



A Practical Self-Shadowing Algorithm for Interactive Hair Animation

Florence Bertails, Clément Ménier, Marie-Paule Cani

► To cite this version:

Florence Bertails, Clément Ménier, Marie-Paule Cani. A Practical Self-Shadowing Algorithm for Interactive Hair Animation. Graphics Interface, May 2005, Victoria, Canada. inria-00519449

HAL Id: inria-00519449

<https://inria.hal.science/inria-00519449>

Submitted on 20 Sep 2010

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

A Practical Self-Shadowing Algorithm for Interactive Hair Animation

Florence Bertails

Clément Ménier

Marie-Paule Cani

GRAVIR - IMAG/INRIA, Grenoble, France

{florence.bertails, clement.menier, marie-paule.cani}@imag.fr



Figure 1: A dynamic hair without self-shadowing (left) and shaded with our algorithm (right). The 3D light-oriented map storing hair density and transmittance (middle). The whole simulation (including animation and our rendering) is running interactively on a standard CPU.

Abstract

This paper presents a new fast and accurate self-shadowing algorithm for animated hair. Our method is based on a 3D light-oriented density map, a novel structure that combines an optimized volumetric representation of hair with a light-oriented partition of space. Using this 3D map, accurate hair self-shadowing can be interactively processed (several frames per second for a full hairstyle) on a standard CPU. Beyond the fact that our application is independent of any graphics hardware (and thus portable), it can easily be parallelized for better performance. Our method is especially adapted to render animated hair since there is no geometry-based precomputation and since the density map can be used to optimize hair self-collisions. The approach has been validated on a dance motion sequence, for various hairstyles.

Key words: Hair self-shadowing, interactive rendering, hair simulation.

1 Introduction

Self-shadowing is of great relevance to the realistic appearance of hair as it contributes to the impression of volume (see Figure 1). Considering the high number of thin, translucent fibers composing a human hair, this phenomenon is difficult to reproduce both accurately and efficiently.

Our work was especially motivated by the need for a simple, fast and accurate technique to render animated

sequences involving dynamic hair. Recently, much effort has been made to achieve interactive frame rates in the simulation of dynamic hair [3, 22, 1]. But most of the time, these good performances only include the cost for animation while realistic hair rendering is done offline.

Approaches targeting interactive self-shadowing are very recent and mostly rely on advanced GPU's capabilities [17, 12]. Though successful, these methods are highly dependent on the hardware architecture, and remain difficult to implement. This paper investigates an alternative solution based on the CPU which turns out to be simpler to implement, more flexible, and which still yields interactive frame rates.

1.1 Previous Work

Realistic rendering of human hair requires the handling of both local and global hair properties. Local hair properties describe the way individual hair fibers are illuminated and then represented in the image space, whereas global properties define how the hair fibers interact together. Global hair properties especially include hair self-shadowing which plays a crucial role in volumetric hair appearance and which is the main focus of this paper.

Local Illumination

To render an individual hair strand, Kajiya and Kay earlier proposed a reflectance model [9] that has been widely used subsequently. Their model is composed of a lambertian diffuse component and an anisotropic specular com-

ponent. Many later approaches have subsequently proposed further refinements to this model [13, 2, 7, 10]. Recently, Marschner *et al.* [16] measured the scattering from real individual hair fibers and proposed a physical-based scattering model accounting for subtle scattering effects (such as multiple specular highlights) observed in their experiments. In our approach, we use Kajya-Kay’s reflectance model but our self-shadowing technique could be combined with any other local illumination model.

Self-Shadowing

Two main techniques are generally used to cast self-shadows into volumetric objects¹: shadow maps and ray casting through volumetric densities.

In basic depth-based shadow maps, the scene is rendered from the light point of view, and the depth of every visible surface is stored into a 2D *shadow map*. A point is shadowed if the distance between the point and its projection to the light’s camera is greater than the depth stored in the shadow map. This algorithm is not adapted to render semi-transparent objects such as hair because it only stores a single depth per pixel. To handle the self-shadowing of semi-transparent objects, Lokovic *et al.* proposed an extension to the traditional shadow maps: the *deep shadow map* [15]. For each pixel of the map, the method stores a *transmittance* function (also called *visibility function*) that gives the fraction of light penetrating at every sampled depth along a ray casted from the pixel.

