

A Statistical Approach to Volume Data Quality Assessment

Chaoli Wang, *Member, IEEE*, and Kwan-Liu Ma, *Senior Member, IEEE*

Abstract—Quality assessment plays a crucial role in data analysis. In this paper, we present a reduced-reference approach to volume data quality assessment. Our algorithm extracts important statistical information from the original data in the wavelet domain. Using the extracted information as feature and predefined distance functions, we are able to identify and quantify the quality loss in the reduced or distorted version of data, eliminating the need to access the original data. Our feature representation is naturally organized in the form of multiple scales, which facilitates quality evaluation of data with different resolutions. The feature can be effectively compressed in size. We have experimented with our algorithm on scientific and medical data sets of various sizes and characteristics. Our results show that the size of the feature does not increase in proportion to the size of original data. This ensures the scalability of our algorithm and makes it very applicable for quality assessment of large-scale data sets. Additionally, the feature could be used to repair the reduced or distorted data for quality improvement. Finally, our approach can be treated as a new way to evaluate the uncertainty introduced by different versions of data.

Index Terms—Quality assessment, reduced reference, wavelet transform, statistical modeling, generalized Gaussian density, volume visualization.

1 INTRODUCTION

LEVERAGING the power of supercomputers, scientists can now simulate many things from galaxy interaction to molecular dynamics in unprecedented details, leading to new scientific discoveries. The vast amounts of data generated by these simulations, easily reaching tens of terabytes, however, present a new range of challenges to traditional data analysis and visualization. A time-varying volume data set produced by a typical turbulent flow simulation, for example, may contain thousands of time steps with each time step having billions of voxels and each voxel recording dozens of variables. As supercomputers continue to increase in size and power, petascale data is just around the corner.

A variety of data reduction methods have been introduced to make the data movable and enable interactive visualization, offering scientists options for studying their data. For instance, subsets of the data may be stored at a reduced precision or resolution. Data reduction can also be achieved with transform-based compression methods. A popular approach is to generate a multiresolution representation of the data such that a particular level of details is selected according to the visualization requirements and available computing resources. In addition, data may be altered in other fashions. Furthermore, it could be desirable to smooth the data or enhance a particular aspect of the data before rendering. Finally, the data may be distorted or corrupted during the transmission over a network.

Research has been conducted to evaluate the quality of rendered images *after* the visualization process [6], [29]. However, few studies focus on analyzing the data quality *before* the visualization actually takes place. It is clear that the original volume data may undergo various changes due to quantization, compression, sampling, filtering, and transmission. If we assume that the original data has full quality, all these changes made to the data may incur quality loss, which may also affect the final visualization result. In order to compare and possibly improve the quality of the reduced or distorted data, it is important for us to identify and quantify the loss of data quality. Unequivocally, the most widely used data quality metrics are *mean square error* (MSE) and *peak signal-to-noise ratio* (PSNR). Although easy to compute, they do not correlate well with perceived quality measurement [18]. Moreover, these metrics require access to the original data and are *full-reference* methods. They are not applicable to our scenario, since the original data may be too large to acquire or compare in an efficient way. Therefore, it is highly desirable to develop a data quality assessment method that does not require a full access of the original data.

In this paper, we introduce a *reduced-reference* approach to volume data quality assessment. We consider the scenario where a set of important statistical information is first extracted from the original data. For example, the extraction process could be performed at the super-computer centers where the large-scale data are produced and stored or, ideally, in situ when the simulation is still running. We then compress the feature information to minimize its size. This makes it easy to transfer the feature to the user as “carry-on” information for volume data quality assessment, eliminating the need to access the original data again. Our feature representation not only serves as the criterion for data quality assessment but also could be used as quality improvement to repair the reduced or distorted data. This is achieved by matching some of its

- The authors are with the Visualization and Interface Design Innovation (VIDI) Research Group, Department of Computer Science, University of California, Davis, 2063 Kemper Hall, One Shields Avenue, Davis, CA 95616. E-mail: {wangcha, ma}@cs.ucdavis.edu.

Manuscript received 18 Aug. 2007; revised 23 Oct. 2007; accepted 19 Nov. 2007; published online 3 Dec. 2007.

For information on obtaining reprints of this article, please send e-mail to: tvccg@computer.org, and reference IEEECS Log Number TVCGSI-2007-08-0111.

Digital Object Identifier no. 10.1109/TVCG.2007.70628.

feature components with those extracted from the original data. We have tested our algorithm on scientific and medical data sets with various sizes and characteristics to demonstrate its effectiveness.

2 BACKGROUND AND RELATED WORK

Unlike the Fourier transform with sinusoidal basis functions, the wavelet transform is based on small waves, called *wavelets*, of varying frequency and limited duration [7]. Wavelet transforms provide a convenient way to represent localized signals simultaneously in space and frequency. The particular kind of dual localization makes many functions and operators using wavelets “sparse” when transformed into the wavelet domain. This sparseness, in turn, brings us a number of useful applications such as data compression, feature detection, and noise removal.

Besides sparseness, wavelets have many other favorable properties, such as multiscale decomposition structure, linear time, and space complexity of the transformations, decorrelated coefficients, and a wide variety of basis functions. Studies of the *human visual system* (HVS) support a multiscale analysis approach, since researchers have found that the visual cortex can be modeled as a set of independent channels, each with a particular orientation and spatial frequency tuning [4], [21]. Therefore, wavelet transforms have been extensively used to model the processing in the early stage of biological visual systems. They have also gained much popularity and have become the preferred form of representation for image processing and computer vision algorithms.

In volume visualization, Muraki [15] introduced the idea of using the wavelet transform to obtain a unique shape description of an object, where a 2D wavelet transform is extended to three-dimensional (3D) and applied to eliminate wavelet coefficients of lower importance. Over the years, many wavelet-based techniques have been developed to compress, manage, and render 3D [8], [11] and time-varying volumetric data [12], [23]. They are also used to support fast access and interactive rendering of data at runtime. In this paper, we employ the wavelet transform to generate multiscale decomposition structures from the input data for feature analysis.

The wavelet transform on a one-dimensional signal can be regarded as filtering the signal with both the scaling function (a low-pass filter) and the wavelet function (a high-pass filter), and downsampling the resulting signals by a factor of two. The extension of the wavelet transform to a higher dimension is usually achieved using separable wavelets, operating on one dimension at a time. The 3D wavelet transform on volume data is illustrated in Fig. 1. After the first iteration of wavelet transform, we generate one low-pass filtered wavelet subband (LLL_1) with $1/8$ of the original size and seven high-pass filtered subbands (HLL_1 , LHL_1 , HHL_1 , LLH_1 , HLH_1 , LHH_1 , and HHH_1). We can then successively apply the wavelet transform to the low-pass filtered subband, thus creating a multiscale decomposition structure (a good introduction of wavelets for computer graphics can be found in [24]). In our experiment, the number of decomposition levels is usually between three and five, depending on the size of original data.

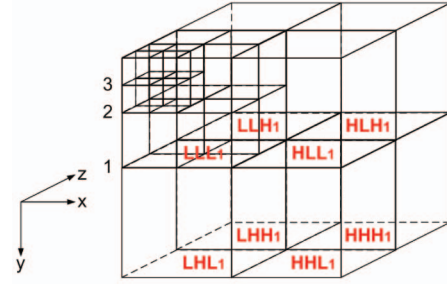


Fig. 1. Multiscale wavelet decomposition of a 3D volumetric data. L = low-pass filtered; H = high-pass filtered. The subscript indicates the level and a larger number corresponds to a coarser scale (lower resolution). An example of three levels of decomposition is shown here.

There is a wealth of literature on quality assessment and comparison in the field of image and video processing. Details on this are beyond the scope of this paper, and we refer interested readers to [2] for a good survey. Here, we specifically review some related work in the field of graphics and visualization. Jacobs et al. [9] proposed an image querying metric for searching in an image database using a query image. Their metric makes use of multiresolution wavelet decompositions of the query and database images and compares how many significant wavelet coefficients the query has in common with potential targets. In [6], Gaddipati et al. introduced a wavelet-based perceptual metric that builds on the subband coherent structure detection algorithm. The metric incorporates aspects of the HVS and modulates the wavelet coefficients based on the contrast sensitivity function. Sahasrabudhe et al. [18] proposed a quantitative technique that accentuates differences in images and data sets through a collection of partial metrics. Their spatial domain metric measures the lack of correlation between the data sets or images being compared. Recently, Wang et al. [26] introduced an image-based quality metric for interactive level-of-detail selection and rendering of large volume data. The quality metric design is based on an efficient way to evaluate the contribution of multiresolution data blocks to the final image.

In [29], Zhou et al. performed a study of different image comparison metrics that are categorized into spatial domain, spatial-frequency domain, and perceptually based metrics. They also introduced a comparison metric based on the second-order Fourier decomposition and demonstrated favorable results against other metrics considered. In our work, we use the wavelet transform to partition the data into multiscale and oriented subbands. The study on volume data quality assessment is thus conducted in the spatial-frequency domain rather than the spatial domain.

