

Chapter 4

Shape Quality Assessment for 3D Watermarking Attacks

4.1 Introduction

The precedent Chapter has presented the varied factors that influence the perception of three-dimensional shapes rendered on 2D screens. In this Chapter, we propose an image-based metric to assess the perceived distortion of 3D models in the context of watermarking attacks. Watermarking is the topic of the next Part of the thesis and, in a few words, relates to hiding data on 3D models in a robust, imperceptible and secure way. Watermarking attacks are modifications of the 3D models that must be imperceptible and that can erase the data hidden by a watermarking scheme. In both cases, it is very important to assess whether a watermarking algorithm or watermarking attack produces perceptible distortions or not. The precedent Chapter has illustrated the poor results obtained by current 3D metrics as well as the emergence of perceptive metrics. Here, we propose a new image-based metric for the assessment of distortion perception for the most common watermarking attacks : noise addition, smoothing and simplification. The goal of our research is to design a metric able to follow human perception as well as defining perceptibility thresholds.

This Chapter is organized as follows. Section 4.2 presents our image-based metric composed of the Mutual Information criterion for images and the automatic generation of 2D views. Then Section 4.3 is dedicated to the psychovisual experiment needed to validate the ability of the metric to follow human visual perception of shape distortions. Finally, Section 4.4

presents the validation of the metric and compares the results of the experiment for our metric and most used 3D metrics in Computer Graphics.

4.2 Visual similarity metric through 2D projections

Image-based metrics have already been proposed by Lindstrom and Turk [89] who use a set of 2D views and compare them with a root mean square error metric. In this Chapter we also use automatically generated 2D views but compare them with another criterion. This criterion is the *Mutual Information* which is known to better capture the statistical similarity of images than the root mean square metric. The other advantage of image-based metrics on 3D-based metrics relates to the fact that there is no limitation on the representation of the 3D data. Indeed, image-based metrics can deal with meshes, NURBS, implicit surfaces etc. The major drawback of image-based metrics relates to animations. Rogowitz and Rushmeier [114] have demonstrated that, despite their performance for assessing perceived distortions on still 3D meshes, image-based metric do not behave well for animated meshes when compared with 3D-based metrics.

4.2.1 Mutual information

The *Mutual Information* (MI) measures the statistical dependance between two scalar features and has therefore become a really popular similarity measure in image processing. The Mutual Information between Random Variables (RV) X and Y is defined as :

$$MI(X, Y) = H(X) + H(Y) - H(X, Y), \quad (4.1)$$

where $H(\cdot)$ and $H(\cdot, \cdot)$ are the marginal and joint entropies ($H(X) = -E\{\log(p(X))\}$ if p is the probability density function of the RV X). In Eq. 4.1, X and Y denote the luminance of both images. Therefore, the key point of all Mutual Information based algorithms is the estimation of a probability density function (PDF). Viola [128] implemented an estimation of the PDFs by using Parzen Windows basis functions. In this work, the PDFs are estimated using histograms, following the implementation of the Insight Toolkit [62]. The normalized MI is commonly used and can be computed by:

$$NMI(X, Y) = \frac{MI(X, Y)}{\frac{H(X) + H(Y)}{2}} = \frac{MI(X, Y)}{\frac{MI(X, X) + MI(Y, Y)}{2}}, \quad (4.2)$$

with NMI always in the interval $[0, 1]$. If X and Y are equal then $NMI = 1$. Otherwise, if X and Y are statistically independent, $NMI = 0$. Now that we have a measure of similarity, we need a measure of dissimilarity. Several choices are possible, the most obvious is to take $1 - NMI(X, Y)$ as a measure of difference.

4.2.2 Generating 2D views

Assessing the quality of a 3D shape through interactions is somewhat natural but quite difficult to quantify. The usual procedure is exactly the opposite task since 3D modeling from 2D projections (or *silhouettes*) is a well explored field [96]. A possible approach to measure the distortion between a mesh M_1 and its modified version M_2 is to estimate the difference between multiple corresponding views. Obviously, rendering conditions (see Chapter 3) must be exactly the same for corresponding views of M_1 and M_2 . Exploring this issue, Rogowitz and Rushmeier [114] have performed several experiments to assess the influence of rendering conditions on perceived distortion for mesh simplification. They have shown that simple rendering conditions are optimal for shape distortion perception. We have followed their conclusions and selected these conditions:

- Only one white light source behind the camera position.
- Gouraud shading and diffusive reflection.
- no color and no texture added on the 3D models.
- screen of 1280×600 pixels and display windows of 600×600 .

Indeed, if texture has always been used to enhance quality and speed of 3D rendering techniques [105], it is however proved that human vision is also very sensitive to geometric features and silhouettes differences [55]. Sander et al. [116] have even demonstrated that silhouette clipping can produce renderings of similar quality to high-resolution meshes in less rendering time. Distortions of silhouettes and feature lines can be measured by several well-known techniques such as the Turning or signature functions, Hausdorff or Frechet distances, but these informations are by nature captured by the NMI criterion.