Kim and Neumann proposed a practical implementation of this approach, called the *opacity shadow maps* [11], and applied it to hair rendering. In their method, the hair volume is uniformly sliced along the light rays and each hair volume comprised between two consecutive slices is rendered from the light’s point of view into the alpha buffer, leading to an opacity shadow map. Final rendering is done by interpolating the different opacity shadow maps. This method has been used together with hair representations at different LODs for the interactive animation of hair [23].

The opacity shadow maps technique was recently exploited by other authors [17, 12] to get a fast rendering by using recent GPU’s capabilities. Koster *et al.* achieved real-time results by accelerating the implementation of the opacity shadow maps and by making some assumptions about the geometry of hair. Mertens *et al.* used an adaptive clustering of hair fragments instead of an uniform slicing, which enabled them to interactively build a more accurate transmittance function.

Volume rendering is a common approach for visualizing datasets that come on 3D grids [8]. To render

semi-transparent volumetric objects with shadows, a first step usually consists of casting rays from the light source into the volume, and storing the light attenuation function (in voxels for instance). Then, actual rendering is done by ray tracing from the viewpoint. Such methods can accurately render both volumetric data such as clouds or smoke [8] as well as data with fine geometry such as hair [9]. Ray-tracing-based approaches often yield good quality results, but they are usually very expensive in terms of time and memory. More recently, splatting approaches have been used in order to achieve interactive shadowing and rendering of volumetric dataset [18, 25]. Billboard splatting has been successfully applied to the rendering of clouds [6]. In the case of hair, this method is still efficient but it does not seem to be really adapted to render the fine geometry of hair [1].

Ray tracing-based methods can often be very prohibitive in terms of rendering time, as they require the calculation and the sorting of multiple intersections between the rays and the objects that need to be shadowed. Conversely, the key benefit of the shadow map-based approaches is the light-oriented sampling of geometry, which makes the computation of accumulative transmittance straightforward. Actually, our method is inspired by both. Combining a volumetric representation of density with a light-oriented sampling allows us to define a practical and interactive self-shadowing algorithm.

1.2 Overview

Our goal is to provide an easy, accurate and efficient way of casting shadows inside hair. Our method has to be flexible enough to handle and accelerate simulations that involve animated hair.

Our main contribution is to propose a new algorithmic structure called *3D light-oriented shadow map* that is inspired by both traditional 3D density volumes and more recent 2D shadow maps as it combines an optimized volumetric representation of hair with a light-oriented partition of space. This voxel structure stores the light attenuation through the hair volume, and it is used to compute the final color attributed to each hair drawing primitive (hair segment for instance).

The main advantages of our method are the following:

- Our application is portable, simple to implement and can render a whole hairstyle composed of thousands of hair strands interactively on a standard CPU. Furthermore, it can easily be parallelized to increase performance.
- The approach is especially adapted to animated hair since the algorithmic structures that we used are efficiently updated at each time step. Moreover we

¹Please refer to [24] for a complete survey on shadowing methods.

show that our data structures provide an inexpensive way of processing hair self-collisions.

- Our technique does not make any assumption about the geometry of hair, and thus can be applied to render any hairstyle. It has been validated on various hairstyles, either static or animated with different kinds of motion.

Section 2 describes our 3D light-oriented shadow map structure. Section 3 explains how the self-shadowing process can be efficiently done by using this new structure. Section 4 deals with the two extensions of the method: on the one hand, we show that our 3D map is very helpful to process hair self-collisions efficiently; on the other hand we provide a parallelized version of our algorithm that improves the global performance of the simulation. The last two sections discuss results before concluding.

2 3D Light-Oriented Shadow Map

Our 3D shadow map is a uniform cubic voxel grid that associates to each voxel (or cell) a *density* value and a *transmittance* value.

The different hair models that we want to render are composed of a set of segments, but our algorithm could also apply to other kinds of geometry such as polygonal surfaces for example.

2.1 A Light-Oriented Local Frame

In our method, the light rays are assumed to be parallel (ie. coming from an infinitely distant source), which is a reasonable assumption for handling common lighting conditions like sun light. This point will be discussed in conclusion.

