



**IRWIN AND JOAN JACOBS**  
**CENTER FOR COMMUNICATION AND INFORMATION TECHNOLOGIES**

# **On the Statistics of Wide-Sense Markov Random Fields and its Application to Texture Interpolation**

**Shira Nemirovsky and Moshe Porat**

**CCIT Report # 682**  
**January 2008**

■ ■ ■ ■ Electronics  
■ ■ ■ ■ Computers  
■ ■ ■ ■ Communications

**DEPARTMENT OF ELECTRICAL ENGINEERING**  
**TECHNION - ISRAEL INSTITUTE OF TECHNOLOGY, HAIFA 32000, ISRAEL**



## **On the Statistics of Wide-Sense Markov Random Fields and its Application to Texture Interpolation**

Shira Nemirovsky and Moshe Porat  
Department of Electrical Engineering  
Technion – Israel Institute of Technology  
Haifa 32000, Israel

### **Abstract**

Random field models are one of the most common classes of image models. Such models, as suggested in the literature, attempt to characterize the correlation among neighboring pixels in the image. Since images are defined on 2D grids, it is more appropriate to model them as realizations of 2D time series or random field models. By assuming a separable correlation function for a 2D autoregressive (AR) model, a straightforward generalization of the 1D time series AR model to 2D is obtained. This simple model and its extensions have found many applications in image restoration, image compression, and texture classification and segmentation. In this work we explore the statistical properties of wide-sense Markov random fields with a separable correlation function. We analyze the effect of interpolation on the statistics of such images, and develop corresponding mathematical relations. Motivated by these results and relations, we propose a new method for texture interpolation, based on an orthogonal decomposition model for textures. Experiments with natural texture images and comparison with presently available interpolation methods demonstrate the advantages of the proposed method.

## 1. A Short Introduction

Wide-sense Markov random fields constitute a useful 2D auto-regressive image model, which is highly applicable to image restoration, image compression and texture classification and segmentation [1], [2]. In this work we focus on sampling and interpolation of images that may be modeled (or approximated) by wide-sense Markov fields in particular, and by 2D AR models in general. We develop useful mathematical relations between the statistical features of low- and high-resolution versions of such images.

We proceed to associating the above results with texture interpolation, a problem of interest since many natural images may be described as a collage of various textures [3], [4], [5].

We propose a new interpolation method for textures, based on an orthogonal decomposition model for textures by Francos et al. [6], [7]. Experiments on Brodatz texture images [3] demonstrate the advantages of the proposed method compared to presently available interpolation methods.

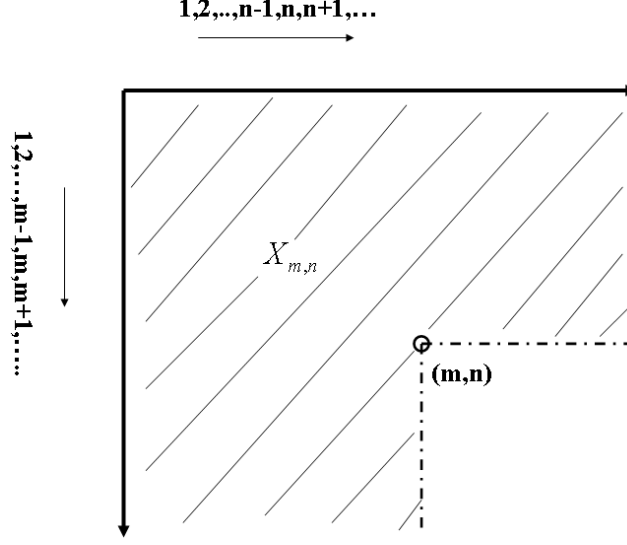
This paper is organized as follows. Section 2 introduces the wide sense Markov model and the new developed relations for its high and low resolution statistics. In Section 3 the application of texture interpolation, as well as a texture model, are reviewed. Then, the new method for texture interpolation is proposed. Section 4 summarizes the work and suggests future research.

## 2. Wide-Sense Markov Random Fields

### 2.1 Representation of images by wide-sense Markov random fields

We denote by  $X_{m,n}$  all the pixels that belong to an L-shaped region of an image matrix as shown in Figure 2.1, that is:

$$X_{m,n} = \{(i, j) | i < m \text{ or } j < m\}. \quad (2.1)$$



**Figure 2.1** – The L-shaped region formed by the dashed lines is  $X_{m,n}$ . Note that the pixel at the location  $(m,n)$  is not included in  $X_{m,n}$ .

Given a discrete random field  $f(m,n)$  on an  $M \times N$  array of points, and assuming that the gray level  $f(i,j)$  at all points within  $X_{m,n}$  is known, we wish to estimate the gray level at  $(m,n)$  under the constraint that this estimate would be a linear function of the gray levels in  $X_{m,n}$ . In other words, denoting by  $\hat{f}(m,n)$  the optimal estimator for  $f(m,n)$  based on the gray levels at  $X_{m,n}$ , one may write

$$\hat{f}(m,n) = \sum_{\substack{i,j \\ \text{s.t.} \\ (m-i,n-j) \in X_{m,n}}} c_{i,j} f(m-i,n-j), \quad (2.2)$$

where the coefficients  $c_{i,j}$  are chosen such that the mean-square estimation error

$$e_{m,n} = E \left\{ \left[ f(m,n) - \hat{f}(m,n) \right]^2 \right\} \quad (2.3)$$

is minimized.  $\hat{f}(m,n)$  is termed the "linear least square estimate of  $f(m,n)$ ".

By substituting (2.2) in (2.3), differentiating with respect to each  $c_{i,j}$ , and setting each derivative equal to zero, the following set of equations is obtained for the unknown coefficients  $\{c_{i,j}\}$ :

$$E \left\{ \left[ f(m,n) - \hat{f}(m,n) \right] f(i,j) \right\} = 0, \quad \text{for all } (i,j) \in X_{m,n}. \quad (2.4)$$

The last result is merely the orthogonality principle in linear least squares estimation, which states that the optimal estimation coefficients must be such that the estimation error

$f(m, n) - \hat{f}(m, n)$  is statistically orthogonal to each  $f(i, j)$  that is used in forming the linear estimator.

We now proceed to the definition of wide-sense Markov fields. Let  $D$  represent the following pairs of  $(i, j)$ :

$$D = \{(0, 1), (1, 1), (1, 0)\}. \quad (2.5)$$

A random field will be called **wide-sense Markov** [1] if the coefficients  $c_{i,j}$  in (2.2) are such that  $\hat{f}$  may be written as

$$\hat{f}(m, n) = \sum_{(i,j) \in D} c_{i,j} f(m-i, n-j). \quad (2.6)$$

In other words, the least squares estimate of  $f(m, n)$  in terms of  $X_{m,n}$  is the same as the estimate in terms of only the three nearest neighbors on the left and above corner. Substituting (2.6) in (2.4), the following conditions that a wide-sense Markov random field must satisfy are obtained:

$$E \left\{ \left[ f(m, n) - \sum_{(i,j) \in D} c_{i,j} f(m-i, n-j) \right] f(p, q) \right\} = 0 \quad (2.7)$$

for all  $(p, q) \in X_{m,n}$ , i.e., for  $(p, q) = \{(m-1, n), (m-1, n-1), (m, n-1)\}$ . Inserting these values of  $(p, q)$  in (2.7) yields the following set of conditions, which can be solved for  $c_{1,0}, c_{1,1}, c_{0,1}$ :

$$c_{1,0} R_{ff}(0, 0) + c_{1,1} R_{ff}(0, 1) + c_{0,1} R_{ff}(-1, 1) = R_{ff}(-1, 0), \quad (2.8a)$$

$$c_{1,0} R_{ff}(0, 1) + c_{1,1} R_{ff}(0, 0) + c_{0,1} R_{ff}(-1, 0) = R_{ff}(-1, 1), \quad (2.8b)$$

$$c_{1,0} R_{ff}(1, -1) + c_{1,1} R_{ff}(1, 0) + c_{0,1} R_{ff}(0, 0) = R_{ff}(0, -1), \quad (2.8c)$$

where

$$R_{ff}(\alpha, \beta) = E \{ f(m, n) f(m + \alpha, n + \beta) \} \quad (2.9)$$

is the auto-correlation matrix of the Random field  $f(m, n)$ , and where the expectation value  $E \{ \cdot \}$  is taken over all the pairs  $f(m, n), f(m + \alpha, n + \beta)$ .

By now it is obvious that a wide-sense Markov random field is described by the following difference equation:

$$f(m, n) - \sum_{(i,j) \in D} c_{i,j} f(m-i, n-j) = \xi(m, n), \quad (2.10)$$

where  $\xi(m, n) \triangleq f(m, n) - \hat{f}(m, n)$  is the difference at each point.

The following theorem characterizes the difference  $\xi(m, n)$ .

**Theorem 2.1:** A discrete random field  $f(m, n)$  is wide-sense Markov if and only if it satisfies the following difference equation for all points  $(m, n)$  with  $m > 1$  and  $n > 1$ :

$$f(m, n) - \sum_{(i,j) \in D} c_{i,j} f(m-i, n-j) = \xi(m, n), \quad (2.11)$$

and for the points on the topmost row and the leftmost columns

$$f(1, 1) = \xi(1, 1) \quad (2.12a)$$

$$f(m, 1) - Af(m-1, 1) = \xi(m, 1), \quad m > 1 \quad (2.12b)$$

$$f(1, n) - Bf(1, n-1) = \xi(1, n), \quad n > 1, \quad (2.12c)$$

where the  $c_{i,j}$  are the solutions of (2.8), and

$$A = \frac{R_{ff}(1, 0)}{R_{ff}(0, 0)}, \quad B = \frac{R_{ff}(0, 1)}{R_{ff}(0, 0)}, \quad (2.13)$$

where  $\xi(m, n)$  is a discrete random field of orthogonal random variables, i.e.,

$$E\{\xi(m, n)\xi(p, q)\} = 0, \quad m \neq p \text{ or } n \neq q. \quad (2.14)$$

**Proof:** We first point out that (2.12b) and (2.12c) are consistent with (2.11) in the sense of the dependence of the estimator on points in  $X_{m,n}$ . Specifically, the least squares estimate  $\hat{f}(m, 1)$  of  $f(m, 1)$  in terms of all the points in  $X_{m,1}$  is the same as that in terms of  $f(m-1, 1)$  only, that is

$$\hat{f}(m, 1) = Af(m-1, 1). \quad (2.15)$$

By arguments similar to those leading to (2.7), the coefficient  $A$  must be such that

$$E\{[f(m, 1) - Af(m-1, 1)]f(p, 1)\} = 0, \quad \text{for all } (p, 1) \in X_{m,1}. \quad (2.16)$$

For  $p = m-1$ , (2.16) yields

$$A = \frac{R_{ff}(-1, 0)}{R_{ff}(0, 0)} = \frac{R_{ff}(1, 0)}{R_{ff}(0, 0)}. \quad (2.17)$$

Thus,

$$f(m, 1) = Af(m-1, 1) + \xi(m, 1), \quad (2.18)$$

where  $\xi(m,1) = f(m,1) - \hat{f}(m,1)$  is the error.

In a similar manner it is easy to show that (2.12c) holds, where the coefficient  $B$  is given by (2.13). We now proceed to the proof of the theorem. Assuming that the field is wide-sense Markov, we want to show that (2.14) is true. From (2.7), (2.16) we conclude that each  $\xi(m,n)$  is orthogonal to  $f(p,q)$  for all  $(p,q) \in X_{m,n}$ . It is easy to see that for any  $(k,l) \in X_{m,n}$ ,  $\xi(k,l)$  is a linear combination of  $f(p,q)$ 's in  $X_{m,n}$ . Therefore,  $\xi(m,n)$  must be orthogonal to  $\xi(k,l)$  for all  $(k,l) \in X_{m,n}$ . This must hold for every point  $(m,n)$  in the image, and thus

$$E\{\xi(p,q)\xi(m,n)\} = 0, \text{ for } p \neq m \text{ or } q \neq n, \quad (2.19)$$

which is (2.14).

