u,v,w} = s(u, v, w, p_1, p_2, \dots, p_n)$ and the distribution is totally characterized by the parameters p_i , $i = 1, \dots, n$. Such a function should satisfy the following requirements:

- Approach the empirical zonal distribution as closely as possible for a certain set of the parameters.
- The parameters should allow an easy control over the distribution.

The idea behind these requirements are twofold: ease in the specification and adaptivity considerations. Varying the parameters should have a direct impact on the entropy of the coefficients (and thus the compression ratio) and the image quality. It should be possible to select the parameters to get the best quality for a given entropy.

Parametrical Function

Different types of functions, varying in complexity, were investigated. The major goal is not to fit perfectly the zonal distribution but rather to find a function that fits reasonably well the low and medium coefficients, and that allows one to alter the shape of the distribution in an easily controllable way. The following function satisfied these requirements:

$$S_{u,v,0} = a + b_1(u+1)^{c_1} b_2(v+1)^{c_2} \quad \text{for } u = 0, \dots, 15, v = 0, \dots, 15. \quad (3.3)$$

In order to keep a symmetrical distribution, $b_1 = b_2 = b$, $c_1 = c_2 = c$. This finally gives

$$S_{u,v,0} = a + b^2(u+1)^c(v+1)^c \quad \text{for } u = 0, \dots, 15, v = 0, \dots, 15. \quad (3.4)$$

Depending on the values of a , b and c , a reasonable fit of the zonal distribution is achieved in the low and medium frequencies. Most important, the crescent shape is preserved. The exponent c controls the shape in the low and medium range. A one-dimensional example is given below to show what can be achieved by adjusting the parameters. Let $S_u = a + b(u+1)^c$. Examples of S_u are shown in Fig. 3.4.

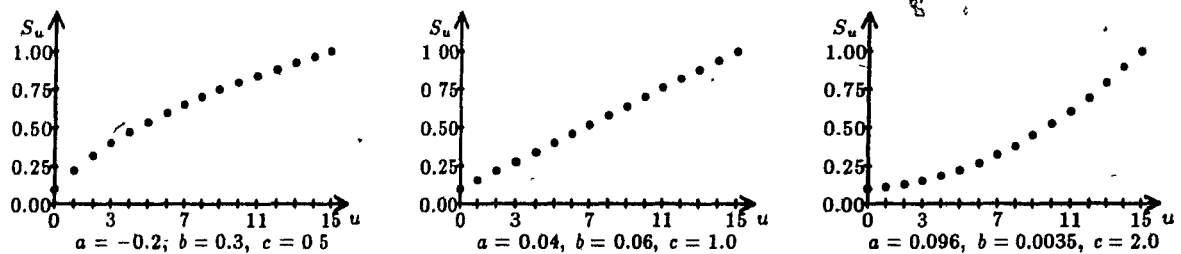


Fig. 3.4 Examples of one-dimensional distributions, S_0 and S_{15} are fixed.

From a practical viewpoint, the value of $S_{0,0,0}$ cannot be raised without introducing perceptible block effects. On the other hand, the value of $S_{15,15,0}$ and high frequency coefficients are not of critical importance. Substituting the set $\{ S_{0,0,0}, S_{15,15,0} \}$ for $\{ a, b \}$ † yields a set of three parameters that are more meaningful, namely, $\{ S_{0,0,0}, S_{15,15,0}, c \}$. These can be thought of as the lower bound, upper bound and 'curvature' of the Q-step distribution respectively. By fixing the value of the upper and lower bounds, the distribution becomes a function of only one parameter. If needed, it is always possible to override certain values generated by the function. An example is given below. The contour plot of Fig. 3.5

† $b^2 = \frac{S_{15,15,0} - S_{0,0,0}}{16^{2c} - 1}$ and $a = S_{0,0,0} - b^2$

gives a general idea on the shape of the distribution and the order of magnitude of the quantization step distribution. It is useful when comparing two distributions. The quantization step size table in Table 3.4 gives the exact values of the step sizes. It is more precise.

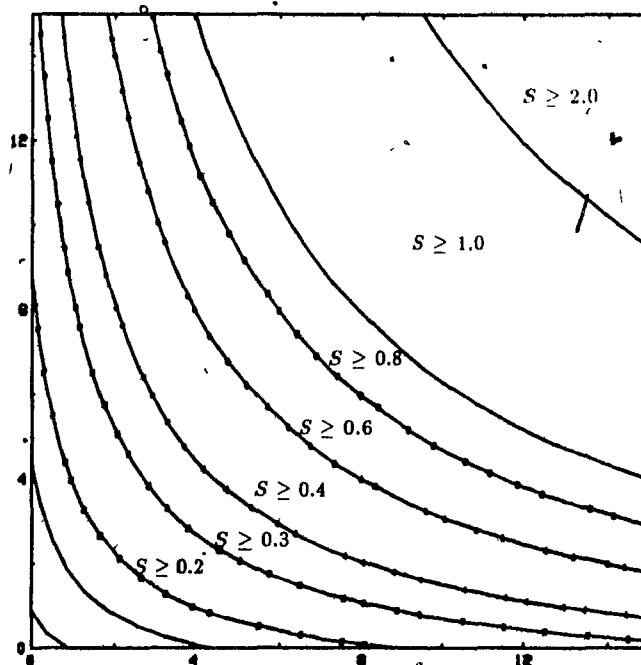


Fig. 3.5 Contour plot of a block quantizer step size distribution. The values of the parameters are: $S_{0,0,0} = 0.1$, $S_{15,15,0} = 3.0$ and $c = 1.0$

Having a function of only one parameter is interesting because it permits an easy specification of a block quantizer. Subjective tests can be conducted to find the value of c at the threshold of perception.

$v \backslash h$	0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
0	0.10	0.11	0.12	0.13	0.15	0.16	0.17	0.18	0.19	0.20	0.21	0.23	0.24	0.25	0.26	0.27
1	0.11	0.13	0.16	0.18	0.20	0.23	0.25	0.27	0.29	0.32	0.34	0.36	0.38	0.41	0.43	0.45
2	0.12	0.16	0.19	0.23	0.26	0.29	0.33	0.36	0.40	0.43	0.46	0.50	0.53	0.57	0.60	0.63
3	0.13	0.18	0.23	0.27	0.32	0.36	0.41	0.45	0.50	0.54	0.59	0.63	0.68	0.73	0.77	0.82
4	0.15	0.20	0.26	0.32	0.37	0.43	0.49	0.54	0.60	0.66	0.71	0.77	0.83	0.88	0.94	1.00
5	0.16	0.23	0.29	0.36	0.43	0.50	0.57	0.63	0.70	0.77	0.84	0.91	0.98	1.04	1.11	1.18
6	0.17	0.25	0.33	0.41	0.49	0.57	0.65	0.73	0.81	0.88	0.96	1.04	1.12	1.20	1.28	1.36
7	0.18	0.27	0.36	0.45	0.54	0.63	0.73	0.82	0.91	1.00	1.09	1.18	1.27	1.36	1.45	1.54
8	0.19	0.29	0.40	0.50	0.60	0.70	0.81	0.91	1.01	1.11	1.21	1.32	1.42	1.52	1.62	1.73
9	0.20	0.32	0.43	0.54	0.66	0.77	0.88	1.00	1.11	1.23	1.34	1.45	1.57	1.68	1.79	1.91
10	0.21	0.34	0.46	0.59	0.71	0.84	0.96	1.09	1.21	1.34	1.46	1.59	1.71	1.84	1.97	2.09
11	0.23	0.36	0.50	0.63	0.77	0.91	1.04	1.18	1.32	1.45	1.59	1.73	1.86	2.00	2.14	2.27
12	0.24	0.38	0.53	0.68	0.83	0.98	1.12	1.27	1.42	1.57	1.71	1.86	2.01	2.16	2.31	2.45
13	0.25	0.41	0.57	0.73	0.88	1.04	1.20	1.36	1.52	1.68	1.84	2.00	2.16	2.32	2.48	2.64
14	0.26	0.43	0.60	0.77	0.94	1.11	1.28	1.45	1.62	1.79	1.97	2.14	2.31	2.48	2.65	2.82
15	0.27	0.45	0.63	0.82	1.00	1.18	1.36	1.54	1.73	1.91	2.09	2.27	2.45	2.64	2.82	3.00

Table 3.4 Table of a block quantizer step size distribution. The values of the parameters are: $S_{0,0,0} = 0.1$, $S_{15,15,0} = 3.0$ and $c = 1.0$

3.4 Thresholds for Time Varying Images(2D-DCT)

3.4.1 Extension of the Results Obtained with Still Images

For time varying images, the previous study gives a good starting point since at the limit, a still image may be considered as the repetition of the same frame. However, when viewing a sequence, the subjective perceptibility of the added quantization noise is different and temporal masking plays an important role. Because experiments on time varying images are more time consuming and involve more computer resources (CPU time and storage), a subset of four sequences was used. These sequences are EIA, QUILT, OSCAR and ALBERT.

Preliminary studies indicated that the ranges $0.1 \leq S_{0,0,0} \leq 0.15$ and $S_{15,15,0} \leq 6$ were adequate. The values $S_{0,0,0} = 0.1$ and $S_{15,15,0} = 3$ were chosen.

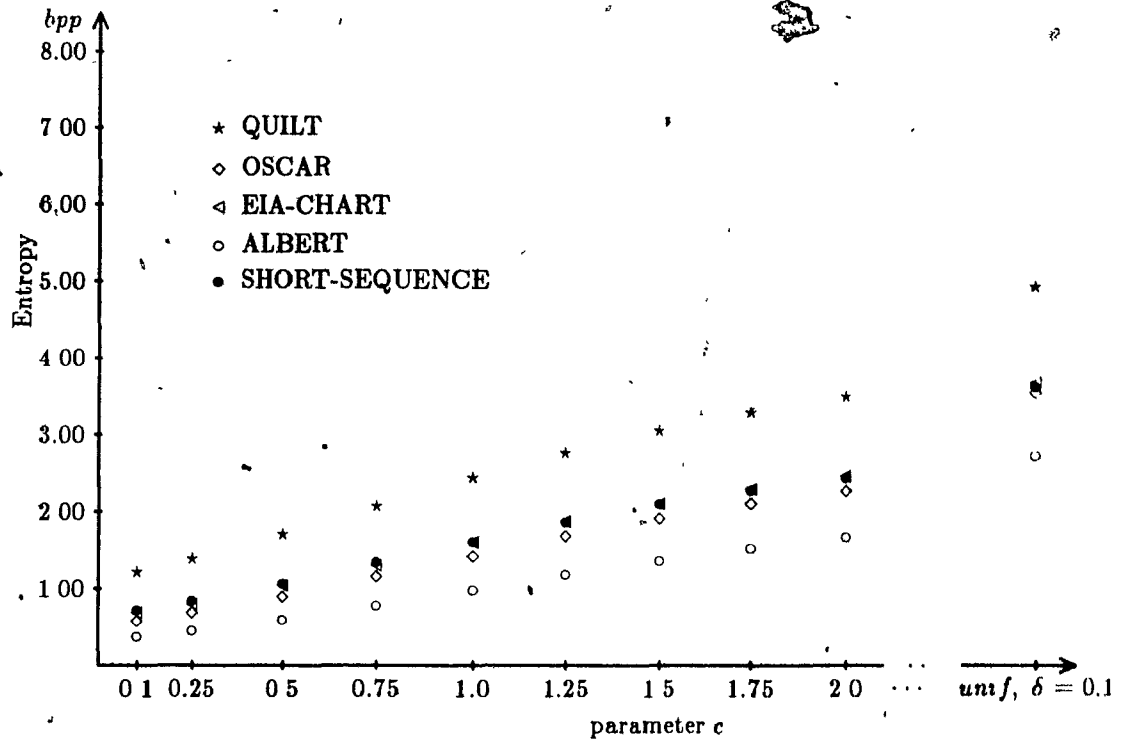


Fig. 3.6 Entropy vs c for the sequences ALBERT, EIA, OSCAR, QUILT, SHORT-SEQUENCE.

In that case, varying c from 2.0 to 0.1 gives a coded images ranking from 'no distortion perceptible' to 'annoying distortion'. When c is lowered, the degradation is smooth and graceful and the entropy decreases substantially. Fig. 3.6 shows the variation of the entropy for the four test sequences and SHORT-SEQUENCE which is composed of eleven frames from different images (Appendix A). The entropy is defined by

$$H_{avg} = \frac{1}{NMK} \sum_{u=0}^{N-1} \sum_{v=0}^{M-1} \sum_{w=0}^{K-1} H(u, v, w) \quad (3.5)$$

where

$$H(u, v, w) = - \sum_{l \in L_{u,v,w}} p_l \log_2(p_l) \quad (3.6)$$

p_l is the relative frequency of element l and $L_{u,v,w}$ is the set of possible values for the equivalent integer of $\hat{F}(u, v, w)$.

3.4.2 Subjective Tests

In the previous section, it was shown that a block quantizer can be completely specified by a single parameter c . Varying c has both an effect on the compression ratio and the level of distortion introduced by the coder. It is not claimed that the parametrical function or the parameter chosen are the best ones. However they are simple and give good results.

The objective in this section is to study the effect of varying c on the subjective judgement of a group of viewers concerning the quality of the coding. Procedures for subjective tests are outlined in report 405-4 [26] and recommendation 500 [27] of the CCIR and also by Sallio & Kretz [28] and Allnatt [29]. The EBU[†] method using the 5-grade impairment scale is best suited.

Setup

A population of 15 viewers participated in the tests, of whom 10 were experienced and 5 were inexperienced viewers. They were all male between 24 and 34 years of age. Although not tested for visual acuity, they all claimed to have a good vision. The population was divided into 5 groups of 3. Each group attended two viewing sessions of approximately 30 minutes each. The viewing distance was four times the screen height (about 1.10m). Three images were chosen for the tests, namely, ALBERT, EIA and QUILT. Due to the limitation in the duration time of a session, only eight levels of distortion can be tested. One level must correspond to the original (no distortion) for anchoring the results. The seven others must be chosen to range between a level corresponding to 5 in the impairment scale to a level corresponding to 1. The average level must be close to 3.

[†] European Broadcasting Union

A sequence is twelve frames of an image. A presentation is the test of a particular image and consists of the:

- display of the original sequence for 10 seconds,
- display of a neutral grey for 5 seconds,
- display of the coded sequence for 10 seconds,
- display of a neutral grey for 10 seconds (voting period).

The total time of a presentation is thus 35 seconds. A session consists of 18 presentations of the three test images for a total time of 31'10". The first two presentations are just examples and not accounted for. The next 16 presentations correspond to the actual test. Each level of distortion (including 'no distortion') is presented twice. The order of presentation is different for each image and each session. A pseudo-random order is necessary to eliminate adaptation effects. Details can be found in Appendix C. The following levels of distortion were selected:

d_1 correspond to no distortion at all

d_2 correspond to a block quantizer with $c = 2.0$

d_3 correspond to a block quantizer with $c = 1.5$

d_4 correspond to a block quantizer with $c = 1.0$

d_5 correspond to a block quantizer with $c = 0.75$

d_6 correspond to a block quantizer with $c = 0.50$

d_7 correspond to a block quantizer with $c = 0.25$

d_8 correspond to a bad coding, $S_{0,0,0} = S_{0,1,0} = S_{1,0,0} = 0.1$, $S_{u,v,0} = 3.0$ elsewhere

Appendix B gives the contour plots corresponding to these levels of distortion and the first frame of each coded image.