We use 18 to 364 views (Figure 4.1) by regularly discretizing ϕ and θ angles on the bounding sphere to compare meshes with the NMI and root

mean square criteria. This strategy is very similar to the subdivision of an initial octahedron and its projection on the sphere frequently used for silhouette clipping. Results naturally present a better accuracy if the number of views increases. The optimal tradeoff between accuracy and speed of course depends of the aimed application.

Before comparing pairs of views, we must register both models. We assume that there is no non-uniform scaling of the models and that both models share the same scale. Iterative Closest Point (ICP) methods are the most widely used to synchronize meshes using rotations and translations [21]. Their performance strongly depends on the initial position of both meshes and oscillations may occur in some cases. Another strategy is to compute the NMI criterion between a first view of M_1 and *all* views of M_2 . Selecting the pair that produces the best similarity score gives the rotation (since we know the position of the camera for each model view) needed to register both models.

Knowing whether some views of a model are more important than others for human perception remains an open issue. Cognition researchers have pointed out that protrusions (see Chapter 9) as well as sharp boundary features are the key regions of a model that influence human vision [60]. Computer graphics researchers have also shown that salient parts of a model influence user viewpoints when interacting with it [137]. Since no other evidence on that particular point has been published yet, we prefer to select views in an isotropic way.

When each pair of corresponding views has been compared by the NMI criterion, we combine these scores by a simple mean:

$$D_{NMI}(M_1, M_2) = \frac{1}{N} \sum_{i=1}^N NMI(M_{1i}, M_{2i}) \quad (4.3)$$

where M_{xi} is the i -th view of model M_x with $x = 1, 2$ and N is the total number of views of each model. Since we use a mean, the number of views does not affect much the result. Taking other combinations such as the median or the maximum have not given interesting results.

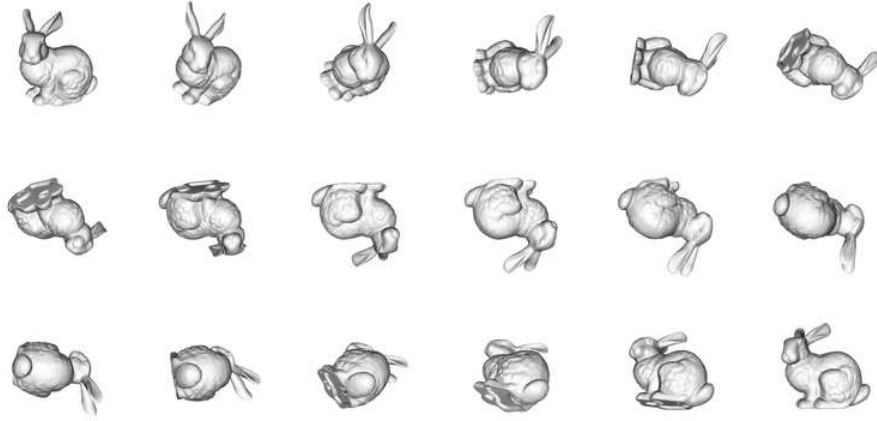


Figure 4.1: Stanford Bunny model views generated to compute the error between two meshes by comparing pairs of corresponding 2D projections.

4.3 Subjective evaluation experiment

In order to validate our metric, we have followed a protocol very close to the experiments of Pan et al. [105] as well as of Corsini et al. [38]. Rendering conditions are the same as those selected in the previous Section. The statistical analysis of the data have also been inspired by ITU-R [2] and ITU-T [1] recommendations. The models used for the subjective experiment are represented in Figure 4.2. They have been displayed on a 17 inch-TFT monitor with participants centered in front of it and at a distance of ≈ 40 cm. The experiment has been run with 21 people (7 females and 14 males) aged between 21 and 35. Users can compare the original model and the distorted one in neighboring windows of the same size and with the same rendering conditions. Cameras are synchronized so that mouse interactions with one model are also operated on the other to enable a perfect comparison (see Fig. 4.3).

The script of the experiment is presented in the next subsection. Inspired by the protocol of Corsini et al. [38], it is composed of several steps that help the user to get familiar with 3D model interactions and with the notion of perceived distortion. In the experiment, the user gives subjective scores between 0 and 10 (0 if no distortion is perceived, 10 if the distortion

is maximal). The user answers are first calibrated by showing him/her six maximal distortions and helping him/her to evaluate and track distortions on six exercise models that are not taken into account for the statistical analysis. Then the user evaluates 40 samples of model distortions without any interaction with the supervisor of the experiment. These samples were sorted by model type but randomly scrambled in the type of attack and intensity to avoid any learning influences in the perception of the distortion.

