

# RGL semester project report on the FLIP algorithm

Juliette Parchet

Referent : Sébastien Speirer

Professor : Wenzel Jakob

Fall 2021

# Contents

|          |                                                                                                 |           |
|----------|-------------------------------------------------------------------------------------------------|-----------|
| <b>1</b> | <b>Introduction</b>                                                                             | <b>2</b>  |
| 1.1      | Motivation behind the $\mathcal{F}$ LIP tool . . . . .                                          | 2         |
| 1.2      | Human Vision System . . . . .                                                                   | 2         |
| <b>2</b> | <b><math>\mathcal{F}</math>LIP details</b>                                                      | <b>3</b>  |
| 2.1      | Color Pipeline . . . . .                                                                        | 3         |
| 2.2      | Feature Pipeline . . . . .                                                                      | 4         |
| 2.3      | $\mathcal{F}$ LIP difference value . . . . .                                                    | 4         |
| <b>3</b> | <b>Results, Comparisons, and Performance</b>                                                    | <b>6</b>  |
| 3.1      | Results of Mitsuba implementation . . . . .                                                     | 6         |
| 3.2      | How to use Mitsuba $\mathcal{F}$ LIP . . . . .                                                  | 7         |
| 3.3      | Comparison of the Mitsuba implementation with the Numpy implementation of $\mathcal{F}$ LIP . . | 7         |
| 3.4      | Performance of Mitsuba against the Numpy implementation of $\mathcal{F}$ LIP . . . . .          | 8         |
| 3.5      | Performance of $\mathcal{F}$ LIP against the L1, L2 errors used currently . . . . .             | 8         |
| <b>4</b> | <b>Conclusion</b>                                                                               | <b>10</b> |
| 4.1      | Why it could it be helpful for the Mitsuba framework . . . . .                                  | 10        |
| 4.2      | What I have learned from this bachelor's project . . . . .                                      | 10        |

# Chapter 1

## Introduction

### 1.1 Motivation behind the FLIP tool

Measuring the image quality of a rendering algorithm is essential in the graphic rendering domain. A standard technique used in computer graphics is flipping between two similar images to reveal their differences (a *reference* image, obtained by a trusted way, representing the authentic image; and a *test* image, an approximated rendering of the reference image). The FLIP algorithm aims to automate this strenuous method and render a new image containing the amplitude of the perceived distance between the reference and test image per pixel. As explained in the FLIP paper<sup>1</sup>, the algorithm is developed on the evaluation of error maps of the manual flipping method. Indeed, the FLIP tool aims at rendering a differential image between test and reference images by taking into account numerous aspects of the human visual system. This way, it would highlight the differences humans perceive and mask some that humans do not recognize.

### 1.2 Human Vision System

FLIP especially pays attention to the viewing conditions of the images. It is expressed with the parameter  $p$  - the number of pixels per degree (which considers the monitor size, the resolution of the screen, and the distance from the screen). This parameter  $p$  is an optional argument in the `flip(...)` function<sup>2</sup>, and is set by default to the value 67 (for a 4k-resolution monitor of size 0.69 x 0.39 m<sup>2</sup>, of resolution 3840 x 2160 pixels, with the observer at a distance 0.7 m of the screen).

Another particular point targeted by FLIP is the Hunt effect, which states that the chromatic differences between images with higher brightness appear larger. To counter this effect, we use the L\*a\*b Hunt color space, which lowers the chromatic distance between two colors if the brightness is low by applying a linear function to the  $a$  and  $b$  channels.

Furthermore, FLIP takes into account the edge and point-based differences - such as fireflies, with the use of the feature pipeline. We will discuss more the feature pipeline further in the report.