Results

Table 3.5 gives the mean rating and the standard deviation for each image. These results have been filtered according to the CCIR procedure ([26] [27]) to eliminate inconsistent opinions or inconsistent viewers.

Test	EIA		ALBERT		QUILT	
	mean	st dev	mean	st dev	mean	st dev
d_1	4.8	0.4	4.8	0.4	4.9	0.3
d_2	4.8	0.4	4.9	0.3	4.9	0.2
d_3	4.7	0.5	4.8	0.4	4.9	0.3
d_4	4.4	0.5	4.8	0.4	4.6	0.5
d_5	4.0	0.6	4.7	0.4	4.4	0.5
d_6	3.3	0.6	4.3	0.6	3.9	0.5
d_7	2.6	0.6	3.3	0.9	3.1	0.6
d_8	1.0	0.1	1.1	0.3	1.2	0.5

Table 3.5 Subjective opinion rating

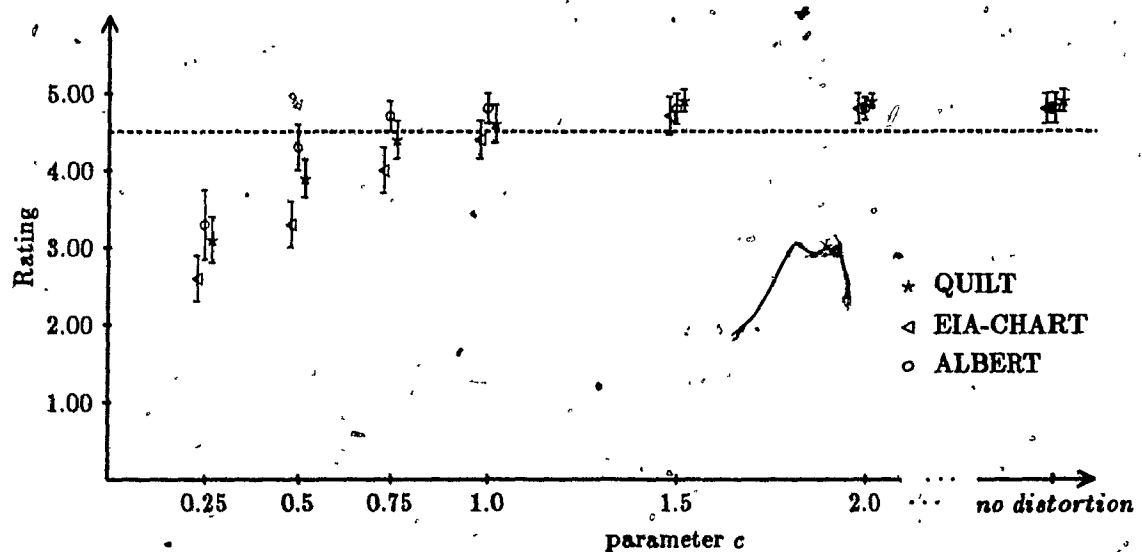


Fig. 3.7 Subjective rating vs c

3.4.3 Discussion

As expected, the level of perceptible distortion depends much on the image itself. The more critical the image, the higher is the value of the parameter c corresponding to the threshold of perception. The threshold is defined as the point where the distortion is imperceptible for 50% of the viewers and corresponds approximately to a mean subjective rating of 4.5 on the impairment scale.

For ALBERT, the threshold is around $c = 0.7$ whereas for QUILT it is a little more than $c = 1.0$. For EIA, the threshold is a little less than $c = 1.0$. For broadcast television, the complexity of the bulk of the images lies between the complexity of ALBERT and the complexity of QUILT and hence, the threshold will be between $c = 1.0$ and $c = 0.75$. For the rest of the study, it will be safely assumed that the threshold corresponds to $c = 1.0$.

An interesting observation concerns the rating of the original image with no distortion. Within the error interval represented by the standard deviation, it has the same rating as the levels of distortion corresponding to $c = 2.0$ and $c = 1.5$. This suggests that some (or all) viewers are very critical in their ratings. In general, ratings depend on the particular image but only to a limited extent. The graphs for the three images are pretty close.

There are no obvious correlations between the subjective ratings and the NSNR for the set of images used. For simple images with not much activity such as ALBERT, only a low level of RMS error can be tolerated before distortion becomes perceptible. For QUILT, distortion becomes perceptible at much higher values of the RMS error. This is probably due to the specific manner in which the error is introduced by the DCT. Other studies with different kinds of distortion [30] have found a logistic function that linked perceptual ratings to the RMS error.

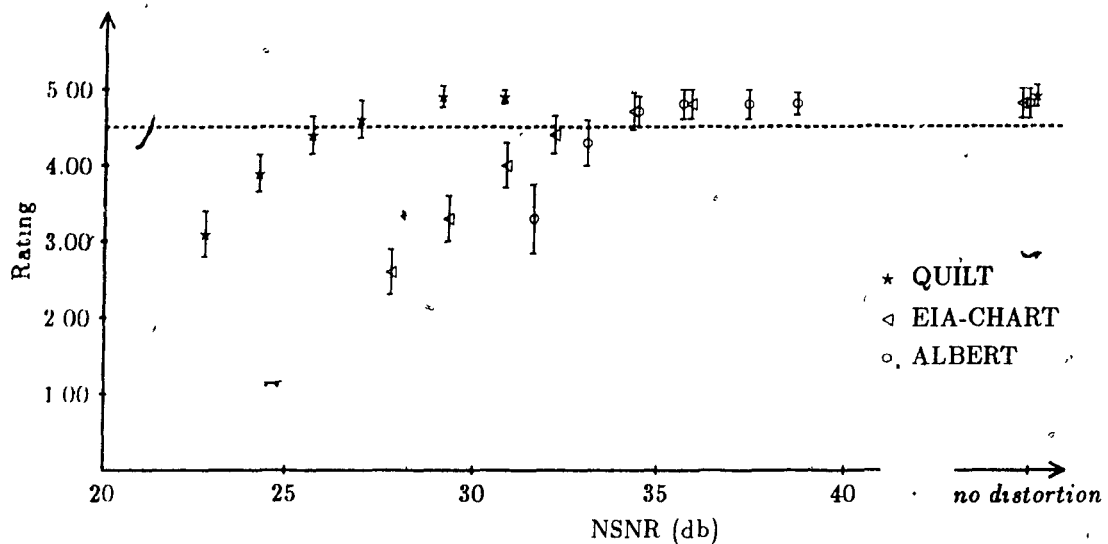


Fig. 3.8 Subjective rating vs NSNR

3.5 Non Uniform Quantization

A particular class of non-uniform quantizers is studied here, their main characteristic is that the step size is an increasing function of the input value. The idea behind this requirement is that a large value for a coefficient generally indicates a block with activity. Activity implies that the masking property of the eye tolerates more distortion. More quantization noise is introduced with wider step sizes for large input values. It is hoped that this will increase the compression ratio.

As seen earlier, $S_{u,v,w}^{(n+1)} = \min(\text{slope}_{u,v,w}^n, \text{saturation}_{u,v,w}) S_{u,v,w}$. In other words, until the saturation level is reached, the step size is $S_{u,v,w}^{(n+1)} = \text{slope}_{u,v,w} \times S_{u,v,w}^{(n)}$. As a starting basis, the block quantizer with $c = 1.0$ was chosen since it corresponds to the threshold of perception for uniform quantizers. In the first set of experiments, the saturation level was set to be very high corresponding in fact to no saturation at all. The slope was varied from 1.0 (uniform quantizer) to 2.0 for all quantizers except $Q_{0,0,0}$ which stayed uniform. It appeared that quality degraded very fast with an increasing slope without a substantial decrease in entropy. This

suggests that when a few high valued coefficients are very coarsely quantized, much distortion is introduced. A saturation level of around 3 seemed to cure this problem. Varying the slope between 1.0 and 1.5 did not introduce additional distortion. However the gain in entropy was very low, on average less than 5%. For a slope higher than 1.5, distortion began to appear gradually without much gain in compression ratio. A higher saturation level brings less than 1% decrease in entropy.

The conclusion is that non-uniform quantization does not permit appreciable bandwidth compression for the added complexity.

3.6 Extension to 3D-DCT

It was found in the preceding study that appreciable bandwidth compression can be achieved with a 2-D DCT. At the threshold of perception, according to the entropy and depending on the images, compression ratios from 3:1 (QUILT) to 10:1 (ALBERT) were achieved. By taking advantage of the temporal correlation, higher compression ratios can be expected. An extensive study of block quantizers for the 3-D DCT would take too much time to conduct in the framework of this thesis. Also it is unlikely that a simple function of one parameter can be found that will fit a large class of images and most importantly, types of motions. For the above mentioned reasons, the study of the 3-D DCT is not as thorough as for the 2-D DCT and is mainly done to get an idea of the gain over 2-D DCT that can be achieved.

3.6.1 Advantage over Two-Dimensional Processing

Let the block size be $N \times M \times K$. For the particular case of images with absolutely no motion, the first transform frame is given by:

$$F(u, v, 0) = \frac{1}{N} \frac{1}{M} \sum_{n=0}^{N-1} \sum_{m=0}^{M-1} \sum_{k=0}^{K-1} C(u)C(v)C(0) \cdot f(n, m, 0) \cdot \cos\left(\frac{\pi(2n+1)u}{2N}\right) \cos\left(\frac{\pi(2m+1)v}{2M}\right) \cos(0) \quad (3.7)$$

with $C(0) = 1$ and $C(l) = \sqrt{2}$ for $l \neq 0$.

This is the same as the 2-D DCT of one frame since $f(n, m, k) = f(n, m, 0)$ and

$$F(u, v, w) = 0 \text{ for } w > 0 \quad (3.8)$$

since there is no temporal energy.

The average entropy per coefficient is thus $H_3 = \frac{1}{K} H_2$, where H_3 is the average entropy when the image is processed with a 3-D DCT and H_2 with a 2-D DCT. For $K = 2$, the bandwidth is halved while the image quality remains unchanged. Higher values for K yield even more savings.

On the other hand, if the temporal correlation is null (successive frames are totally different and uncorrelated), then no savings are achieved by 3-D DCT processing over the 2-D DCT. In real life, the gain lies somewhere between these two extremes and depends to a great extent on the types of images processed. However it suffices to scan a few television channels to notice that a fair part of the material consists of scenes with fixed or slowly moving backgrounds.

3.6.2 Experiments

A block size of $16 \times 16 \times 4$ was chosen and statistics on the real valued transform coefficients were gathered. The goal is to find a block quantizer that yields images comparable to those obtained at the threshold of perception with two-dimensional processing. A transform block can be seen as four frames of 16×16 coefficients. Each frame is specified independently of others. The parametrical function developed

earlier is used to generate the quantization step distribution for a frame. A block quantizer is thus completely specified by the values c_0 , c_1 , c_2 and c_3 that are the parameters for the frames.

In addition to the distortion present with two-dimensional processing, three-dimensional processing introduces 'temporal' distortion. If temporal coefficients are too coarsely quantized, motion is not well reconstructed in the sense that moving object may be displaced with respect to the original frame. This phenomenon is easily perceptible when viewing the images frame by frame. In real time display, the temporal masking property of the human visual system hides much of the distortion.

Since thorough experiments with many types of motion are limited by the computer resources (and time also), the values of c_0 , c_1 , c_2 and c_3 chosen are approximate. They give results comparable to the threshold of perception and a margin was allowed to account for faster types of motion (temporal coefficients are quite finely quantized). It should be stressed that these experiment are really meant to see if the gain with respect to 2D DCT is significant or not. The values chosen are: $c_0 = 2.0$ and $c_1 = c_2 = c_3 = 1.0$.

Chapter 4 Coding of the DCT Coefficients

In the previous chapter, it was shown that appreciable bandwidth compression can be achieved by the use of DCT coding and proper design of the block quantizer. An approximation to the lower bound of the average number of bits per pels (or bits per coefficient) was given by the entropy H_{avg} . It was found that the entropy depended both on the image itself and on the particular block quantizer used. In practice however, the compression ratio will depend on the technique used to code the quantized coefficients. To design an efficient block coder, a thorough knowledge of the source [†] statistics is needed.

The coders used here are inspired by that of Chen and Pratt [7], namely, they use Huffman coding for coding non zero coefficients and a combination of run-length and Huffman coding for coding zero valued coefficients.

The concepts of Huffman coding and run-length coding are introduced and followed by a brief look at the quantized block characteristics. Different coders are then simulated and the performances compared. The chapter concludes with a look at some 'post filtering' that can be embedded in the coders.

[†] The source consists of the blocks of coefficients.

4.1 A Brief Review of Coding Theory, Huffman and Run-Length Coding

In a general context, let S be a source of cardinality M with source symbols s_i , $i = 1, \dots, M$ having probabilities p_i . It is assumed to be stationary. Let A be a code alphabet of cardinality D and source symbols a_i . The source emits a sequence $\{s_{j_0} s_{j_1} \dots s_{j_n} \dots\}$. A coder is a mapping from a sequence of source symbols to a sequence of code symbols. A first order coder is a mapping of each of the source symbol s_i into a sequence c_i of code letters called a code word

$$s_i \mapsto c_i = \{a_{i1} a_{i2} \dots a_{il_i}\} \quad (4.1)$$

where l_i is the length of the codeword c_i . In a vector coder, codewords are assigned to vectors which are a finite length sequence of source symbols.

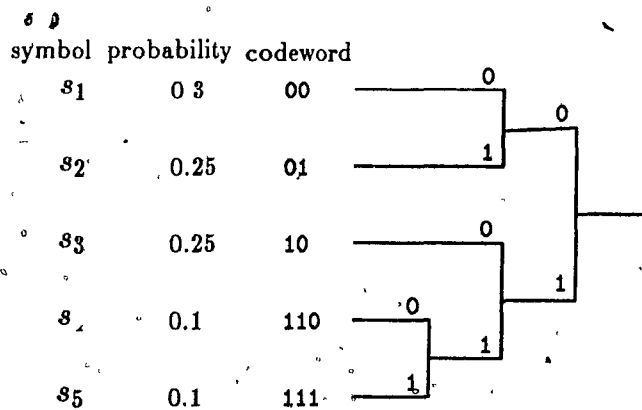
$$\{s_{j_0} s_{j_1} \dots s_{j_n}\} \mapsto \{c_{j_0} c_{j_1} \dots c_{j_n}\} \mapsto \{a_{j_0 1} \dots a_{j_0 l_0} \dots a_{j_n k}\} \quad (4.2)$$

There is no indication in the code symbol sequence where a codeword begins and ends. This information must be determined from the structure of the code itself, assuming however that the start of the code sequence is known.