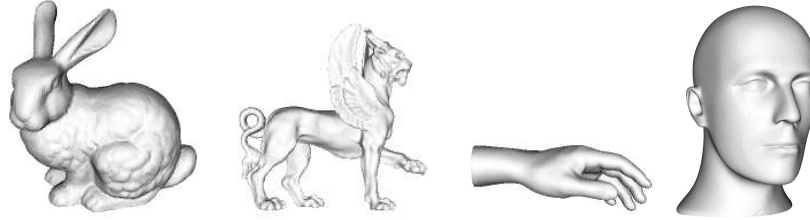


Figure 4.2: *Models used for the experiment. From left to right Stanford Bunny, Feline, Hand and Head models.*

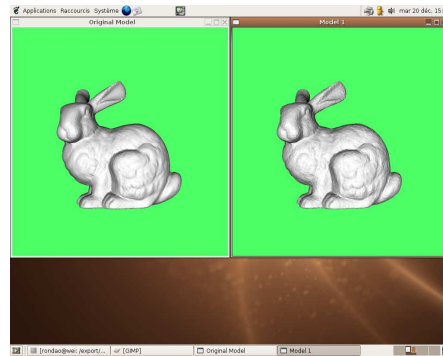


Figure 4.3: *Snapshot of the experiment environment. The original model (on the left) and the tested model (on the right) are represented with exactly the same rendering conditions and with synchronized cameras. Interactions are made by the use of a mouse.*

4.3.1 Experiment script

The script of the experiment has been inspired by works of Corsini et al. [38] and Pan et al. [105]. The first step is to position the user centered in front of the screen and about 40 cm from it. Then the supervisor explains the subject of this test i.e. the evaluation of the distortions introduced by watermarking attacks on the shape of a 3D model.

What is a 3D model? A 3D model is a digital media that represents a 3D shape. A sphere, a cylinder, a cube... are very simple examples of such shapes. Much more complex 3D models are used in animation movies and video games that nowadays get more and more popular.

What is the purpose of the test? Here, we deal with 3D model manipulations that produce some distortions that should not be perceived by the human eye. These manipulations consist in reducing the number of polygons used to represent the 3D model, smoothing the surface of the 3D model or modifying the roughness of the surface by producing small and random displacements of each point of the surface. The purpose of this test is the evaluation of the visual impact of these distortions. The test is quite simple. You will interact with some 3D models and you will have to indicate whether you see a certain kind of distortion. Some examples of these distortions will help you to understand what you have to evaluate.

3D Models and Perceived Distortion During the test you will interact with the models (by zooming and rotating them) and you will have to indicate how much you perceive such distortion by giving a score that is proportional to its distortion value. First you will see the four original models without any distortion. (*Show the original models*).

Before the experiment, you will see how 3D distortions look like. In a few words the roughness of the surface of the model is modified in some way. It is important that you focus on the smoothness of the parts of each model. (*Show the attacked models*.)

Now you will learn how to interact with the 3D model by using a mouse. To rotate the model, simply push the left button and move the mouse. When you want to stop rotating the model, release the left button. For zooming in the model, push the right button and move the mouse ahead

or back to zoom in or to zoom out respectively. Release the mouse button to stop zooming. (*Interaction trial - Try to rotate the model. Try to zoom it*).

Experiment Here are some examples that help you to numerically evaluate how much you perceive the distortions. You have to choose the worst case and mentally assign a value of 10 to it. This will give you an idea of the distortions that you will see. The distortions can be present or not on the model. So, you assign a score of 10 to the most evident distortions. If the perception of the distortions during the test is half the worst examples you chose, give it 5; if it is 1/10th as bad give it 1. Remember that the question is how much you perceive, how much you notice such kind of distortions. (*Show worst cases*).

Before we start the experiment you will have six practical exercises to be sure that you understand the task. You will respond in these trials just like you will in the main experiment. You will first interact with the original and the possibly distorted model. Then you will have to give a score to indicate how evident the distortion is. (*Practice Trials*).

You will be shown 40 samples during the test. This takes about 30 minutes. You can take the time you want to interact with the model and write your answer, but wait to hit $< q >$ until you are ready to go on. If by mistake, you close a window before being able to score the model, please call me. (*Experiment without any help of the supervisor*).

4.3.2 Statistical Analysis of the Experimental Results

The set of scores entered by each subject for each sample of the experiment is named by *subjective scores*. These subjective scores must be condensed by statistical techniques to yield results that summarize the characteristics of the system under test. The summarized score values are referred to as *Mean Opinion Score* (MOS) and represent the perceived distortion on an attacked 3D model. The subjective scores are very sensitive to scale and bias for each subject. Even if maximally distorted samples, compared to their original version, have been presented to the subjects of the experiment in order to tune their scale, some have globally given scores higher than the mean of the others. We present here the methods used to normalize the scale and bias of each subject. We also show how

to combine subjective scores and how to evaluate the accuracy of the estimates.