Conversely, if (2.14) is true for  $\xi(m,n)$  as defined by (2.11), (2.12), we need to show that  $f(m,n)$  is a wide-sense Markov field. This is done as follows. We observe that  $f(2,1) = A\xi(1,1) + \xi(2,1)$  and that  $f(1,2) = B\xi(1,1) + \xi(1,2)$ , using (2.12b) and (2.12c), respectively. Inserting these equalities in (2.11), we obtain

$$f(2,2) = (c_{1,1} + Bc_{1,0} + Ac_{0,1})\xi(1,1) + c_{1,0}\xi(1,2) + c_{0,1}\xi(2,1) + \xi(2,2).$$

In a similar manner, any  $f(m,n)$  may be expressed as a linear combination of  $\xi(p,q)$ 's with  $(p,q) \in Y_{m,n}$ , as defined in Figure 2.2. Using (2.14), this implies that  $f(m,n)$  is orthogonal to any  $\xi(k,l)$  for which  $k > m$  or  $l > n$ . From this we may conclude that the estimation error  $\xi(k,l) = f(k,l) - \hat{f}(k,l)$  is orthogonal to all  $f(p,q)$  with  $(p,q) \in X_{k,l}$ , which, combined with (2.6), is exactly (2.7). Therefore, the random field  $f(m,n)$  is wide-sense Markov, which completes the proof of the theorem.

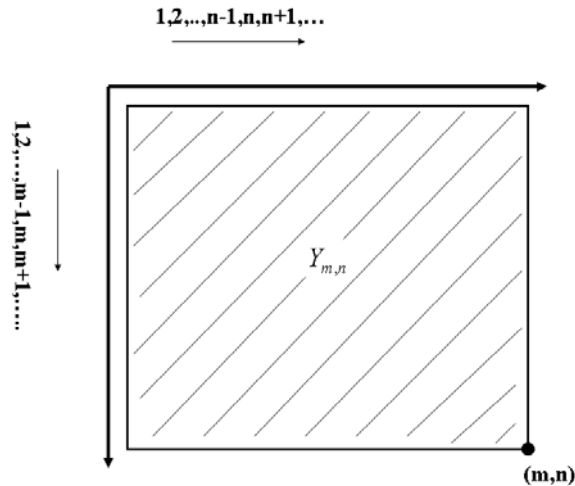


Figure 2.2 – The region  $Y_{m,n}$ . Note that here, the pixel in  $(m,n)$  is included in the region.

We proceed now to an interesting special case of a wide-sense Markov field, for which the autocorrelation function is

$$R_{ff}(\alpha, \beta) = \exp(-c_1 |\alpha| - c_2 |\beta|) . \quad (2.20)$$

Note that for this autocorrelation  $R_{ff}(0,0)=1$ , which means that the original image is first normalized by dividing it by  $\sqrt{R_{ff}(0,0)}$ . By this normalization we avoid the need to insert a multiplicative constant to  $R_{ff}(\alpha, \beta)$ .

The expression in (2.20) may be rewritten as

$$R_{ff}(\alpha, \beta) = \rho_h^{|\alpha|} \rho_v^{|\beta|} , \quad (2.21)$$

where  $\rho_h = e^{-c_1}$  and  $\rho_v = e^{-c_2}$  are measures of the horizontal and vertical correlation, respectively. Substituting (2.20) in (2.8), the following solution is obtained:

$$c_{1,0} = \rho_h, \quad c_{0,1} = \rho_v, \quad c_{1,1} = -\rho_h \rho_v . \quad (2.22a)$$

Hence, the autocorrelation in (2.21) may be rewritten as

$$R_{ff}(\alpha, \beta) = c_{10}^{|\alpha|} c_{01}^{|\beta|} . \quad (2.22b)$$

Also, inserting (2.21) into (2.13) yields:

$$A = c_{10}, \quad B = c_{01} . \quad (2.22c)$$

In the following sections we explore the special characteristics of this wide-sense random Markov random field. Specifically, we will examine the statistical properties of such an image under sampling. We first turn to a description of the sampling operation.

## 2.2 Sampling operation

In this section we introduce the sampling operation, which later on will be applied to a given Markov random field image.

For simplicity, we discuss a square image, whose dimensions are  $(2M) \times (2M)$ . We denote this image by  $f$  (the larger image or the high-resolution image).

Sampling  $f$  every other pixel yields a low-resolution version of this image, which we denote by  $f^*$  (the small image). Specifically, if  $k, l$  are the row and column indices of  $f^*$  respectively, then

$$f^*(k, l) = f(2k-1, 2l-1). \quad (2.23)$$



It is easy to see that the low-resolution version  $f^*$  is of dimensions  $M \times M$ , and thus, the high-resolution version  $f$  consists of  $3M^2$  more pixels than  $f^*$ . Figures 2.3 and 2.4 demonstrate the relation between  $f, f^*$ .

|          |          |          |          |          |          |
|----------|----------|----------|----------|----------|----------|
| $f_{11}$ | $f_{12}$ | $f_{13}$ | $f_{14}$ | $f_{15}$ | $f_{16}$ |
| $f_{21}$ | $f_{22}$ | $f_{23}$ | $f_{24}$ | $f_{25}$ | $f_{26}$ |
| $f_{31}$ | $f_{32}$ | $f_{33}$ | $f_{34}$ | $f_{35}$ | $f_{36}$ |
| $f_{41}$ | $f_{42}$ | $f_{43}$ | $f_{44}$ | $f_{45}$ | $f_{46}$ |
| $f_{51}$ | $f_{52}$ | $f_{53}$ | $f_{54}$ | $f_{55}$ | $f_{56}$ |
| $f_{61}$ | $f_{62}$ | $f_{63}$ | $f_{64}$ | $f_{65}$ | $f_{66}$ |

**Figure 2.3 – High resolution image**

|          |          |          |
|----------|----------|----------|
| $f_{11}$ | $f_{13}$ | $f_{35}$ |
| $f_{31}$ | $f_{33}$ | $f_{35}$ |
| $f_{51}$ | $f_{53}$ | $f_{55}$ |

**Figure 2.4 – Low resolution image (sampled image)**

## 2.3 Statistical properties of sampled Markov random field image

We now turn to examining the statistical properties of the high resolution and low resolution Markov random field images, given that the autocorrelation of the high resolution image is as in (2.21) and (2.22), i.e., has an exponential form. For simplicity, we avoid general indices  $i, j$ , and instead use numerical indices, as the ones that appear in Figures 2.3 and 2.4. This indices convention is valid due to the global nature of the definition of a wide-sense Markov field (see (2.6)), or the fact that the coefficients  $c_{01}, c_{10}, c_{11}$  are constant and that the white noise  $\xi$  is assumed to have the same variance throughout the whole image.

**Theorem 2.2:** Let  $f$  be a wide-sense Markov random field with an autocorrelation of the form:

$$R_{ff}(\alpha, \beta) = \rho_h^{|\alpha|} \rho_v^{|\beta|},$$

where

$$c_{1,0} = \rho_h, \quad c_{0,1} = \rho_v, \quad c_{1,1} = -\rho_h \rho_v.$$

Then, the sampled version  $f^*$  is also wide-sense Markov and its coefficients are given by

$$c_{10}^* = c_{10}^2, \quad c_{01}^* = c_{01}^2, \quad c_{11}^* = -c_{01}^* c_{10}^*. \quad (2.24)$$

**Proof:** We aim to express  $f_{55}$  using  $f_{33}, f_{35}, f_{53}$ , and then show that the pixels at these points are linearly related through constant coefficients that are the wide-sense Markov linear coefficients of the low-resolution image.

Using the Markov property of  $f$ , note that

$$f_{44} = c_{01}f_{43} + c_{11}f_{33} + c_{10}f_{34} + \xi_{44}, \quad (2.25a)$$

$$f_{45} = c_{01}f_{44} + c_{11}f_{34} + c_{10}f_{35} + \xi_{45}, \quad (2.25b)$$

$$f_{54} = c_{01}f_{53} + c_{11}f_{43} + c_{10}f_{44} + \xi_{54}, \quad (2.25c)$$

$$f_{55} = c_{01}f_{54} + c_{11}f_{44} + c_{10}f_{45} + \xi_{55}, \quad (2.25d)$$

where  $\xi_{mn}$  is the white noise generating the image  $f$ .

Inserting (2.25a), (2.25b) and (2.25c) into (2.25d), we obtain:

$$\begin{aligned} f_{55} = & c_{01} [c_{01}f_{53} + c_{11}f_{43} + c_{10}f_{44} + \xi_{54}] + \\ & c_{11} [c_{01}f_{43} + c_{11}f_{33} + c_{10}f_{34} + \xi_{44}] + \\ & c_{10} [c_{01}f_{44} + c_{11}f_{34} + c_{10}f_{35} + \xi_{45}] + \xi_{55}. \end{aligned} \quad (2.26)$$

Opening the brackets and rearranging so that the noise-dependant terms will appear at the end of the expression, we get

$$\begin{aligned} f_{55} = & c_{01}^2 f_{53} + c_{01} c_{11} f_{43} + c_{01} c_{10} f_{44} + c_{11} c_{01} f_{43} + \\ & + c_{11}^2 f_{33} + c_{11} c_{10} f_{34} + c_{10} c_{01} f_{44} + c_{10} c_{11} f_{34} + c_{10}^2 f_{35} + \\ & + c_{01} \xi_{54} + c_{11} \xi_{44} + c_{10} \xi_{45} + \xi_{55}. \end{aligned} \quad (2.27)$$

Inserting (2.25a) into (2.27),

$$\begin{aligned} f_{55} = & c_{01}^2 f_{53} + c_{01} c_{11} f_{43} + c_{01} c_{10} [c_{01} f_{43} + c_{11} f_{33} + c_{10} f_{34} + \xi_{44}] + \\ & + c_{11} c_{01} f_{43} + c_{11}^2 f_{33} + 2c_{11} c_{10} f_{34} + c_{10} c_{01} [c_{01} f_{43} + c_{11} f_{33} + c_{10} f_{34} + \xi_{44}] + c_{10}^2 f_{35} + \\ & + c_{01} \xi_{54} + c_{11} \xi_{44} + c_{10} \xi_{45} + \xi_{55}. \end{aligned} \quad (2.28)$$

Rearranging the last expression yields

$$f_{55} = c_{01}^2 f_{53} + c_{10}^2 f_{35} + (c_{11}^2 + 2c_{10}c_{01}c_{11})f_{33} + (2c_{10}^2c_{01} + 2c_{11}c_{10})f_{34} + (2c_{01}^2c_{10} + 2c_{11}c_{01})f_{43} + N, \quad (2.29)$$

where

$$N \triangleq 2c_{01}c_{10}\xi_{44} + c_{01}\xi_{54} + c_{11}\xi_{44} + c_{10}\xi_{45} + \xi_{55}, \quad (2.30)$$

Observing that in the last equality the terms  $f_{34}, f_{43}$  do not belong to the sampled image  $f^*$ , we proceed to express them in terms of the pixels that belong to  $f^*$ , that is:

$$f_{43} = c_{01}f_{42} + c_{11}f_{32} + c_{10}f_{33} + \xi_{43} \quad (2.31a)$$

$$f_{34} = c_{01}f_{33} + c_{11}f_{23} + c_{10}f_{24} + \xi_{34}. \quad (2.31b)$$