In practice, the cardinality of the code alphabet is two, $\{0, 1\}$. The efficiency of a coder is measured by its average word length $\bar{l} = \sum_{i=1}^M p_i l_i$, $\bar{l}$ has the same dimension as the entropy and is measured in bits per symbol. Among first order coders, two classes are widely used, namely, fixed length coders and prefix variable length coders. Fixed length coders are simple but give poor results when source probabilities are not uniform. Prefix codes (also called tree codes) are less simple but perform better. These coders will always have a code length of at least one bit per source symbol. For sources with lower entropy, vector coding should be considered[31].

4.1.1 Huffman Coding

For a given source, it is not always possible to find a first order coder whose average word length equals the source entropy. However the Huffman code is optimum in the sense that no other first order coder can outperform it. The Huffman code is tree like, or prefix code. Each codeword is represented by a leaf in the tree and is uniquely decodable. The code is built from the source symbol probabilities. The length of a codeword is given by the number of nodes between the leaf and the root of the tree. Short codewords are assigned to most probable elements. An example is given below for a five elements source. The value of the entropy is 2.1855 while the average codeword length is 2.2.



4.1.2 Run-Length Coding

Run-length coding is primarily used for the transmission of binary images. The document is scanned horizontally resulting in alternating sequences of black and white pels. Compression is achieved by transmitting the values of the length of the alternating sequences rather than the values of the pels themselves. The source then becomes the possible values for the sequence lengths. A first order coder can further be used for coding the elements of this new source. This technique is efficient when sequences are long, as is usually the case for typed documents.

4.1.3 Discussion on the Estimation of the Entropy of Quantized Blocks.

In section 3.4.1, the average entropy was defined as the mean of the entropy of the individual coefficients (Eq. 3.5, 3.6). In other words, a block was considered as NMK different sources. However in the context of DCT coding, a block should be considered as a unique entity. The entropy of a block would be H_B defined as

$$H_B = - \sum_{b \in B} p_b \text{Log}_2(p_b) \quad (4.3)$$

where B is the set of all possible blocks. An example of a block is given in Table 4.1. From a practical viewpoint, the set B would be so large that any computation of H_B would be virtually impossible. In the rest of the thesis, 'entropy' will actually mean H_{avg} , which is more tractable. It should however be noted that it is only an approximation and $H_{avg} \geq \frac{1}{NMK} H_B$. The equality is achieved if the transform coefficients are independent.

4.2 Some Observations on the Quantized Coefficients

Table 4.1 shows a typical block. It is taken from the first frame of image ALBERT and covers about two thirds of his right eye.

A striking observation is that nearly half of the coefficients are zero valued. Moreover, they form a rather uniform big cluster. Since ALBERT moves his head while talking, 'vertical' coefficients have quite large values compared to 'horizontal' ones. This is due to the line interlace effect of frame coding. Other sample blocks from different parts of the image, and different images, reveal the same structure, namely, many zeros. In the uniform background in ALBERT only a few coefficients are not zeros. On the other hand, there are less zeros in QUILT because of the important high frequency content.

$u \backslash h$	0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
0	768	112	9	4	1	1	5	0	5	-1	1	0	0	0	0	0
1	94	-1	-2	2	-4	1	-2	-1	-3	1	-2	0	0	0	0	0
2	127	18	-12	-6	-2	-1	0	0	1	0	0	0	0	0	0	0
3	-17	-27	9	0	4	0	0	-1	0	0	0	0	0	0	0	0
4	95	-7	-7	-7	-7	-1	-1	0	-1	0	0	0	0	0	0	0
5	-56	-11	-1	7	3	2	1	0	1	0	0	0	0	0	0	0
6	27	16	2	-5	-3	-1	-1	0	-1	0	0	0	0	0	0	0
7	-4	-9	-5	2	2	1	0	0	0	0	0	0	0	0	0	0
8	-1	1	2	0	0	0	0	0	0	0	0	0	0	0	0	0
9	6	3	1	-1	-2	-1	0	0	0	0	0	0	0	0	0	0
10	-11	-5	0	2	1	0	0	0	0	0	0	0	0	0	0	0
11	12	3	0	-2	-1	-1	0	0	0	0	0	0	0	0	0	0
12	-11	-1	1	1	1	0	0	0	0	0	0	0	0	0	0	0
13	1	2	0	-1	0	0	0	0	0	0	0	0	0	0	0	0
14	-6	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0
15	-3	-1	1	0	0	0	0	0	0	0	0	0	0	0	0	0

Table 4.1 A typical quantized transform coefficient bloc. Taken from ALBERT in the region of his right eye.
Parameter $c = 1.0$.

The second noteworthy observation concerns the probability density function (PDF) of individual coefficients. Figs 4.1 and 4.2 show that the shape of the PDF differs depending on the position of the coefficient within a block. The distribution gets narrower with increasing spatial frequency. Also the probability that the coefficient is zero increases quickly, from around 20% for $F(2,2,0)$ to nearly 90% for $F(7,7,0)$. This explains the large number of zeros in a block. Looking at the PDF without taking into account the value zero shows that the symbols have very different probabilities. Note that the coefficient $F(0,0,0)$, being the average luminance value of the pels in a block, has a more uniform PDF (not shown here). These statistics are from LONG-SEQUENCE quantized with a block quantizer at the threshold of perceptibility ($c = 1.0$). Since this image really consists of a variety of images, these results are quite general and robust.

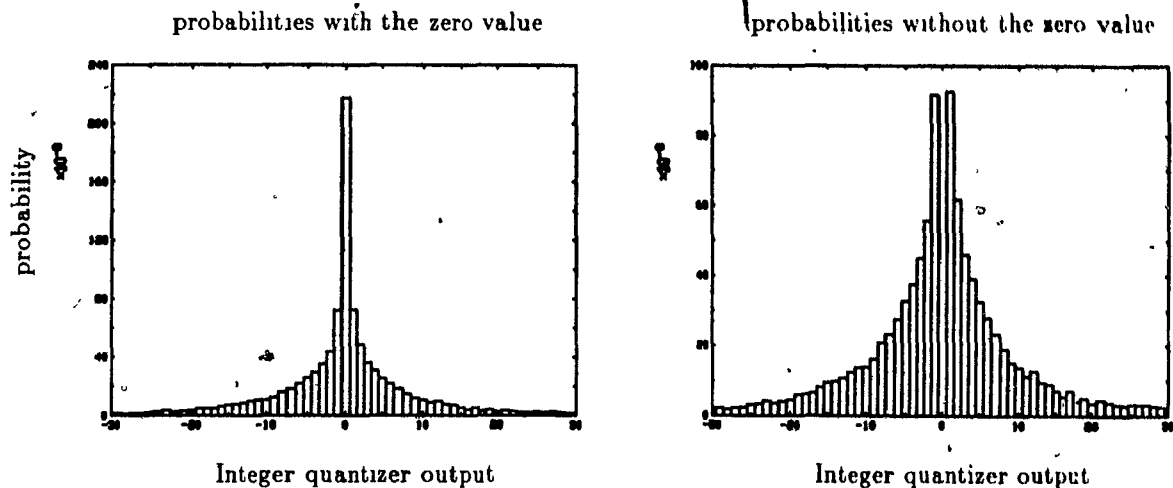


Fig. 4.1 Probability distribution functions for $F(2, 2, 0)$ quantized, with, and without, the zero value. Different vertical scales used.

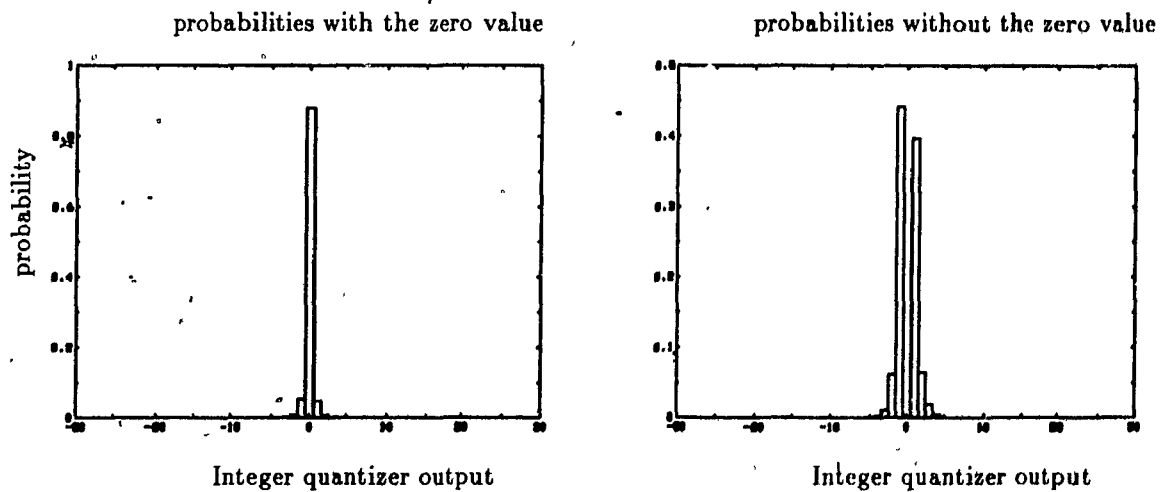


Fig. 4.2 Probability distribution functions for $F(7, 7, 0)$ quantized, with, and without, the zero value. Different vertical scales used.

4.3 Dual Huffman and Run-Length Coders

The design of a straight block coder, namely, a mapping of every possible block

to a codeword, is very impractical due to the high cardinality of the source. An efficient alternative using Huffman and run-length coding is presented.

The block is mapped into a one-dimensional vector by the use of a scanning pattern. The vector thus consists of NMK coefficients in a known order. Taking the example of Table 4.1 and using a zigzag scanning pattern starting from $F(0,0,0)$ would result in:

```

      768 112 94 127 -1 9 4 -2 18 -17 95 . . .
      . . . 1 1 0 -2 0 0 0 1 -1 2 0 1 0 . . .
      . . . 0 0 1 0 0 . . . . . 0 0 0 0 0

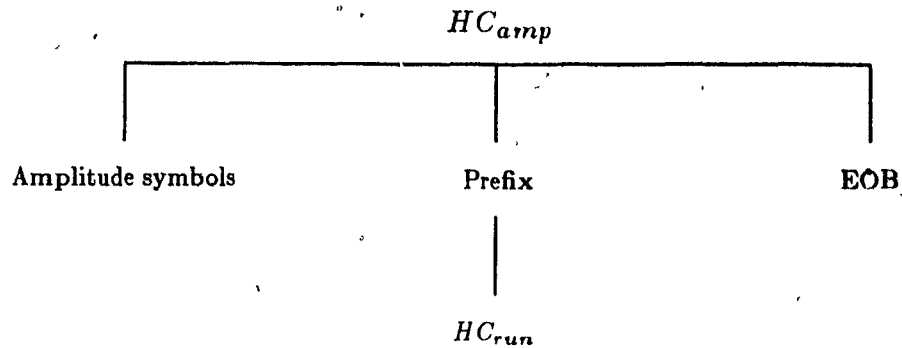
```

The first coefficients have a rather large variance and are rarely zero. Gradually, the variance decreases and long successions of zeros start to appear. The last third of the vector consists of zeros only.

To code the non-zero amplitudes, a Huffman code HC_{amp} is used. Another Huffman code HC_{run} is used to code the runs of consecutive zeros. The following sequence:

3 4 0 0 2 0 0 0 0 0 0 -1 0 0 0 0 0 5 0 . . . 0 0 0 would be coded $HC_{amp}(3)$
 $HC_{amp}(4)$ $HC_{run}(2)$ $HC_{amp}(2)$ $HC_{run}(7)$ $HC_{amp}(-1)$ $HC_{run}(5)$ $HC_{amp}(5)$ EOB.
 Knowing the length of a run of zeros permit one to compute the position of the next non-zero coefficient. A special codeword, EOB (End Of Block), is used to say that all of the remaining coefficients are zero. EOB reduces the occurrence of long runs.

The statistics used to design the two previous coders are gathered as follows: A histogram of non-zero amplitudes is generated with two additional source symbols, EOB and *prefix*. Prefix designates the start of a run of zeros. The HC_{amp} code is generated according to these statistics. HC_{run} is generated from the histogram of the length of zero sequences, excluding the one going to the end of block. In the actual coder HC_{run} is a subcode of HC_{amp} with the prefix code *prefix* (ie, the prefix leaf designates that what follows represent the code of a run of zeros).



This coder will be called a *separated* dual coder since HC_{amp} and HC_{run} are really two separate Huffman codes. It follows the idea of Chen and Pratt in their 'scene adaptive coder' [7].

Another method consists of having an *integrated* dual coder. In this case, the source consists of the non-zero amplitudes, the runs of zeros and the EOB symbols. $\{\dots -3, -2, -1, 1, 2, 3, \dots, 1_z, 2_z, 3_z, \dots, EOB\}$ where l_z designates a run of l zeros. The two previously mentioned histograms are merged to produce a single one that is used to generate an integrated Huffman code. This method was also used by Dubois and Moncet for the coding of NTSC pictures [32].

Note that in both the separated and integrated coders, the $F(0,0,0)$ coefficient is treated differently. It is always the first element in the scanning pattern and is assigned a fixed length codeword of twelve bits. It is not taken into account when generating the histograms.

The choice of the scanning pattern is important. One would want to have long runs of zeros. Since zeros are coded in groups, it is possible in theory to get bit rates less than one bit per coefficient.

Block Coding With NMK Coders

This represents an extension of the ideas presented above. Each coefficient has its own set of amplitude and run-length histograms. The amplitude histogram is computed from the values taken only by that particular coefficient. The runs are

those starting at the coefficient in the scanning process. Both separate or integrated codes can be specified for each coefficient except for $F(0,0,0)$. In this coder, the codes are adapted to each coefficient. This arrangement will hopefully yield lower bit rates at the cost of a huge increase in memory requirements and to a lesser extent, complexity. It was however found that the gain was negligible for the effort deployed.

4.4 Results

In this section, three factors affecting the efficiency of a dual Huffman coder are studied:

- The type of coder, namely, integrated or separated.
- The scanning pattern.
- The robustness of the coder.

Robustness

As seen previously, a coder is generated from histograms generated for a particular image. This coder will then be optimal (or adapted) to that image and may be poor for another image. The robustness is a measure of the sensitivity of a coder to different source images. It shows how a particular coder performs with respect to one that is specifically adapted to the source image. In order to test the robustness of the dual Huffman coder, two coders designed from different statistics are tried on a set of images, Table 4.2. The 'general' coder draws its statistics from LONG-SEQUENCE. For each image, an 'individual' coder is generated from the statistics of the particular image only. The general coder will be said robust if it performs relatively close to the individual coders on average.

Scanning Pattern

Due to the fact that the energy is concentrated in low frequency coefficients, the zigzag scanning pattern performs well. A more general method would be to scan the coefficients in a decreasing order of weighted variance. The weighted variance being the variance of the coefficient divided by the step-size of its quantizer. The idea is that it would allow for less runs of zeros and a more efficient use of the EOB codeword.

Results

All experiments were conducted using a block quantizer with the parameter $c = 1.0$ corresponding to the threshold of perception. Table 4.2 summarizes the results.