Normalization

A subject can be susceptible both of random and systematic errors. In order to get rid of the systematic errors, we use a normalizing procedure. Obviously, this step must be performed before combining the measurements across subjects. The unscaled annoyance value m_{ij} obtained from subject i after viewing the test sample j can be represented by the following sample [38]:

$$m_{ij} = g_i a_j + b_i + n_{ij} \quad (4.4)$$

where a_j is the true annoyance value for test sample j in the absence of any error, g_i is the gain factor of subject i , b_i is the bias of subject i and n_{ij} is generally assumed to be a sample from a zero-mean white Gaussian noise.

Since the gain and bias can have large variations across subjects of the experiment, the normalization procedure aims at reducing these systematic errors. In order to estimate whether the variations of gain and bias between subjects are significant, we use a two-way analysis of variance (ANOVA) approach [120]. A two-way ANOVA divides the total variations of a table of data into two parts: the variations attributed to the columns of the data table, and the variations attributed to the rows of the data table. Specifically, the F-test can be used to determine the likelihood that the means of the columns, or the means of the rows, are different. The experimental data m_{ij} is arranged so that each row represents data for one test sample and each column represents all the data from one test subject. The ANOVA assumes that the data can be modeled as [120]:

$$m_{ij} = \mu + \alpha_i + \beta_j + \epsilon_{ij} \quad (4.5)$$

where μ is the overall mean, α_i is the subject effect, β_j is the sample effect and ϵ_{ij} relates to the experiment errors. Table 4.1 illustrates the ANOVA results for the raw data collected during the experiment. The F-values (22.18 for subjects and 90.18 for samples) are large. The right-most column of the table contains the probabilities that the subject effect and the sample effect are constant and that there are no differences among the subjects or among the test samples. Obviously, we do not expect the sample effect to be constant since it depends on the particular test sample with varying

distortions. These probabilities are 0 for Table 4.1. This means that there are significant variations from subject to subject in the subjective scores. In order to check whether these variations are also due to the gain factor, we perform an ANOVA on the logarithm of m_{ij} :

$$\ln(m_{ij}) = \ln(g_i a_j + b_i + n_{ij}) \approx \ln(g_i) + \ln(a_j) + \frac{b_i}{g_i a_j} + \frac{n_{ij}}{g_i a_j}. \quad (4.6)$$

In this equation, $\alpha_i = \ln(g_i)$, $\mu + \beta_j = \ln(a_j)$ and $\epsilon_{ij} \approx \frac{b_i}{g_i a_j} + \frac{n_{ij}}{g_i a_j}$. Here, ϵ_{ij} is no longer independent of the other factors, however, if b_i and n_{ij} are small, this does not matter much in the analysis. Table 4.2 contains the ANOVA results for $\ln(m_{ij} + 1)$ ($m_{ij} + 1$ to avoid numerical problems when $m_{ij} = 0$). Here again the F values are large and probabilities are zero.

Source	SS	df	MS	F	Prob>F
Subjects	665.35	19	35.018	22.18	0
Samples	5552.7	39	142.377	90.18	0
Error	1169.95	741	1.579		
Total	7388	799			

Table 4.1: Anova test for raw data : user mean bias. The probability that the means of the subjects are equal and the probability that the means of the samples are equal, is zero.

Source	SS	df	MS	F	Prob>F
Subjects	31.816	19	1.67455	17.54	0
Samples	319.645	39	8.19601	85.85	0
Error	70.744	741	0.09547		
Total	422.205	799			

Table 4.2: Anova test for raw data : user gain bias. The probability that the means of the gains of each subject are equal and the probability that the means of the gains of each sample are equal, is zero.

This shows that the gain varies from a subject to another and that this requires a correction before the combination of all users' results. The measurements are corrected in the following way:

$$\hat{m}_{ij} = \frac{1}{\hat{g}_i} (m_{ij} - \hat{b}_i) \quad (4.7)$$

where $\hat{g}_i$ is the corrected gain of subject i , $\hat{b}_i$ is the corrected bias of subject i and $\hat{m}_{ij}$ is the normalized score.