---

<sup>1</sup>Andersson, Jim Nilsson, Tomas Akenine-Möller, Magnus Oskarsson, Kalle Åström, and Mark D.Fairchild. 2020.FLIP: A Difference Evaluator for Alternating Images. Proc. ACM Comput. Graph. Interact. Tech. 3, 2, Article 15 (August 2020), 23 pages. <https://doi.org/10.1145/3406183>

<sup>2</sup>or the `flip_pooling(...)` function

## Chapter 2

### TLIP details

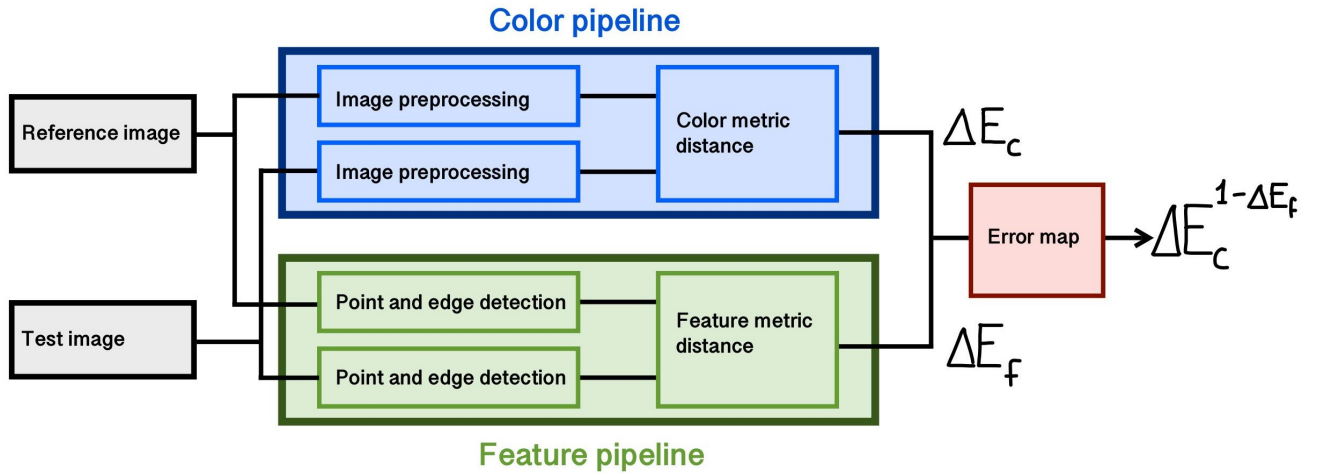


Figure 2.1: the TLIP pipeline

As we can see in Figure 2.1, the TLIP algorithm is divided into two distinct pipelines: color and feature pipelines. Let us take a closer look at these two.

#### 2.1 Color Pipeline

The color pipeline first preprocesses and filters the reference and test images. To void some color shifts during filtering (leading to achromatic and chromatic confusion), TLIP uses the  $YyCx Cz$  color space, a linearized version of CIE  $L^*a^*b$  (obtained via the XYZ tristimulus values) using the standard illuminance D65. Here,  $Yy$  contains the achromatic information,  $Cx$  the red/green information, and  $Cz$  the blue/yellow information. We transform the contrast sensitivity functions from the frequency to the spatial domain, where we perform a gaussian-based convolution of the reference and test images. Finally, the images are transformed into the CIE  $L^*a^*b$  color space, a perceptually uniform color space. The idea behind the CIE  $L^*a^*b$  color space is that the numeric distances between colors agree with the perceived distance between them. Then as discussed before, to consider the Hunt effect, we adjust the  $a$  and  $b$  channels and transform the images into the  $L^*a^*b$  Hunt color space.

The color pipeline then computes the color metric when the preprocessing is done for the reference and test images. As the color difference is often significant in rendering applications, TLIP uses the HyAB distance metric instead of a more common one like the Euclidian distance. Indeed, the HyAB distance metric, a hybrid model between the Euclidian and the city-block models, is designed to handle large color



distances. Finally, a last correction of the result is made: to reduce the gap between significant differences in luminance versus large differences in chrominance,  $\mathcal{F}$ LIP offers to compress large distances into the top of the  $[0, 1]$  range and map smaller distances to the remaining part of it. This process is shown in Figure 2.2.

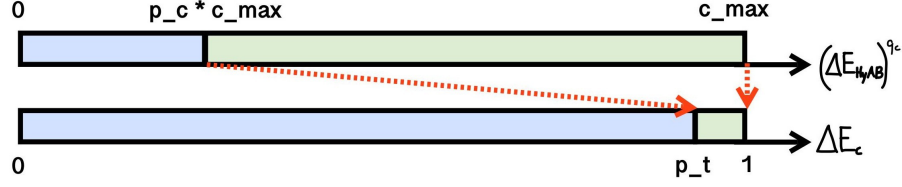


Figure 2.2: Error mapping compression of the color distance

Ultimately, the color pipeline yields this mapped color distance, denoted by  $\Delta E_c$ , resulting from preprocessing of the reference and test images, and their color metric distance computation.