Inserting (2.31a), (2.31b) in (2.29) yields:

$$\begin{aligned} f_{55} = & c_{01}^2 f_{53} + c_{10}^2 f_{35} + (c_{11}^2 + 2c_{10}c_{01}c_{11})f_{33} + \\ & + (2c_{10}^2c_{01} + 2c_{11}c_{10})(c_{01}f_{33} + c_{11}f_{23} + c_{10}f_{24} + \xi_{34}) + \\ & + (2c_{01}^2c_{10} + 2c_{11}c_{01})(c_{01}f_{42} + c_{11}f_{32} + c_{10}f_{33} + \xi_{43}) + N \end{aligned} \quad (2.32)$$

Which, after rearranging becomes

$$f_{55} = c_{01}^2 f_{53} + c_{10}^2 f_{35} + (c_{11}^2 + 6c_{10}c_{01}c_{11} + 4c_{10}^2c_{01}^2)f_{33} + N_{total} \quad (2.33)$$

Where

$$\begin{aligned} N_{total} = & N + (2c_{10}^2c_{01} + 2c_{11}c_{10})(c_{11}f_{23} + c_{10}f_{24}) + \\ & + (2c_{01}^2c_{10} + 2c_{11}c_{01})(c_{01}f_{42} + c_{11}f_{32}). \end{aligned} \quad (2.34)$$

Inserting the expression for  $N$  from (2.30), we obtain

$$\begin{aligned} N_{total} = & 2c_{01}c_{10}\xi_{44} + c_{01}\xi_{54} + c_{11}\xi_{44} + c_{10}\xi_{45} + \xi_{55} + \\ & + (2c_{10}^2c_{01}c_{11} + 2c_{11}^2c_{10})f_{23} + (2c_{10}^3c_{01} + 2c_{11}c_{10}^2)f_{24} + \\ & + (2c_{01}^3c_{10} + 2c_{11}c_{01}^2)f_{42} + (2c_{01}^2c_{10}c_{11} + 2c_{11}^2c_{01})f_{32}. \end{aligned} \quad (2.35)$$

Our goal now is to show that  $N_{total}$  has the form of a white noise, which is not dependant on  $f_{24}, f_{42}, f_{23}, f_{32}$ . To obtain this result, we note that (2.22) implies that the wide-sense Markov field coefficients for the case of exponential autocorrelation are related by

$$c_{11} = -c_{10}c_{01}. \quad (2.36)$$

Using this last result, it is easy to observe that the terms that are dependant on  $f_{24}, f_{42}, f_{23}, f_{32}$  vanish:

$$2c_{10}^2 c_{01} c_{11} + 2c_{11}^2 c_{10} = -2c_{10}^3 c_{01}^2 + 2c_{10}^3 c_{01}^2 = 0 \quad (2.37a)$$

$$2c_{10}^3 c_{01} + 2c_{11} c_{10}^2 = 2c_{10}^3 c_{01} - 2c_{10}^3 c_{01} = 0 \quad (2.37b)$$

$$2c_{01}^3 c_{10} + 2c_{11} c_{01}^2 = 2c_{01}^3 c_{10} - 2c_{10} c_{01}^3 = 0 \quad (2.37c)$$

$$2c_{01}^2 c_{10} c_{11} + 2c_{11}^2 c_{01} = -2c_{01}^3 c_{10}^2 + 2c_{10}^2 c_{01}^3 = 0. \quad (2.37d)$$

Inserting (2.36) and (2.37) into (2.35), we obtain:

$$\begin{aligned} N_{total} &= 2c_{01} c_{10} \xi_{44} + c_{01} \xi_{54} - c_{01} c_{10} \xi_{44} + c_{10} \xi_{45} + \xi_{55} = \\ &= c_{01} c_{10} \xi_{44} + c_{01} \xi_{54} + c_{10} \xi_{45} + \xi_{55}. \end{aligned} \quad (2.38)$$

Returning to (2.33) and using (2.36) and (2.38) yields:

$$\begin{aligned} f_{55} &= c_{01}^2 f_{53} + c_{10}^2 f_{35} + (c_{01}^2 c_{10}^2 - 6c_{10}^2 c_{01}^2 + 4c_{10}^2 c_{01}^2) f_{33} + N_{total} = \\ &= c_{01}^2 f_{53} + c_{10}^2 f_{35} - c_{01}^2 c_{10}^2 f_{33} + N_{total}. \end{aligned} \quad (2.39)$$

Keeping in mind that  $f_{33}, f_{35}, f_{53}, f_{55}$  all belong to the sampled image  $f^*$  (see Figure 2.4), and recalling that  $f(m, n)$  is orthogonal to any  $\xi(k, l)$  for which  $k > m$  or  $l > n$  (see proof of Theorem 2.1), and also that  $\xi$  is a white noise, it is straightforward to see that (2.39) is no other than the difference equation of a wide-sense Markov field (see (2.10)). In other words, (2.39) may be rewritten as:

$$f_{55} = c_{01}^* f_{53} + c_{10}^* f_{35} + c_{11}^* f_{33} + \xi_{55}^*, \quad (2.40)$$

where:

$$c_{01}^* = c_{01}^2, \quad c_{10}^* = c_{10}^2, \quad c_{11}^* = -c_{10}^* c_{01}^* \quad (2.41a)$$

$$N_{total} = \xi_{55}^*. \quad (2.41b)$$

From (2.40),  $\xi^*$  is a white noise generating the wide-sense Markov field of the sampled image  $f^*$ .

To complete the proof, the boundaries of the sampled image  $f^*$  should be considered. To do so, we express the pixels that belong to the leftmost column of the image  $f$ :

$$f_{21} = A f_{11} + \xi_{21}, \quad (2.42a)$$

$$f_{31} = A f_{21} + \xi_{31}, \quad (2.42b)$$

which, combined together, yield

$$f_{31} = A(Af_{11} + \xi_{21}) + \xi_{31} = A^2 f_{11} + A\xi_{21} + \xi_{31}. \quad (2.43)$$

Inspection of (2.43) reveals that the pixels  $f_{31}, f_{11}$  (see Figure 2.4) on the leftmost boundary of the sampled image  $f^*$  are related in a similar manner to the pixels on the leftmost boundary of the high-resolution image  $f$  (see (2.12b)). In other words, (2.43) may be rewritten as

$$f_{31} = A^* f_{11} + \xi_{31}^*, \quad (2.44)$$

where

$$A^* = A^2 \quad (2.45a)$$

$$\xi_{31}^* = A\xi_{21} + \xi_{31}, \quad (2.45b)$$

and  $\xi^*$  is the white noise generating the leftmost boundary of the sampled image  $f^*$ .

It is easy to see that (2.45a) is in good agreement with (2.22b) and (2.41a), and indeed

$$A^* = A^2 = c_{10}^2 = c_{10}^*. \quad (2.46)$$

In a similar manner it is easy to show that  $B^* = B^2$  is the coefficient that corresponds to the generation of the topmost boundary of  $f^*$ .

Finally, recalling our convention that we use numerical indices for simplicity, and noting that the same procedure may be done for all the pixels that belong to  $f^*$ , the proof is completed.

□

Using Theorem 2.2, it is possible to derive relations between the statistics of the original and sampled images, and also between their generating white noise statistics (expectation value and variance). We start with the noise and then proceed to handling the statistical properties of the images.

**Theorem 2.3:** Let  $f$  be a wide-sense Markov random field, and let  $f^*$  be its sampled version.

- (a) Denoting by  $\sigma_{\xi}^2, \sigma_{\xi^*}^2$  the variances of the generating white noise of  $f, f^*$  respectively, the following relation holds:

$$\sigma_{\xi^*}^2 = (c_{01}^2 c_{10}^2 + c_{01}^2 + c_{10}^2 + 1) \sigma_{\xi}^2. \quad (2.47a)$$

- (b) Denoting by  $\mu, \mu^*$  the expectation values of the white noise of  $f, f^*$  respectively,

$$\mu_{\xi^*} = (c_{01} c_{10} + c_{01} + c_{10} + 1) \mu_{\xi}. \quad (2.47b)$$

**Proof:** We recall that  $\xi$  is a white noise, i.e., that  $\xi(m_1, n_1), \xi(m_2, n_2)$  are statistically independent for  $m_1 \neq m_2$  or  $n_1 \neq n_2$ . Using the well-known result that the variance of the sum of two independent random variables is simply the sum of their variances, we obtain from (2.38):

$$\sigma_{\xi^*}^2 = (c_{01} c_{10})^2 \sigma_{\xi}^2 + c_{01}^2 \sigma_{\xi}^2 + c_{10}^2 \sigma_{\xi}^2 + \sigma_{\xi}^2. \quad (2.48)$$

Simple rearranging yields (2.47a), which is the desired result.

In order to prove (2.47b), we take the expectation value  $E\{\cdot\}$  from both sides of (2.38), which gives

$$\mu_{\xi^*} = c_{01} c_{10} \mu_{\xi} + c_{01} \mu_{\xi} + c_{10} \mu_{\xi} + \mu_{\xi}. \quad (2.49)$$

(2.49) is identical to (2.47b), and this completes the proof of the theorem.  $\square$

**Theorem 2.4:** The **first moments** (i.e., the expectation value) of the high resolution image  $f$  and the low-resolution image  $f^*$  are equal.

**Proof:** We first formulate an expression to the expectation value of a wide-sense Markov random field. Taking the expectation value on both sides of the difference equation (2.11), we obtain:

$$E \left\{ f(m, n) - \sum_{(i,j) \in D} c_{i,j} f(m-i, n-j) \right\} = E \{ \xi(m, n) \}, \quad (2.50)$$

where the expectation is taken over all the points that belong to  $f$ . Using the linearity of the expectation operator  $E\{\cdot\}$ , denoting  $\mu = E\xi$  as the expectation value of the white noise  $\xi$  and  $F = E\{f\}$  as the expectation of the image  $f$ , we obtain:

$$F = \frac{\mu}{1 - c_{10} - c_{01} - c_{11}}. \quad (2.51)$$

In a similar manner, we may express the expectation of the sampled image  $F^* = E\{f^*\}$ :

$$F^* = \frac{\mu^*}{1 - c_{10}^* - c_{01}^* - c_{11}^*} . \quad (2.52)$$

Now, inserting the relations between the coefficients of the high-resolution image and the sampled image (see (2.41a)), and also the relation between the expectation value of the white noise in the original and the sampled image (see (2.47b)), we get

$$F^* = \frac{(c_{01}c_{10} + c_{01} + c_{10} + 1)\mu}{1 - c_{10}^2 - c_{01}^2 + c_{10}^2c_{01}^2} . \quad (2.53)$$

Using (2.51) to express  $\mu$  in terms of  $F$  and putting it in (2.52) yields:

$$F^* = \frac{(1 + c_{10} + c_{01} + c_{10}c_{01})(1 - c_{10} - c_{01} + c_{01}c_{10})F}{1 - c_{10}^2 - c_{01}^2 + c_{01}^2c_{10}^2} . \quad (2.54)$$

Noting that

$$\begin{aligned} & (1 + c_{01}c_{10} + c_{10} + c_{01})(1 + c_{01}c_{10} - (c_{10} + c_{01})) = \\ & = (1 + c_{01}c_{10})^2 - (c_{10} + c_{01})^2 = 1 + 2c_{01}c_{10} + c_{10}^2c_{01}^2 - c_{10}^2 - 2c_{10}c_{01} - c_{01}^2 = \\ & = 1 + c_{10}^2c_{01}^2 - c_{10}^2 - c_{01}^2, \end{aligned} \quad (2.55)$$

we obtain from (2.54):

$$F = F^* , \quad (2.56)$$

which completes the proof of the theorem.  $\square$

**Theorem 2.5:** Let  $f$  be a wide-sense Markov image with an exponential autocorrelation  $R_{ff}(\alpha, \beta) = c_{10}^{|\alpha|}c_{01}^{|\beta|}$  (see (2.22b)). Then, the autocorrelation of the sampled image  $f^*$  is given by sampling the autocorrelation of  $f$ , i.e.,

$$R_{ff}^*(\alpha, \beta) = R_{ff}(2\alpha, 2\beta). \quad (2.57)$$

**Proof:** Recalling the relation between the coefficients of the high-resolution image  $f$  and the low-resolution image  $f^*$  (see (2.41a)), we use (2.22b) to obtain

$$R_{ff}^*(\alpha, \beta) = (c_{10}^*)^{|\alpha|} (c_{01}^*)^{|\beta|} = (c_{10}^2)^{|\alpha|} (c_{01}^2)^{|\beta|} = (c_{10})^{2|\alpha|} (c_{01})^{2|\beta|} = R_{ff}(2\alpha, 2\beta), \quad (2.58a)$$

which is the desired result.  $\square$

A result of Theorem 2.5 is that  $f$  and  $f^*$  have equal second order moments, or in other words equal variances, since the second moment is by definition the value of the autocorrelation function at the origin:

$$Ef^2 = R_{ff}(\alpha, \beta) \Big|_{\alpha, \beta=0} = R_{ff}(2\alpha, 2\beta) \Big|_{\alpha, \beta=0} = Ef^{*2}. \quad (2.58b)$$

Note that the proof of Theorem 2.5 holds also for the case where the autocorrelation function is not normalized, i.e.,  $R_{ff}(0, 0) = \text{const} \neq 1$ . In such a case, the only modification is that  $R_{ff}(2\alpha, 2\beta), R_{ff}^*(\alpha, \beta)$  are multiplied by the constant  $\text{const}$ .