The bias is simply estimated by the mean of all measurements made by each subject:

$$\hat{b}_i = \frac{1}{J} \sum_{j=1}^J m_{ij} - \mu \quad (4.8)$$

where $\mu = \frac{1}{J} \sum_{j=1}^J \sum_{i=1}^N m_{ij}$, $J = 40$ is the total number of samples and $N = 21$ is the total number of subjects. Table 4.3 shows the results after correcting the bias. The means of the subjects are now the same with a F-value equal to zero and a probability equal to 1. In order to adjust the

Source	SS	df	MS	F	Prob>F
Subjects	0	19	0	0	1
Samples	3553.76	39	91.1221	93.17	0
Error	724.74	741	0.9781		
Total	4278.5	799			

Table 4.3: Anova test for normalized data : user mean bias. The user mean bias is now corrected as the probability that the means of each subject are equal, is equal to one.

gain, we estimate the range of each subject by computing the standard deviation of the values. The gain factor becomes:

$$\hat{g}_i = \frac{4\delta_i}{N} \quad (4.9)$$

where δ_i is the standard deviation of all the N values recorded by the i -th subject. Table 4.4 illustrates that the data now only depend on the models as the means of bias and gain are the same for each subject. This fact indicates that the experiment is well designed, i.e. the experimental methodology is not affected by any systematic error.

Source	SS	df	MS	F	Prob>F
Subjects	0.432	19	0.02271	0.52	0.9527
Samples	132.803	39	3.4052	78.61	0
Error	32.097	741	0.04332		
Total	165.331	799			

Table 4.4: Anova test for normalized data : user gain bias. The probability that the gain of each subject is constant, is now equal to .9527.

Overall Scores and Screening

Now that the data have been normalized, the subjective scores can be combined into a single overall score for each test model using the sample mean. The subjective MOS is given by:

$$\mu_j = \frac{1}{N} \sum_{i=1}^N m_{ij} \quad (4.10)$$

Confidence intervals have been calculated by using Student's T-distribution [120]. This distribution is appropriate when only a small number of samples is available. The sample standard deviation s_j is calculated for each sequence j using μ_j and the limits are calculated by:

$$l_j = t_{(0.05, N)} s_j \quad (4.11)$$

where $t_{(0.05, N)}$ is the t -value associated with a probability of 0.95 and N degrees of freedom. As the number of observations N increases, the confidence interval decreases. After this step, one can perform the subject screening procedure. An expected range of values is calculated for each model. If a subject is erratic on both sides of the expected range, we prune this subject. In our experiment, after normalization and screening, we have discarded only one subject.

4.4 Validation of the proposed metric

This Section presents the comparison between state-of-the art metrics, our proposed metric and the subjective evaluations obtained during the experiment. Figures 4.4 to 4.5 compare our 2D-projections-based metric with usual metrics found in the literature and presented in Chapter 3 and the results of the experiment presented in Section 4.3 for the three more common watermarking attacks: Gaussian noise addition, Laplacian smoothing and vertex decimation.

These results show that the NMI criterion (referred as mutual information) is usually close to the subjective visual perception of the 3D models. The confidence intervals of the collected statistical data are also represented in these figures. Least Squares methods [89], Hausdorff distance, VSNR (a.k.a. RMSE), the Geometric Laplacian do not capture very well

the curve dynamic of the subjective results. Notice that, for the simplification attack, the VSNR and the Geometric Laplacian cannot be computed because the models do not share the same connectivity anymore.

4.4.1 Noise addition

Figures 4.4 and 4.5 show that the human eye is very sensitive to noise addition for the *Bunny* and *Head* models, respectively. All metrics underestimate the perception curve but NMI is the closest in our experiments. The distortions related to Fig. 4.4 can be viewed in Fig. 4.12. Noise perception is not the same for both models. It seems that the users of the experiment have perceived the noise addition distortion on the human face as much more noticeable than for the other model. This may be partly explained by the fact that the *Head* model is smoother than the *Bunny* model. The other possible origin of this result is related to the fact that the human eye is much more trained to distinguish human faces than bunnies.

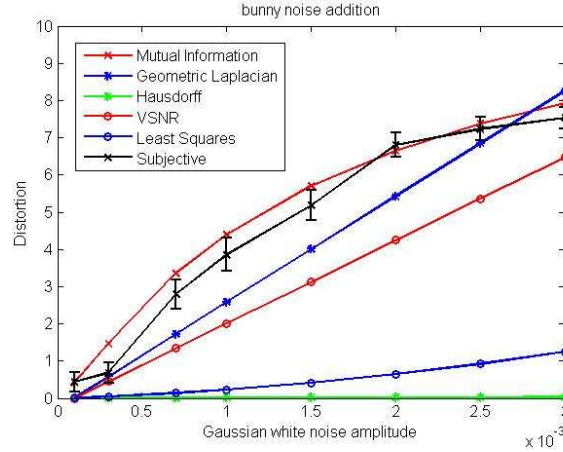


Figure 4.4: Noise attack for the *Bunny* model for increasing force in terms of the amplitude factor α . VSNR and Geometric Laplacian linearly evolve with the noise factor. These two metrics as well as the NMI metric are the closest to the perceptive curve.