## 2.2 Feature Pipeline

By simply examining the color distance between two images, we would miss significant (and very strongly detected by humans) differences: the edge and point detections. This is why, in addition to the color pipeline,  $\mathcal{F}$ LIP developed a feature pipeline that detects edges and points in the images and then computes the feature metric distance between the reference and test images. This result will amplify the color metric distance computed by the color pipeline (as shown in Figure 2.1).

The feature detection (edge and point detection) is done in the  $YyCxCz$  color space. It only analyzes the achromatic  $Yy$  channel, which contains the majority of the high-frequency spatial data. The feature detection is done via the convolution of the image, and the kernels are computed based on the first gaussian derivative for the edge detection and the second gaussian derivative for the point detection.

After the feature screening, the feature metric distance is calculated. Considering that edges and points never occur in the same pixels, this distance is calculated by taking the maximum absolute distance between the edge and point magnitude of the reference and test images. Finally, this distance is normalized in  $[0, 1]$  and taken to the power of  $q_f = 1/2$  to increase the incidence of the feature difference on the final error map.  $\Delta E_f$  denotes this final value.

We can observe in Figure 2.3 the significance of the feature pipeline: the perceptually noticeable difference in the structure of the stool feet is enhanced by the feature pipeline.



Figure 2.3: 1) The reference image of the tool, 2) the test image of the stool, 3)  $\Delta E_c$  : the color pipeline only, 4) the entire  $\mathcal{F}$ LIP pipeline

## 2.3 $\mathcal{F}$ LIP difference value

When the per-pixel color metric and feature metric distances are computed, the  $\mathcal{F}$ LIP difference value (denoted by  $\Delta E = (\Delta E_c)^{1-\Delta E_f}$ ) can be calculated. It results in an Enoki TensorXf of shape  $H \times W \times 1$  (with

H the image height, W the image width, and 1 indicate that  $\Delta E$  only has one channel). As both  $\Delta E_c$  and  $\Delta E_f$  are in  $[0, 1]$ ,  $\Delta E$  is also in  $[0, 1]$ . We can remark that if  $\Delta E_c = 0$  then  $\Delta E = 0$ , independently of  $\Delta E_f$ . Likewise, if  $\Delta E_f = 1$ , then  $\Delta E = 1$  is maximal, regardless of  $\Delta E_c$ . Furthermore, if  $\Delta E_f = 0$ , then  $\Delta E$  is entirely determined by  $\Delta E_c$ .

## Chapter 3

# Results, Comparisons, and Performance

The results and performance analysis are performed with a ThinkPad T480s, with the UserBenchmarks analysis<sup>1</sup>: CPU: Intel Core i7-8550U - 47.9%, GPU: Intel UHD Graphics 620 (Mobile Kaby Lake R) - 5.7%, SSD: Samsung PM981 NVMe PCIe M.2 512GB - 135.1%, RAM: Samsung M471A1K43BB1-CRC M471A1K43CB1-CRC 16GB - 67.4%, MBD: Lenovo 20L7CTO1WW.

### 3.1 Results of Mitsuba implementation

As discussed earlier, the Mitsuba  $\mathcal{F}$ LIP implementation can yield an error map  $\Delta E$ . Nevertheless, we may want to condense the error map found to a smaller set of values. That is the role of the `flip_pooling` function, which renders a histogram containing the values of  $\Delta E$  stored in several bins. It is also possible to extract valuable values as the mean or the median of  $\Delta E$ . Let us see here these different results.