3 ALGORITHM OVERVIEW

From a mathematical standpoint, we can treat volume data as 3D arrays of intensity values with locally varying statistics that result from different combinations of abrupt features like boundaries and contrasting homogeneous regions. In line with this consideration, we advocate a statistical approach for volume data quality assessment. Given a volumetric data set, a first attempt may lead us to examine its statistics in the original spatial domain.

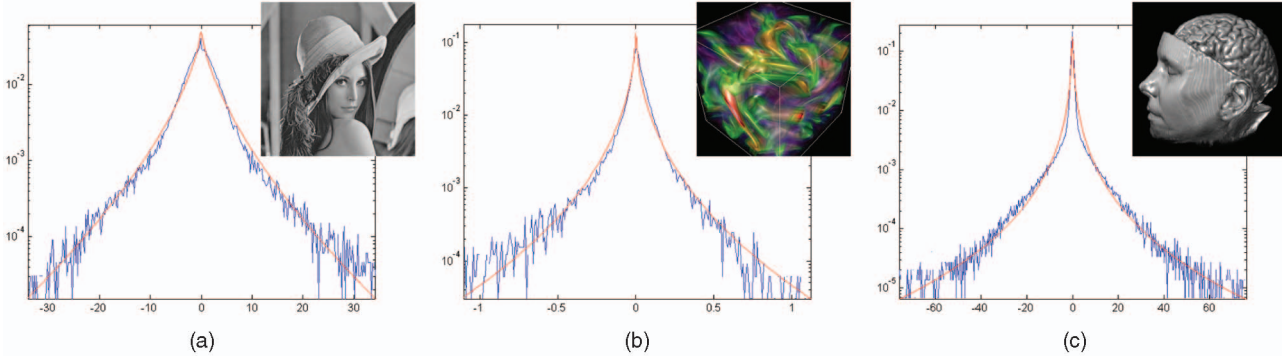


Fig. 2. (a)-(c) Three wavelet subband coefficient histograms (blue curves) fitted with a two-parameter GGD model (red curves) for the Lena image, the vortex data set, and the brain data set, respectively. The estimated parameters (α, β) are $(1.2661, 0.6400)$, $(5.8881e - 003, 0.4181)$, and $(7.1276e - 003, 0.2709)$ for (a), (b), and (c), respectively. The overall fitting is good for all three examples. (a) Lena image, HL_1 subband. (b) Vortex data set, HHL_2 subband. (c) Brain data set, HLL_1 subband.

However, even first-order statistics such as histograms would vary significantly from one portion of data to another and from one data set to another. This defies simple statistical modeling over the entire data set, as well as the subsequent quality assessment.

Instead of spatial domain analysis, we can transform the volume data from the spatial domain to the spatial-frequency domain using the wavelet transform and analyze its frequency statistics. Since frequency is directly related to rate of change, it is intuitive to associate frequencies in the wavelet transform with patterns of intensity variations in the spatial data. Furthermore, the wavelet transform allows us to analyze the frequency statistics at different scales. This will come in handy when we evaluate the quality of reduced or distorted data with different resolutions. Compared with the statistics of data in the spatial domain, the local statistics of different frequency *subbands* are relatively constant and easily modeled. This is realized using *generalized Gaussian density* (GGD) to model the marginal distribution of wavelet coefficients at different subbands and scales (Section 4.1). We also record information about selective wavelet coefficients (Section 4.3) and treat the low-pass filtered subband (Section 4.4) as part of our feature representation.

Our feature thus consists of multiple parts, and each part corresponds to certain essential information in the spatial-frequency domain. Note that this data analysis and feature extraction process can be performed when we have the access to the original data or, ideally, in situ where a simulation is running. Once we extract the feature from the data, we are able to use it for quality assessment without the need to access the original data. Given a reduced or distorted version of data, we compare its feature components with those derived from the original data using predefined distance functions. This gives us an indication of quality loss in relation to the original data. We can also use the feature to perform a cross-comparison of data with different reduction or distortion types to evaluate the uncertainty introduced in different versions of data. Moreover, by forcing some of its statistical properties to match those of the original data, we may repair the reduced or distorted data for possible quality improvement.

4 WAVELET SUBBAND ANALYSIS

In the multiscale wavelet decomposition structure, the low-pass filter subband corresponds to average information that represents large structures or overall context in the volume data. After several iterations of wavelet transforms, the size of the low-pass filter subband is small compared with the size of original data (already less than 0.2 percent for a three-level decomposition). Thus, we can directly treat it as part of the feature. On the other hand, the high-pass filtered subbands correspond to detail information that represents abrupt features or fine characteristics in the data. They spread across all different scales with an aggregate size nearly equal to the size of original data. The key issues are how to extract important feature information from these high-pass filtered subbands and how to compress the feature.

4.1 Wavelet Subband Statistics

Studies on natural image statistics reveal that the histogram of wavelet coefficients exhibits a marginal distribution at a particular high-pass filtered subband. An example of the Lena image and the coefficient histogram of one of its wavelet subbands is shown in Fig. 2a. The y -axis is on a log scale in the histogram. As we can see, the marginal distribution of wavelet coefficients creates a sharp peak at zero and more extensive tails than the Gaussian density. The intuitive explanation of this is that natural images usually have large overall structures consisting of smooth areas interspersed with occasional abrupt transitions, such as edges and contours. The smooth areas lead to near-zero coefficients, and the abrupt transitions give large-magnitude coefficients. In [14], Mallat shows that such a marginal distribution of the coefficients in individual wavelet subbands can be well-fitted with a two-parameter GGD model:

$$p(x) = \frac{\beta}{2\alpha\Gamma(\frac{1}{\beta})} \exp\left(-\left(\frac{|x|}{\alpha}\right)^\beta\right), \quad (1)$$

where Γ is the Gamma function, that is, $\Gamma(z) = \int_0^\infty e^{-t} t^{z-1} dt$, $z > 0$.

In the GGD model, α is the *scale* parameter that describes the standard deviation of the density, and β is the *shape* parameter that is inversely proportional to the decreasing rate of the peak. As an example, the plots of the GGD distribution

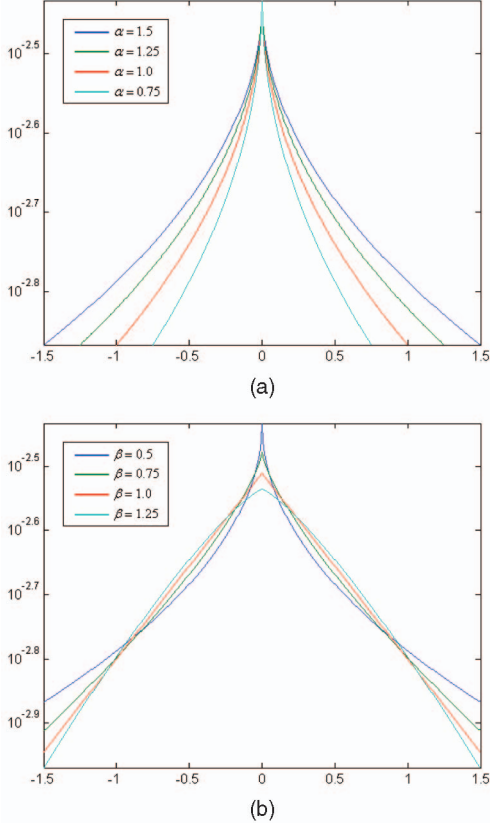


Fig. 3. The GGD distribution and its model parameters (α, β) . (a) α varies when $\beta = 0.5$. (b) β varies when $\alpha = 1.5$. The figure shows the sensitivity of the shape of GGD plots with respect to the model parameters.

under varied (α, β) values are illustrated in Fig. 3. The model parameter (α, β) can be estimated using the *moment matching* method [25] or the *maximum likelihood* rule [16]. Numerical experiments in [25] show that 98 percent of natural images satisfy this property. Even for the remaining two percent, the approximation of the real density by a GGD is still acceptable. Note that the GGD model includes the Gaussian and the Laplacian distributions as special cases with $\beta = 2$ and $\beta = 1$, respectively. The GGD model provides a very efficient way for us to summarize the coefficient histograms of an image, as only two parameters are needed for each subband. This model has been used in previous work for noise reduction [22], image compression [3], texture image retrieval [5], and quality encoding [28]. In this paper, we use the moment matching method, which takes the mean and the variance of wavelet coefficients in a subband to compute its GGD model parameters (α, β) (see the Appendix for the implementation detail). In Fig. 2a, the red curve is the GGD function with parameters estimated using the moment matching method. The result fits the original wavelet coefficient distribution quite well.

We extend this statistical model to 3D volume data since many scientific and medical data share the same intrinsic characteristics as natural images, that is, homogeneous regions mixed with abrupt transitions. Moreover, the rate or proportion of homogeneous regions and abrupt transitions is also similar for image and volume: In 2D, we have an area of homogeneous regions versus the edge length of abrupt

transitions, and in 3D, we have a volume of homogeneous regions versus the surface area of abrupt transitions. Initial experiments on two small data sets give very promising results. Figs. 2b and 2c show one example of wavelet subband coefficient histograms for each data set and their respective well-fitted GGD curves. Thus, with only two GGD parameters, we are able to capture the marginal distribution of wavelet coefficients in a subband that otherwise would require at least hundreds of parameters using a histogram. We shall see in Section 5 that this GGD model works well for larger data sets too. Next, we discuss the distance measure for wavelet subband statistics.