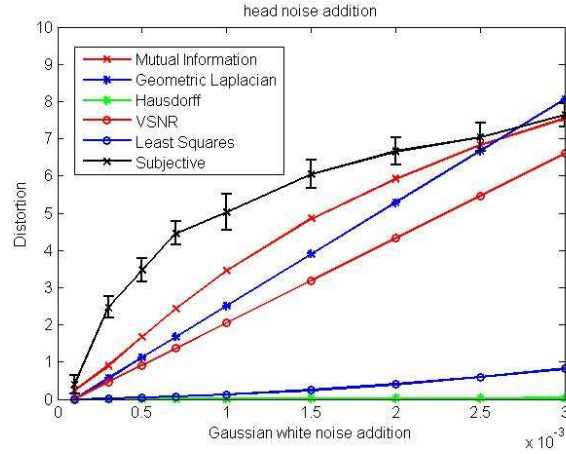


Figure 4.5: Gaussian noise attack for the Head model for increasing force in terms of α . All metrics underestimate the perception of noise on the surface, maybe because the model represents a human face. Here again, NMI is closer to the perceived distortion.

4.4.2 Vertex simplification

In the case of the simplification attack on the *Bunny* and *Hand* models, the NMI is very close to the perception curve and nearly always in the confidence area around the latter (see Fig. 4.6 and 4.7).

4.4.3 Smoothing

The smoothing attack is also well captured for the *Feline* model in Figure 4.8. NMI behaves well compared to state-of-the-art metrics. Some corresponding 3D meshes are represented in Fig. 4.10. The fact that smoothing is not so noticeable for 3D meshes is based on the fact that we test models which are represented by already smooth surfaces. It is then more difficult to discriminate a fair and a «fairer» version of the same surface. Furthermore, as seen in the previous Chapter, rendering techniques usually smooth normal angle differences in order to give the perception of a continuous surface. This second step of smoothing can explain why smoothness differences are not well distinguished by human users.

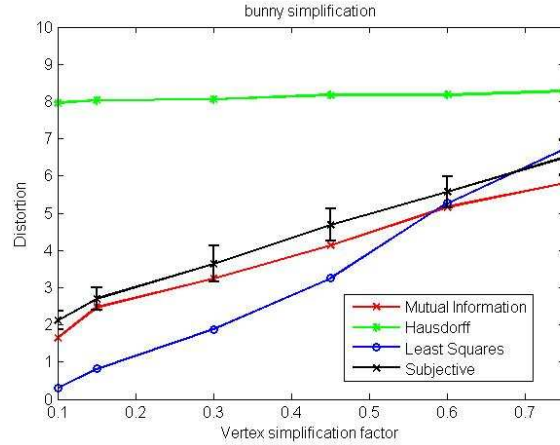


Figure 4.6: *Simplification attack for the Bunny model for increasing force in terms of percentage of decimated points. VSNR and Geometric Laplacian cannot be computed since the connectivity of the models differ. Image-based metrics are closer to perception than the Hausdorff distance.*

Finally, the smoothing attack undergone by model *Head* is represented in Fig. 4.9 with some views in figure Fig. 4.11. Here again, the perceived noise distortions are underestimated by all metrics while smoothing is well captured by NMI. This model has been differently perceived when compared to other models for the same attacks. This may be partially explained by the fact that human are very acquainted with human faces. It would however be impossible to design a 3D metric able to follow the content-dependent variations of human eye perception. It may be then useful to adapt the metric to the targeted application in case of very specific 3D models such as human faces.

4.4.4 Image-based metrics and view variance

It may also be interesting to consider the variance of the views contributions for image-based metrics. Intuitively, smoothing and noise addition attacks produce global distortions while vertex simplification minimizes local distortions. Our tests (figures 4.13 to 4.15) confirm this intuition and also show that the NMI metric shows lower standard deviation variance than the Least Squares metric. Figure 4.15 shows that for the case of the

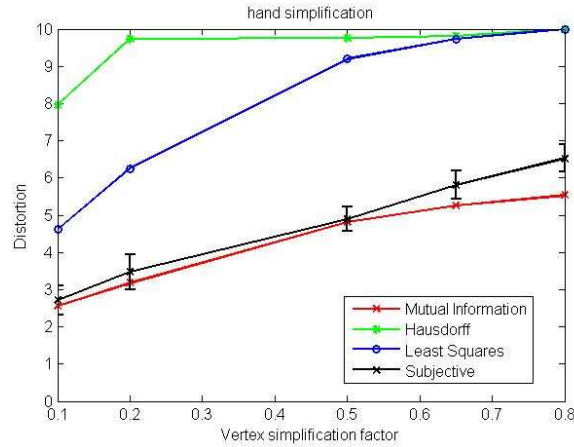


Figure 4.7: *Vertex simplification attack for the Hand model for increasing force in terms of iterations. Image-based metrics are closer to human perception than Hausdorff distance. Here again VSNR and Geometric Laplacian cannot be computed.*

simplification attack, the standard deviation does not monotonically increase with the amplitude of the attack.