Instead of having a fixed-oriented structure like in previous approaches, our map is always aligned with the light direction. More precisely, the map is placed in a local frame $\mathcal{R} = (O, \mathbf{X}_{map}, \mathbf{Y}_{map}, \mathbf{Z}_{map})$ where $\mathbf{X}_{map}$ coincides with the normalized light vector $\mathbf{L}$ and O is the origin of the map (see Figure 3).

As we shall see in Section 3.2, this configuration is very helpful for computing the accumulated transparencies efficiently. Note that for non-animated data requiring “dynamic” lighting (ie. a moving light), this choice would not be appropriate since the material geometry is to be recomputed each time the light moves. But in our case, the geometry of hair needs to be updated at each time step, so the moving light case does not yield extra cost for us.

2.2 Object Space to Map Space

To occupy a limited memory, our data structure exploits the fact that during animation, the hair volume is always located inside a bounding box of constant spatial dimension. Indeed hair always remains attached to a scalp, and

hair strands are assumed to be inextensible. Storing hair elements can thus be done inside a bounded structure, provided we build a mapping function from the 3D object space to this 3D bounding space.

The spatial dimension of the map is thus fixed and only depends on the maximal length l_{max} of a hair strand. If the dimension of the map is superior or equal to $2 \times l_{max} + h_{max}$, where h_{max} is the maximal dimension of the head, it is ensured that the grid will always represent a bounding volume for the hair at any time step. Of course, the best choice for the dimension of the map is the minimal number satisfying the constraint above.

The size (or resolution) of the map (ie. the number of cells it contains) depends on the desired accuracy of self-shadowing. Some tests have been made in Section 5 to compare results using different map resolutions.

In the remainder of the paper, $NCELLS$ will denote the number of cells in each direction $\mathbf{X}_{map}$, $\mathbf{Y}_{map}$ and $\mathbf{Z}_{map}$ of the map frame $\mathcal{R}$, and ds will represent the step of the map, ie. the spatial dimension of a cell (see Figure 2).

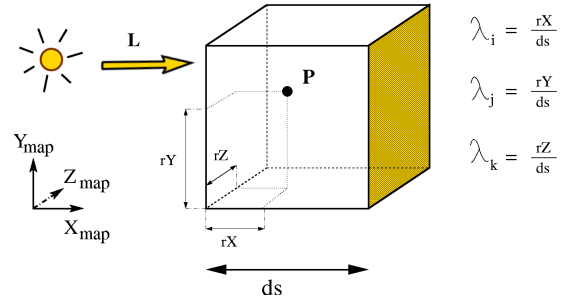


Figure 2: One cell of the map containing a point P . The λ_i parameters give the location of P inside the cell, and will be useful for the filtering process (see Section 3.3). By convention, each cell will store the quantity of light received by its back side (yellow side).

To find the index of the cell corresponding to a point $P(x, y, z)$, the coordinates of P are first expressed in the map frame $\mathcal{R}$ as $(x_{map}, y_{map}, z_{map})$, and the following mapping function then applied:

$$\Psi : \mathbb{R}^3 \longrightarrow [0..NCELLS]^3$$

$$\begin{bmatrix} x_{map} \\ y_{map} \\ z_{map} \end{bmatrix} \longmapsto \begin{bmatrix} \left\lfloor \frac{x_{map}}{ds} \right\rfloor \bmod NCELLS \\ \left\lfloor \frac{y_{map}}{ds} \right\rfloor \bmod NCELLS \\ \left\lfloor \frac{z_{map}}{ds} \right\rfloor \bmod NCELLS \end{bmatrix}$$

Figure 3 shows the mapping between the object space and the map space.

Thanks to the mapping function Ψ , access to elements of the map is done in constant time, which greatly contributes to the efficiency of the method.