**Remark:** Although the relations for the variance of the noise in the original and the sampled image were obtained for pixels that are not located on the boundaries (leftmost column and topmost row) of the image, the results that were presented above do not change for these pixels. This may be easily seen by observing that (2.12) may be obtained from (2.11) by  $A = c_{10}, B = 0$  for (2.12b) and  $A = 0, B = c_{01}$  for (2.12c). Hence it follows that the pixels on the boundaries should obey the same statistical properties that were obtained for inner pixels. For example, inspection of (2.45) shows that  $\sigma_{\xi^*}^2 = (A^2 + 1)\sigma_{\xi}^2 = (c_{10}^2 + 1)\sigma_{\xi}^2$ , which is exactly (2.48) for  $c_{01} = 0$ .

To demonstrate Theorem 2.5, we created a wide-sense Markov image in Matlab, with the following parameters: linear coefficients  $c_{01} = c_{10} = 0.9$ , and variance of the generating white noise  $\sigma_{\xi}^2 = 0.1$  (see Figure 2.5). Then, we calculated the autocorrelation of the image, both theoretically (according to (2.22b)) and experimentally (see Figure 2.6). For the experimental calculation we used the following well-known formula, from the theory of nonparametric spectrum estimation:

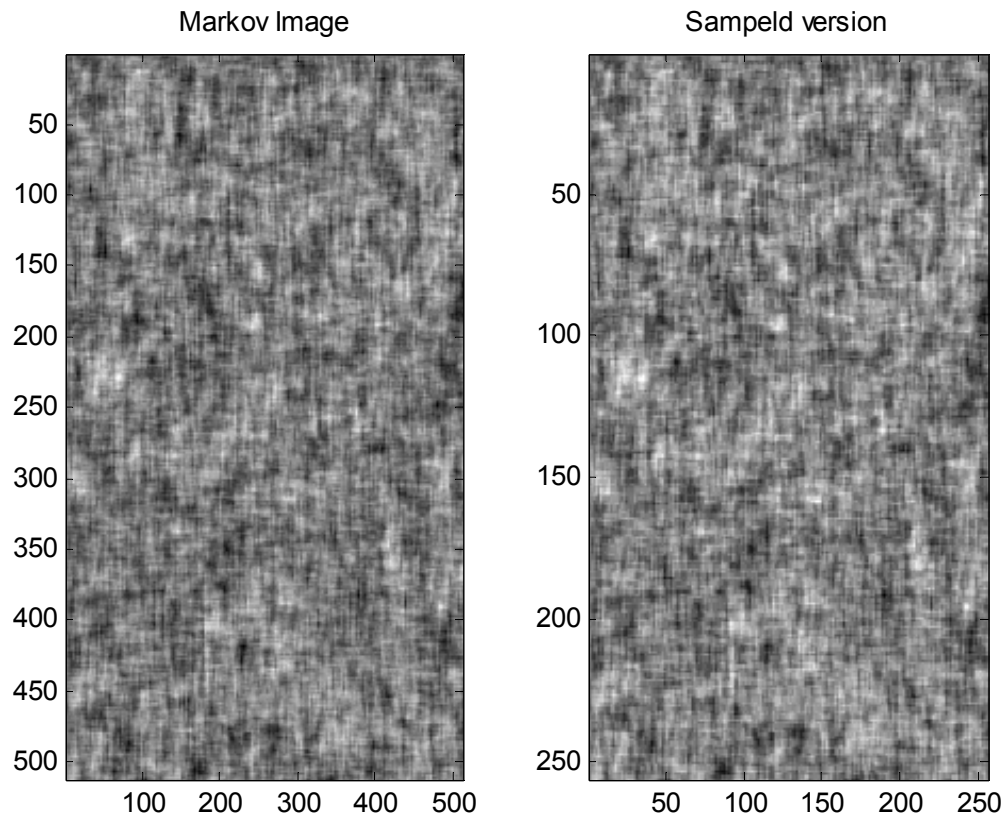
$$\hat{R}_{ff}(\alpha, \beta) = \begin{cases} \frac{1}{(M - |\alpha|)(N - |\beta|)} \sum_{m=0}^{M-|\alpha|-1} \sum_{n=0}^{N-|\beta|-1} f(m, n) f(m + |\alpha|, n + |\beta|) & , \text{sign}(\alpha \cdot \beta) > 0 \\ \frac{1}{(M - |\alpha|)(N - |\beta|)} \sum_{m=0}^{M-|\alpha|-1} \sum_{n=|\beta|+1}^N f(m, n) f(m + |\alpha|, n - |\beta|) & , \text{else.} \end{cases}$$

Note that this expression yields:  $\hat{R}_{ff}(\alpha, \beta) = \hat{R}(-\alpha, -\beta)$ .

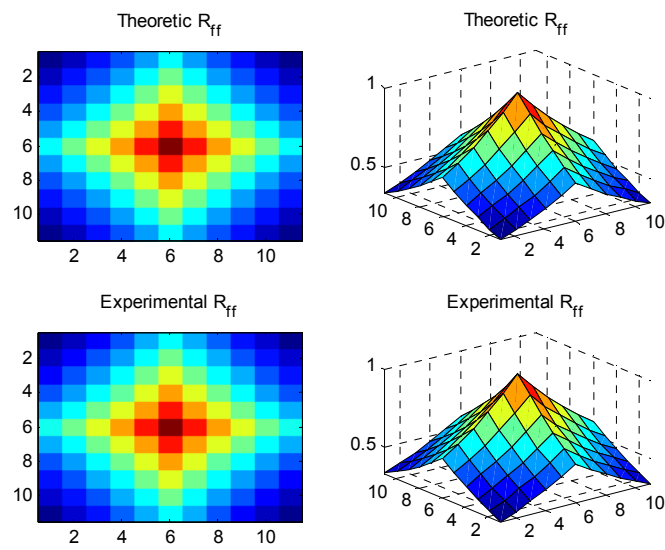
We then sampled the image to obtain a low-resolution version (shown in Figure 2.5), and calculated (experimentally) the autocorrelation of the sampled image (Figure 2.7). Finally, we



compared the autocorrelation of the sampled image to the sampled version of the autocorrelation of the original (high-resolution) image (also shown in Figure 2.7).

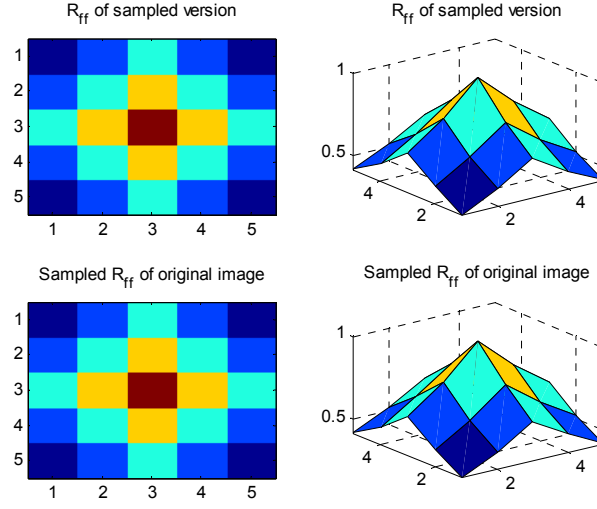


**Figure 2.5 – high and low resolutions of wide-sense Markov image**



**Figure 2.6 – Theoretic and experimental autocorrelation of high resolution image**

Numerical comparison between the theoretically and experimentally calculated autocorrelation shows that the deviation between the two is less than 5%.



**Figure 2.7 – autocorrelation of sampled image vs. sampled autocorrelation**

Numerical comparison of the sampled autocorrelation and the autocorrelation of the sampled image yields a deviation of less than 1% between the two. Thus, this experiment is in agreement with the theoretical result of Theorem 2.5.

## 2.4 High-order moment statistics for Gaussian Markov Random Fields

In the previous section we saw that for a wide-sense Markov field with an exponential autocorrelation, the two first moments (expectation value + variance) are invariant under sampling. This result was obtained under the sole assumption that the noise  $\xi$  is white. No further assumptions regarding the noise statistics (i.e., uniformly or normally distributed) were made. In this section we discuss the question of invariance of higher-order moments, or the equivalent question of invariance of the probability distribution of the image under sampling.

We first review a few major terms from probability theory, which will help us in developing our results.

**Definition 2.1:** Let  $X$  be a random variable with a probability density function  $f(x)$ . The **characteristic function** of  $X$  is given by

$$\Phi(\omega) = \int_{-\infty}^{+\infty} f(x) e^{j\omega x} dx. \quad (2.59)$$

This function has a maximum at the origin  $\omega = 0$  since  $f(x) \geq 0$ :

$$|\Phi(\omega)| \leq \Phi(0) = 1. \quad (2.60)$$

If  $j\omega$  is changed to  $s$ , the obtained integral

$$\Phi(s) = \int_{-\infty}^{+\infty} f(x) e^{sx} dx \quad (2.61)$$

is termed the **moment generating function** of  $X$ .

Note: we use the same notation,  $\Phi$ , for the characteristic function and for the moment generating function for simplicity. The distinction between the two functions is through the argument -  $s$  or  $\omega$ .

From (2.59), (2.61) it is clear that

$$\Phi(\omega) = E\{e^{j\omega x}\}, \quad \Phi(s) = E\{e^{sx}\}. \quad (2.62)$$

(2.62) leads to the following result:

$$\text{If } Y = aX + b, \text{ then } \Phi_Y(\omega) = e^{jb\omega} \Phi_X(a\omega). \quad (2.63)$$

A well known characteristic function is for the case of a normal (Gaussian) random variable, with mean  $\mu$  and variance  $\sigma^2$ , which equals

$$\Phi(\omega) = e^{j\eta\omega} e^{-\sigma^2\omega^2/2}. \quad (2.64)$$

(2.64) may be easily obtained by using (2.63) and the following integral:

$$\frac{1}{\sigma\sqrt{2\pi}} \int_{-\infty}^{+\infty} e^{-x^2/2\sigma^2} e^{j\omega x} dx = e^{-\sigma^2\omega^2/2}. \quad (2.65)$$

**Inversion Formula:** Inspection of (2.59) reveals that  $\Phi(-\omega)$  is the Fourier transform of  $f(x)$ , which means that the characteristic functions basically possess the same properties of

the Fourier transform. Moreover, the probability density distribution  $f(x)$  is actually the Fourier transform of the characteristic function  $\Phi(\omega)$ :

$$f(x) = \frac{1}{2\pi} \int_{-\infty}^{+\infty} \Phi(\omega) e^{-j\omega x} d\omega. \quad (2.66)$$

**Definition 2.2:** The **n-th moment** of a random variable  $X$  is

$$m_n = E\{X^n\} \quad (2.67)$$

**Moment Theorem:** The derivatives of  $\Phi(s)$  at the origin equal the moments of  $X$ , i.e.,

$$\Phi^{(n)}(0) = m_n. \quad (2.68)$$

**Proof:** Differentiating (2.61)  $n$  times with respect to  $s$  gives

$$\frac{\partial^n \Phi(s)}{\partial s^n} = \frac{\partial^n}{\partial s^n} \left( \int_{-\infty}^{+\infty} f(x) e^{sx} dx \right) = \int_{-\infty}^{+\infty} f(x) x^n e^{sx} dx. \quad (2.69)$$

Evaluating (2.69) at  $s = 0$ , we obtain

$$\left. \frac{\partial^n \Phi(s)}{\partial s^n} \right|_{s=0} = \left[ \int_{-\infty}^{+\infty} f(x) x^n e^{sx} dx \right]_{s=0} = \int_{-\infty}^{+\infty} f(x) x^n dx = m_n, \quad (2.70)$$

which is the desired result.