4.4.5 Final observations

Several observations can be made on the results of our experiment. Users seem to be very sensitive to noise addition, sensitive to simplification and almost not sensitive to smoothing. RMSE (VSNR) and Geometric Laplacian criteria fail at describing that fact because smoothing causes as large values as noise addition. This is a good point of our metric which solves the problem presented in Chapter 3 (see Fig. 3.6).

Even though we did not perform a perceptibility threshold experiment (such as Just-Noticeable Difference tests [32]), we can say that for the four models tested, noise addition attacks become perceptible starting from 0.002 (amplitude factor), smoothing attacks starting from 100 iterations and simplification starting from 10 pc. This is important for the next part of the thesis which is dedicated to 3D watermarking.

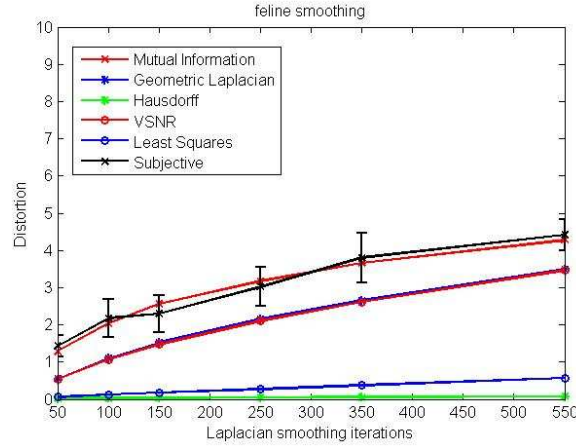


Figure 4.8: *Smoothing attack for the Feline model for increasing force in terms of iterations. NMI (and up to a small bias, VSNR and Geometric Laplacian) capture(s) well the perceived distortion.*

4.5 Conclusion

This Chapter has presented a new perceptual metric for the visual quality assessment of 3D objects from 2D projections in the particular context of watermarking attacks. This metric has been compared to the state-of-the-art and validated by a psychovisual experiment. If image-based metrics are not new by themselves, our contribution relates to the use the Mutual Information criterion to estimate distortions of 3D watermarking attacks. Future work will extend our experiments to measure the perceived artifacts caused by a set of 3D watermarking schemes. Furthermore, it would be interesting to compare the metric of Corsini et al. [38] with our works.

This work has been published in [4, 7].

The remainder of this thesis is now dedicated to 3D watermarking which is introduced in the next Chapter.

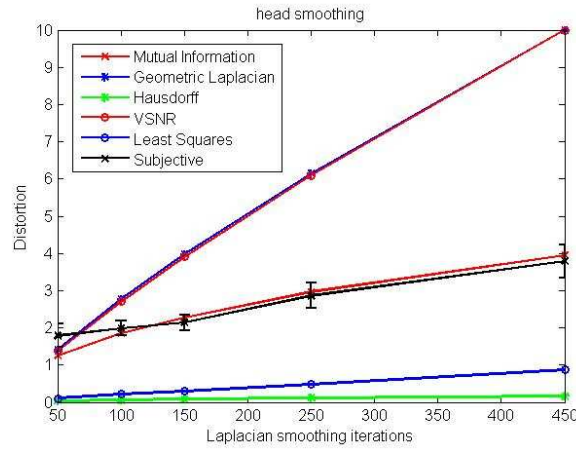


Figure 4.9: *Laplacian smoothing attack for the Head model for increasing force in terms of iterations. VSNR and Geometric Laplacian linearly evolve with the number of iterations while perception does not significantly change. NMI is close to the perception curve.*

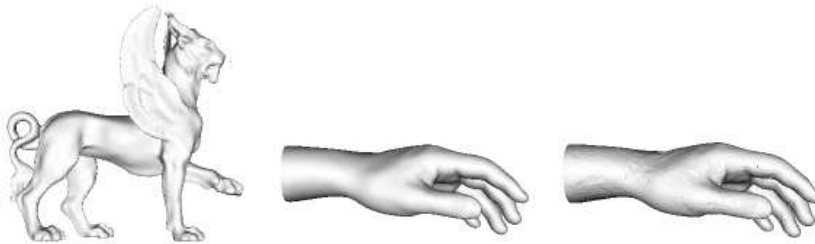


Figure 4.10: *From left to right, Feline model attacked by smoothing with 350 iterations and Hand model attacked by simplification of 20 and 80 pc., respectively.*

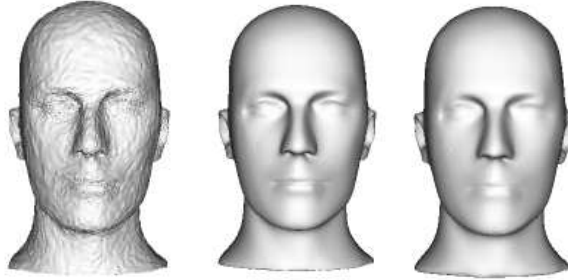


Figure 4.11: From left to right, Head model represented for the noise attack with $\alpha = 0.0025$ and smoothing attack with 150 and 450 iterations.