Let $p(x)$ and $q(x)$ denote the probability density functions of the wavelet coefficients in the same subband of the original and distorted data, respectively. Here, we assume the coefficients to be independently and identically distributed. Let $\mathbf{x} = \{x_1, x_2, \dots, x_N\}$ be a set of randomly selected coefficients. The log-likelihoods of $\mathbf{x}$ being drawn from $p(x)$ and $q(x)$ are

$$l(p) = \frac{1}{N} \sum_{i=1}^N \log p(x_i) \quad \text{and} \quad l(q) = \frac{1}{N} \sum_{i=1}^N \log q(x_i), \quad (2)$$

respectively. Based on the law of large numbers, when N is large, the log-likelihoods ratio between $p(x)$ and $q(x)$ asymptotically approaches the *Kullback-Leibler distance* (KLD) (also known as the *relative entropy* of p with respect to q):

$$d(p||q) = \int p(x) \log \frac{p(x)}{q(x)} dx. \quad (3)$$

Although the KLD is not a true metric, that is, $d(p||q) \neq d(q||p)$, it satisfies many important mathematical properties. For example, it is a convex function of p . It is always nonnegative and equals zero only if $p(x) = q(x)$. In this paper, we use the KLD to quantify the difference between wavelet coefficient distributions of the original and distorted data. This quantity is evaluated numerically as follows:

$$d(p||q) = \sum_{i=1}^M P(i) \log \frac{P(i)}{Q(i)}, \quad (4)$$

where $P(i)$ and $Q(i)$ are the normalized heights of the i th histogram bin, and M is the number of bins in the histogram. Note that the coefficient histogram Q is computed directly from the distorted data, whereas the coefficient histogram P is approximated using its GGD parameters (α, β) extracted from the original data.

Finally, the KLD between the distorted and original data over all subbands is defined as

$$D_1 = \log \left(1 + \sum_{i=1}^B d(p^i||q^i) \right), \quad (5)$$

where B is the total number of subbands analyzed, p^i and q^i are the probability density functions of the i th subbands in the original and distorted data, respectively, and $d(p^i||q^i)$ is the estimated KLD between p^i and q^i .

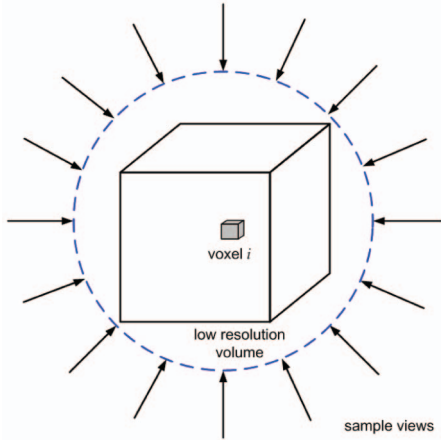


Fig. 4. A voxel's visual importance in the low-resolution volume is the multiplication of its opacity and average visibility. The average visibility is calculated using a list of evenly sampled views along the volume's bounding sphere.

4.2 Voxel Visual Importance

At runtime, a transfer function is applied to the input volume where the scalar data values are mapped to optical quantities such as color and opacity, and the volume is projected into 2D images. To capture the visualization-specific contribution for each voxel, we define a voxel's visual importance ω as follows:

$$\omega(i) = \alpha(i) \cdot \bar{\nu}(i), \quad (6)$$

where $\alpha(i)$ is the opacity of voxel i , $\bar{\nu}(i)$ is its average visibility. As sketched in Fig. 4, given an original large volume data, we use its low-resolution form (for practical and performance concern) to calculate visual importance values of all voxels within the volume. In (6), $\alpha(i)$ and $\bar{\nu}(i)$ account for the emission and attenuation of voxel i , respectively.

To calculate the average visibility, we consider a list of evenly sampled views along the bounding sphere that encloses the volume and take the average of the visibility from those sample views. Given a view along the bounding sphere, the visibility for each voxel in the low-resolution data is acquired in this way: we render the low resolution data by drawing front-to-back view-aligned slices and evaluate the visibility of all the voxels during the slice drawing. The visibility of a voxel is computed as $(1 - \alpha)$ right before the slice containing the voxel is to be drawn, where α is the accumulated opacity at the voxel's screen projection. This process repeats for each sample view. Finally, for each voxel in the volume, we use the average of its visibility from all sample views to calculate its visual importance. Essentially, the visual importance indicates the average contribution of a voxel in association with a given input transfer function. This visualization-specific term is then normalized and incorporated into the following wavelet coefficient selection.

4.3 Wavelet Coefficient Selection

The GGD model captures the marginal distribution of wavelet coefficients at each individual subband. Using the distance defined in (5), we are able to know how close the coefficient distributions of distorted data are in relation to

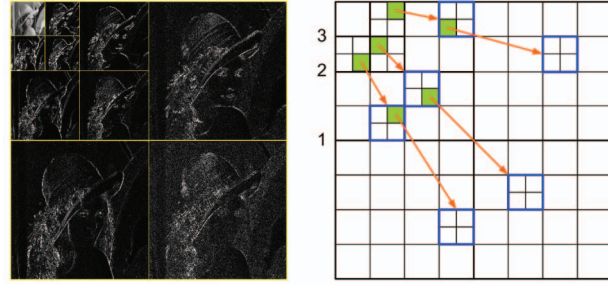


Fig. 5. The parent-child dependencies of wavelet coefficients in different subbands. In this 2D example, a coefficient in a coarse scale has four child coefficients in the next finer scale of similar orientation. The arrow points from the subband of the parents to the subband of the children.

the original data. Nevertheless, the histogram itself does not tell the spatial-frequency positions of wavelet coefficients. This limits our ability to compare the data in finer detail and to possibly repair the distorted data. Therefore, along with the *global* GGD parameters per wavelet subband, we also need to record *local* information about wavelet coefficients for data quality assessment and improvement.

An important observation is that, although the coefficients of wavelet subbands are approximately decorrelated, they are not statistically independent. For example, Fig. 5 shows a three-level decomposition of the Lena image. It can be seen that coefficients of large magnitude (bright pixels) tend to occur at neighboring spatial-frequency locations and also at the same relative spatial locations of subbands at adjacent scales and orientations. Actually, in 2D, a coefficient c in a coarse scale has four child coefficients in the next finer scale. Each of the four child coefficients also has four child coefficients in the next finer scale. Furthermore, if c is insignificant with respect to some threshold ϵ , then it is likely that all of its descendant coefficients are insignificant too. This coefficient dependency has been exploited in several image compression algorithms, such as the embedded zerotree wavelet (EZW) encoding [20] and a following image codec based on set partitioning in hierarchical trees (SPIHT) [19]. These algorithms have also been extended to 3D volumetric image compression in medical application [13]. In this paper, we utilize the coefficient dependency to store selective wavelet coefficients in an efficient manner.

There are two categories of wavelet coefficients that are of importance for the purpose of quality assessment and improvement. One category is the coefficients of large magnitudes that correspond to abrupt features like edges or boundaries. As we can see in Fig. 2, they are along the tails of the marginal coefficient distribution where the perceptually significant coefficients generally reside. The other category is neighboring near-zero coefficients, which correspond to homogeneous regions. They are close to the zero peak of the distribution and are important indications of data regularity. Taking into account the visualization-related factor, we modulate the wavelet coefficients with the voxel visual importance values (Section 4.2) at their nearest spatial-frequency locations. In this case, a wavelet coefficient is large only if it has both large magnitude and high voxel visual importance; a wavelet coefficient is near

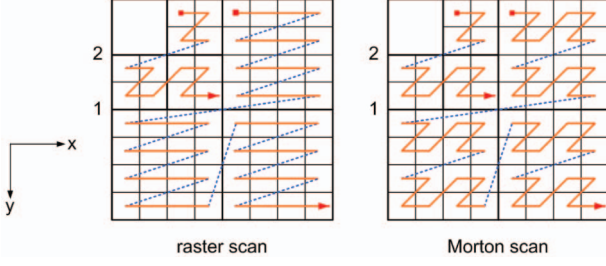


Fig. 6. Scan of wavelet coefficients in the raster order and the Morton order. The blue dashed line segments indicate discontinuities in the scan. Compared with the raster order, the Morton order preserves the spatial-frequency locality better.

zero if it has either near-zero magnitude or near-zero voxel visual importance.