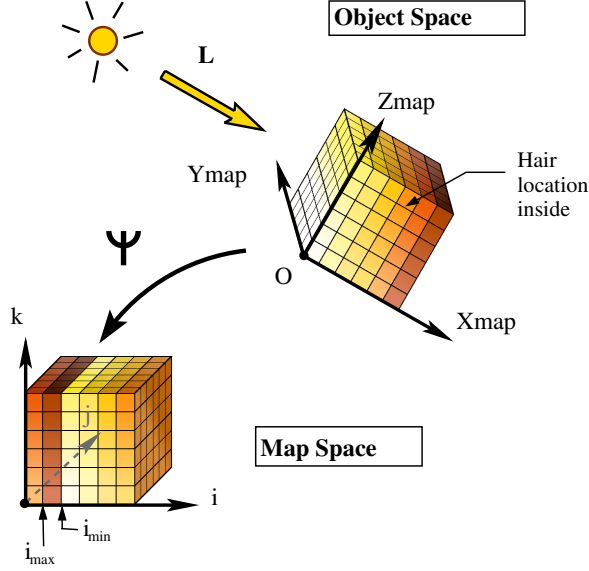


Figure 3: Correspondence between the object space and the map space. Because of the modulo operator in the mapping function Ψ , the first slice of the map (in light order) does not necessarily have the lowest index. The first slice and the last slice have consecutive indexes.

3 Self-Shadowing Algorithm

Our self-shadowing algorithm is composed of three main steps: hair density filling (1), transmittance computation (2), and filtering (3). Initially, each cell of the map has a null density (and we call it an *empty* cell).

The following figure summarizes the whole rendering pipeline².

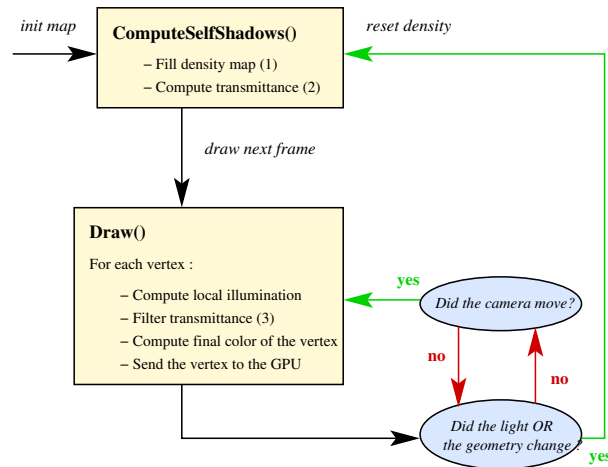


Figure 4: The rendering pipeline.

²In our case, each hair strand is drawn as an OpenGL line strip

3.1 Filling Hair Density into the Map

The first step of the algorithm consists of filling the map with hair density. This is simply done by traversing the geometry of hair and doing the following operations:

- Each hair strand s_i is sampled using a Catmull-Rom spline into $nSmooth$ points P_k^i ;
- For each point P_k^i , the density of the cell $\Psi(P_k^i)$ is incremented.

Of course the resulting density value obtained for one cell only makes sense relative to values of the other cells. Indeed, each isolated density value is arbitrary, and especially depends on the number of sample points used for each strand. Assuming that hair sampling is uniform, which is a reasonable assumption, the relative density multiplied by a scaling factor f approximates the light fall off through the corresponding cell. This quantity is commonly called the *extinction* parameter [15].

In practice, our hair sampling is the same as the one that is used at the final drawing stage, in order to ensure that each drawn vertex belongs to a non-empty cell.

3.2 Computing Transmittance

The fraction of light that penetrates to a point P of space can be written as [15]:

$$\tau(p) = \exp\left(-\int_0^l \kappa(l') dl'\right) \quad (1)$$

where l is the length of the path from the light to the point, and κ is the extinction function along the path.

The function τ is called the *transmittance* function. A sampled evaluation of τ can be done by accumulating transparencies of sampled regions along the light direction.

In our case, we need to evaluate the transmittance function at each cell of the map. To do this, we compute the transparency of each cell (i, j, k) as:

$$t(i, j, k) = \exp(-\kappa_{i,j,k} ds) \quad (2)$$

where the extinction coefficient $\kappa_{i,j,k}$ is computed using the density value of the cell (i, j, k) , as explained before in Section 3.1: $\kappa_{i,j,k} = f \times d_{i,j,k}$ where $d_{i,j,k}$ is the density of cell (i, j, k) and f is a scaling factor.