Note: the term "**moment generating function** of  $X$ " is justified by the result of this theorem.

Using the above results enables the formulation of the following relation between the moment generating function and the moments of a random variable. Expanding  $\Phi(s)$  into a series near the origin and using (2.68), one obtains

$$\Phi(s) = \sum_{n=0}^{\infty} \frac{m_n}{n!} s^n, \quad (2.71)$$

where (2.71) is valid only if all moments are finite and the series converges absolutely near  $s = 0$ . Since  $f(x)$  may be determined in terms of  $\Phi(s)$  (see (2.66)), the conclusion from (2.71) is that (under the stated convergence conditions) the density of a random variable is uniquely determined if all its moments are known.

**Note:** The above discussion and results were developed for the continuous case, but may be easily expanded to the discrete case; for example, suppose that  $X$  is a discrete random variable which takes the values  $x_i$  with probabilities  $p_i$ . The characteristic function in this case is of the following form:

$$\Phi(\omega) = \sum_i p_i e^{j\omega x_i}. \quad (2.72)$$

This expansion of the preceding results to the discrete case enables us to treat the gray level of the pixels of a Markov random field as a probability distribution function and discuss its properties in terms of the tools developed so far. We now proceed to doing so for the case of a wide-sense Markov random field with an exponential autocorrelation, generated by normally distributed white noise.

**Theorem 2.6:** Let  $f$  be a wide-sense Markov random field with an exponential autocorrelation function, generated by a normally distributed white noise. Let  $f^*$  be its low-resolution version. Then, the moments of  $f^*$  up to any order are equal to those of  $f$ ; i.e., the probability distribution of the gray levels of  $f$  is equal to the probability distribution of  $f^*$ .

**Proof:** We first show that the gray level of  $f$  is normally distributed. To see that, we observe that each pixel in  $f$  may be expressed as a linear combination of samples of the white Gaussian noise  $\xi$ . For example,

$$f_{11} = \xi_{11} \quad (2.73)$$

$$\begin{aligned} f_{22} &= c_{10}f_{12} + c_{01}f_{21} + c_{11}f_{11} + \xi_{22} = \\ &= c_{10}[c_{01}f_{11} + \xi_{12}] + c_{01}[c_{10}f_{11} + \xi_{21}] + c_{11}f_{11} + \xi_{22} = \\ &= 2c_{10}c_{01}f_{11} - c_{10}c_{01}f_{11} + c_{10}\xi_{12} + c_{01}\xi_{21} + \xi_{22} = \\ &= c_{10}c_{01}\xi_{11} + c_{10}\xi_{12} + c_{01}\xi_{21} + \xi_{22}. \end{aligned}$$

Now, since  $\xi$  is normally distributed **white** noise, the last result implies that each pixel in  $f$  is a linear combination of statistically independent normally distributed random variables, and thus each pixel in  $f$  is also a normally distributed random variable. Furthermore, if we take a linear combination of pixels in  $f$ , and substitute each pixel with its linear combination of samples of white noise (as in (2.73)), we finally obtain a linear combination of samples of

white noise as well, which, due to the normal distribution of the noise, is also normally distributed. Now, we recall the well-known result from probability theory which states that if **any** linear combination of normally distributed random variables is a normally distributed random variable as well, then the random variables have a **joint statistics** which is **normally distributed**. Using this result, we conclude that the joint statistics of the gray levels of pixels in  $f$  is normally distributed.

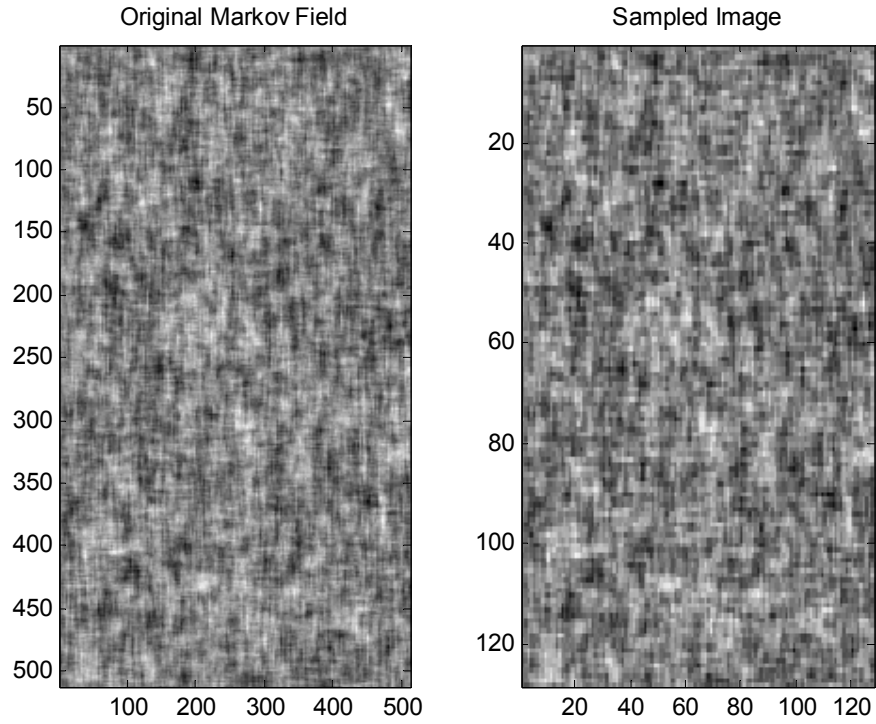
Furthermore, the low-resolution image  $f^*$  is also normally distributed, since as we saw in previous sections,  $f^*$  is also wide-sense Markov and the same arguments regarding the normal distribution may be applied to it.

We now use Theorems 2.4 and 2.5, according to which the first two moments of a wide-sense Markov random field  $f$  with an exponential autocorrelation are invariant under sampling. Since  $f$  and  $f^*$  are normally distributed, and since the characteristic function of a normally distributed random variable is dependent only on its first two moments (see (1.64)), we conclude that  $f$  and  $f^*$  have **equal characteristic functions**, and thus **equal generating moment functions** and **equal probability distribution functions**.  $\square$

To conclude Theorem 2.6, we proved that for a wide-sense random Markov field, with an exponential autocorrelation and normally distributed generating white noise, the **probability distribution is invariant under sampling**.

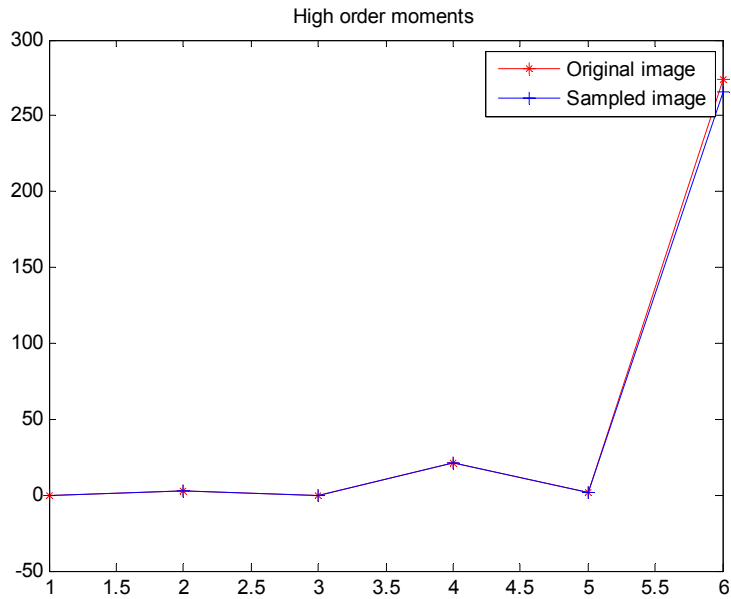
The following example, simulated in Matlab, demonstrates the results obtained so far.

In Figure 2.8 a wide-sense Markov field and its sampled (low-resolution) version are shown. The high-resolution (original) field was generated with the following parameters: variance and mean value of the normally distributed white noise:  $\sigma_\xi^2 = 0.1, \mu_\xi = 0$ ; vertical and horizontal correlation coefficients:  $c_{01} = c_{10} = 0.9$ .



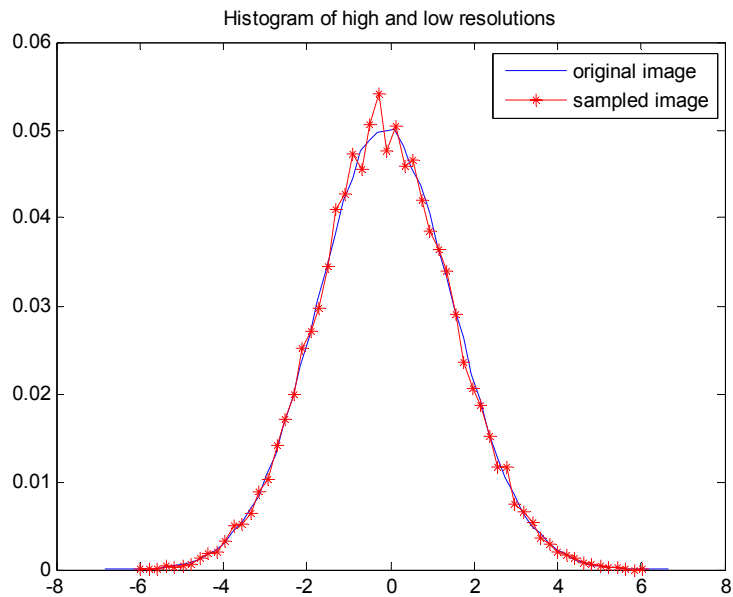
**Figure 2. 8 – High resolution and low resolution wide-sense Markov random field.**

In Figure 2.9 a comparison between the moments up till order 6 is presented. It is clear that the moments are almost identical. Note also that, as expected, the odd order moments are nearly zero – which is the case for normally distributed random variables with zero expectation value. Also, for the even order moments, the discrepancies between the 2, 4, 6 order moments of the high- and low-resolutions were 0.6% , 1.3% and 2.9% , respectively. We thus observe good fitting between theory and experimental results.



**Figure 2.9 – moments up till order 6 for the low and high resolutions.**

In Figure 2.10 the histograms of the high- and low-resolutions are shown. It is easy to see that the two histograms, which correspond to the probability distributions of the gray levels, are nearly identical. The small discrepancy stems from the fact that the histogram is merely a discrete approximation of the continuous probability density function.