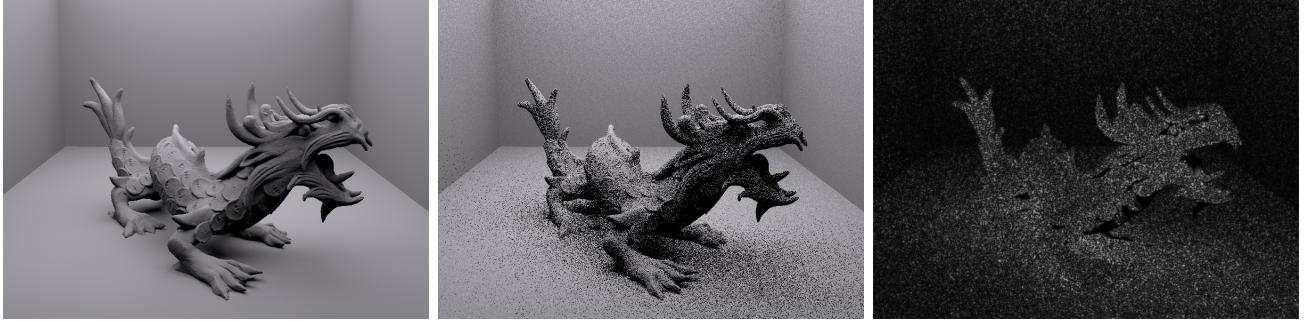


Figure 3.1: 1) the reference dragon image, 2) the test dragon image, 3) the error map  $\Delta E$  of the dragon images (TensorXf type)

In Figure 3.1, we can see the  $\mathcal{F}$ LIP result with the most data: the error map  $\Delta E$ . It is returned by the `flip(...)` method as an Enoki TensorXf of dimension  $H \times W \times 1$  (with  $H$  the height of the image,  $W$  the width of the image and 1 to signify that the TensorXf only has one channel). It is the main aim of  $\mathcal{F}$ LIP: it contains all the data computed by the algorithm and is the automated result of the manual flipping technique.

In Figure 3.2, we can observe a plotted histogram of  $\Delta E$ , in addition to the mean and the median values. These results are obtained with the `flip_pooling(...)` function, which renders the histogram of the error map, and optionally the mean and/or the median. These results could prove helpful, for example, as a machine learning loss function parameter or as an approval threshold for validation of an implemented algorithm.

---

<sup>1</sup><https://www.userbenchmark.com/>

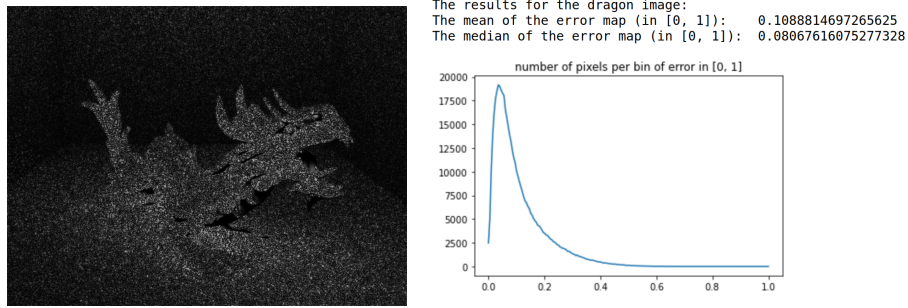


Figure 3.2: 1) the error map  $\Delta E$  of the dragon images (TensorXf type), 2) the plotted histogram, the mean and the median of  $\Delta E$

### 3.2 How to use Mitsuba $\mathcal{FLIP}$

All the functions used for the Mitsuba  $\mathcal{FLIP}$  implementation are written in the `flip_mitsuba.py` python document. The only functions which should be used are

`flip(reference, test, p = 67, p_c = 0.4, p_t = 0.95)` and  
`flip_pooling(reference, test, nb_bins = 100, return_mean = False, return_median = False, p = 67, p_c = 0.4, p_t = 0.95)`

As explained in the documentation, the `flip(...)` function takes two Enoki TensorXf images as parameters, and returns an Enoki TensorXf new image indicating the perceived distance between the reference and test images. Furthermore, it takes some optional arguments. The parameter  $p$ , as we discussed earlier (with its default setup being 67 pixels per degree), is calculated as  $p = d \times \frac{W_p \times \pi}{W_m \times 180}$ , with  $W_p$  the monitor width pixel number,  $W_m$  the monitor width in meters, and  $d$  the distance from the screen in meters. The parameters  $p_c$  and  $p_t$  are used as mentioned above to compress the error mapping of the color pipeline distance and can be changed to optimize the  $\mathcal{FLIP}$  algorithm.