The transparencies are then composited together to get the final transmittance of each cell (i, j, k) :

$$Trans(i, j, k) = \prod_{i'=i_{min}}^i \exp(-d_{i',j,k} f ds) \quad (3)$$

where i_{min} is the index of the map slice that is the closest to the light (see Figure 3).

As we mentioned in the previous section, the novelty of our approach in comparison with previous algorithms using voxel grids is that cells are sorted along the light direction: accumulating transparencies then becomes straightforward:

- A transmittance parameter $prevTrans$ is first initialized to 1 which is the proper value for a transparent and fully illuminated cell;
- The column (j, k) is traversed, starting from slice i_{min} (the closest slice to the light) until slice i_{max} (the furthest slice):
 - If cell (i, j, k) is non-empty, its transmittance is set to $prevTrans \times \exp(-d_{i,j,k} f ds)$ (using Equation 3) and the parameter $prevTrans$ is updated to this value.
 - Otherwise cell (i, j, k) is given the transmittance $prevTrans$.

Note that some empty cells might also be in shadow, since filling densities into the map does not necessary yield a connective set of non-empty cells. Even if only vertices belonging to non-empty cells will be drawn, giving a proper transmittance value to empty cells is important because such cells could be involved in the filtering process, if a non-empty cell has empty neighbors (see next section). The algorithm described above guarantees that every cell of the map has a proper transmittance value.

3.3 Filtering and Composing Colors

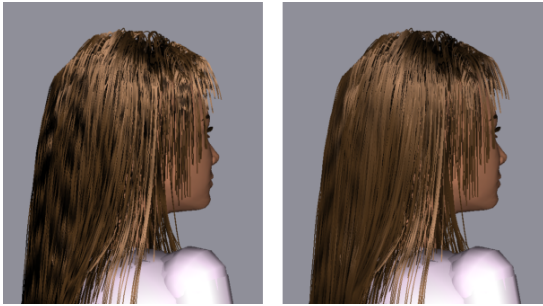


Figure 5: The effect of filtering the transmittance values. Self-shadows without filtering (left): regular patterns aligned with the map are visible. Self-shadows with filtering (right): artefacts have vanished, hair looks coherent.

Before drawing hair primitives, it is necessary to filter transmittance values, otherwise regular patterns aligned with the density map will be quite visible, as shown by Figure 5.

For each point P that has to be sent to the GPU for final drawing:

- We compute the relative position of P $(\lambda_i, \lambda_j, \lambda_k)$ with respect to its corresponding cell $\Psi(P)$ (see Figure 2).
- We compute filtered transmittance at point P by applying a trilinear interpolation as:

$$Trans^f(P) = \sum_{\substack{i' \in \{i-1, \dots, i\} \\ j' \in \{j-1, \dots, j\} \\ k' \in \{k-1, \dots, k\}}} A_{i'} A_{j'} A_{k'} Trans(i', j', k')$$

$$\text{where } A_{i'} = \begin{cases} \lambda_i & \text{if } i' = i \\ (1 - \lambda_i) & \text{otherwise} \end{cases}$$

(similar for $A_{j'}$ and $A_{k'}$)

- Finally, the color Φ_P of vertex P is obtained by the following equation:

$$\Phi_P = \Phi_{Ambient} + Trans^f(P) \times (\Phi_{Diffuse} + \Phi_{Specular}(P))$$

4 Extensions

4.1 Handling Hair Self-Collisions

Because of the high number of hair strands composing a human hairstyle, hair self-collisions represent a difficult and computationally expensive issue in hair animation. In practice, it often takes more than 80% of the simulation time [21].

An acceptable approximation of hair self-interaction consists of considering that internal collisions mainly result into the hair volume [14]. Starting from this assumption, hair density information is very useful: if the density is local over a fixed threshold (corresponding to maximum quantity of hair that can be contained in a cell), the hair strands should undergo constraints that spread them out.

Hair is animated using an approach closed to hair guidance methods [5, 4]. In our case, hair is composed of approximately a hundred wisps where each hair wisp is simulated through three guide hair strands. Each guide hair strand is animated using a fast rigid links simulation [19]. Final rendered hair strands are simply interpolated from the guide hair strands within each wisp.