Starting from the coarsest scale, we scan each wavelet subband and encode coefficients of interest. As illustrated with a 2D example in Fig. 6, we follow the *Morton order* (Z-curve order) as opposed to the ordinary raster order to better utilize the spatial-frequency locality. For neighboring near-zero coefficients (at least eight consecutive coefficients in 3D), we run-length encode their positions (that is, the scan orders). For large-magnitude coefficients, we encode their positions and values as well. In general, most scientific and medical data have a low-pass spectrum. When the data are transformed into a multiscale wavelet decomposition structure, the energy in the subbands decreases from a fine scale (high resolution) to a coarse scale (low resolution). Therefore, the wavelet coefficients are, on the average, smaller in the finer scales than in the coarser scales. Accordingly, we vary the thresholds of near-zero coefficients (ε) and large-magnitude coefficients (τ) for different scales. Finally, information of the selective coefficients is further compressed using the open source *zlib* compressor.

We define the distance for the selective wavelet coefficients as follows:

$$D_2 = \log \left(1 + \sum_{i=1}^B \sqrt{\frac{\sum_{j=1}^{L_i} \left(\frac{c_j - c'_j}{\text{max}_i} \right)^2 + \sum_{k=1}^{Z_i} \left(\frac{c'_k}{\text{max}_i} \right)^2}{L_i + Z_i}} \right), \quad (7)$$

where B is the number of subbands over all the scales, L_i and Z_i are the numbers of large-magnitude coefficients selected and near-zero coefficients selected in the i th subband, respectively. c_j and c'_j are the j th large coefficients selected from the original and distorted data, respectively, and max_i is the largest magnitude (modulated by visual importance) of all coefficients at the i th subband. For near-zero coefficients, we assume the original coefficients $c_k = 0$ and only consider coefficients in the distorted data with $|c'_k| > \varepsilon$ for the calculation.

4.4 Low-Pass Filtered Subband

The low-pass filter subband in the multiscale wavelet decomposition structure corresponds to the average information that represents large structures or overall context in the volume data. Compared with the size of original data, the size of this subband is usually small after several iterations of wavelet decomposition. Therefore, we directly incorporate it as part of the feature. Let b_i and b_j be the low-pass filter

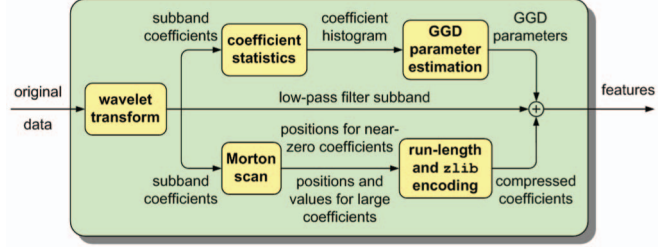


Fig. 7. Our feature representation of the original data in the wavelet domain.

subbands of the original and distorted data, respectively. The similarity between b_i and b_j is defined as

$$S = \frac{\sigma_{ij}}{\sigma_i \sigma_j} \cdot \frac{2\mu_i \mu_j}{\mu_i^2 + \mu_j^2} \cdot \frac{2\sigma_i \sigma_j}{\sigma_i^2 + \sigma_j^2} = \frac{4\sigma_{ij} \mu_i \mu_j}{(\sigma_i^2 + \sigma_j^2)(\mu_i^2 + \mu_j^2)}, \quad (8)$$

where σ_{ij} is the covariance between b_i and b_j , μ_i and μ_j are the mean values of b_i and b_j , respectively, and σ_i and σ_j are their standard deviations. Equation (8) consists of three parts, namely, *loss of correlation*, *luminance distortion*, and *contrast distortion*. Collectively, these three parts capture the structure distortion of the low-pass filtered subband in the distorted data. This similarity measure comes from the image quality assessment literature [27] and has been shown to be consistent with the luminance masking and contrast masking features in the HVS, respectively. The dynamic range of S is $[-1, 1]$. The best value of 1 is achieved when $b_i = b_j$. Hence, we define the distance between b_i and b_j as follows:

$$D_3 = \sqrt{1.0 - (S + 1.0)/2.0}. \quad (9)$$

4.5 Summary

In summary, as shown in Fig. 7, our feature representation of the original data in the wavelet domain includes three parts: the GGD model parameters from wavelet subband statistics, selective wavelet coefficients, and the low-pass filtered subband. Given a reduced or distorted version of data, we analyze the quality loss by calculating its distances to the original data for each of the feature components [(5), (7), and (9)]. Each partial distance indicates some quality degradation with reference to the original data, and the summation of all these partial distances gives the overall degradation. Thus, an overall distance could be computed heuristically as the weighted sum of the three individual distances:

$$D = k_1 D_1 + k_2 D_2 + k_3 D_3, \quad (10)$$

where $k_i > 0$, $i = 1, 2$, and 3 . Note that there is no need for normalizing this overall distance. For the purpose of data quality improvement, it is advantageous to keep each distance separate (or further at the subband level) so that we know which parts cause significant quality degradation. We can then repair accordingly using the feature extracted from the original data.

5 RESULTS

We experimented with our algorithm on six floating-point data sets, as listed in Table 1. Among the six data sets, three of them are from scientific simulation, and the

TABLE 1
The Six Floating-Point Data Sets and Their Feature Sizes

data set	dimension	data size	DL	LP dimension	LP size	HP size	feature size	ratio
turbulent vortex flow	128^3	8.0MB	3	16^3	16.0KB	42.2KB	58.2KB	0.710%
solar plume velocity magnitude	$512^2 \times 2048$	2.0GB	5	$16^2 \times 64$	64.0KB	376.0KB	440.0KB	0.021%
supernova angular momentum	864^3	2.40GB	5	27^3	76.9KB	852.6KB	929.5KB	0.037%
UNC brain	$128^2 \times 72$	4.5MB	3	$16^2 \times 9$	9.0KB	53.7KB	62.7KB	1.361%
head aneurysm	512^3	512MB	4	32^3	128.0KB	198.0KB	326.0KB	0.062%
visible woman	$512^2 \times 1728$	1.69GB	5	$16^2 \times 54$	54.0KB	753.4KB	807.4KB	0.046%
DL = decomposition levels								
LP dimension (size) = low-pass filtered subband's dimension (size); HP size = all high-pass filtered subbands' feature size								

remaining three are from medical application. These six data sets vary greatly in size and exhibit quite different characteristics. For multiscale wavelet decomposition, we specifically restricted our attention to the Daubechies family of orthogonal wavelets, as the evaluation of all possible wavelet transforms is out of the scope of our experiments. The decision for levels of wavelet decomposition is based on the size of input data, as well as the trade-off between the size of feature and the robustness of GGD model parameters.

In our test, the threshold for near-zero wavelet coefficients ε_i at the i th subband was chosen as $\max_i / (2^{L+3})$, where $\max_i$ is the largest magnitude (modulated by visual importance) of all coefficients at the i th subband, and L is the total number of decomposition levels we have. The threshold for large-magnitude wavelet coefficients τ_i at the i th subband was chosen as $\max_i / (2^{s+2})$, where s is the scale in which the i th subband locates. We varied τ_i according to the scale because the wavelet coefficients in a subband became more important as the scale increases. In Table 1, the size of low-pass filtered subband is the uncompressed size of LLL_n , after n levels of decomposition. The feature size of all high-pass filtered subbands includes their respective GGD parameters and selective wavelet coefficients in the compressed form.

From the last column in Table 1, we can see that for all six data sets, the size of feature is small compared with the original data. Note that the size of feature does not increase in proportion to the size of original data. This is mainly due to the increase of decomposition levels for larger data sets, as we can afford to have more levels of wavelet decomposition while still keeping the GGD parameters robust. Our experiment confirms that the GGD model generally performs well when the size of the input data becomes larger. For instance, Fig. 8 shows one of the wavelet subband coefficient histograms for the visible woman and the solar plume data sets, and their respective GGD curves. On the other hand, the feature size is also data and transfer function dependent. For example, the ratio for the brain data set is 1.361 percent, which is relatively high compared with the vortex data set having the same number of decomposition levels. Thus, it follows that we record a higher percentage of high-pass filtered subband feature information for the brain data set.

Next, we report results of volume data quality assessment and improvement using the extracted feature. To compare the quality of rendered images, we used a GPU raycaster for volume rendering. All tests were performed on a 2.33-GHz Intel Xeon processor with 4-Gbytes main memory and an Nvidia GeForce 7900 GTX graphics card with 512-Mbytes video memory.