The `flip_pooling(...)` function takes the same arguments as the `flip(...)` method but additionally takes `nb_bins`, which indicates the number of bins of the histogram's x-axis to construct. In addition, the method also takes `return_mean` and `return_median`, which indicates if the function should return the mean and/or the median of  $\Delta E$ . This function returns (always) the pooled histogram of  $\Delta E$  (an Enoki Float type) with its x-axis going from 0 to 1, and returns (if asked) the mean and/or the median of  $\Delta E$ .

### 3.3 Comparison of the Mitsuba implementation with the Numpy implementation of $\mathcal{FLIP}$

In Figure 3.3 we can observe the result of the Enoki and Numpy  $\mathcal{FLIP}$  implementations.<sup>2</sup> They seem very similar, and indeed, below is the displayed absolute difference between these implementations: the difference is almost zero, with the mean = 0.000. In conclusion, the LDR  $\mathcal{FLIP}$  algorithm seems to be correctly implemented in the Mitsuba framework because almost identical to the Numpy implementation provided by  $\mathcal{FLIP}$ . However, one can see some performance differences.

<sup>2</sup>See the `test_rapport_flip_pipeline.ipynb` Jupyter Notebook for more details

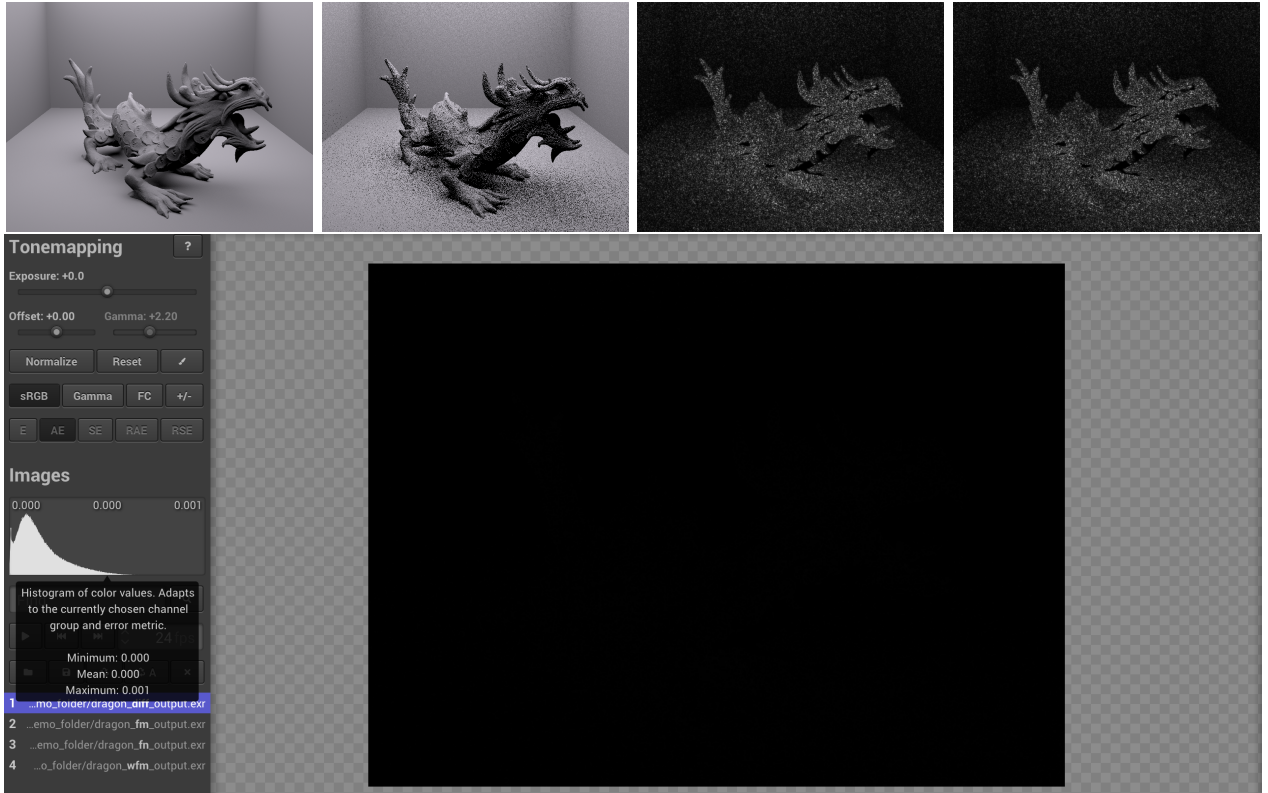


Figure 3.3: 1) the reference dragon image, 2) the test dragon image, 3) Mitsuba TFLIP  $\Delta E$ , 4) Numpy TFLIP  $\Delta E$ , 5) the absolute difference between 3) and 4) displayed in the Tev application