Using the density map at each time step, hair self-collisions are processed by applying repulsive forces from the center of each cell having a too high density. Although this method is extremely simple, it yields convincing results. Furthermore, this is a very cheap way to handle hair self-collisions (it only takes 2.5% of the whole processing time). Please visit our website and watch our videos at <http://www.evasion.imag.fr/Publications/2005/BMC05a/>.

4.2 Parallelization of the Algorithm

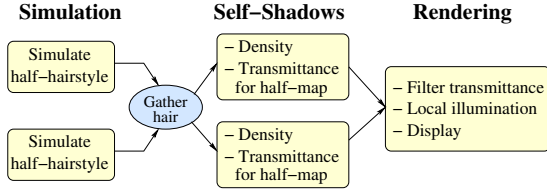


Figure 6: A parallel version of our algorithm.

One advantage of having a CPU-based algorithm is that parallelization can be considered in order to increase its efficiency. As a matter of fact, the described method is very well suited for such a technique. We present here the parallel implementation of the simulation and self-shadowing algorithms.

- **Simulation:** thanks to the use of the density-map for handling self-collisions, each hair wisp can be simulated independently. This allows for a straight-forward parallelization where each processor computes a part of the hair, gathering at the end their partial results.
- **Self-Shadowing:** here again a straight-forward parallelization can be applied thanks to the fact that the map is light-oriented. As described in Section 3.2, the calculations for each column (j, k) can be done independently.

We have tested this implementation on a standard PC cluster and were able, using 3 CPUs, to easily double the frame rate in comparison with the single processor results given in the next section.

When trying to use more CPUs, the network gathering and sending of the vertices to the GPU became the main bottleneck. Sending vertex arrays directly to the GPU should reduce this bottleneck.

5 Results and Discussion

Our algorithm has been applied both to static and dynamic hairstyles. In each case we compare it with existing methods in terms of quality and performance.

5.1 Rendering Static Hair

Figures 1 and 7 show that our self-shadowing algorithm produces good visual results for merely synthetic hairstyles as well as for hairstyles captured from real hair geometry. We can see in Figure 7 that self-shadows make volumetric wisps stand out, whereas no self-shadows flatten the hair.

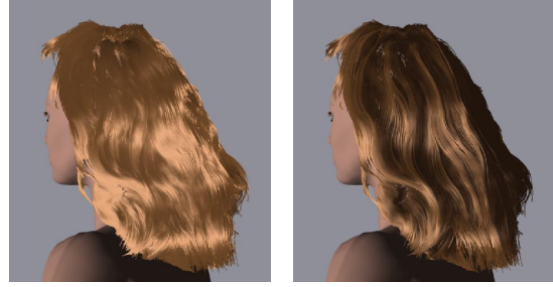


Figure 7: Applying our self-shadowing algorithm to a hairstyle captured from photographs by the method of Paris *et. al* [20]. The hairstyle is composed of 87,500 hair strands (1,123 K segments) and it took 2 seconds to render it.

Figure 8 shows results obtained on curly hair when using different map resolutions. We can notice that for fine resolutions (128×128 or 256×256), curly wisps are properly shadowed and their shape is thus clearly visible, which is not the case for the coarsest resolutions. In practice, we found that a 128×128 resolution was sufficient to account for small shape details of hair.

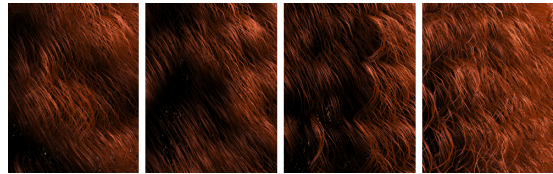


Figure 8: Evaluation of the quality of self-shadowing, using different map resolutions. From left to right: 32×32 with $ds = 0.5$; 64×64 with $ds = 0.2$, 128×128 with $ds = 0.1$ and 256×256 with $ds = 0.05$.

	Map reset + density	Trans	Filter + draw	Total rendering
Smooth	0.038	0.015	0.037	0.09
Curly	0.062	0.015	0.053	0.13

Table 1: Detailed performance of the rendering process (computing density, transmittance, filtering and final drawing) of a smooth hairstyle composed of 100K segments and a curly hairstyle composed of 200K segments. The results are expressed in seconds per frame; they have been obtained using an Intel P4 CPU at 3GHz.