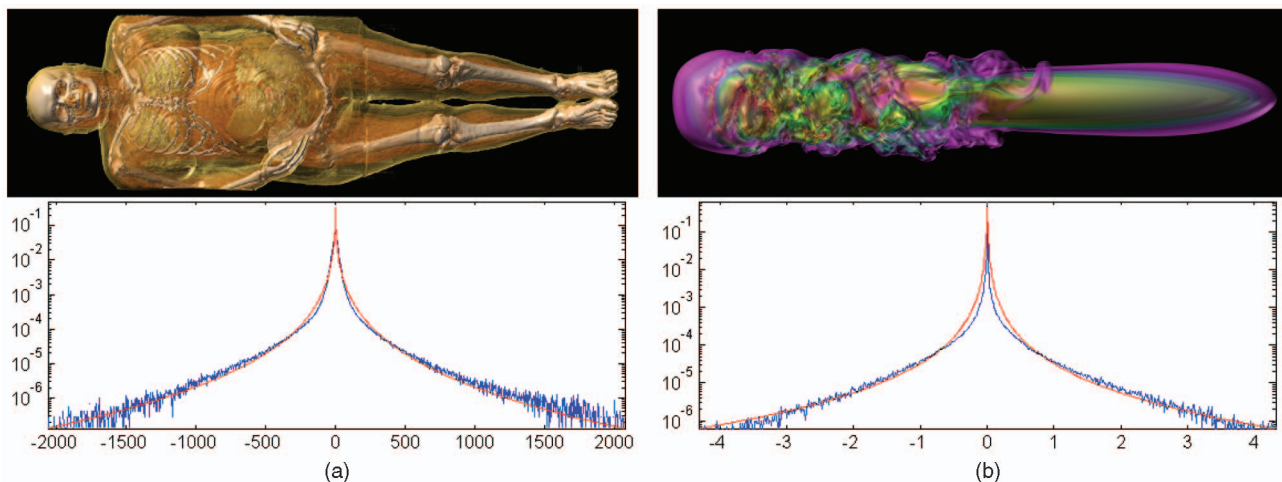


Fig. 8. (a) and (b) Two wavelet subband coefficient histograms (blue curves) fitted with the GGD model (red curves) for the visible woman and the solar plume data sets, respectively. The estimated parameters (α, β) are $(1.6786e - 002, 0.2425)$ and $(1.4922e - 009, 0.1405)$, respectively. In general, the fitting works well for both data sets. (a) Visible woman data set, HLH₂ subband. (b) Solar plume data set, LLH₂ subband.

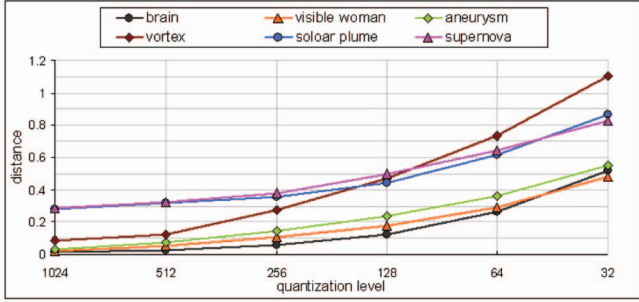


Fig. 9. Quality assessment on six test data sets with six different quantization levels. The data quality gets increasingly worse as the number of quantization levels decreases.

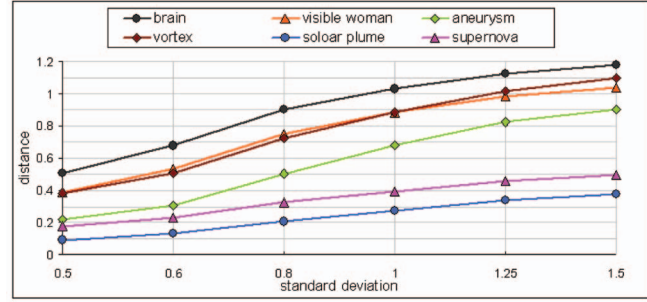


Fig. 10. Quality assessment on six test data sets with 5^3 Gaussian smooth filters of six different standard deviations. The data quality gets worse as the standard deviation increases.

5.1 Quality Assessment

First of all, we experimented with our quality measure on different data sets and observed how the quality of data changes for the same reduction or distortion type. We used the three smaller data sets (brain, vortex, and aneurysm) of their original resolutions and the other three data sets (visible woman, solar plume, and supernova) of their second highest resolutions in the test. To calculate the overall distance, we used (10) with $(k_1, k_2, k_3) = (0.1, 1.0, 1.0)$. Please note that in this paper, we assume that the original data set has full quality. Thus, any changes made to the data would involve possible quality loss, even though the desire is to enhance the data from a certain perspective.

Quantization is a commonly used approach for data reduction. Our first example studies the quality loss under the uniform quantization scheme. Fig. 9 shows the quality assessment result of all six data sets with six different quantization levels. A larger distance indicates a greater degree of quality degradation. More specifically, Table 2 lists all partial and overall distances for the aneurysm data set. Although different data sets have different responses of quality loss due to quantization, the overall trend is fairly obvious: the data quality gets increasingly worse as the number of quantization levels decreases.

Our second example studies the quality loss under the Gaussian smooth filtering. Fig. 10 shows how the data quality changes with six different Gaussian smooth filters for all six data sets. We applied a discrete Gaussian kernel of size 5^3 with different standard deviations. A larger standard deviation indicates a greater degree of smoothing since neighboring voxels carry more weight. Table 3 lists all partial and overall distances for the solar plume data set. It is clear that the data quality gets worse

as the standard deviation increases. Unlike quantization, however, the rate of quality loss decreases gradually in the sequence.

Besides quality assessment of data with the same type of reduction or distortion, the feature also avails us to perform cross-type data quality comparison. For example, Fig. 11 gives quality assessment results on the solar plume and the visible woman data sets under four different distortion types: mean shift (of the data range over 256), voxel misplacement (with two slices of voxels misplaced), averaging filter (using a kernel of size 3^3), and salt-and-pepper noise (with an equal probability of $1/1024$ for the bipolar impulse). Table 4 lists all partial and overall distances, MSE, and PSNR for these four distortion types. For both data sets, we can see that the mean shift introduces the minimum quality loss here, followed by the voxel misplacement. The salt-and-pepper noise incurs the most quality degradation. This result is consistent with perceived quality in rendered images. However, the MSE and the PSNR incorrectly recognize the mean shift as having a larger distortion than the voxel misplacement for both data sets. Note that they also give the opposite results on the averaging filter for the two data sets. This is due to the reason that the MSE and the PSNR metrics are only voxel-based and do not consider the overall structure distortion of the data.

5.2 Quality Improvement

Since the feature captures essential information from the original data, it can be utilized to improve the quality of distorted or corrupted data. In this paper, we do not show examples where the feature is used to construct a higher resolution data from a low-resolution data, as it is the most

TABLE 2

Partial and Overall Distances for the Aneurysm Data Set with Six Different Quantization Levels

level	D_1	D_2	D_3	D
1024	0.1961	0.0071	0.0038	0.0305
512	0.5227	0.0144	0.0077	0.0744
256	0.9933	0.0277	0.0153	0.1423
128	1.5111	0.0551	0.0303	0.2365
64	1.9518	0.1065	0.0594	0.3611
32	2.2407	0.2089	0.1150	0.5480

TABLE 3

Partial and Overall Distances for the Solar Plume Data Set with Six Gaussian Smooth Filters of Different Standard Deviations

σ	D_1	D_2	D_3	D
0.5	0.2687	0.0656	0.0004	0.0929
0.6	0.3793	0.0921	0.0006	0.1306
0.8	0.6325	0.1422	0.0010	0.2065
1.0	0.9225	0.1822	0.0015	0.2760
1.25	1.2039	0.2165	0.0020	0.3389
1.5	1.3548	0.2377	0.0023	0.3755

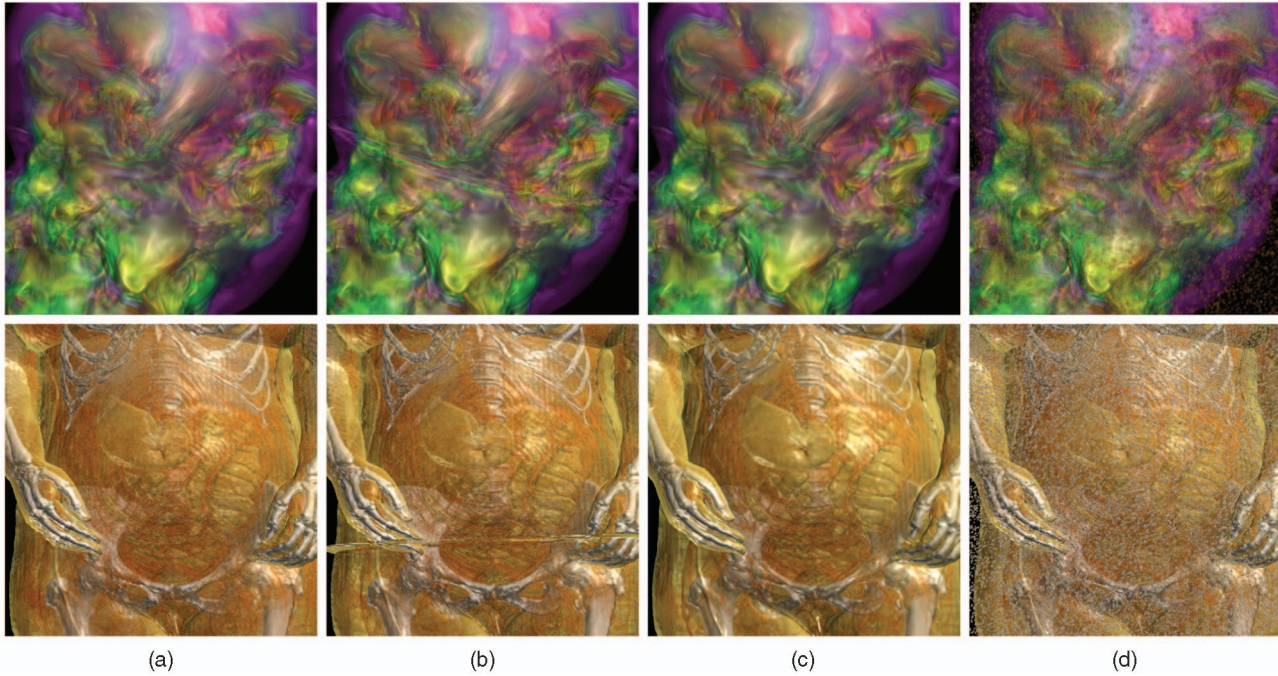


Fig. 11. Cross-type quality assessment on low-resolution solar plume ($256^2 \times 1024$) and visible woman ($256^2 \times 864$) data sets. The data quality degrades as the overall distance (listed in Table 4) increases from (a) to (d). The rendered images are cropped for a closer comparison. (a) Mean shift. (b) Voxel misplacement. (c) Averaging filter. (d) Salt-and-pepper noise.