### 3.4 Performance of Mitsuba against the Numpy implementation of TFLIP

After running 1000 times the TFLIP method both in Mitsuba and Numpy frameworks (with the dragon image: width = 800 pixels, height = 600 pixels, linear RGB,  $p = 67$ ,  $p_c = 0.4$ ,  $p_t = 0.95$ ), we notice that the Mitsuba TFLIP is around 2 times slower than the Numpy TFLIP.<sup>3</sup> For the dragon image, Mitsuba TFLIP takes 0.91 seconds, where Numpy TFLIP takes 0.40 seconds.

The difference was much more significant in the early stages of the project (around ten times slower), and some improvements were made (particularly in the feature pipeline: the `convolve_edge_detection` function was improved by only using simple computations and minimizing the calls to `lookup_offset(...)`).

Some improvements are yet to be made, particularly in the convolution calculation: the `lookup_offset` function uses some `ek.gather(...)`, which are very costly, and it slows down a lot the algorithm execution. Furthermore, we could improve the performance by parallelizing the preprocessing of the images, the edge and point detection, and, more generally, the two pipelines. Indeed, these processes are primarily independent and would be pretty easy to parallelize.

### 3.5 Performance of TFLIP against the L1, L2 errors used currently

As we can remark in figure 4.3, the error map is very different between the TFLIP algorithm and the L1 (city-block or Manhattan) and L2 (Euclidean) techniques. One could argue that the TFLIP algorithm renders more faithfully than the L1 and L2 techniques the errors distinguished between the reference and test images. For example, the fireflies on the ground are way more precise in 3), and on the contrary in

<sup>3</sup>See the `test_rapport_performance.ipynb` Jupyter Notebook for more details



4) and 5) it is blurry. Also, the features of the dragon are mostly well-rendered in the test image, and 3) shows that fact, contrarywise of 4) and 5).

However, another aspect to take into account is the performance, and here  $\mathcal{F}$ LIP is really behind L1 and L2. When running the different methods 1000 times: Mitsuba  $\mathcal{F}$ LIP is approximately 1714 times slower than L1, and 1830 times slower than L2<sup>4</sup>. The Mitsuba  $\mathcal{F}$ LIP runs in approximately 0.86 seconds, the L1 in 0.00051 seconds, and L2 in 0.00047 seconds.