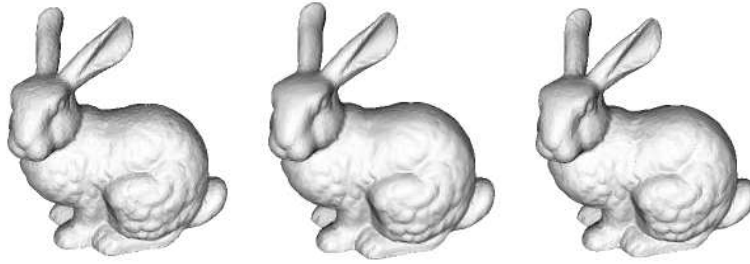


Figure 4.12: From left to right, Bunny model attacked by Gaussian noise attack with $\alpha = .001$ and $.0003$ and a simplification attack of $.45$ pc.

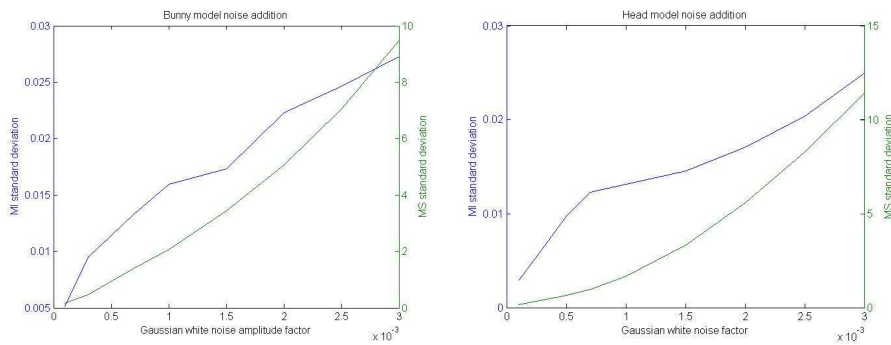


Figure 4.13: From left to right, noise addition and standard deviation of views contributions for the Least Squares (MS) and NMI (MI) metrics tested on Bunny and Head models, respectively.

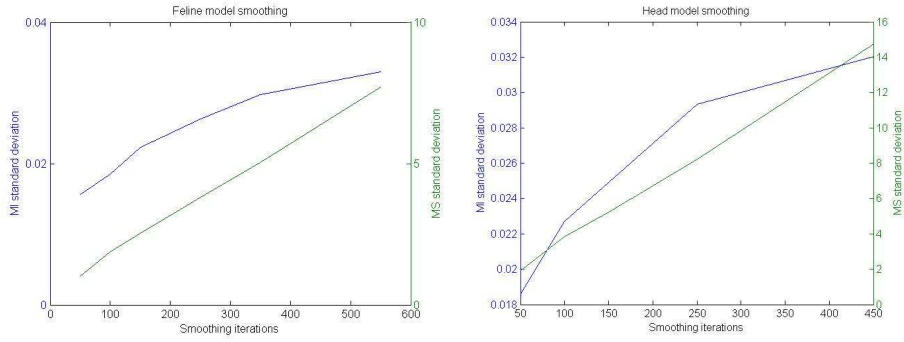


Figure 4.14: From left to right, smoothing addition and standard deviation of views contributions for the Least Squares (MS) and NMI (MI) metrics tested on Feline and Head models, respectively.

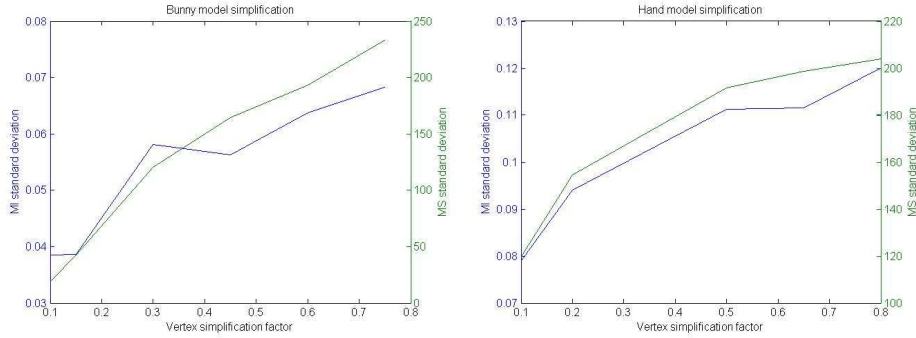


Figure 4.15: From left to right, simplification attack and standard deviation of views contributions for the Least Squares (MS) and NMI (MI) metrics tested on Bunny and Head models, respectively.