Image	Entropy H_{avg}	Individual Coders	Zigzag scanning		Variance scanning	
			integrated	separated	integrated	separated
LONG-SEQ	1.541		1.618	1.693	1.641	1.702
ALBERT	0.987	1.058	1.056	1.092	1.072	1.099
EIA CHART	1.607	1.771	1.767	1.852	1.804	1.870
FLOWERS	1.410	1.576	1.573	1.642	1.602	1.658
TOYS	1.627	1.777	1.791	1.882	1.841	1.917
LETTERS	1.407	1.630	1.635	1.703	1.662	1.710
MEISEL	1.124	1.396	1.395	1.462	1.388	1.450
OSCAR	1.433	1.595	1.592	1.663	1.634	1.690
QUILT	2.455	2.948	2.952	3.111	2.928	3.087
WAGONS	1.338	1.576	1.580	1.646	1.601	1.647
YVON	0.814	0.843	0.843	0.875	0.875	0.896

Table 4.2 Bit rates obtained in *bits per pels* for different coders. The individual coders are 'integrated' and use a zigzag scanning pattern.

As expected, the integrated coder perform better than the separated one yielding a gain of about 4.5% for the zigzag scanning pattern. The zigzag scanning pattern is quite good since it actually outperforms the weighted variance scanning pattern by about 1.5%, which is not significant. The robustness of the coders seem adequate for most images. Except for TOYS which gave a difference of 0.78%, the difference was less than 0.3 % which can be considered negligible. Compared to the average entropy, the bit rate obtained are acceptable. The performance depends much on the image. YVON yielded a bit rate only 3.4% higher than its entropy while MEISEL was 19.5% above entropy for integrated coders. On average, the bit rate is around 10% higher than the entropy. Appendix E gives table of entropies for the coefficients and the average number of bits taken to code each coefficient for the different coders. The average number of bits is solely based on the non-zero values taken by the coefficient. In other words, the share of bits to code its zero values is not taken into account. For that reason, the average number of bit for a coefficient cannot directly be compared to its entropy.

4.5 Post Filtering

Within the clusters of zeros in a block of quantized coefficients, there are often isolated coefficients with non-zero amplitude. Most of the time their amplitude is ± 1 or ± 2 . These coefficients are usually medium or high frequency ones and are not important in the viewpoint of perceptibility. Forcing these isolated coefficients to zero will yield fewer and longer runs of zeros, thus lowering the bit rate. Formally, a coefficient will be said to be isolated if in the scanning pattern, the preceding and the following coefficients are both zero valued.

Furthermore, experiments showed that skipping the isolated ± 1 did not introduce any additional distortion if any was present. In fact no differences were

perceptible except that the filtered image appeared very slightly less noisy than its non filtered counterpart. A noticeable reduction in the bit rate resulted however. Skipping also the ± 2 isolated coefficients had no effect and a negligible bit rate saving. This suggest that the majority of isolated coefficient have an amplitude of ± 1 . Post filtering can be viewed as a type of quantization with memory. In that case, the threshold is dependent on the value of adjacent pels. The results obtained are summarized in Table 4.3.

Image	Entropy H_{avg}	No post filtering		With post filtering	
		Zigzag scanning	Variance scanning	Zigzag scanning	Variance scanning
LONG-SEQ	1.541	1.618	1.641	1.416	1.394
ALBERT	0.987	1.056	1.072	0.859	0.838
EIA CHART	1.607	1.767	1.804	1.552	1.548
FLOWERS	1.410	1.573	1.602	1.353	1.318
TOYS	1.627	1.791	1.841	1.606	1.616
LETTERS	1.407	1.635	1.662	1.454	1.431
MEISEL	1.124	1.395	1.388	1.220	1.190
OSCAR	1.433	1.592	1.635	1.356	1.330
QUILT	2.455	2.952	2.928	2.669	2.617
WAGONS	1.338	1.580	1.601	1.383	1.353
YVON	0.814	0.843	0.875	0.706	0.697

Table 4.3 Influence of skipping isolated (± 1) coefficients on the bit rate. Comparison of scanning patterns for the coders with post-filtering.

A zigzag scanning pattern and an integrated coder were used as the reference. A coder with post filtering is not an entropy coder since the reconstructed block of coefficients at the receiver may be different from the one at the transmitter. This explains that bit rates are sometime below entropy. The gain with respect to coders with no post-filtering is appreciable. Comparing integrated coders with

zigzag scanning pattern show a reduction in the bit rate ranging from 9.6% for QUILT to 18.6% for ALBERT. The LONG-SEQUENCE gave a gain of 12.5%. In the case of post-filtering coders, the variance scanning pattern seems to give slightly better results than the zigzag pattern. The difference is however quite small. Since they do not seem to affect the image quality, coders with post-filtering are thus considered best. For LONG-SEQUENCE, the bit rate obtained of 1.394 bpp using a variance scanning pattern is equivalent to a compression ratio of 5.7.

4.6 Extension to 3D DCT

In this section, the block quantizer described in Section 3.6 is used. The quality of the images obtained is similar with those obtained in the preceding study. However since two-dimensional and three-dimensional DCT do not produce exactly the same impairments in the viewpoint of quality, results cannot be directly compared.

Image	2D DCT		3D DCT	
	No post F.	With post F.	No post F.	With post F.
LONG-SEQ	1.618	1.416	1.107	0.891
ALBERT	1.056	0.859	0.647	0.466
EIA CHART	1.767	1.552	1.245	1.004
FLOWERS	1.573	1.353	1.170	0.926
TOYS	1.791	1.606	0.935	0.752
LETTERS	1.635	1.454	1.450	1.257
MEISEL	1.395	1.220	0.914	0.718
OSCAR	1.592	1.356	0.976	0.734
QUILT	2.952	2.669	2.013	1.683
WAGONS	1.580	1.383	1.242	1.030
YVON	0.843	0.706	0.479	0.344

Table 4.4 Comparison of 2D coding and 3D coding for images of approximately same quality.

An integrated coder with a weighted variance/scanning pattern is used. It is compared to its 'post filtering' counterpart as well as to the two-dimensional coder.

The gain over two dimensional precessing is important and as expected depends on the motion. For example TOYS which is mainly a still picture shows an impressive gain. ALBERT and YVON which exhibit a fixed background also benefit from three-dimensional processing reaching bit rates below 0.5 bits per pel. On average with LONG-SEQUENCE, the bit rate is close to 1.1 bits per pel and post filtering allows a further gain of 20% yielding a bit rate of 0.9 bits per pels.

Chapter 5

Conclusion

This study was aimed at investigating a block quantization method for discrete cosine transform image coding that is adapted to the properties of the human visual system. In the method developed, a block quantizer is generated by a function of a small set of parameters. Individual quantizers have a uniform step size. The parametric function allows, for a given distortion level, to determine a step size for each coefficient based on its overall subjective importance.

A subjective test demonstrated the validity of this approach. Overall image distortion was successfully controlled by one parameter for a set of three images exhibiting various features. The opinion of the viewers correlated well with the parameter value. The value of the parameter at which the distortion is at the threshold of perceptibility depends on the particular image and its 'activity'. However the values for the three test images were close.

The variable rate block coder used here is similar the one developed by Chen and Pratt. Two enhancements have been added, namely, the use of an integrated Huffman coder and post-filtering. These two enhancement permit a further reduction of the bit rate of about 20% with respect to the original method.

For intraframe coding, an average bit rate of 1.4 bits per pel was obtained. However the bit rate depends on the particular image and varied from a high of 2.6 bits per pel for QUILT to a low of 0.7 bits per pel for YVON.

The use of non-uniform quantizers, for which the step size increases with the value of the transform coefficient, was investigated to determine the further reduction in bit rate possible over uniform quantization. However, the gain achieved was not significant enough to justify the added complexity. This is because most coefficients are low valued and only a few are affected by non-uniform quantization.

A better bandwidth compression was obtained with interframe three-dimensional DCT coding. However in that case, the gain over intraframe coding depends on the motion present in the image or more generally its temporal activity. An average bit rate of 0.9 bits per pel was achieved, with a low of 0.35 bits per pel for a head and shoulder scene typical of teleconferencing. The choice of the scanning pattern is more delicate since the energy distribution among the coefficients depends greatly upon the motion present in the image. Further psychovisual experiments are certainly needed to optimize the three-dimensional block quantizers and to study the effect of both the motion present in the images and the temporal size of the blocks.

Adaptive block quantization can be implemented with two potential uses. One may want to adapt the output bit rate to a certain channel bandwidth, or in other words regulate the output bit rate. In that case a simple feedback mechanism derived from some buffer status can be used to control the block quantizer (increase or decrease the parameter 'c' in Eq. 3.4). Another possible use would be to minimize the bit rate while keeping the quality of the image above the threshold of perception. In that case, one will need to do an extensive study to separate images into perceptual classes and, according to some measure of activity, select the most appropriate block quantizer. The overhead added by adaptivity would be minimal in terms of transmission since only one or two descriptive values of the block quantizer should be sent. Computational complexity however would increase since the quantizers should be constantly updated. A simple method would be to have a bank of block quantizers in memory and switch to the appropriate one as needed.

With decreasing hardware cost, this would lower processing overhead.

Appendices

Appendix A. Test Images

This appendix presents the first frame of the various sequences used for the test. A sequence comprises usually twelve frames of the same scene except for SHORT-SEQUENCE and LONG-SEQUENCE.

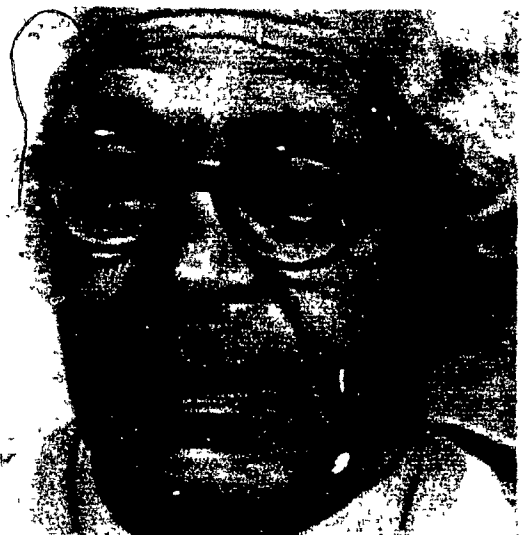
SHORT-SEQUENCE comprises eleven frames: EIA CHART, YVON, FLOWERS, OSCAR MVT, GASTON, TOYS, QUILT, WAGONS, OSCAR, ANNABELLE ALBERT,

LONG-SEQUENCE comprises 144 frames, twelve each from the scenes: YVON, EIA CHART, ALBERT, WAGONS, FLOWERS, QUILT, OSCAR, MEISEL, LETTERS, TOYS, ANNABELLE, OSCAR MVT,

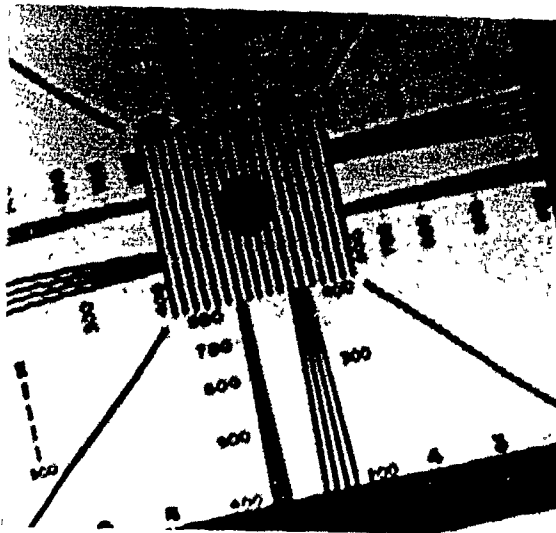
The actual size of an image on the display screen is $136\text{mm} \times 124\text{mm}$. For the subjective test, images are viewed from a distance of approximately one meter which corresponds to four times the picture height of the display screen.