In comparison with [17] our self-shadowing algorithm runs at a higher frame rate (11 FPS instead of 6 FPS for 100K hair segments).

5.2 Rendering Dynamic Hair

Figure 9 shows two snapshots from our hair animations. Our self-shadowing algorithm captures the fine discontinuities observed in real hair during motion, as illustrated in Figure 10.



Figure 9: A smooth brown hairstyle (100 K segments) and a curly red hairstyle (200 K segments) animated with different dance motions and interactively rendered with our algorithm.

	Anim	Hair self-collisions	Rendering	Total simu
Smooth	0.067	0.003	0.09	0.16
Curly	0.254	0.003	0.13	0.557

Table 2: Detailed performance of the simulation (animation, rendering and hair self-collisions) of two hairstyles composed of 134 animated wisps: a smooth hair style (100K rendered segments) and a curly hair style (200 K rendered segments). The results are expressed in **seconds per frame**; they have been obtained using an Intel P4 CPU at 3GHz.

Table 2 gives the detailed performance of the whole simulation, including animation, hair self-collisions and rendering for both smooth and curly hairstyles. Note that the animation time is not the same for the two hairstyles, because it includes the update and the smoothing of the interpolated hair strands.

A hair composed of 3350 hair strands and 100K segments is thus completely simulated at an interactive frame rate of 6 FPS. For aesthetic results, we have implemented hair-body collisions using a standard method based on spheres approximation. Handling such collisions makes the performance fall down to 3.5 FPS for the smooth hairstyle, and 1.5 FPS for the curly hairstyle, but no optimization has been developed yet for that specific problem, considering it was beyond the scope of this paper.

In our approach, the hair volume is properly generated using a repulsive force field based on local densities, as explained in 4.1. However, this method does not account

for hair anisotropy nor wisps interpenetration. This could be done by adding more information to the map, such as hair orientation.



Figure 10: A real shadowed hair (left) and our model (right) with similar lighting conditions.

6 Conclusion and Future Work

We have presented a new hair self-shadowing algorithm based on a 3D-light oriented density map. Our approach can interactively render various hairstyles composed of thousands of hair strands, and yields convincing results. Our algorithm can easily be parallelized to improve the performance. Furthermore, we have shown that our density map is very helpful in accelerating the simulation process, as it can be used to handle self-collisions in an inexpensive way with good visual results. We are planning to use the hair density information again to optimize hair-body collisions.

For simplicity purposes, our approach makes the assumption of an infinitely distant source, which could be a limitation for rendering scenes illuminated by punctual sources. Yet, it seems that we could easily handle the case of punctual sources by only changing our mapping function Ψ . Instead of considering an uniform square space partition, the new mapping function Ψ' should account for an angular space partition starting from the source point, and then sampled normally to the light rays.

Our method could also handle several light sources by simply adding as many light-oriented maps as sources. The final transmittance of a point P would have to be interpolated between the transmittance values obtained from the different sources.

To get a better precision in our computations for a low cost, an interesting idea would be to follow the same approach as Mertens *et. al* [17] who build an adaptive slicing along a light ray and thus get a better approximation of the visibility function than approaches using a uniform slicing.

Acknowledgements

This work was supported by L'Oréal Recherche. Thanks to Thanh Giang for proofreading the paper.