**Figure 2.10 – Histograms of high and low resolutions.**



To further validate that the histograms in Figure 2.10 indeed correspond to normal distribution, we recall the following known formula for moments of even orders of zero-mean normal distribution:

$$E[X^{2k}] = \frac{(2k)!}{2^k \cdot k!} \sigma^{2k}. \quad (2.74)$$

For example, the 2<sup>nd</sup>, 4<sup>th</sup> and 6<sup>th</sup> moments are given by:

$$E[X^2] = \sigma^2, \quad E[X^4] = 3\sigma^4, \quad E[X^6] = 15\sigma^6. \quad (2.75)$$

To compare between the formula and the calculated moments, we first calculate the 2<sup>nd</sup>, 4<sup>th</sup> and 6<sup>th</sup>-order moments of the (high-resolution) Markov field  $f$  in the following manner:

$$\hat{m}_{2k} = \frac{\sum_{i=1}^N \sum_{j=1}^N f_{i,j}^{2k}}{N^2}. \quad (2.76)$$

This calculation yields:  $\hat{m}_2 = 2.6514$ ,  $\hat{m}_4 = 21.0527$ ,  $\hat{m}_6 = 273.7359$ . Now, we first evaluate  $\hat{\sigma}^2$  (the hat represents the fact that this is an estimation for  $\sigma^2$ ) by finding where the histogram of  $f$  falls to  $e^{-1/2}$  of its maximum value; this yields

$$\hat{\sigma}^2 = 2.465, \quad (2.77)$$

which gives 7% relative error with respect to  $\hat{m}_2$ . Then, we calculate:

$$3(\hat{m}_2)^2 = 21.09, \quad 15(\hat{m}_2)^3 = 279.59. \quad (1.78)$$

Clearly, (2.78) is in very good agreement with the theoretical result (2.75), which indicates that the histogram in Figure 2.10 is indeed a good approximation of normal distribution.

### 3. Application: Interpolation of Textures

In Section 2 the wide-sense Markov model was addressed, and theoretical results regarding its distribution were obtained. Specifically, according to Theorem 2.6, the probability distribution of the gray levels of the low-resolution field  $f^*$  is equal to the probability distribution of the high-resolution field  $f$ .

We now turn to an application of that result, which is the problem of **texture interpolation**. First, we give a short definition of texture and briefly address the problem of texture interpolation and its motivation. Then, we describe a model for texture analysis, by which a given texture may be analyzed and synthesized. Based on this model, we propose a new method for texture interpolation. The performance of the proposed method is demonstrated using experimental results on digitized image textures taken from Brodatz [3].

#### 3.1. The Texture Interpolation Problem

**Texture** may be defined as a structure which is made up of a large ensemble of elements that significantly resemble each other, organized according to some kind of ‘order’ in their locations. The elements are organized such that there is not a specific element that attracts the viewer's eyes, but the human viewer gets an impression of uniformity when he looks at the texture [4].

**Image textures** can be specified as: (1) grayscale or color images of natural textured surfaces, and (2) simulated patterns that approach these natural images [5].

The problem of texture interpolation may be stated as follows: given a down-sampled (low-resolution) texture, we would like to obtain an up-sampled (high-resolution) version of the same texture, such that the general visual impression will be preserved.

The problem of image interpolation (not necessarily textures) is very common and appears in many image processing applications. Considering the fact that many natural images may be described as a collage of various textures patches, or at least as a collection of textures and relatively smooth areas, it is easy to see that an efficient texture interpolation scheme may be useful. This is the motivation for the method proposed in this work.

#### 3.2. The Texture Model

In this section we briefly describe a texture model upon which we later base the proposed interpolation method. In this model, proposed by Francos et al. (see [4], [6], [7]), the texture field is assumed to be a realization of a 2D homogeneous random field. Based on a 2D Wold-

like decomposition of homogeneous random field, the texture field is decomposed into a sum of two mutually orthogonal, spatially homogeneous components: (1) the **global structural component**, which is **deterministic** in prediction theory sense; and (2) the **purely stochastic component**, or purely **in-deterministic** component.

The above decomposition may be formulated as follows. Denoting the homogeneous 2D texture field by  $\{y(m, n)\}$ , it can be represented by the orthogonal decomposition

$$y(m, n) = w(m, n) + v(m, n), \quad (3.1)$$

where  $w(m, n)$  is the purely in-deterministic component and  $v(m, n)$  is a deterministic field. It is important to note that the deterministic component  $v(m, n)$  may be further decomposed as

$$v(m, n) = h(m, n) + g(m, n), \quad (3.2)$$

where  $h(m, n)$  is called the **harmonic field**, and  $g(m, n)$  is termed the **generalized evanescent field**. Generally speaking, the harmonic field generates the periodic features of the texture field, while the evanescent component generates global directional features in the texture field. The spectral density function of the harmonic field is a sum of 2D delta functions, while the spectral density function of the generalized evanescent field is a sum of 1D delta functions that are supported on lines of rational slope in the frequency domain. A further discussion of this decomposition exceeds the scope of this paper; for an extensive treatment see [4], [6], [7].

For simplicity, in this paper we assume that  $g(m, n) = 0$ , that is, we treat the deterministic component  $v(m, n)$  as composed of only the harmonic field, i.e.,  $v(m, n) = h(m, n)$ . According to this assumption, we confine our discussion to textures that possess periodic features, and not directional ones. Experiments with many textures from [1] reveal that this assumption is not very restrictive, and in fact enables us to simply treat a wide variety of textures.

Returning to the harmonic field  $h(m, n)$ , it may be expressed as a countable sum of weighted 2D sinusoids, i.e.,

$$h(m, n) = \sum_{k=1}^P \{C_k \cos 2\pi(n\omega_k + mv_k) + D_k \sin 2\pi(n\omega_k + mv_k)\}, \quad (3.3)$$

where  $(\omega_k, \nu_k)$  are the spatial frequencies of the  $k$ 'th harmonic, and for practical use, the  $C_k$ 's and  $D_k$ 's may be treated as constant parameters.

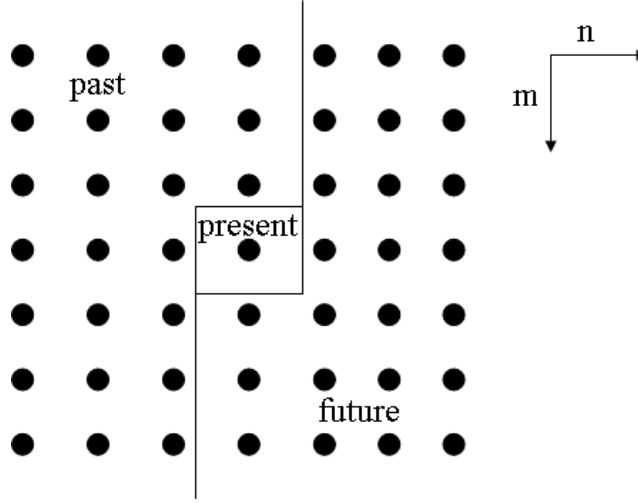
To complete this brief review of the model, we have to refer to the in-deterministic component  $w(m, n)$ . It is shown [4] that this component may be expressed according to a **2D auto-regressive (AR) model** (see [4], [6], [7] for a discussion regarding the auto-regressive modeling), i.e. ,

$$w(m, n) = - \sum_{(k, l) \in \chi_{m, n}} a(k, l) w(m - k, n - l) + u(m, n), \quad (3.4)$$

where  $u(m, n)$  is a 2D **white innovations field**,  $\{a(k, l)\}$  are the auto-regressive model parameters, and  $\chi_{m, n}$  is defined as follows:

$$\chi_{m, n} = \{(k, l) | l = n, k < m\} \cup \{(k, l) | l < n, \text{all } k\}. \quad (3.5)$$

Note that the definition of  $\chi_{m, n}$  in (3.5) differs from its definition in Section 2. Figure (3.1) clarifies this definition.



**Figure 3.1 – Two-dimensional support for the auto-regressive model**

Note that in Fig. (3.1), the pixel at location  $(m, n)$  is referred to as the "present", while the pixels that belong to  $\chi_{m, n}$  are referred to as the "past", and the remaining pixels as the "future". These terms allow us to define a 2D ordering relation, which is an extension to the 1D natural time ordering in the simple 1D auto-regressive model. Since the purely in-deterministic component  $w(m, n)$  in (3.4) is auto-regressive with relation to its "past" samples, it may be viewed as generated by a **causal** auto-regressive model.

In order to estimate the 2D AR model parameters  $\{a(k, l)\}$  (according to the formulation in (3.4)), we seek for a set of parameters that minimizes the innovation field variance, which is also the prediction error variance. Applying the orthogonality principle for optimal linear estimation, the following set of **normal equations** is obtained:

$$r(i, j) + \sum_{(k, l) \in \mathcal{K}(i, j)} a(k, l) r(i - k, j - l) = \begin{cases} \sigma^2, & \text{for } (i, j) = (0, 0) \\ 0, & \text{else} \end{cases}, \quad (3.6)$$

where  $\{r(i, j)\}$  is the autocorrelation sequence of the field  $w(m, n)$ , and  $\sigma^2$  is the variance of the innovation process (or the prediction error variance).

To conclude the review of this texture model, the generality of the above decomposition allows the modeling and parameterization of a wide variety of texture types. Finally, it is worth noting that this model is motivated by findings about human vision [4], [11].

### 3.3. New Method for Texture Interpolation

In this section we present a new method for texture interpolation. The method relies on the theoretical results that were obtained in Section 2, and uses the texture model of Francos et al., as reviewed in the previous section. The main idea behind the method is the following. Given a down-sampled (low-resolution) texture, we wish to extract its purely in-deterministic component, evaluate its optimal (low-resolution) auto-regressive model parameters, and use these parameters to generate the high-resolution purely in-deterministic component.

The harmonic component of the down-sampled texture is also extracted, and by filtering and zero padding in the frequency domain is perfectly interpolated to create its up-sampled version. The generated deterministic (i.e., harmonic) and in-deterministic components are then combined, to yield the high-resolution output texture.

Prior to formulating and elaborating on the stages of this method, we describe its motivation. Findings about the human visual system suggest that second-order statistics represent the textural information of homogenous purely in-deterministic random textures [4], [11]. The normal equations (3.6) imply that the auto-correlation function strongly depends on the AR model parameters, and thus we may conclude that ‘similar’ AR parameters correspond to visually similar textures.

We now recall that in Section 2 the wide-sense Markov field was addressed. This random field is defined such that each pixel is a linear combination of its three upper-left corner nearest neighbors, plus some additive white noise. Furthermore, for the case of a wide-sense

Markov field with an exponential auto-correlation, we saw that the probability distributions of its high and low resolution were equal (see Theorem 2.6).

Now, note that the wide sense Markov field is a special case of the auto-regressive model, and as such may be treated as a special case of an in-deterministic texture component. It is thus motivating to assume that the probability distributions of the low and high versions of the in-deterministic components of general auto-regressive causal model (not necessarily wide-sense Markov), will also be approximately equal and will have similar visual features, and as such – also similar autocorrelations. Thus, it is motivating to approximate the AR model parameters of the high-resolution version by the parameters of the low resolution, and in this manner obtain an interpolation from low to high resolution for the in-deterministic component.

The proposed interpolation method is applied as follows:

**Input:** An  $N \times M$  low resolution texture image, denoted by  $y^*$ .

**Output:** An  $2N \times 2M$  high resolution texture image, denoted by  $y$ .

**Step 1:** Extraction and interpolation of the harmonic component:

**1.1** Calculation of the DFT and **periodogram** (i.e., squared absolute value of the DFT) of  $y^*$ .

**1.2** Finding the peaks (2D delta functions) of the periodogram from step 1.1, and creating a frequency domain filter (mask) whose value is 1 at the locations of those peaks and 0 elsewhere.

**1.3** Filtering the DFT of  $y^*$  with the frequency domain mask from step 1.2, to obtain the DFT of the harmonic component of  $y^*$ .

**1.4** Zero padding the filtered DFT of  $y^*$  from step 1.3 in order to obtain the DFT which corresponds to the harmonic component of the high resolution  $y$ .