ALBERT



MEISEL



EIA CHART



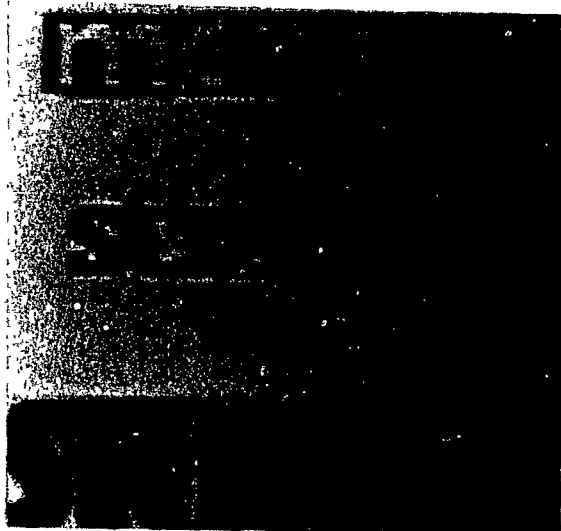
WAGONS



FLOWERS



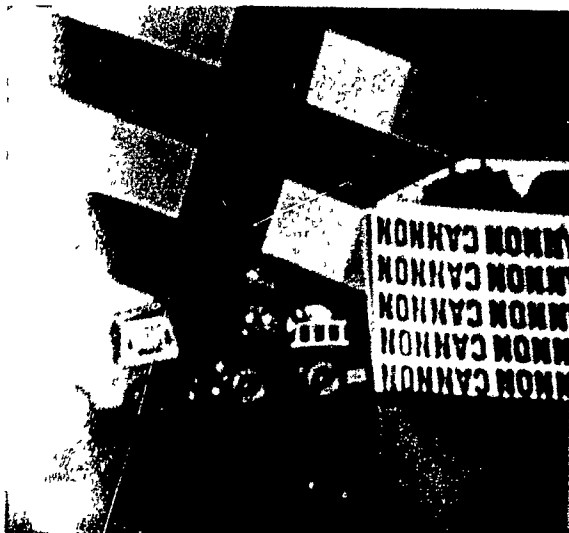
OSCAR



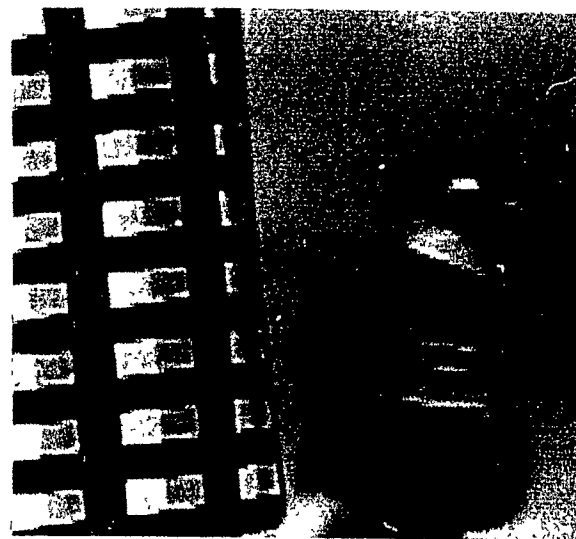
LETTERS



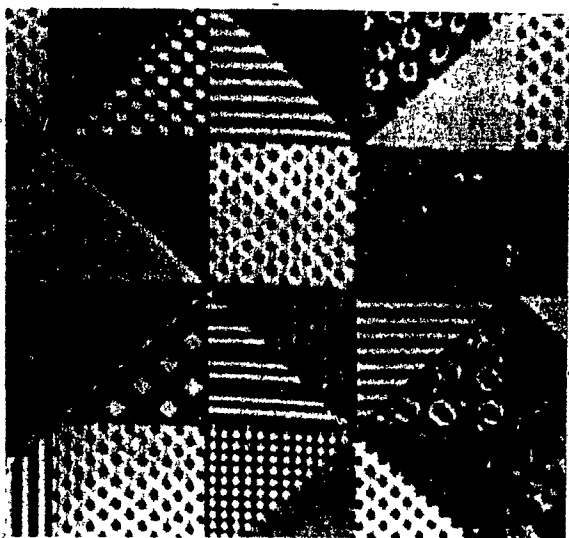
ANNABELLE



TOYS



OSCAR MVT



QUILT



YVON



GASTON

Appendix B. Subjective Tests, pictures and contour plots

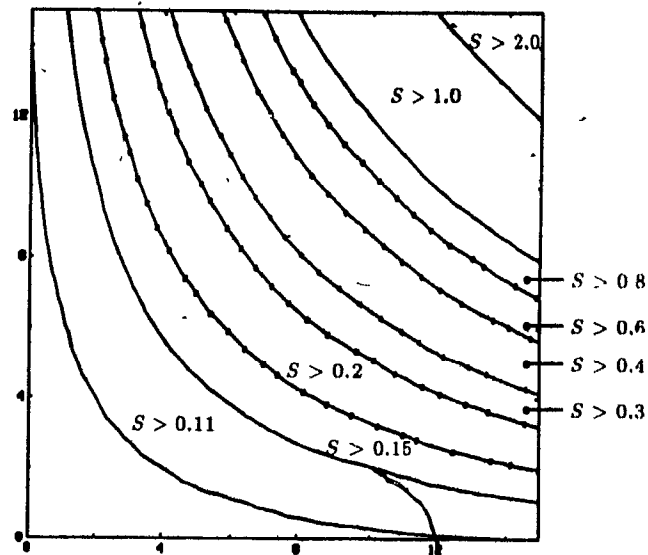


Fig. B.1 Level of distortion d_2 corresponding to $c = 2.0$

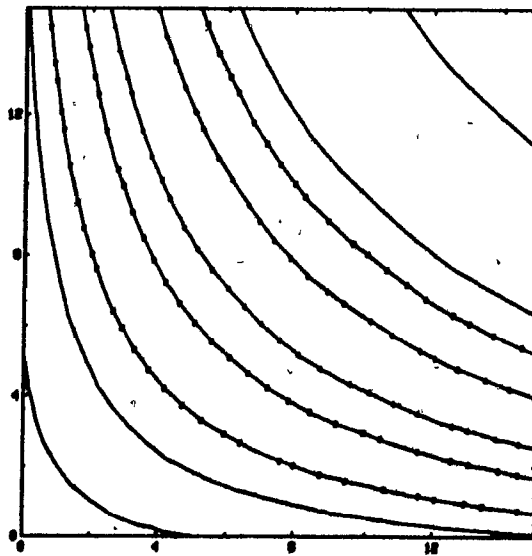


Fig. B.2 Level of distortion d_3 corresponding to $c = 1.5$

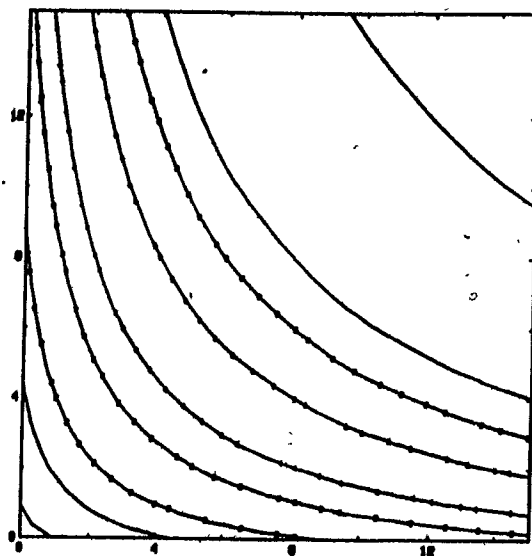


Fig. B.3 Level of distortion d_4 corresponding to $c = 1.0$

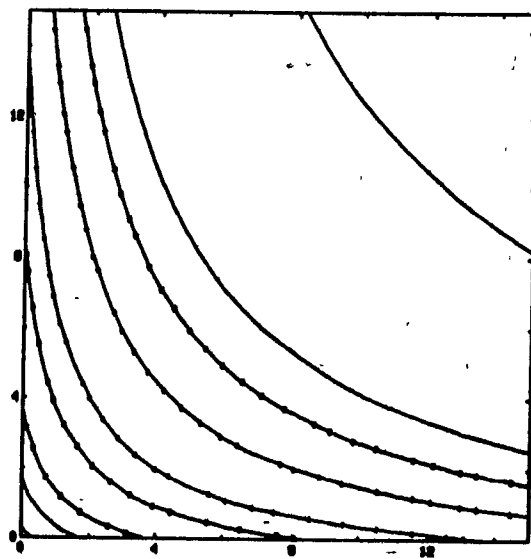


Fig. B.4 Level of distortion d_5 corresponding to $c = 0.75$

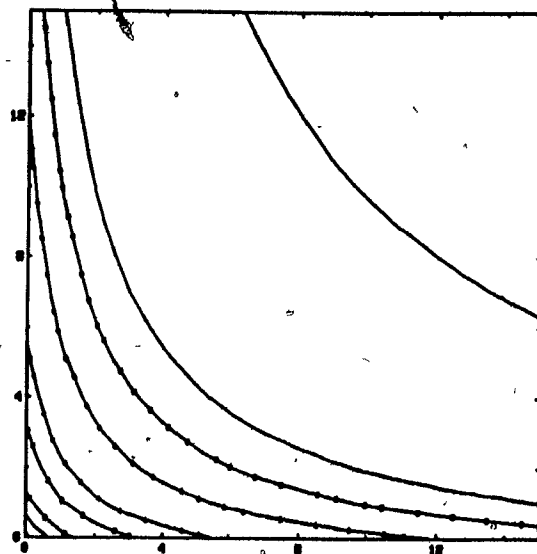


Fig. B.5 Level of distortion d_6 corresponding to $c = 0.5$

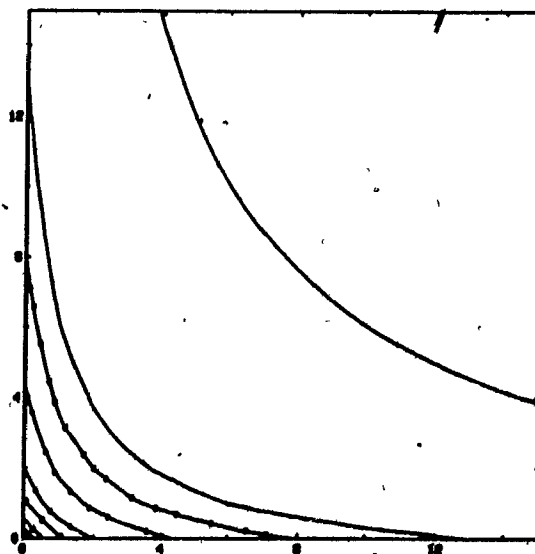
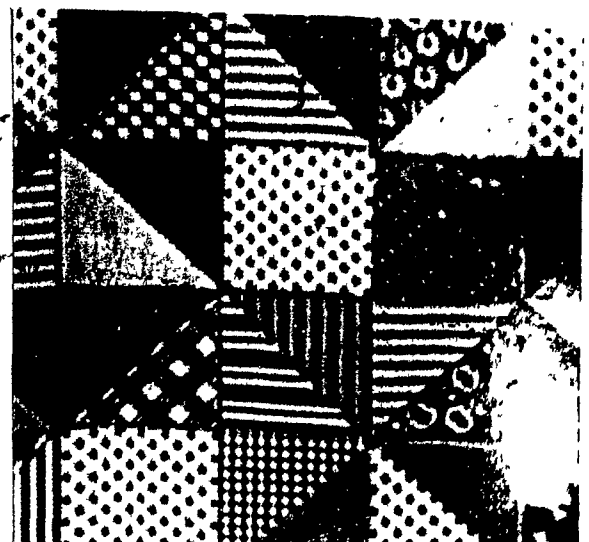
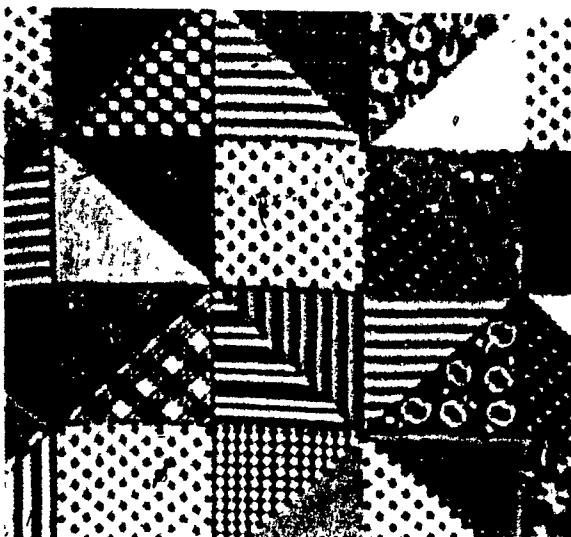
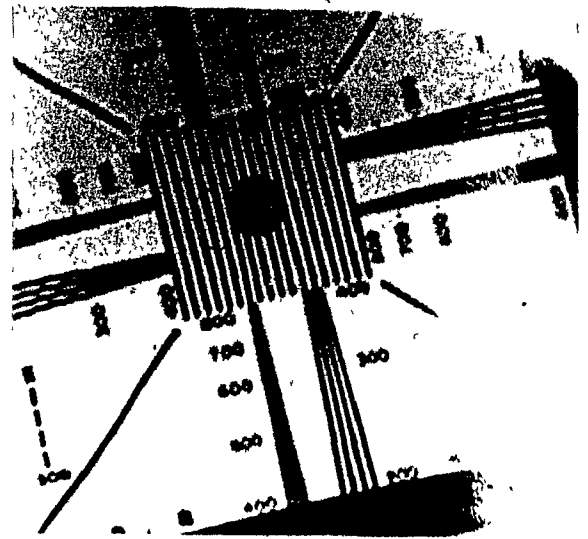
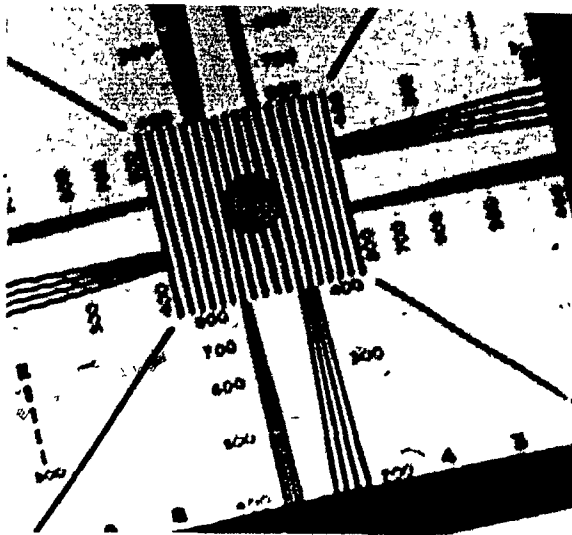
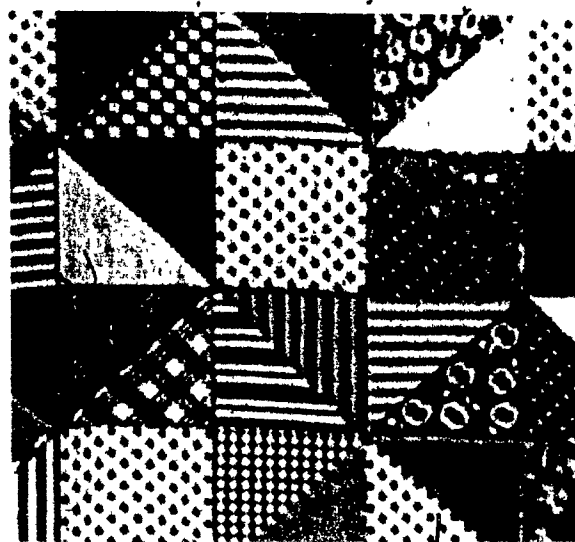
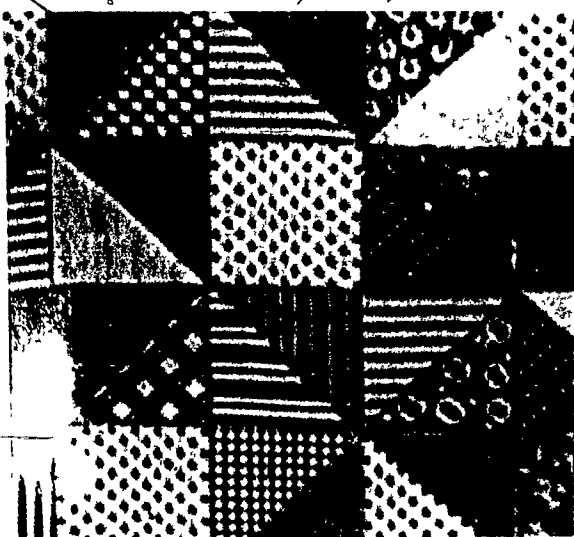
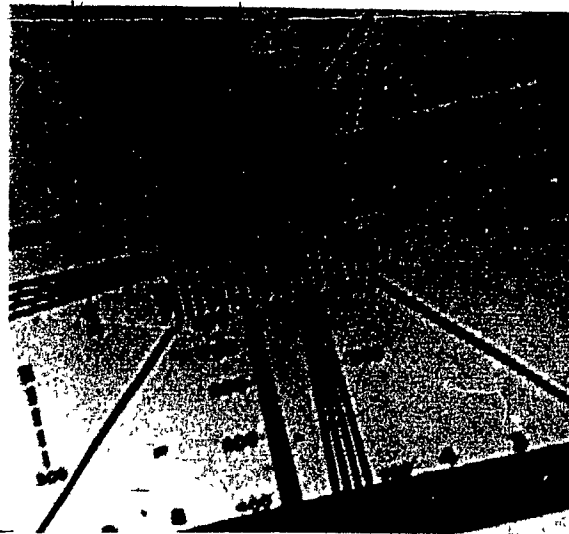
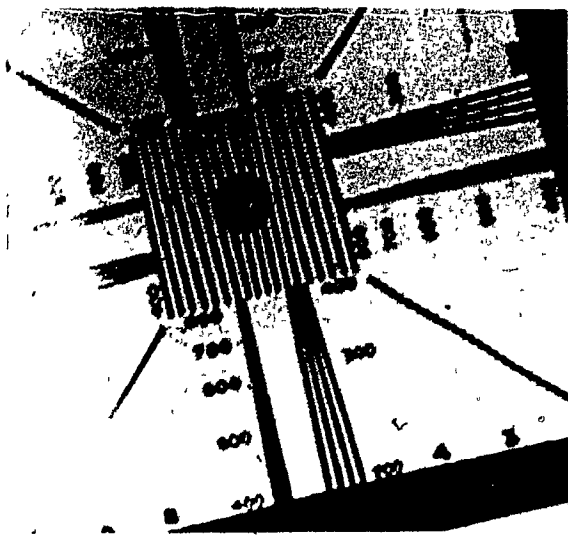