References

- [1] Y. Bando, B-Y. Chen, and T. Nishita. Animating hair with loosely connected particles. *Computer Graphics Forum*, 22(3):411–418, 2003. Proceedings of Eurographics’03.
- [2] D. Banks. Illumination in diverse codimensions. In *Proceedings of ACM SIGGRAPH’94*, Computer Graphics Proceedings, Annual Conference Series, pages 327–334, 1994.
- [3] F. Bertails, T-Y. Kim, M-P. Cani, and U. Neumann. Adaptive wisp tree - a multiresolution control structure for simulating dynamic clustering in hair motion. In *ACM SIGGRAPH Symposium on Computer Animation*, pages 207–213, July 2003.
- [4] J. Chang, J. Jin, and Y. Yu. A practical model for hair mutual interactions. In *ACM SIGGRAPH Symposium on Computer Animation*, pages 73–80, July 2002.
- [5] A. Daldegan, N. M. Thalmann, T. Kurihara, and D. Thalmann. An integrated system for modeling, animating and rendering hair. *Computer Graphics Forum*, 12(3):211–221, 1993.
- [6] Y. Dobashi, K. Kaneda, H. Yamashita, T. Okita, and T. Nishita. A simple, efficient method for realistic animation of clouds. In *SIGGRAPH ’00: Proceedings of the 27th annual conference on Computer graphics and interactive techniques*, pages 19–28. ACM Press/Addison-Wesley Publishing Co., 2000.
- [7] D. Goldman. Fake fur rendering. In *Proceedings of ACM SIGGRAPH’97*, pages 127–134, 1997.
- [8] J. Kajiya and B. Von Herzen. Ray tracing volume densities. In *Proceedings of the 11th annual conference on Computer graphics and interactive techniques*, pages 165–174. ACM Press, 1984.
- [9] J. Kajiya and T. Kay. Rendering fur with three dimensional textures. In *Proceedings of ACM SIGGRAPH’89*, Computer Graphics Proceedings, Annual Conference Series, pages 271–280, 1989.
- [10] T-Y. Kim. *Modeling, Rendering and Animating Human Hair*. PhD thesis, University of Southern California, 2002.
- [11] T-Y. Kim and U. Neumann. Opacity shadow maps. In *Rendering Techniques 2001*, Springer, pages 177–182, July 2001.
- [12] M. Koster and H-P. Seidel. Real-time rendering of human hair using programmable graphics hardware. In *Computer Graphics International (CGI)*, pages 248–256, June 2004.
- [13] A. M. LeBlanc, R. Turner, and D. Thalmann. Rendering hair using pixel blending and shadow buffers. *The Journal of Visualization and Computer Animation*, 2(3):92–97, – 1991.
- [14] D-W. Lee and H-S. Ko. Natural hairstyle modeling and animation. *Graphical Models*, 63(2):67–85, March 2001.
- [15] T. Lokovic and E. Veach. Deep shadow maps. In *Proceedings of the 27th annual conference on Computer graphics and interactive techniques*, pages 385–392. ACM Press/Addison-Wesley Publishing Co., 2000.
- [16] S. Marschner, H. Jensen, M. Cammarano, S. Worley, and P. Hanrahan. Light scattering from human hair fibers. *ACM Transactions on Graphics (Proceedings of the SIGGRAPH conference)*, 22(3):281–290, July 2003.
- [17] T. Mertens, J. Kautz, P. Bekaert, and F. Van Reeth. A self-shadow algorithm for dynamic hair using density clustering. In *Proceedings of Eurographics Symposium on Rendering*, 2004.
- [18] M. Nulkar and K. Mueller. Splatting with shadows. *Volume Graphics*, pages 35–50, 2001.
- [19] C. Van Overveld. An iterative approach to dynamic simulation of 3-D rigid-body motions for real-time interactive computer animation. *The Visual Computer*, 7:29–38, 1991.
- [20] S. Paris, H. Briceño, and F. Sillion. Capture of hair geometry from multiple images. *ACM Transactions on Graphics (Proceedings of the SIGGRAPH conference)*, 2004.
- [21] E. Plante, M-P. Cani, and P. Poulin. Capturing the complexity of hair motion. *Graphical Models (Academic press)*, 64(1):40–58, January 2002.
- [22] K. Ward and M. Lin. Adaptive grouping and subdivision for simulating hair dynamics. In *Proceedings of Pacific Graphics’03*, September 2003.
- [23] K. Ward, M. Lin, J. Lee, S. Fisher, and D. Macri. Modeling hair using level-of-detail representations. In *International Conference on Computer Animation and Social Agents (CASA)*, May 2003.
- [24] A. Woo, P. Poulin, and A. Fournier. A survey of shadow algorithms. *IEEE Computer Graphics and Applications*, 10(6):13–32, 1990.
- [25] C. Zhang and R. Crawfis. Shadows and soft shadows with participating media using splatting. *IEEE Transactions on Visualization and Computer Graphics*, 9(2):139–149, 2003.