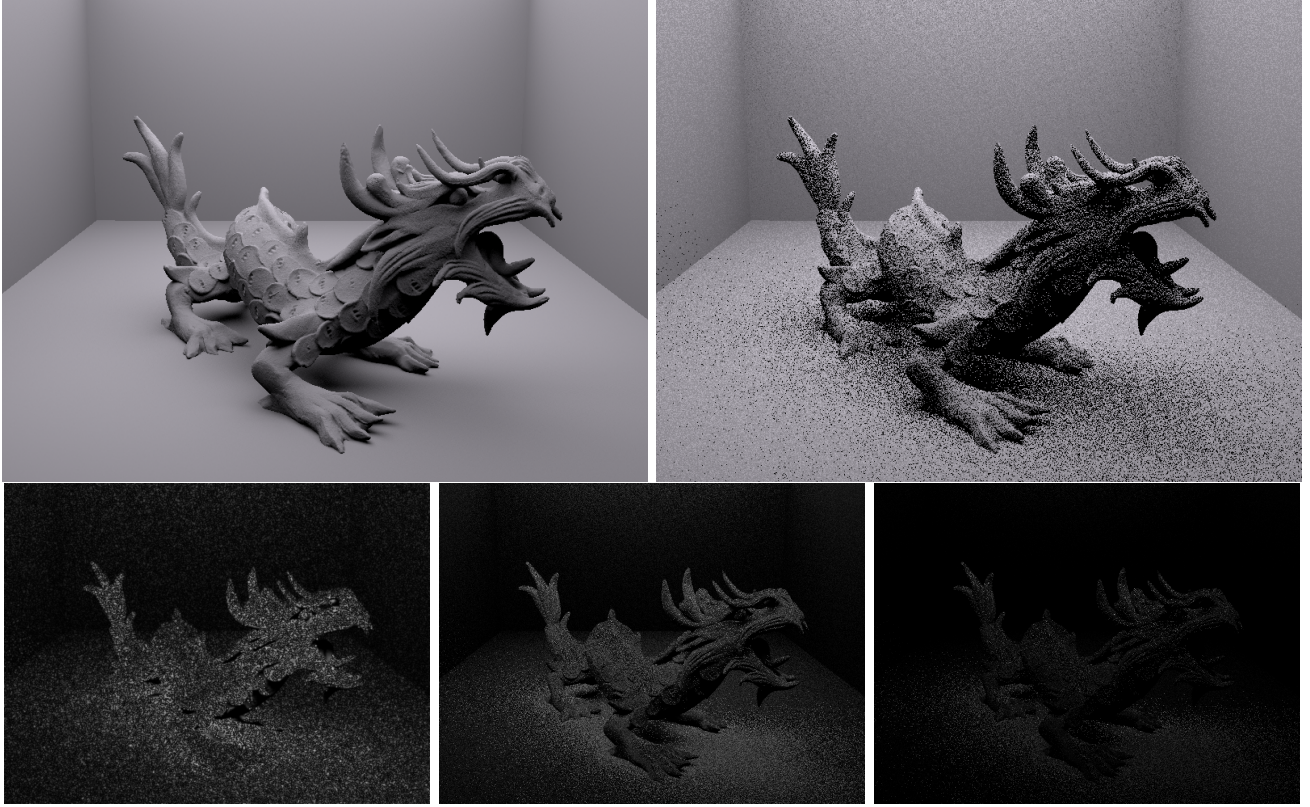


Figure 3.4: Error map between the 1) reference and 2) test images for 3) the Mitsuba  $\mathcal{F}$ LIP, 4) the L1 (Manhattan or city-block) distance, 5) the L2 (Euclidean) distance

---

<sup>4</sup>See the `test_rapport.L1.L2.comparison.ipynb` Jupyter Notebook for more details

# Chapter 4

## Conclusion

### 4.1 Why it could it be helpful for the Mitsuba framework

As we discussed in the introduction, currently used methods to assess the quality of a rendering algorithm (by comparing a reference and test images) can be: manual flipping, or simple distance computations (as L1 or L2 methods). I believe that the manual flipping between two images is very time-consuming, and that the L1/L2 methods, even if very fast, do not represent as well as the  $\mathcal{F}$ LIP method the perceived differences between two images. Furthermore,  $\mathcal{F}$ LIP does not only render the new image  $\Delta E$ , but can also render a histogram or single values such as the mean and the median. Because of the performance issue, Mitsuba  $\mathcal{F}$ LIP may not yet be used with large data sets, but I believe that for smaller ones,  $\mathcal{F}$ LIP could prove very useful and precise.

### 4.2 What I have learned from this bachelor's project

This semester, I learned more about the computer graphic domain and the complexity of evaluating the quality of a rendered image. By analyzing the  $\mathcal{F}$ LIP paper (and related works), I broadened my vision of the computer graphics field. I also had to assimilate an unknown and not fully documented library: Enoki. It was quite a challenge, as I never had to deal with a similar issue during my studies. Hopefully, my referent was here to guide me through this new situation and explained to me how to use some of the most useful methods of Enoki, like `ek.gather(...)`, or `ek.scatter(...)`. I also learned more about Enoki types as `Array3f`, `Bitmap`, or `TensorXf`, which differs from the python types I have worked with. Finally, I experienced running semester-length project outside of a course for the first time, which allowed me to have a more flexible calendar and project direction. In conclusion, working with the Realistic Graph Laboratory was a unique, engaging, and stimulating experience.

All the work done (the  $\mathcal{F}$ LIP implementation, the images used in the report, the tests performed,...) can be found in Github at the web address: <https://github.com/Jucifer06/Flip>