Fig. B.6 Level of distortion d_7 corresponding to $c = 0.25$



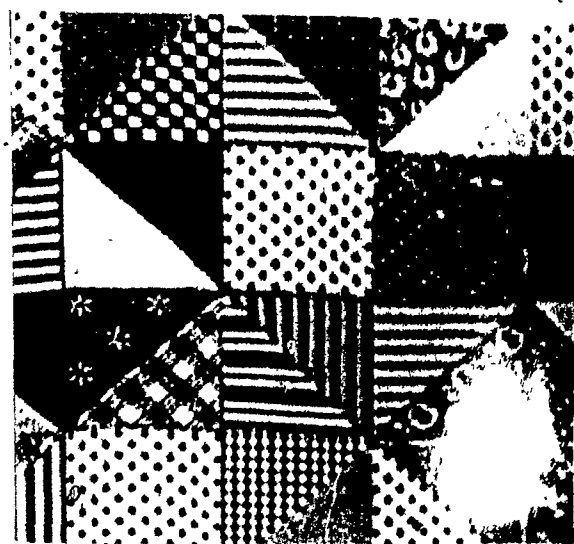
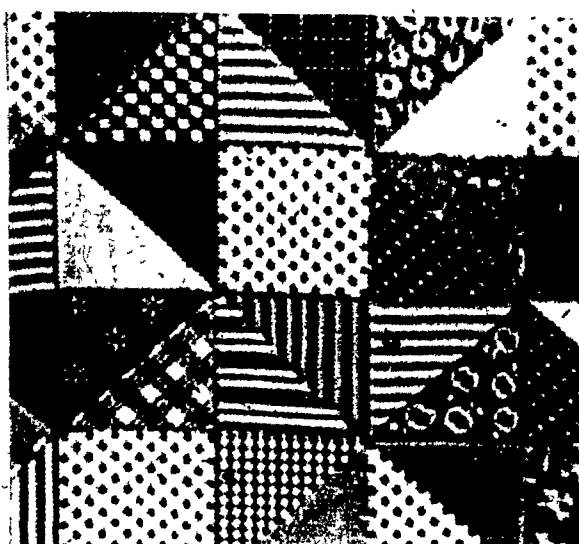
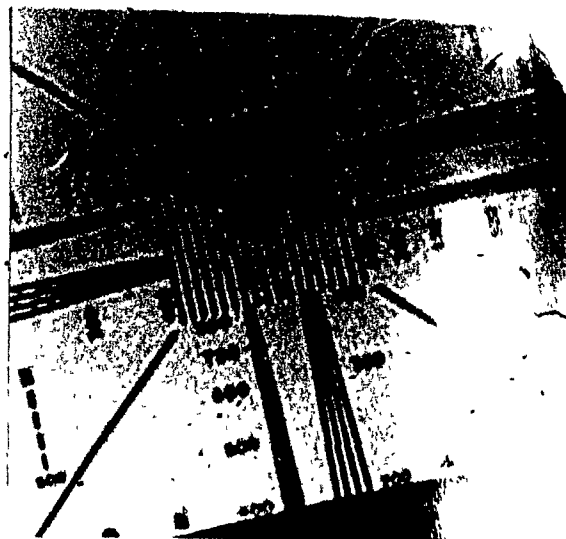
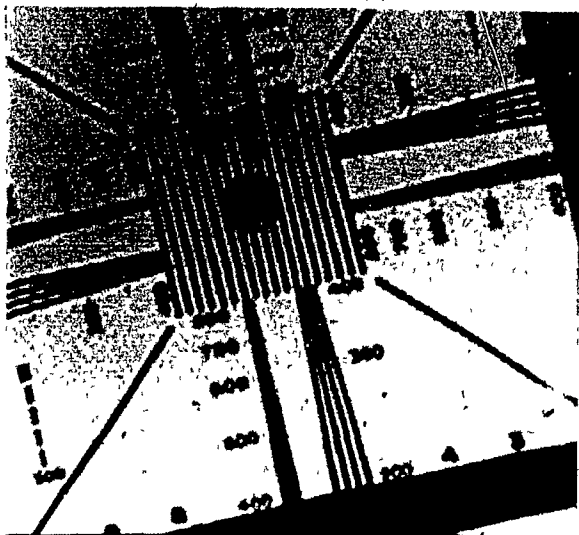
Distortion d_2

Distortion d_3



Distortion d_4

Distortion d_5



Distortion d_6

Distortion d_7

Appendix C. Subjective Tests Procedure

The subjective tests are comprised of two sessions of half an hour each. In order to eliminate any bias, a pseudo random order of presentation is necessary. The first two presentations are examples chosen to show the kind of impairments being tested. The next 16 presentations are actual tests and can be separated into four quadrants, namely 3 to 6, 7 to 10, 11 to 14 and 15 to 18. Since each level of distortion is presented four times, it should be spread over the four quadrants. Also, in a presentation, EIA-CHART, ALBERT and QUILT should display different levels of distortion. Levels of distortion were first assigned in a random order and reordered to meet the above mentioned restrictions. Table C.1 and C.2 give the order used for the first and second session respectively.

PRESENTATION	CHART	ALBERT	QUILT
1	d_7	d_5	d_7
2	d_5	d_7	d_5
3	d_6	d_5	d_3
4	d_4	d_3	d_2
5	d_3	d_1	d_8
6	d_1	d_8	d_7
7	d_5	d_6	d_4
8	d_8	d_7	d_5
9	d_3	d_2	d_1
10	d_2	d_4	d_6
11	d_7	d_8	d_5
12	d_6	d_7	d_1
13	d_4	d_3	d_7
14	d_2	d_1	d_4
15	d_1	d_4	d_2
16	d_8	d_5	d_6
17	d_7	d_2	d_8
18	d_5	d_6	d_7

Table C.1 First session order of presentation

PRESENTATION	CHART	ALBERT	QUILT
1	d_8	d_6	d_8
2	d_6	d_8	d_6
3	d_5	d_4	d_1
4	d_2	d_6	d_4
5	d_7	d_5	d_6
6	d_8	d_7	d_6
7	d_6	d_8	d_7
8	d_1	d_5	d_2
9	d_7	d_1	d_3
10	d_4	d_3	d_8
11	d_8	d_5	d_2
12	d_1	d_4	d_8
13	d_3	d_2	d_6
14	d_5	d_6	d_3
15	d_4	d_3	d_1
16	d_3	d_8	d_5
17	d_6	d_1	d_7
18	d_2	d_7	d_4

Table C.2 Second session order of presentation

Viewers were given a blank table to fill up and the following explicative note:

SUBJECTIVE TEST

SESSION 1

Impairment Scale

5	Imperceptible
4	Perceptible, but not annoying
3	Slightly annoying
2	Annoying
1	Very annoying

SEQUENCE	CHART	ALBERT	QUILT
1			
2			
3			
4			
5			
6			
7			
8			
9			
10			
11			
12			
13			
14			
15			
16			
17			
18			

your NAME please : _____

SUBJECTIVE TEST

DESCRIPTION

We have asked you to come along to help us assess the effects of impairments that can occur in digital coding of TV. The complete test will consist of two separate sessions of which this is the first.

We are going to show a series of presentations, each of which consists of :

- the original TV sequence for 10 seconds.
- a neutral grey for about 5 seconds.
- the coded TV sequence for 10 seconds.
- a neutral grey for 10 seconds, voting period.

We would like your opinion of the overall impairment level of the coded TV sequence with respect to the original. Your opinion should be expressed in terms of the 5-grade impairment scale ranging from 5, *imperceptible*, to 1, *very annoying*, which is shown at the top of the report form provided.

You do not need to use all the grades and you may use a given grade as often as you wish. Please record your opinion in the report form during the 10 second voting period. At the beginning of a sequence, the HDVS system may display for a fraction of a second some frames that do not belong to the sequence, please ignore them.

First, we are going to show you six examples of the impairments we are testing. We would like you to record your opinion of these in the first two rows, but these results will not be included in the main analysis.

The session will last for about half an hour. Remember, all that is required is your personal opinion.

Any questions ?

thank you.

Appendix B. Statistics on the transform coefficients

The mean, standard deviation, minimum and maximum values of all coefficients were computed for both $16 \times 16 \times 1$ and $16 \times 16 \times 4$ blocks. Tables of the results for the standard deviation are given below. As for the mean values, except for $F(0,0,0)$, they are all zero and thus are not included. The image used to generate these statistics is LONG-SEQUENCE. The standard deviation tables give an idea of the energy distribution among the coefficients.

$u \backslash h$	0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
0	27.41	6.13	3.38	2.53	2.67	2.39	2.27	1.94	1.60	1.13	0.61	0.49	0.41	0.38	0.20	0.14
1	8.50	3.62	2.36	1.90	1.62	1.54	1.81	1.55	0.99	0.77	0.45	0.36	0.29	0.23	0.17	0.11
2	5.02	2.83	2.27	1.90	1.44	1.23	1.49	1.35	0.94	0.70	0.45	0.37	0.29	0.22	0.16	0.10
3	3.56	2.25	2.29	2.21	1.52	1.02	0.98	0.77	0.66	0.58	0.44	0.32	0.26	0.21	0.16	0.09
4	3.48	2.01	2.03	1.98	1.31	0.89	0.75	0.58	0.58	0.46	0.36	0.28	0.23	0.19	0.14	0.09
5	2.93	1.73	1.34	1.17	0.89	0.69	0.57	0.53	0.45	0.34	0.27	0.23	0.20	0.17	0.13	0.09
6	2.36	1.43	1.08	0.92	0.72	0.60	0.48	0.39	0.35	0.27	0.22	0.22	0.20	0.15	0.11	0.07
7	2.25	1.21	0.92	0.74	0.58	0.54	0.46	0.36	0.29	0.23	0.20	0.23	0.22	0.16	0.10	0.06
8	1.53	0.98	0.80	0.58	0.44	0.38	0.37	0.31	0.23	0.20	0.16	0.21	0.24	0.17	0.10	0.07
9	0.92	0.75	0.70	0.53	0.38	0.33	0.31	0.27	0.24	0.21	0.19	0.21	0.21	0.16	0.11	0.07
10	0.78	0.60	0.63	0.54	0.40	0.31	0.27	0.23	0.24	0.23	0.20	0.21	0.19	0.15	0.11	0.07
11	0.63	0.53	0.52	0.48	0.40	0.34	0.33	0.29	0.28	0.27	0.24	0.21	0.19	0.16	0.12	0.08
12	0.59	0.54	0.57	0.57	0.44	0.33	0.31	0.28	0.31	0.34	0.29	0.23	0.20	0.18	0.14	0.09
13	0.61	0.60	0.60	0.58	0.47	0.44	0.49	0.42	0.37	0.36	0.29	0.26	0.24	0.20	0.16	0.09
14	0.61	0.65	0.57	0.51	0.44	0.49	0.64	0.53	0.37	0.33	0.27	0.26	0.25	0.20	0.14	0.09
15	0.92	0.98	0.83	0.72	0.78	0.81	0.84	0.63	0.46	0.42	0.33	0.35	0.40	0.32	0.22	0.11

Table D.1 Standard deviation of 2D-DCT coefficients, the test image is LONG-SEQUENCE.

$v \backslash h$	0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
0	26.85	5.79	3.02	1.94	1.80	1.28	1.21	1.46	1.42	0.93	0.42	0.30	0.20	0.16	0.11	0.09
1	8.20	3.34	2.04	1.52	1.22	0.96	0.87	0.94	0.77	0.60	0.30	0.22	0.16	0.12	0.09	0.07
2	4.73	2.52	1.89	1.47	1.14	0.80	0.71	0.82	0.70	0.49	0.28	0.24	0.17	0.13	0.10	0.06
3	3.25	1.94	1.67	1.70	1.05	0.73	0.66	0.50	0.43	0.35	0.28	0.19	0.15	0.11	0.09	0.05
4	3.20	1.73	1.47	1.45	0.85	0.66	0.57	0.43	0.41	0.25	0.21	0.17	0.12	0.10	0.08	0.05
5	2.77	1.50	1.05	0.84	0.56	0.48	0.43	0.43	0.35	0.22	0.16	0.12	0.11	0.09	0.07	0.05
6	2.18	1.21	0.86	0.67	0.47	0.40	0.31	0.27	0.26	0.18	0.12	0.10	0.09	0.08	0.06	0.04
7	2.10	1.02	0.72	0.57	0.42	0.38	0.28	0.23	0.20	0.12	0.11	0.09	0.09	0.07	0.06	0.04
8	1.40	0.79	0.61	0.44	0.31	0.25	0.23	0.21	0.15	0.12	0.09	0.09	0.09	0.08	0.05	0.04
9	0.78	0.51	0.47	0.35	0.24	0.18	0.15	0.13	0.13	0.10	0.09	0.08	0.08	0.08	0.06	0.04
10	0.66	0.41	0.40	0.35	0.25	0.17	0.15	0.11	0.11	0.10	0.09	0.08	0.08	0.07	0.05	0.04
11	0.49	0.33	0.30	0.28	0.21	0.16	0.14	0.13	0.11	0.10	0.09	0.08	0.08	0.07	0.06	0.04
12	0.43	0.28	0.34	0.31	0.25	0.16	0.14	0.12	0.11	0.11	0.09	0.09	0.09	0.07	0.06	0.04
13	0.38	0.29	0.34	0.32	0.23	0.17	0.15	0.12	0.12	0.11	0.10	0.10	0.10	0.08	0.07	0.04
14	0.39	0.31	0.31	0.26	0.21	0.17	0.16	0.15	0.12	0.11	0.10	0.10	0.11	0.09	0.06	0.05
15	0.54	0.42	0.42	0.37	0.32	0.24	0.22	0.17	0.15	0.13	0.11	0.13	0.16	0.12	0.07	0.06

Table D.2 Standard deviation of 3D-DCT coefficients, frame 0, the test image is LONG-SEQUENCE.