TABLE 4
Partial and Overall Distances, MSE, and PSNR for the Solar Plume and the Visible Woman Data Sets with Four Different Distortion Types

data set	type	D_1	D_2	D_3	D	rank	MSE	PSNR	rank
solar plume	mean shift	$2.0556e-4$	$4.5864e-7$	$5.8988e-2$	0.0590	4	$5.5252e-2$	48.1648	2
	misplacement	$1.3126e-1$	$9.4561e-2$	$5.5746e-3$	0.1133	3	$2.4239e-2$	51.7431	3
	averaging	$5.2393e-1$	$1.2551e-1$	$8.1216e-4$	0.1787	2	$5.2299e-3$	58.4033	4
	noise	$3.1198e+0$	$1.8137e+0$	$2.8900e-2$	2.1546	1	$6.7305e+0$	27.3078	1
visible woman	mean shift	$8.8428e-5$	$6.9770e-7$	$2.3914e-2$	0.0239	4	$1.8691e+3$	48.1648	3
	misplacement	$1.5366e-2$	$1.1612e-1$	$4.6497e-3$	0.1223	3	$1.5397e+3$	49.0066	4
	averaging	$1.6449e+0$	$5.4139e-1$	$1.4596e-3$	0.7073	2	$1.9289e+5$	28.0279	1
	noise	$1.7343e+0$	$7.8530e-1$	$9.7468e-3$	0.9685	1	$1.2152e+4$	40.0347	2

common way of using the wavelet transform and compression.

Our first example deals with missing data. Fig. 12a shows the rendering of a low-resolution supernova data set with $1/8$ (that is, an octant) of data missing. The missing of data could result from incomplete data transmission, or even a bug in the data reduction source code. Recall that we keep each partial distance separate (and actually at the subband level). This helps us identify which parts introduce the dramatic change (in this case, the subband at the same orientation as the missing portion) and then repair accordingly using the feature information.

The repairing scheme works as follows: First, a multi-scale wavelet decomposition structure is built from the corrupted data, where the size of low-pass filtered subband at the coarsest scale equals the size of low-pass filtered subband recorded in the feature (Section 4.4). Note that for the repairing purpose, we keep the low-pass filtered subbands in all scales. Then, we improve the high-pass filtered subbands across all scales by replacing the wavelet coefficients with their corresponding coefficients in the

feature; that is, those large-magnitude and near-zero wavelet coefficients selected (Section 4.3). Next, starting from the coarsest scale (the lowest resolution), we reconstruct the low-pass filtered subband at the next finer scale using the low-pass filtered subband recorded in the feature and improved high-pass filtered subbands. The reconstructed low-pass filtered subband is used to correct the missing part in the same-scale low-pass filtered subband decomposed from the corrupted data. The corrected low-pass filtered subband is then used to reconstruct the next finer scale in an iterative manner. In this way, we are able to automatically repair the missing portion in the corrupted data scale by scale. Finally, an optional median filter is applied to the corrected portion of data at the finest scale in order to suppress potential noise and produce a better match with the original GGD model parameters. Fig. 12b shows the result after this automatic repairing process. It is clear that the data quality improves as the overall distance decreases.

We can also apply a similar repairing process for noise reduction. For example, Fig. 13a shows the rendering of the

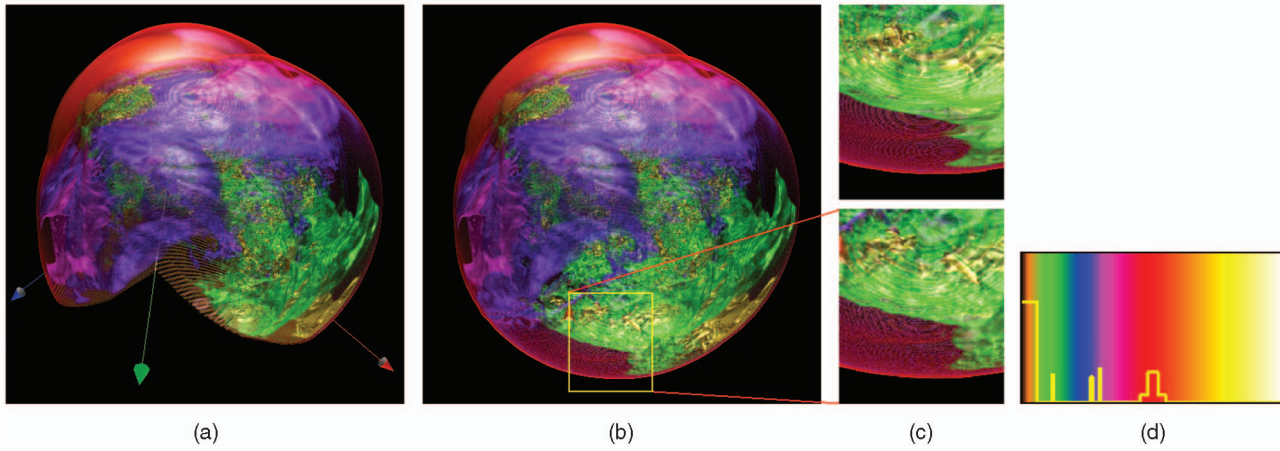


Fig. 12. Quality improvement on a low-resolution supernova data set (432^3) with 1/8 of data missing, as shown in (a). (b) The result after an automatic repairing process using the feature information. In (c), a portion of (b) is zoomed in for comparison with the reference image displayed on the top. (a) Before, $D = 1.3346$. (b) After, $D = 0.5536$. (c) Comparison. (d) Transfer function.

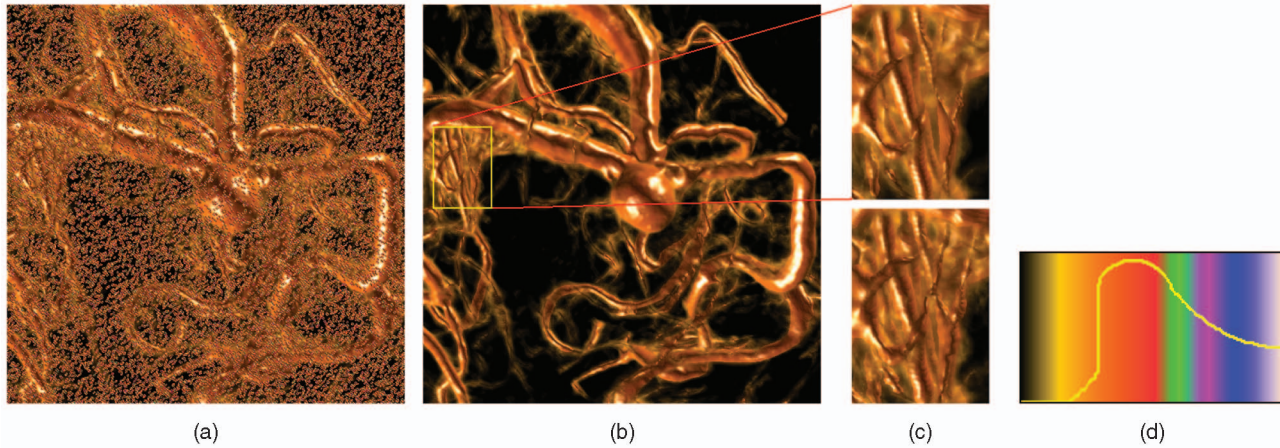


Fig. 13. Quality improvement on the aneurysm data set (512^3) distorted by random noise, as shown in (a). (b) The result after an automatic repairing process using the feature information. In (c), a portion of (b) is zoomed in for comparison with the reference image displayed on the bottom. (a) Before, $D = 2.3094$. (b) After, $D = 0.7188$. (c) Comparison. (d) Transfer function.

aneurysm data set distorted by random noise. This kind of distortion can be detected through the observation of a sequence of sudden spikes appearing in the wavelet coefficient subband histograms. The denoising process also follows a coarse-to-fine manner as usual, but there is no need to keep the low-pass filtered subband at every scale in the wavelet decomposition structure. Another difference is that for each high-pass filtered subband i , we first set large-magnitude wavelet coefficients to zero (if they are larger than the threshold τ_i) before improving them with their corresponding coefficients in the feature. Fig. 13b shows the result after this repairing process. As we can see, the noise is eliminated, whereas the fine structure of the blood vessels is preserved.