**1.5** Applying the inverse DFT on the output of 1.4, to obtain the estimated harmonic component of the high resolution, denoted by  $\hat{h}$ .

**Step 2:** Extraction and interpolation of the purely in-deterministic component:

**2.1** Filtering the DFT of  $y^*$  with a mask that is the negative of the mask in Step 1.2, to obtain the DFT of the purely in-deterministic component of the low resolution.

**2.2** Applying inverse DFT on the output of Step 2.1, to obtain an estimation of the (space domain) purely in-deterministic component in low resolution, denoted by  $\hat{w}^*$ .

**2.3** Choosing the number of AR model parameters (AR order), or in other words the size of the past support that is used in the model.

**2.4** Estimating the optimal AR model parameters according to the minimum innovation process variance criterion, using a least squares solution. These AR model parameters are estimated from the low-resolution version.

**2.5** Re-arranging the pixels of  $\hat{w}^*$  in a  $2N \times 2M$  matrix, such that sampling of this matrix yields again the low resolution  $\hat{w}^*$  (as in Figures 3.3, 3.4).

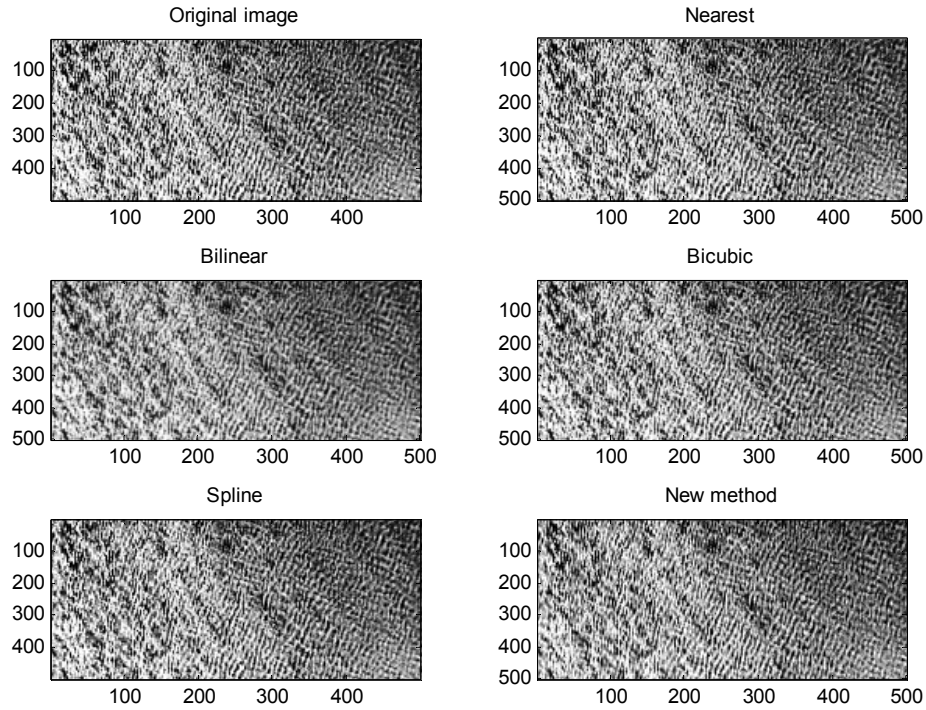
**2.6** Calculating the gray levels at the locations of missing pixels (i.e., that are not populated by the original low-resolution pixels' values) according to the AR model (3.4), and using the model parameters that were obtained from the low resolution (in Step 2.4). The innovation process  $u$  that drives the model is white Gaussian noise with some pre-determined variance  $\sigma^2$ . Pixels whose support neighborhood exceeds the boundaries of the image are initialized simply by a white noise with the same distribution like the innovation process.

**2.7** The output of Step 2.6 is the estimated purely in-deterministic component in the high resolution, denoted by  $\hat{w}$ .

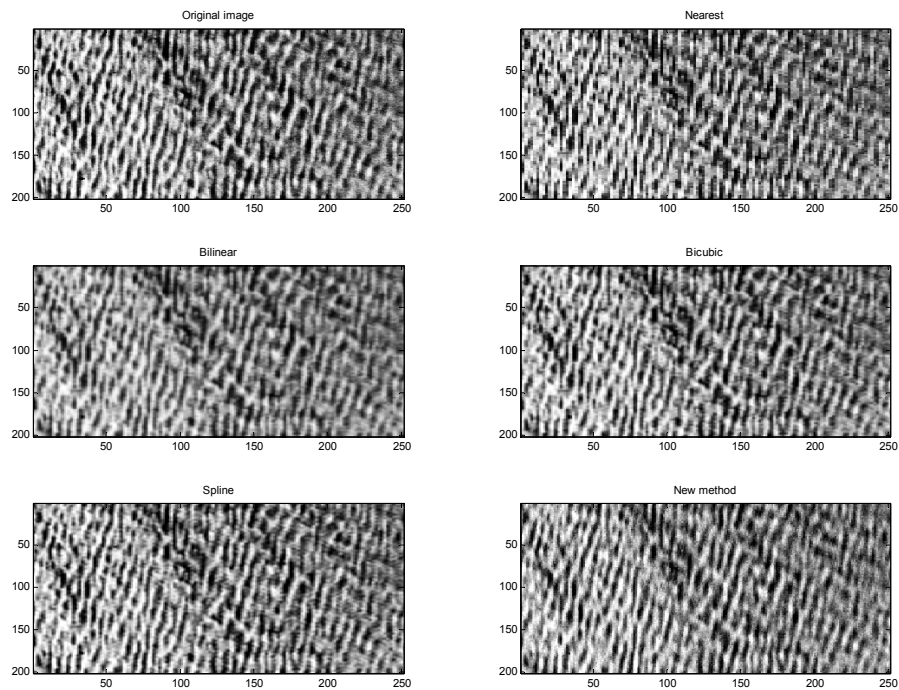
**Step 3:** Combining the estimated harmonic and purely in-deterministic components from Steps 1 and 2, to obtain the estimated texture image in high resolution, denoted by  $\hat{y}$ , i.e.,  $\hat{y} = \hat{h} + \hat{w}$ .

## 4.4 Experimental Results

In this section experimental results of the proposed method are shown. The textures were taken from the Brodatz album. Each original  $2N \times 2M$  texture was sampled such that an  $N \times M$  texture was obtained. Then, various interpolation methods were applied on the sampled texture and compared to the original (high-resolution) texture in terms of visual impression and PSNR. The texture ‘water’ and its interpolation using various interpolation methods (nearest neighbor, bilinear, bicubic and spline) is shown in Figure 3.2. A close look on a patch from the texture is shown in Figure 3.3.



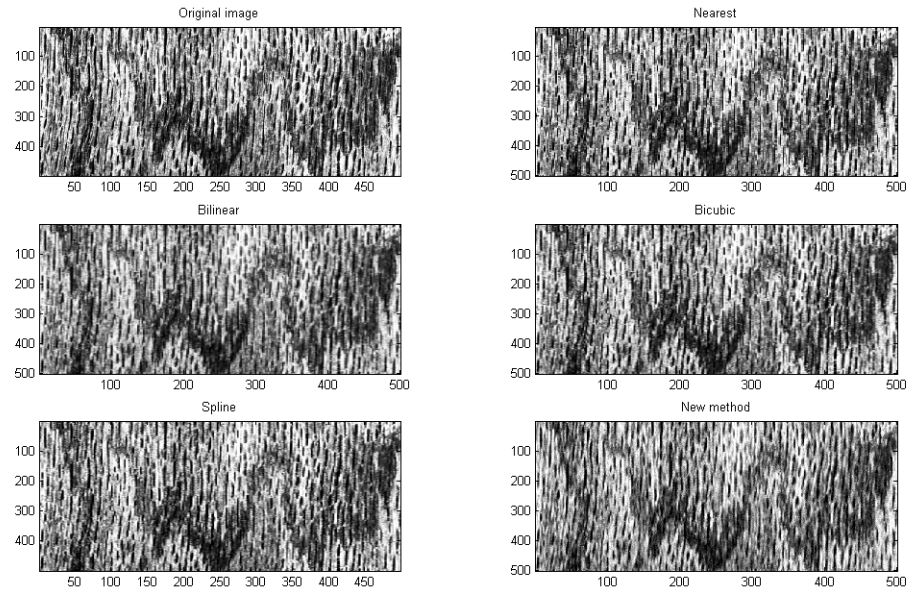
**Figure 3.2 – Interpolation results for the texture 'water'.**



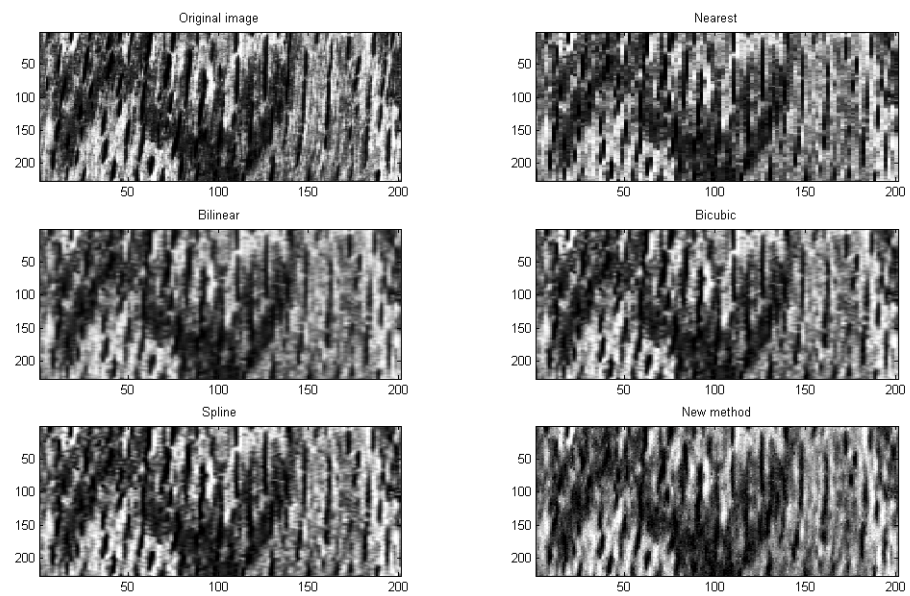
**Figure 3.3 – A close up on a patch from the textures in Fig. 3.2 (water).**

Similar figures (3.4-3.7) are shown below, for the textures 'wood grain' and 'herring bone weave'.