$v \backslash h$	0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
0	4.29	1.57	1.24	1.30	1.53	1.42	1.37	0.72	0.36	0.30	0.19	0.16	0.14	0.14	0.09	0.06
1	1.81	1.10	0.99	0.93	0.86	0.90	1.18	0.67	0.38	0.28	0.17	0.14	0.11	0.10	0.08	0.05
2	1.39	1.04	1.05	1.00	0.71	0.71	0.97	0.75	0.38	0.26	0.19	0.14	0.11	0.09	0.07	0.05
3	1.22	0.95	1.38	1.15	0.94	0.53	0.53	0.39	0.27	0.21	0.16	0.12	0.11	0.09	0.07	0.04
4	1.10	0.82	1.24	1.13	0.85	0.45	0.35	0.23	0.21	0.17	0.14	0.11	0.09	0.08	0.06	0.04
5	0.91	0.68	0.69	0.68	0.57	0.38	0.26	0.20	0.16	0.13	0.10	0.09	0.08	0.07	0.06	0.04
6	0.70	0.59	0.53	0.51	0.43	0.33	0.23	0.17	0.13	0.10	0.08	0.08	0.08	0.06	0.05	0.03
7	0.64	0.51	0.47	0.36	0.31	0.27	0.23	0.16	0.10	0.09	0.07	0.08	0.08	0.07	0.05	0.03
8	0.46	0.44	0.41	0.29	0.23	0.20	0.19	0.13	0.09	0.08	0.07	0.07	0.08	0.06	0.05	0.03
9	0.35	0.39	0.40	0.30	0.22	0.18	0.16	0.12	0.09	0.08	0.07	0.07	0.07	0.07	0.05	0.03
10	0.29	0.29	0.38	0.33	0.23	0.17	0.13	0.10	0.09	0.08	0.07	0.07	0.07	0.06	0.05	0.03
11	0.26	0.25	0.28	0.29	0.25	0.20	0.17	0.14	0.11	0.10	0.08	0.08	0.07	0.07	0.05	0.03
12	0.26	0.25	0.31	0.38	0.26	0.19	0.16	0.13	0.12	0.12	0.10	0.09	0.08	0.07	0.07	0.04
13	0.31	0.28	0.32	0.36	0.28	0.27	0.28	0.24	0.16	0.14	0.10	0.09	0.09	0.08	0.07	0.04
14	0.29	0.30	0.29	0.29	0.25	0.31	0.41	0.32	0.17	0.14	0.11	0.10	0.09	0.08	0.07	0.05
15	0.44	0.45	0.42	0.42	0.49	0.51	0.45	0.34	0.19	0.16	0.12	0.12	0.13	0.11	0.10	0.05

Table D.3 Standard deviation of 3D-DCT coefficients, frame 1, the test image is LONG-SEQUENCE.

$u \backslash h$	0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
0	3.01	1.06	0.74	0.86	1.17	1.33	1.15	0.94	0.51	0.42	0.25	0.20	0.18	0.13	0.09	0.05
1	1.15	0.70	0.58	0.60	0.57	0.75	0.98	0.80	0.42	0.30	0.20	0.16	0.13	0.10	0.07	0.05
2	0.83	0.65	0.59	0.60	0.47	0.56	0.83	0.70	0.41	0.33	0.21	0.16	0.12	0.10	0.07	0.05
3	0.74	0.55	0.68	0.78	0.52	0.42	0.44	0.39	0.33	0.29	0.22	0.13	0.11	0.09	0.07	0.04
4	0.73	0.52	0.61	0.69	0.49	0.34	0.30	0.25	0.27	0.24	0.17	0.12	0.10	0.08	0.06	0.04
5	0.57	0.45	0.42	0.42	0.37	0.30	0.24	0.20	0.18	0.15	0.12	0.10	0.09	0.07	0.06	0.04
6	0.49	0.40	0.33	0.33	0.30	0.27	0.24	0.18	0.15	0.12	0.10	0.09	0.09	0.07	0.05	0.03
7	0.44	0.34	0.27	0.26	0.23	0.24	0.26	0.19	0.13	0.10	0.09	0.09	0.10	0.07	0.04	0.03
8	0.35	0.29	0.26	0.20	0.17	0.18	0.19	0.15	0.11	0.09	0.08	0.10	0.10	0.08	0.05	0.03
9	0.29	0.29	0.27	0.22	0.18	0.17	0.19	0.15	0.13	0.11	0.09	0.10	0.09	0.07	0.05	0.03
10	0.24	0.25	0.24	0.22	0.18	0.16	0.15	0.13	0.13	0.09	0.10	0.08	0.07	0.05	0.03	0.03
11	0.24	0.24	0.23	0.21	0.20	0.18	0.21	0.17	0.16	0.16	0.12	0.10	0.08	0.07	0.05	0.04
12	0.23	0.25	0.25	0.25	0.21	0.18	0.18	0.17	0.18	0.20	0.13	0.12	0.09	0.08	0.06	0.04
13	0.28	0.30	0.28	0.26	0.24	0.24	0.32	0.26	0.23	0.21	0.15	0.13	0.11	0.10	0.07	0.05
14	0.27	0.33	0.26	0.25	0.23	0.28	0.41	0.33	0.22	0.19	0.14	0.12	0.11	0.10	0.07	0.04
15	0.44	0.50	0.41	0.38	0.42	0.48	0.60	0.41	0.27	0.23	0.17	0.16	0.16	0.19	0.09	0.06

Table D.4 Standard deviation of 3D-DCT coefficients, frame 2, the test image is LONG-SEQUENCE.

$u \backslash h$	0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
0	1.65	0.68	0.48	0.42	0.45	0.54	0.71	0.49	0.41	0.38	0.30	0.29	0.27	0.28	0.12	0.06
1	0.63	0.46	0.33	0.29	0.27	0.30	0.41	0.34	0.27	0.25	0.21	0.19	0.18	0.13	0.09	0.05
2	0.46	0.38	0.31	0.25	0.24	0.25	0.32	0.30	0.27	0.25	0.22	0.18	0.16	0.12	0.08	0.05
3	0.38	0.34	0.29	0.30	0.22	0.20	0.22	0.23	0.25	0.29	0.21	0.18	0.15	0.12	0.09	0.05
4	0.36	0.33	0.27	0.26	0.21	0.17	0.16	0.18	0.23	0.25	0.20	0.16	0.14	0.11	0.09	0.05
5	0.33	0.28	0.24	0.19	0.16	0.16	0.14	0.13	0.16	0.17	0.15	0.15	0.12	0.09	0.07	0.04
6	0.27	0.27	0.20	0.16	0.15	0.14	0.13	0.13	0.14	0.13	0.12	0.15	0.14	0.09	0.06	0.04
7	0.26	0.23	0.18	0.14	0.13	0.14	0.14	0.13	0.12	0.13	0.11	0.16	0.16	0.10	0.06	0.03
8	0.24	0.22	0.18	0.12	0.11	0.11	0.12	0.10	0.10	0.11	0.09	0.14	0.18	0.11	0.06	0.04
9	0.20	0.24	0.19	0.14	0.12	0.12	0.12	0.12	0.13	0.13	0.12	0.15	0.16	0.10	0.06	0.04
10	0.17	0.24	0.18	0.14	0.12	0.11	0.11	0.11	0.13	0.14	0.13	0.15	0.14	0.09	0.06	0.04
11	0.21	0.24	0.21	0.14	0.13	0.13	0.13	0.13	0.17	0.17	0.16	0.14	0.13	0.10	0.07	0.04
12	0.20	0.28	0.23	0.15	0.14	0.13	0.13	0.14	0.20	0.22	0.23	0.15	0.13	0.12	0.09	0.05
13	0.23	0.33	0.25	0.20	0.17	0.17	0.18	0.19	0.22	0.22	0.21	0.17	0.17	0.13	0.10	0.05
14	0.26	0.37	0.26	0.21	0.18	0.19	0.21	0.21	0.20	0.20	0.18	0.18	0.17	0.13	0.09	0.05
15	0.42	0.58	0.40	0.31	0.30	0.31	0.31	0.30	0.29	0.28	0.24	0.26	0.30	0.20	0.16	0.06

Table D.5 Standard deviation of 3D-DCT coefficients, frame 3, the test image is LONG-SEQUENCE.

Appendix E. Entropy and Bit Rates for Individual Coefficients

This appendix presents results concerning the coefficients individually. Table E.1 contains their entropy computed from the LONG-SEQUENCE image and coded with a quantization step distribution generated with the parametrical function defined in Chapter 3 and parameter value $c = 1.0$. The bit rates of Table E.2 are those resulting of the integrated coder using a zigzag scanning pattern, no post filtering, and the same block quantizer as above.

The bit rates of Table E.3 to E.6 result, also of a three-dimensional integrated coder and the block quantizer as defined in Sect. 3.5.2

As mentioned in Chapter 3, bit rates really account for the non-zero values of the coefficients. The amount of bit coded with run length coding is not considered.

$v \backslash h$	0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
0	10.01	7.32	6.11	5.49	4.95	4.61	4.24	3.88	3.44	3.16	2.66	2.39	2.12	1.88	1.36	0.81
1	7.69	6.28	5.44	4.79	4.27	3.81	3.49	3.13	2.74	2.39	1.92	1.62	1.36	1.07	0.68	0.27
2	6.72	5.67	4.97	4.33	3.78	3.26	2.95	2.59	2.22	1.90	1.49	1.21	0.98	0.72	0.40	0.13
3	6.16	5.14	4.54	3.91	3.37	2.78	2.44	2.06	1.75	1.49	1.16	0.86	0.69	0.45	0.24	0.07
4	5.78	4.78	4.14	3.50	2.92	2.34	1.95	1.56	1.35	1.05	0.82	0.62	0.45	0.29	0.13	0.03
5	5.33	4.34	3.65	3.02	2.45	1.94	1.54	1.19	0.99	0.73	0.54	0.39	0.26	0.17	0.07	0.03
6	4.94	3.96	3.22	2.63	2.04	1.60	1.20	0.90	0.73	0.51	0.33	0.26	0.20	0.08	0.03	0.01
7	4.56	3.59	2.86	2.25	1.66	1.30	1.00	0.74	0.51	0.34	0.22	0.19	0.15	0.07	0.01	0.00
8	4.14	3.11	2.45	1.86	1.30	0.96	0.74	0.52	0.33	0.22	0.11	0.14	0.13	0.06	0.00	0.00
9	3.62	2.72	2.16	1.64	1.11	0.80	0.60	0.43	0.30	0.20	0.13	0.13	0.11	0.05	0.00	0.00
10	3.29	2.41	1.92	1.46	1.03	0.68	0.46	0.31	0.26	0.19	0.12	0.10	0.07	0.02	0.00	0.00
11	2.96	2.20	1.70	1.35	0.94	0.68	0.53	0.38	0.28	0.22	0.15	0.10	0.06	0.02	0.00	0.00
12	2.74	2.01	1.61	1.30	0.90	0.60	0.45	0.31	0.30	0.26	0.17	0.09	0.04	0.03	0.01	0.00
13	2.67	1.97	1.56	1.25	0.90	0.71	0.60	0.47	0.33	0.26	0.16	0.11	0.07	0.03	0.01	0.00
14	2.53	1.87	1.42	1.12	0.82	0.71	0.61	0.49	0.32	0.23	0.13	0.09	0.07	0.03	0.00	0.00
15	2.87	2.01	1.49	1.25	0.91	0.85	0.74	0.58	0.34	0.26	0.18	0.16	0.15	0.09	0.05	0.00

Table E.1. Entropy table for the coefficients quantized with the parameter $c = 1.0$ (2D DCT) and generated from the image LONG-SEQUENCE.

$v \backslash h$	0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
0	12.00	8.99	6.69	5.68	4.80	4.29	3.77	3.33	2.83	2.50	1.98	1.66	1.41	1.19	0.76	0.37
1	9.68	7.04	5.67	4.63	3.88	3.23	2.84	2.43	2.02	1.66	1.22	0.96	0.75	0.55	0.30	0.09
2	7.74	6.07	4.93	3.95	3.24	2.59	2.21	1.83	1.48	1.19	0.85	0.64	0.48	0.33	0.15	0.04
3	6.79	5.23	4.25	3.36	2.72	2.05	1.70	1.33	1.07	0.84	0.60	0.40	0.31	0.18	0.08	0.02
4	6.11	4.66	3.71	2.84	2.19	1.61	1.23	0.90	0.74	0.53	0.38	0.27	0.18	0.10	0.04	0.01
5	5.41	4.04	3.12	2.33	1.72	1.23	0.89	0.62	0.48	0.33	0.22	0.14	0.09	0.05	0.02	0.01
6	4.80	3.52	2.60	1.92	1.32	0.93	0.63	0.43	0.32	0.21	0.12	0.09	0.06	0.02	0.01	0.00
7	4.26	3.03	2.19	1.54	0.99	0.71	0.49	0.34	0.21	0.13	0.07	0.06	0.04	0.02	0.00	0.00
8	3.72	2.45	1.72	1.17	0.71	0.47	0.33	0.21	0.12	0.08	0.03	0.04	0.04	0.02	0.00	0.00
9	3.10	2.05	1.46	0.99	0.58	0.37	0.25	0.16	0.10	0.06	0.04	0.04	0.03	0.01	0.00	0.00
10	2.67	1.71	1.22	0.83	0.52	0.31	0.19	0.11	0.09	0.06	0.04	0.03	0.02	0.00	0.00	0.00
11	2.29	1.50	1.04	0.74	0.46	0.30	0.22	0.14	0.10	0.07	0.04	0.03	0.01	0.00	0.00	0.00
12	2.00	1.31	0.95	0.70	0.43	0.26	0.18	0.11	0.11	0.09	0.05	0.03	0.01	0.01	0.00	0.00
13	1.96	1.25	0.91	0.67	0.43	0.32	0.26	0.19	0.12	0.09	0.05	0.03	0.02	0.01	0.00	0.00
14	1.80	1.16	0.79	0.58	0.38	0.32	0.26	0.19	0.11	0.08	0.04	0.02	0.02	0.01	0.00	0.00
15	2.18	1.28	0.84	0.66	0.43	0.40	0.33	0.24	0.12	0.09	0.06	0.05	0.04	0.03	0.01	0.00

Table E.2 Average bit rate for the coefficients quantized with the parameter $c = 1.0$ (2D DCT) and generated from the image LONG-SEQUENCE.

$v \backslash h$	0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
0	12.00	9.53	7.03	5.99	4.99	4.40	3.74	3.35	2.88	2.53	2.09	1.83	1.60	1.38	0.96	0.66
1	10.49	7.83	6.42	5.47	4.61	3.84	3.33	2.92	2.53	2.10	1.71	1.36	1.13	0.88	0.60	0.24
2	8.58	7.05	6.00	5.10	4.31	3.53	2.98	2.52	2.14	1.74	1.35	1.04	0.85	0.60	0.34	0.11
3	7.70	6.35	5.54	4.69	3.94	3.03	2.57	2.02	1.68	1.35	1.03	0.69	0.51	0.33	0.16	0.05
4	7.08	5.95	5.07	4.28	3.39	2.57	2.06	1.59	1.30	0.90	0.62	0.43	0.29	0.15	0.07	0.03
5	6.36	5.37	4.51	3.69	2.91	2.10	1.58	1.11	0.86	0.55	0.36	0.22	0.14	0.09	0.03	0.02
6	5.81	4.85	4.00	3.24	2.36	1.65	1.14	0.74	0.55	0.34	0.20	0.12	0.07	0.04	0.01	0.00
7	5.23	4.33	3.45	2.76	1.86	1.23	0.81	0.54	0.36	0.19	0.11	0.06	0.03	0.02	0.01	0.00
8	4.52	3.71	2.98	2.24	1.44	0.87	0.57	0.35	0.21	0.12	0.06	0.04	0.03	0.01	0.00	0.00
9	3.89	3.10	2.55	1.85	1.14	0.61	0.37	0.21	0.15	0.06	0.03	0.02	0.02	0.01	0.00	0.00
10	3.45	2.72	2.18	1.55	0.98	0.49	0.25	0.14	0.09	0.04	0.02	0.01	0.00	0.00	0.00	0.00
11	3.03	2.40	1.89	1.34	0.82	0.43	0.23	0.11	0.06	0.03	0.01	0.00	0.00	0.00	0.00	0.00
12	2.74	2.15	1.73	1.26	0.76	0.34	0.19	0.07	0.05	0.03	0.01	0.00	0.00	0.00	0.00	0.00
13	2.81	2.08	1.67	1.13	0.63	0.32	0.18	0.07	0.04	0.02	0.00	0.00	0.00	0.00	0.00	0.00
14	2.59	2.01	1.48	0.96	0.54	0.27	0.17	0.07	0.03	0.01	0.00	0.00	0.00	0.00	0.00	0.00
15	3.36	2.21	1.54	1.00	0.54	0.32	0.18	0.08	0.04	0.01	0.00	0.00	0.00	0.00	0.00	0.00

Table E.3 Average bit rate per coefficients, 3D DCT, frame 0, and generated from the image LONG-SEQUENCE.