6 DISCUSSION

6.1 Choice of Wavelets

To extract essential information from the original data, we decomposed the data into multiple scales using wavelet basis functions localized in spatial position, orientation, and spatial frequency. We used the Daubechies family of orthogonal wavelets in our experiment because they

provide a good trade-off between performance and complexity [5], [6]. Moreover, we found that the choice for the number of scaling and wavelet function coefficients has a little effect on assessment accuracy. Therefore, we specifically used the Daubechies D4 transform for efficiency. Other separable wavelets (such as the Gabor wavelets) or redundant transforms (such as the steerable pyramid transform) could also be used in our algorithm. For example, the steerable pyramid transform decomposes the data into several spatial-frequency bands and further divides each frequency band into a set of orientation bands. It can thus help to minimize the amount of aliasing within each subband. However, they are more expensive to compute and require more storage space.

6.2 Timing Performance

The timing of wavelet analysis on the original data includes the time for multiscale wavelet decomposition, GGD parameters estimation, and subband wavelet coefficients selection. This one-time preprocess may take anywhere from seconds to a total of several minutes on a single PC, depending on the size of input data. For data sets that could not be loaded into memory simultaneously, we employed a blockwise

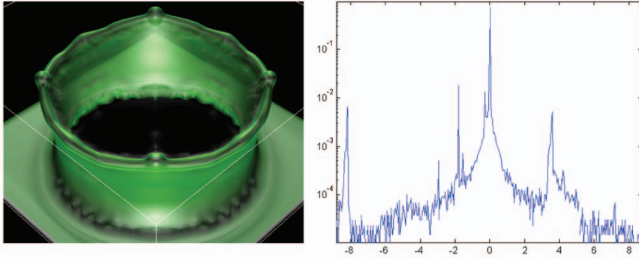


Fig. 14. The “milk crown” physical simulation data set ($512 \times 256 \times 512$) and its LHL_3 subband coefficient histogram, which does not exhibit the marginal distribution.

wavelet transform process and handled boundaries of neighboring blocks to guarantee seamless results. The timing of quality assessment on different versions of data includes the time for wavelet decomposition and distance calculation. At runtime, it usually takes less than one minute on a single PC to evaluate data with the size up to 512 Mbytes. For larger gigabytes data, the time to perform wavelet transforms becomes dominant in the quality assessment process. In the worst case, the assessment time would be similar to the preprocess time if the data we evaluate has the same size as the original data.

6.3 GGD Model

We note that as an approximation, the GGD model introduces a prediction error at each wavelet subband with respect to the corresponding wavelet coefficient distribution. For example, the fit near the center of the histogram in Fig. 8b is not good. This error can be calculated as the KLD between the model histogram and the histogram of wavelet subband coefficients from the original data. Let $d(p_m^i \| p^i)$ denote the prediction error at the i th subband. Accordingly, we use $d(p^i \| q^i) = |d(p_m^i \| q^i) - d(p_m^i \| p^i)|$ to calculate the overall KLD (5). That is, we actually subtract the prediction error from the KLD between the model histogram and the histogram of wavelet subband coefficients from the reduced or distorted data (denoted as $d(p_m^i \| q^i)$) and use the absolute value in the calculation.

On the other hand, our experiment shows that the GGD model generally works well on scientific and medical data sets with different sizes. However, there are cases where this model fails to give good results. Such an example is shown in Fig. 14. For these failed cases, we can store the actual wavelet subband coefficient histograms (per scale) at the expense of increasing the storage or fit each coefficient histogram with splines to smooth out the irregularities. Although there is a need of further research on why these cases fail, the apparent reason is that those data sets do not fall into the category of natural statistics. For this same reason, we cannot partition the original data into blocks in an octree fashion and analyze the individual blocks using the GGD model (each block does not necessarily exhibit the marginal coefficient distribution even though the whole data set does). Therefore, our solution is a multiscale, not a true multiresolution approach.

6.4 Transfer Function

In this work, different versions of a data set were rendered using the same transfer function for the purpose of

subjective data quality comparison. Since our focus was on data quality assessment, we chose to fix rendering parameters so that the possible difference or uncertainty introduced by the visualization process could be minimized. Our current algorithm explicitly takes the input transfer function into consideration by modulating wavelet coefficients with voxel visual importance values at their nearest spatial-frequency positions. The voxel visual importance values were precomputed offline with a given transfer function. If the transfer function changes at runtime, the calculation can be performed online (in this case, we need to keep the low-resolution data).

Our solution is a coarse approximation of voxel contribution to the visualization. The accuracy of voxel visual importance values depends on the resolution of data used and the number of sample views taken for the average visibility calculation. There is a trade-off between the update speed and the accuracy of visual importance values. In practice, we can update the visual importance values within seconds for a low-resolution volume of size around 64^3 with 16 sample views. In our solution, voxel visual importance values are only used to modulate and select wavelet coefficients (Section 4.3). An improvement of our implementation is to store selective wavelet coefficients offline by only considering their magnitudes. At runtime, when the transfer function changes, the visual importance values are calculated and used to further pick visually important coefficients from stored coefficients.

Besides our current algorithm, another way to possibly improve wavelet subband analysis is to apply the idea presented in [1] that classifies the voxels into core, gradient, and unimportant voxels and assigns weight functions for wavelet coefficients accordingly. Alternatively, the users can also provide their own voxel visual importance volume, derived from volume classification or segmentation, for example, to modulate wavelet coefficients. Nevertheless, we understand that this voxel-based approach is not an optimal solution for large data analysis in terms of both efficiency and effectiveness. A better solution could be using some shape functions to approximate the volume data and capture the visual importance aspect. Another direction is to perform a more rigorous study on data quality comparison in association with direct volume rendering algorithm specifications [10].

7 CONCLUSION AND FUTURE WORK

We introduce a reduced-reference approach to volume data quality assessment. A multiscale wavelet representation is first built from the original data, which is well-suited for the subsequent statistical modeling and feature extraction. As shown in Section 5, we extract minimum feature information in the wavelet domain. Using the feature and predefined distance measures, we are able to identify and quantify the quality loss in the reduced or distorted version of data. Quality improvement on distorted or corrupted data is achieved by forcing some of their feature components to match those from the original data. Finally, our approach can be treated as a new way to evaluate the uncertainty introduced by reduced or distorted data.

Our algorithm is flexible with data sets of different sizes, ranging from megabytes to gigabytes in the experiment. We believe that the general approach presented in this paper can be applied to quality assessment and improvement on larger scale data. As we move into the era of petascale computing, our work can help scientists perform in situ processing so that low-resolution data together with a set of features are saved to disk, which greatly reduces storage requirement and facilitates subsequent data analysis, quality assessment, and visualization.

Our current scheme is based on the GGD model, which generally works well on data sets that exhibit natural statistics. We will investigate where and how well the GGD model works for different volume data. Furthermore, we can improve this model by augmenting it with a set of hidden random variables that govern the GGD parameters [17]. Such hidden Markov models may encompass a wider variety of data sets and yield better quality assessment results. On the other hand, we need to conduct a user study to suggest that the visual quality perceived by the users conforms to the quality assessment results obtained from our algorithm. In the future, we also would like to extend this reduced-reference approach to quality assessment of time-varying multivariate data.

APPENDIX

THE CALCULATION OF GGD PARAMETERS

The key MATLAB functions for calculating the GGD parameters (α, β) and for returning the GGD function values are provided as follows:

```
function f = fbeta (x)
% FBETA: an auxiliary function that computes
% beta

f = exp (2 * gammaln (2 ./ x) - gammaln (3 ./ x)
        - gammaln (1 ./ x));
% GAMMALN: logarithm of Gamma function

function [alpha,beta,K] = sbpdf(mean,variance)
% SBPDF: estimate generalized Gaussian
% probability density function of a
% wavelet subband using the moment
% matching method

F = sprintf('fbeta (x) - %g',mean^2 / variance);

try
    beta = fzero(F,[0.01,5]);
    % FZERO: find zero of a function of one
    % variable
catch
    warning('(mean^2 / variance) is out of the range');
    if (mean^2 / variance) > fbeta (5)
        beta = 5;
    else
        beta = 0.01;
    end
end
```

```
alpha=mean * exp(gammaln (1/beta)-gammaln(2/beta));

if (nargout > 2)
    K = beta / (2 * alpha * gamma(1/beta));
    % GAMMA: Gamma function
end

function y = ggpdf (x,alpha,beta,K)
% GGPDF: return generalized Gaussian
% probability density function with
% parameters alpha and beta at the
% value in x

if (alpha <= 0|beta <= 0)
    tmp = NaN;
    y = tmp (ones(size(x)));
    % ONES: create an array of all ones
else
    y = K * exp (-abs(x).^beta ./ (alpha^beta));
    y = y ./ sum (y);
end
```

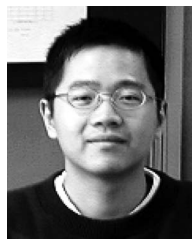
ACKNOWLEDGMENTS

This research was supported in part by the US National Science Foundation through Grants CCF-0634913, CNS-0551727, OCI-0325934, and CCF-9983641 and the US Department of Energy through the SciDAC program with Agreement DE-FC02-06ER25777, DOE-FC02-01ER41202, and DOE-FG02-05ER54817. The aneurysm data set was provided by Michael Meißner and made available by Dirk Bartz at <http://www.volvis.org/>. The visible woman data set was courtesy of the US National Library of Medicine. The solar plume data set was provided by John Clyne (NCAR). The supernova data set was provided by John M. Blondin (NCSU) and Anthony Mezzacappa (ORNL). The milk crown data set was courtesy of Kenji Ono (RIKEN, Japan). Chaoli Wang would like to thank Hongfeng Yu and Shinho Kang for their helps in GPU and MATLAB programming, respectively. Finally, the authors would like to thank the reviewers for their helpful comments.