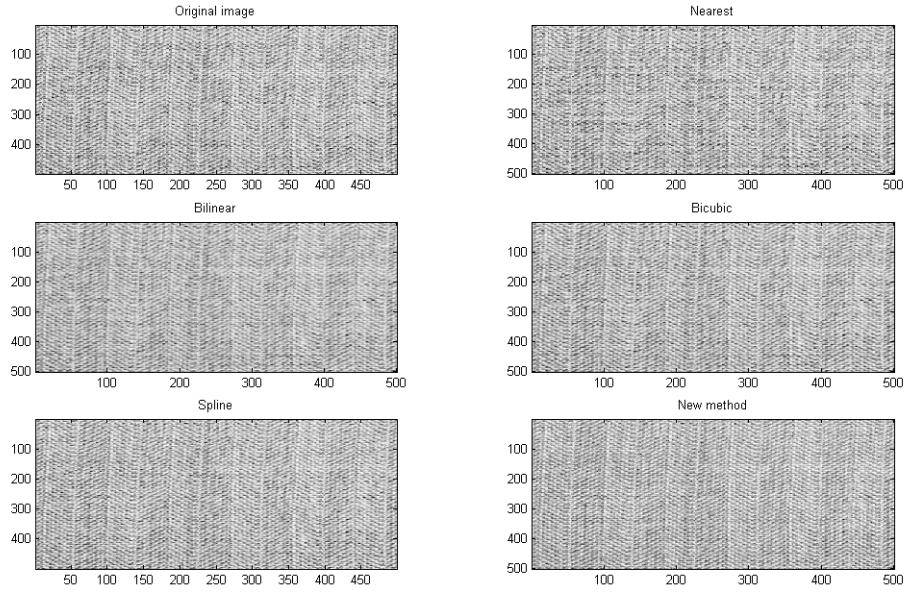




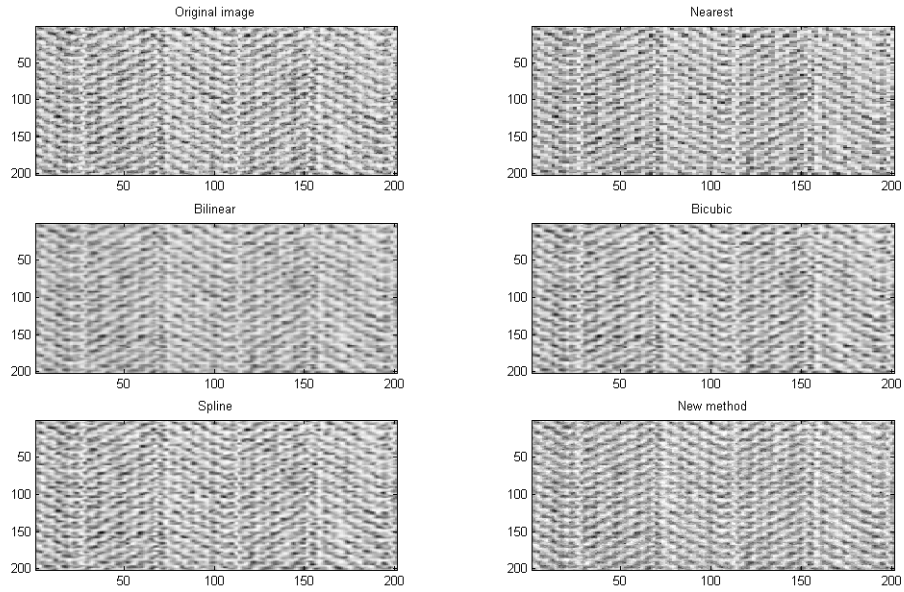
**Figure 3.4 - Interpolation results for the texture ‘Wood grain’**



**Figure 3.5 - A close up on a patch from the textures in Figure 4.4 (wood grain).**



**Figure 3.6 - Interpolation results for the texture "Herringbone weave".**



**Figure 3.7 - A close up on a patch from the textures in Figure 3.6 (herringbone weave).**

Inspection of Figures 3.2-3.7 reveals that the proposed texture interpolation method achieves better interpolation results in the sense that it is less blurry and preserves the fine features of the original texture. The other interpolation methods have a smoothing effect, which tends to blur out the details of the textures.

PSNR values (in dB) for the different methods on the three textures (water, wood grain and herringbone weave) are shown in Table 3.1. The PSNR was calculated according to the following formula:

$$PSNR = 10 \log_{10} \left( \frac{(\max(f))^2}{MSE(f, \hat{f})} \right), \quad (3.7)$$

where  $f$  is the original (high resolution)  $(2N) \times (2M)$  texture image,  $\max(f)$  is its highest gray level value, and  $MSE(f, \hat{f})$  is the mean square error of the interpolated image  $\hat{f}$  and the original image, which is calculated as

$$MSE(f, \hat{f}) = \frac{\sum_{i=1}^{2N} \sum_{j=1}^{2M} (f - \hat{f})^2}{(2N) \cdot (2M)}. \quad (3.8)$$

|                   | Nearest | Bilinear | Bicubic | Spline | New Method |
|-------------------|---------|----------|---------|--------|------------|
| Water             | 14.81   | 16.44    | 16.36   | 18.88  | 16.79      |
| Wood Grain        | 14.95   | 15.99    | 15.83   | 17.9   | 16.3       |
| Herringbone Weave | 19.39   | 20.6     | 20.46   | 23.52  | 21.52      |

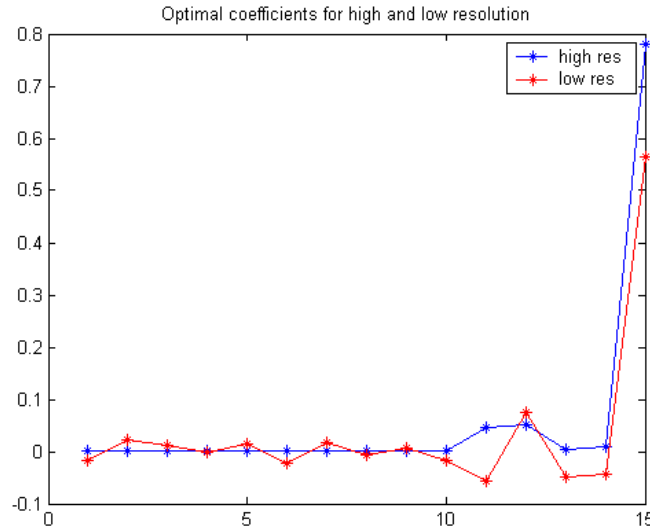
**Table 3.1 – PSNR in dB for different interpolation methods, including the proposed method.**

Inspection of Table 3.1 shows that the new (proposed) method achieves significantly better PSNR results than nearest neighbor, bilinear and bicubic interpolations. It is worth noting that although the spline interpolation achieves higher PSNR values than the proposed method, it still suffers from the blurriness effect, and the proposed method achieves better visual results (see Figures 3.3, 3.5 and 3.7 for a better visual comparison).

The fact that the proposed method achieves better visual results, and still a lower PSNR value in comparison with the spline interpolation may be explained by the realization of the proposed method, which generates the high-resolution in-deterministic component using additive white noise (the 2D AR innovation process). This inserts a random component into the interpolated texture, which results in a lower PSNR value (compared to spline interpolation), but also with a better visual resemblance to the original texture, which is successfully modeled as the sum of a deterministic (harmonic) and in-deterministic components. In other words, since the proposed method is adjusted to the specific model of the texture, it is able to interpolate it successfully.

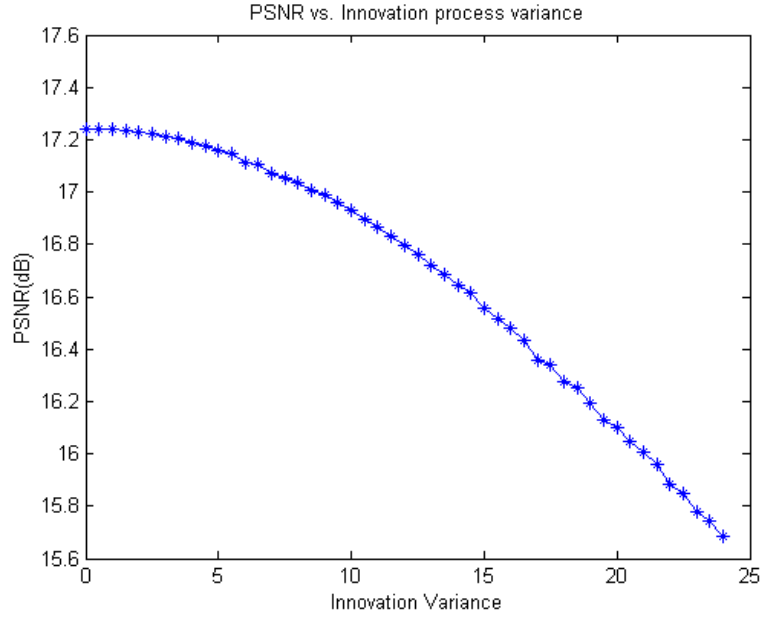
Note also that the PSNR values in Table 3.1 are relatively small. This is due to the noisiness of the in-deterministic component of the texture, which is smoothed by the interpolation and thus increases the interpolation error with respect to the original texture image.

To complete this section, we refer to Figures 3.8 and 3.9. In Figure 3.8 the calculated optimal 2D AR model parameters are shown, for the high and low resolutions of the texture ‘Wood grain’. It is obvious that the two parameter sets resemble each other, which is in correspondence with the discussion of Section 3.3.



**Figure 3.8 – Optimal in-deterministic AR model parameters for low and high resolution, calculated on ‘Wood grain’ texture.**

In Figure 3.9 the PSNR (in dB) of the proposed method versus the innovation process variance is shown, for the case of the ‘water’ texture. As expected, and as discussed above, the PSNR increases as the variance of the innovation process decreases. This may imply that the best choice one can do (in terms of PSNR) is to set the variance of the innovation process to be zero, but it is important to note that visually, this choice doesn't always yield the best visual resemblance to the original texture. In practice, a good rule of thumb is to choose about 0.7-0.8 of the estimated innovation process variance of the low resolution (i.e., the input texture). According to this, the PSNR values of the proposed method (shown in Table 3.1) may be further improved by setting the innovation variance to zero, but we chose to present the PSNR values which correspond to better visual performance, and those were calculated according to the above rule of thumb.



**Figure 3.9 – PSNR vs. Innovation process variance for the case of 'water' texture.**

## 4. Summary and Discussion

In this work we developed new mathematical relations between the high and low resolutions of wide-sense Markov random fields. In particular, we proved that under certain conditions, the probability distribution functions of the low and high resolutions are equal. This result motivated us to propose a new method for texture interpolation, based on an orthogonal decomposition texture model.

The proposed method performance was tested on texture images taken from the Brodatz album. The experimental results demonstrated the advantages of the proposed method over existing interpolation methods. It was shown that the new method achieves better visual performance than nearest neighbor, bilinear, bicubic and spline interpolations, and in addition higher PSNR values than nearest neighbor, bilinear and bicubic interpolations.

Future research will focus on interpolation of color textures (using an extended model for color textures) and further analysis of the relations between the high and low resolution statistics of a general 2D AR model. Specifically, it would be desirable to obtain explicit mathematical relations between the optimal model parameters for the high and low resolutions.

## References

- [1] A. Rosenfeld and A. C. Kak, *Digital Picture Procesing*, Academic Press., London, 1981.
- [2] R. Chellapa and A. Jain, *Markov Random Fields: Theory and Application*, Academic Press, 1991.
- [3] Brodatz, P., *Textures: A Photographic Album for Artists and Designers*, Dover Publications, New York, 1966.
- [4] J. M. Francos, A. Z. Meiri and B. Porat, *A Unified Texture Model Based on a 2-D Wold-Like Decomposition*, IEEE Transactions on Signal Processing, Vol. 41, No. 8, August 1993.
- [5] G. L. Gimel'farb, *Image Textures and Gibbs Random Fields*, Klumer Academic Publishers, 1999.
- [6] J. M. Francos, A. Z. Meiri and B. Porat, *A 2-D Autoregressive, Finite Support, Causal Model for Texture Analysis and Synthesis*, IEEE, 1989.
- [7] J. M. Francos, A. Z. Meiri and B. Porat, *On a Wold-Like Decomposition of 2-D Discrete Random Fields*, IEEE, 1990.
- [8] B. Porat, *A Course in Digital Signal Processing*, John Wiley & Sons, 1997.
- [9] B. Porat, *Digital Processing of Random Signals*, Prentice-Hall Inc, 1994.
- [11] B. Julesz, *A theory of Preattentive Texture Discrimination Based on First-Order Statistics of Textons*, Biological Cybernetics, 1981, pp. 131-138.
- [12] A. Papoulis, *Probability, Random variables and Stochastic Processes*, second edition, McGraw Hill, 1984.
- [13] R. Chellapa and R. L. Kashyap, *Texture Synthesis Using 2-D Noncausal Autoregressive Models*, IEEE Transactions on Acoustics, Speech and Signal Processing, Vol. ASSP-33, No. 1, 1985, pp 194-203.