$u \backslash h$	0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
0	6.34	4.95	4.03	3.52	2.98	2.64	2.31	1.85	1.41	1.20	0.77	0.62	0.49	0.41	0.20	0.07
1	4.98	4.15	3.42	2.83	2.41	1.99	1.66	1.27	0.88	0.65	0.38	0.28	0.18	0.13	0.05	0.02
2	4.34	3.63	2.92	2.43	1.93	1.51	1.22	0.87	0.62	0.41	0.25	0.15	0.10	0.05	0.02	0.01
3	3.85	3.10	2.56	2.05	1.56	1.16	0.88	0.60	0.37	0.26	0.14	0.07	0.04	0.02	0.01	0.00
4	3.55	2.74	2.23	1.65	1.25	0.81	0.58	0.32	0.24	0.14	0.08	0.04	0.01	0.01	0.00	0.00
5	3.05	2.29	1.78	1.29	0.93	0.58	0.35	0.21	0.12	0.06	0.02	0.01	0.00	0.00	0.00	0.00
6	2.66	1.91	1.38	0.97	0.63	0.41	0.24	0.13	0.06	0.02	0.00	0.00	0.00	0.00	0.00	0.00
7	2.24	1.55	1.12	0.73	0.46	0.31	0.17	0.07	0.02	0.01	0.00	0.00	0.00	0.00	0.00	0.00
8	1.86	1.21	0.85	0.49	0.27	0.17	0.11	0.04	0.01	0.00	0.00	0.00	0.00	0.00	0.00	0.00
9	1.49	0.98	0.71	0.43	0.22	0.13	0.07	0.03	0.01	0.00	0.00	0.00	0.00	0.00	0.00	0.00
10	1.25	0.83	0.60	0.39	0.22	0.09	0.04	0.01	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00
11	1.10	0.73	0.51	0.36	0.20	0.11	0.07	0.03	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00
12	1.05	0.66	0.50	0.37	0.19	0.09	0.04	0.02	0.01	0.00	0.00	0.00	0.00	0.00	0.00	0.00
13	1.03	0.67	0.48	0.36	0.20	0.15	0.10	0.07	0.02	0.01	0.00	0.00	0.00	0.00	0.00	0.00
14	0.97	0.63	0.42	0.30	0.18	0.16	0.13	0.08	0.02	0.00	0.00	0.00	0.00	0.00	0.00	0.00
15	1.09	0.74	0.50	0.39	0.23	0.23	0.17	0.11	0.02	0.01	0.00	0.00	0.00	0.00	0.00	0.00

Table E.4 Average bit rate per coefficients, 3D DCT, frame 1, and generated from the image LONG-SEQUENCE.

$u \backslash h$	0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
0	4.07	3.40	2.87	2.62	2.31	2.20	1.95	1.73	1.34	1.19	0.81	0.67	0.51	0.40	0.20	0.08
1	3.28	2.84	2.37	2.03	1.76	1.59	1.35	1.15	0.86	0.65	0.43	0.32	0.22	0.13	0.04	0.01
2	2.85	2.44	2.03	1.71	1.38	1.15	0.96	0.81	0.58	0.44	0.27	0.20	0.11	0.06	0.01	0.00
3	2.59	2.10	1.76	1.44	1.13	0.86	0.71	0.55	0.40	0.31	0.21	0.09	0.06	0.02	0.01	0.00
4	2.38	1.83	1.47	1.20	0.87	0.61	0.46	0.32	0.25	0.18	0.10	0.05	0.02	0.01	0.00	0.00
5	2.06	1.55	1.19	0.90	0.62	0.43	0.31	0.20	0.14	0.08	0.04	0.02	0.01	0.00	0.00	0.00
6	1.80	1.27	0.90	0.64	0.47	0.33	0.23	0.12	0.08	0.04	0.02	0.01	0.01	0.00	0.00	0.00
7	1.53	1.06	0.72	0.48	0.31	0.26	0.18	0.10	0.04	0.02	0.01	0.00	0.01	0.00	0.00	0.00
8	1.22	0.81	0.52	0.33	0.18	0.14	0.10	0.07	0.02	0.01	0.00	0.01	0.01	0.00	0.00	0.00
9	1.10	0.72	0.47	0.28	0.16	0.13	0.09	0.05	0.03	0.01	0.00	0.01	0.00	0.00	0.00	0.00
10	0.97	0.61	0.40	0.26	0.15	0.09	0.05	0.03	0.02	0.02	0.00	0.00	0.00	0.00	0.00	0.00
11	0.87	0.57	0.39	0.24	0.14	0.09	0.09	0.05	0.03	0.03	0.01	0.00	0.00	0.00	0.00	0.00
12	0.83	0.53	0.39	0.24	0.15	0.08	0.05	0.04	0.04	0.04	0.01	0.00	0.00	0.00	0.00	0.00
13	0.82	0.56	0.38	0.27	0.16	0.14	0.14	0.09	0.05	0.03	0.01	0.01	0.00	0.00	0.00	0.00
14	0.70	0.50	0.32	0.23	0.14	0.14	0.15	0.10	0.04	0.02	0.00	0.00	0.00	0.00	0.00	0.00
15	0.87	0.65	0.43	0.33	0.20	0.22	0.21	0.13	0.05	0.04	0.01	0.01	0.01	0.01	0.00	0.00

Table E.5 Average bit rate per coefficients, 3D DCT, frame 2, and generated from the image LONG-SEQUENCE.

$v \backslash h$	0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
0	3.03	2.36	1.90	1.66	1.49	1.45	1.45	1.24	1.08	1.02	0.78	0.67	0.61	0.49	0.24	0.08
1	2.16	1.83	1.41	1.17	0.99	0.92	0.89	0.78	0.62	0.54	0.40	0.33	0.27	0.18	0.09	0.02
2	1.79	1.48	1.14	0.88	0.73	0.64	0.60	0.51	0.42	0.36	0.29	0.21	0.18	0.10	0.03	0.01
3	1.59	1.24	0.93	0.71	0.54	0.43	0.40	0.33	0.29	0.27	0.20	0.14	0.12	0.06	0.02	0.00
4	1.41	1.05	0.75	0.53	0.40	0.29	0.22	0.18	0.17	0.18	0.14	0.10	0.07	0.03	0.01	0.00
5	1.25	0.88	0.59	0.37	0.26	0.22	0.13	0.10	0.11	0.10	0.07	0.06	0.03	0.01	0.00	0.00
6	1.11	0.74	0.46	0.27	0.18	0.15	0.10	0.07	0.07	0.06	0.04	0.04	0.03	0.01	0.00	0.00
7	1.06	0.66	0.36	0.20	0.12	0.11	0.08	0.06	0.04	0.04	0.02	0.03	0.02	0.01	0.00	0.00
8	1.24	0.52	0.30	0.13	0.06	0.05	0.04	0.03	0.01	0.02	0.01	0.02	0.03	0.01	0.00	0.00
9	0.76	0.52	0.27	0.14	0.07	0.05	0.03	0.03	0.02	0.03	0.01	0.02	0.02	0.00	0.00	0.00
10	0.64	0.44	0.24	0.11	0.06	0.03	0.02	0.02	0.02	0.02	0.01	0.02	0.01	0.00	0.00	0.00
11	0.60	0.41	0.25	0.10	0.06	0.05	0.03	0.03	0.04	0.03	0.02	0.01	0.01	0.00	0.00	0.00
12	0.57	0.41	0.25	0.10	0.07	0.03	0.02	0.02	0.04	0.05	0.03	0.01	0.00	0.00	0.00	0.00
13	0.57	0.45	0.26	0.13	0.07	0.06	0.05	0.04	0.04	0.04	0.03	0.02	0.01	0.00	0.00	0.00
14	0.57	0.42	0.25	0.13	0.07	0.06	0.06	0.04	0.03	0.03	0.02	0.01	0.01	0.00	0.00	0.00
15	0.73	0.55	0.35	0.21	0.12	0.11	0.11	0.08	0.05	0.04	0.03	0.03	0.03	0.01	0.01	0.00

Table E.6 Average bit rate per coefficients, 3D DCT, frame 3, and generated from the image LONG-SEQUENCE.

References

- [1] C F Hall, E L Hall, "A Nonlinear Model for the Spatial Characteristics of the Human Visual System," *IEEE Trans Syst Man Cyb*, vol smc 7, No 3, pp 161-170, March 1977
- [2] W K Pratt, *Digital Image Processing* Wiley-Interscience, 1978
- [3] H L Lohscheller, "A Subjectively Adapted Image Communication System," *IEEE Trans Commun.*, vol com 32, No 12, pp. 1316-1322, December 1984
- [4] S. Ericsson, "Frequency Weighted Interframe Hybrid Coding," Report No. TRITA-TTT-8401, Royal Institute of Technology, Stockholm, Sweden, January 1984
- [5] N.C. Griswold, "Perceptual Coding in the Cosine Transform Domain," *Optical Engineering*, vol 19, No.3, pp. 306-311, May-June 1980.
- [6] J.L. Mannos, D.J. Sakrison, "The Effects of a Visual Fidelity Criterion on the Encoding of Images," *IEEE Trans. Inform. Theory*, vol. IT 20, No 4, pp. 525-536, July 1974.
- [7] W H Chen, W.K. Pratt, "Scene Adaptive Coder," *IEEE Trans. Commun*, vol. com 32, No 3, pp. 225-232, March, 1984
- [8] J A. Saghri, A G. Tesher, "Adaptive Transform Coding Based On Chain Coding Concepts," *IEEE Trans Commun.*, vol. com 34, No.2, pp. 112-117, February 1986.
- [9] E. Dubois, "The Sampling and Reconstruction of Time-varying Imagery," *Proc. IEEE*, vol. 73, No.4, pp. 502-522, April 1985.
- [10] F Kretz, "Subjectively Optimal Quantization of Pictures," *IEEE Trans. Commun.*, vol. com-23, No 11, pp. 1288-1292, November 1975.
- [11] F.X.J. Lukas, Z.L. Budrikis, "Picture Quality Prediction Based on a Visual Model," *IEEE Trans. Commun.*, vol. com 30, No.4, pp. 1679-1692, July 1982.
- [12] D.J. Sakrison, "On the role of the Observer and a Distorsion Measure in Image Transmission," *IEEE Trans. Commun.*, vol. com 25, No.11, pp. 1251-1266, November 1977.

- [13] A N Netravali, J.O Limb, "Picture Coding A Review," *Proc IEEE*, vol 68, No 3, pp 366-407, March 1980.
- [14] M. Kunt, A. Ikononopoulos, M Kocher, "Second Generation Image-Coding Techniques," *Proc. IEEE*, vol. 73, No 4, pp. 549-574, April 1985.
- [15] H.G. Muamann, P. Pirsh, H.J Grallert, "Advances in Picture Coding," *Proc IEEE*, vol. 73, No.4, pp. 523-548, April 1985.
- [16] N Ahmed, K. Rao, *Orthogonal Transform for Digital Signal Processing*, Springer-Verlag, 1975.
- [17] A Segall, "Bit Allocation and Encoding for Vector Sources," *IEEE Trans. Inform Theory*, vol IT 22, No.2, pp 162-169, March 1976
- [18] W H Chen, C H. Smith, S C Fralick, "A Fast Computational Algorithm for the Discrete Cosine Transform," *IEEE Trans. Commun*, vol com 25, No.9, pp. 1004-1007, September 1977
- [19] M Vetterli, "Fast 2-D Discrete Cosine Transform," in *ICASSP 84 IEEE International Conference On Acoustic, Speech, And Signal Processing*, 1985, pp. 48.6.1-48.6.4
- [20] E Arnould, JP Dugre, "Real Time Discrete Cosine Transform, an Original Architecture," in *ICASSP 85 IEEE International Conference On Acoustic, Speech, And Signal Processing*, 1984, pp. 40.8.1-40.8.4.
- [21] J Max, "Quantizing for Minimum Distortion," *IRE Trans. Inform. Theory*, vol. IT-6, No.1, pp. 7-12, March 1960.
- [22] P.J. Ready and P.A. Wintz, "Multispectral Data Compression Through Transform Coding and Block Quantization," Information note 050572, Purdue University, Laboratory for Application of Remote Sensing, May 1972.
- [23] J.J.Y Huang and P.M. Schutheiss, "Block Quantization of Correlated Gaussian Random Variable," *IEEE Trans. Commun.*, vol. CS-11, No.3, pp. 289-296, September 1963.
- [24] M.D. Levine, *Vision in Man and Machine*. New York: McGraw-Hill, 1985.
- [25] N.B. Nill, "A Visual Model Weighted Cosine Transform for Image Compression and Quality Assessment," *IEEE Trans. Commun.*, vol. 33, No.6, pp. 551-557, June 1985.
- [26] C. C. I. R., "Subjective Assessment of the Quality of Television Pictures," Report No. 405-4, 1966-1970-1974-1978-1982.
- [27] C. C. I. R., "Method for the Subjective Assessment of the Quality of Television Pictures," Recommendation 500-2, 15th Plenary Assembly, 1982.
- [28] P. Sallio and F. Kretz, "Qualité subjective en télévision numérique, Première partie: Méthodologie de son évaluation," *Revue de radiodiffusion-télévision*, vol. 52, pp. 1-7, 1978.

[29] J. Allnatt, *Transmitted-picture Assessment* Wiley, 1983.

[30] E. Dubois, P. Faubert, "The Effects on Picture Quality of Errors in a Digital Video Tape Recorder," Report no. 85-30, INRS-Télécommunications, September 1985.

[31] R.G. Gallager, *Information Theory and Reliable Communication*. Wiley, 1968.

[32] E. Dubois and J.L. Moncet, "Encoding and Progressive Transmission of Still Pictures in NTSC Composite Format Using Transform Domain Methods," *IEEE Trans. Commun.*, vol. com-34, No 3, pp. 310-319, March 1986.