REFERENCES

- [1] C.L. Bajaj, S. Park, and I. Ihm, "Visualization-Specific Compression of Large Volume Data," *Proc. Pacific Graphics (PG '01)*, pp. 212-222, 2001.
- [2] A.C. Bovik, *Handbook of Image and Video Processing*, second ed. Academic Press, 2005.
- [3] R.W. Buccigrossi and E.P. Simoncelli, "Image Compression via Joint Statistical Characterization in the Wavelet Domain," *IEEE Trans. Image Processing*, vol. 8, no. 12, pp. 1688-1701, 1999.
- [4] J.G. Daugman, "Two-Dimensional Spectral Analysis of Cortical Receptive Field Profiles," *Vision Research*, vol. 20, pp. 847-856, 1980.
- [5] M.N. Do and M. Vetterli, "Wavelet-Based Texture Retrieval Using Generalized Gaussian Density and Kullback-Leibler Distance," *IEEE Trans. Image Processing*, vol. 11, no. 2, pp. 146-158, 2002.
- [6] A. Gaddipati, R. Machiraju, and R. Yagel, "Steering Image Generation with Wavelet Based Perceptual Metric," *Computer Graphics Forum*, vol. 16, no. 3, pp. 241-251, 1997.
- [7] R.C. Gonzalez and R.E. Woods, *Digital Image Processing*, second ed. Prentice Hall, 2002.

- [8] S. Guthe, M. Wand, J. Gonser, and W. Straßer, "Interactive Rendering of Large Volume Data Sets," *Proc. IEEE Visualization Conf. (VIS '02)*, pp. 53-60, 2002.
- [9] C.E. Jacobs, A. Finkelstein, and D.H. Salesin, "Fast Multiresolution Image Querying," *Proc. ACM SIGGRAPH '95*, pp. 277-286, 1995.
- [10] K. Kim, C.M. Wittenbrink, and A. Pang, "Extended Specifications and Test Data Sets for Data Level Comparisons of Direct Volume Rendering Algorithms," *IEEE Trans. Visualization and Computer Graphics*, vol. 7, no. 4, pp. 299-317, Oct.-Dec. 2001.
- [11] T.-Y. Kim and Y.G. Shin, "An Efficient Wavelet-Based Compression Method for Volume Rendering," *Proc. Pacific Graphics (PG'99)*, pp. 147-157, 1999.
- [12] L. Linsen, V. Pascucci, M.A. Duchaineau, B. Hamann, and K.I. Joy, "Hierarchical Representation of Time-Varying Volume Data with $\sqrt{2}$ Subdivision and Quadrilinear B-Spline Wavelets," *Proc. Pacific Graphics (PG '02)*, pp. 346-355, 2002.
- [13] J. Luo and C.W. Chen, "Coherently Three-Dimensional Wavelet-Based Approach to Volumetric Image Compression," *J. Electronic Imaging*, vol. 7, no. 3, pp. 474-485, 1998.
- [14] S. Mallat, "A Theory for Multiresolution Signal Decomposition," *IEEE Trans. Pattern Analysis and Machine Intelligence*, vol. 11, no. 7, pp. 674-693, July 1989.
- [15] S. Muraki, "Volume Data and Wavelet Transforms," *IEEE Computer Graphics and Applications*, vol. 13, no. 4, pp. 50-56, July/Aug. 1993.
- [16] H.V. Poor, *An Introduction to Signal Estimation and Detection*, second ed. Springer, 1994.
- [17] J. Portilla, V. Strela, M.J. Wainwright, and E.P. Simoncelli, "Image Denoising Using Scale Mixtures of Gaussians in the Wavelet Domain," *IEEE Trans. Image Processing*, vol. 12, no. 11, pp. 1338-1351, 2003.
- [18] N. Sahasrabudhe, J.E. West, R. Machiraju, and M. Janus, "Structured Spatial Domain Image and Data Comparison Metrics," *Proc. IEEE Visualization Conf. (VIS '99)*, pp. 97-104, 1999.
- [19] A. Said and W.A. Pearlman, "A New, Fast, and Efficient Image Codec Based on Set Partitioning in Hierarchical Trees," *IEEE Trans. Circuits and Systems for Video Technology*, vol. 6, no. 3, pp. 243-250, 1996.
- [20] J.M. Shapiro, "Embedded Image Coding Using Zerotrees of Wavelet Coefficients," *IEEE Trans. Signal Processing*, vol. 41, no. 12, pp. 3445-3462, 1993.
- [21] R. Shapley and P. Lennie, "Spatial Frequency Analysis in the Visual System," *Ann. Rev. of Neuroscience*, vol. 8, pp. 547-583, 1985.
- [22] E.P. Simoncelli and E.H. Adelson, "Noise Removal via Bayesian Wavelet Coring," *Proc. Int'l Conf. Image Processing (ICIP '96)*, pp. 16-19, 1996.
- [23] B.-S. Sohn, C.L. Bajaj, and V. Siddavanahalli, "Feature Based Volumetric Video Compression for Interactive Playback," *Proc. IEEE Symp. Volume Visualization*, pp. 89-96, 2002.
- [24] E.J. Stollnitz and D.H. Salesin, *Wavelets for Computer Graphics: Theory and Applications*. Morgan Kaufmann, 1996.
- [25] G. Van de Wouwer, P. Scheunders, and D. Van Dyck, "Statistical Texture Characterization from Discrete Wavelet Representations," *IEEE Trans. Image Processing*, vol. 8, no. 4, pp. 592-598, 1999.
- [26] C. Wang, A. Garcia, and H.-W. Shen, "Interactive Level-of-Detail Selection Using Image-Based Quality Metric for Large Volume Visualization," *IEEE Trans. Visualization and Computer Graphics*, vol. 13, no. 1, pp. 122-134, Jan./Feb. 2007.
- [27] Z. Wang, A.C. Bovik, H.R. Sheikh, and E.P. Simoncelli, "Image Quality Assessment: From Error Visibility to Structural Similarity," *IEEE Trans. Image Processing*, vol. 13, no. 4, pp. 600-612, 2004.
- [28] Z. Wang, G. Wu, H.R. Sheikh, E.-H. Yang, and A.C. Bovik, "Quality-Aware Images," *IEEE Trans. Image Processing*, vol. 15, no. 6, pp. 1680-1689, 2006.
- [29] H. Zhou, M. Chen, and M.F. Webster, "Comparative Evaluation of Visualization and Experimental Results Using Image Comparison Metrics," *Proc. IEEE Visualization Conf. (VIS '02)*, pp. 315-322, 2002.



He is a member of the IEEE.



He is a senior member of the IEEE.

► For more information on this or any other computing topic, please visit our Digital Library at www.computer.org/publications/dlib.

Chaoli Wang received the BE and ME degrees in computer science from Fuzhou University, China, in 1998 and 2001, respectively, and the PhD degree in computer and information science from Ohio State University in 2006. Currently, he is a postdoctoral researcher in the Department of Computer Science, University of California, Davis. His research interests are computer graphics and visualization, with a focus on large-scale data analysis and visualization. He is a member of the IEEE.

Kwan-Liu Ma received the PhD degree in computer science from the University of Utah in 1993. He is a professor of computer science at the University of California, Davis, and he directs the DOE SciDAC Institute for Ultrascale Visualization. His research interests include visualization, high-performance computing, and user interface design. He received the US National Science Foundation (NSF) PECASE Award in 2000, Schlumberger Foundation Technical Award in 2001, and the UC Davis College of Engineering's Outstanding Mid-Career Research Faculty Award in 2007. He is the paper chair for the IEEE Visualization 2008 Conference, IEEE Pacific Visualization 2008 Symposium, and Eurographics Parallel Graphics and Visualization 2008 Symposium. He also serves on the editorial boards of the *IEEE Computer Graphics and Applications* and the *IEEE Transactions on Visualization and Graphics*. He is a senior member of the IEEE.