



An Efficient and Practical Near and Far Field Fur Reflectance Model

LING-QI YAN, University of California, Berkeley

HENRIK WANN JENSEN, University of California, San Diego

RAVI RAMAMOORTHY, University of California, San Diego

Physically-based fur rendering is difficult. Recently, structural differences between hair and fur fibers have been revealed by Yan et al. (2015), who showed that fur fibers have an inner scattering medulla, and developed a double cylinder model. However, fur rendering is still complicated due to the complex scattering paths through the medulla. We develop a number of optimizations that improve efficiency and generality without compromising accuracy, leading to a practical fur reflectance model. We also propose a key contribution to support both near and far-field rendering, and allow smooth transitions between them.

Specifically, we derive a compact BCSDf model for fur reflectance with only 5 lobes. Our model unifies hair and fur rendering, making it easy to implement within standard hair rendering software, since we keep the traditional R , TT , and TRT lobes in hair, and only add two extensions to scattered lobes, TT^s and TRT^s . Moreover, we introduce a compression scheme using tensor decomposition to dramatically reduce the precomputed data storage for scattered lobes to only 150 KB, with minimal loss of accuracy. By exploiting piecewise analytic integration, our method further enables a multi-scale rendering scheme that transitions between near and far field rendering smoothly and efficiently for the first time, leading to 6 – 8 $\times$ speed up over previous work.

CCS Concepts: • **Computing methodologies** → **Reflectance modeling**;

Additional Key Words and Phrases: fur, multi-scale, reflectance

ACM Reference format:

Ling-Qi Yan, Henrik Wann Jensen, and Ravi Ramamoorthy. 2017. An Efficient and Practical Near and Far Field Fur Reflectance Model. *ACM Trans. Graph.* 36, 4, Article 67 (July 2017), 13 pages.

DOI: <http://dx.doi.org/10.1145/3072959.3073600>

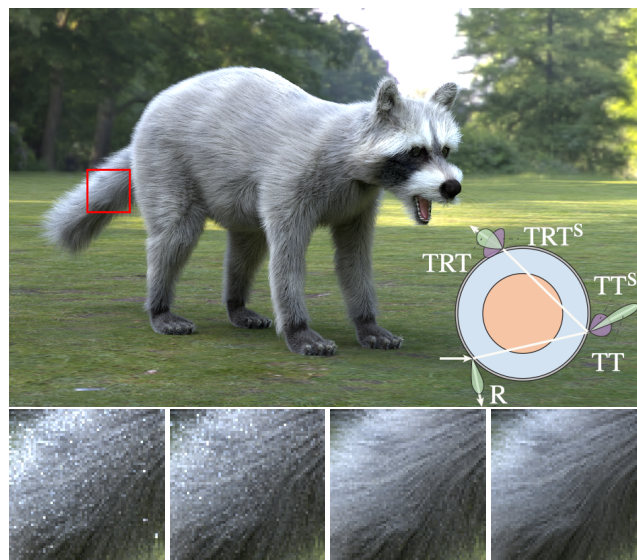
1 INTRODUCTION

Recently, computer-generated virtual animal characters are increasingly used in both films and video games. However, photo-realistic fur rendering is still a problem due to the complexity of light transport within a fur fiber. Recently, Yan et al. (2015) have developed a comprehensive physically-based model for fur reflectance, revealing significant structural differences between hair and fur fibers. Scattering within the central region of a fur fiber, known as the *medulla*, results in complicated light paths. Moreover, this complexity directly leads to significant pre-computation, and limits fur rendering to be near-field only. Even for hair rendering, efficient far field integration schemes are lacking. State of the art methods either assume that the azimuthal section of hair fibers are perfectly smooth (Marschner et al. 2003) so that the far-field integration can be solved, or resort

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than ACM must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.

© 2017 ACM. 0730-0301/2017/7-ART67 \$15.00

DOI: <http://dx.doi.org/10.1145/3072959.3073600>



(Yan et al. 2015) Ours (near field) Ours (multi-scale) (Yan et al. 2015)
ES, 256spp, 11.2s 256spp, 7.4s 256spp, 9.2s EQ, 1800spp, 75.8s

Fig. 1. (Top row) A rendering of a raccoon model using our practical 5-lobe reflectance model (illustrated bottom right) and our multi-scale rendering scheme with 1K samples per pixel. (Bottom row) Insets rendered using different methods. We use 256 samples per pixel for equal sample (ES) comparison, and show equal quality (EQ) comparison with Yan et al. (2015). Our multi-scale model performs more than 8 $\times$ faster for equal quality while being practical and efficient. Even our near field model has significantly less noise than previous work.

to numerical integrations as well as pre-computation (d'Eon et al. 2011).

Motivated by these observations, we aim to improve the efficiency and practicality of fur rendering, and provide a reflectance model that is simple to implement in modern rendering systems. Specifically, we describe a near field fur reflectance model in Sec. 4, focusing on simplicity and accuracy as compared to Yan et al. (2015). In Sec. 5, we illustrate how our near-field model integrates to far-field, and propose a novel multi-scale rendering scheme, focusing on efficiency. Overall, our major contributions are:

Simple reflectance model: Our local illumination model builds upon the double cylinder model representing the cuticle-cortex-medulla structure of fur fibers (Fig. 3 (b)). We *unify the cortex and the medulla's indices of refraction (IORs)*, removing most of the complicated types of light interactions between them. This simplification finally results in only 5 lobes in our model (Fig. 3 (c)), compared to 11 lobes previously. In particular, we keep the R , TT , and TRT lobes in hair models (with intensities modified slightly because of attenuation by the medulla) and only add two new scattered lobes, TT^s and

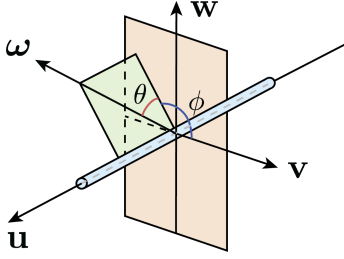


Fig. 2. Longitudinal-azimuthal parameterization for hair/fur fibers. A direction ω is parameterized into θ in the plane spanned by ω and the cylinder axis u , and ϕ orthogonal to the plane.

TRT^s. In this way, our reflectance model (BCSDF) is fully analytic, as opposed to Yan et al. (2015) which requires implicit ray tracing. The simplicity of our model also benefits importance sampling (Sec. 6), leading to faster convergence (Figs. 15 and 16). Our BCSDF can be included in existing hair rendering systems with minimal extra effort, and unifies hair and fur rendering. We are motivated by observations in the literature on fur fibers (Fig. 6), and we verify this simplification by the accuracy of fits to real measurements (Fig. 11).

Improved accuracy and practicality: We improve the accuracy of the model by considering different roughnesses of azimuthal and longitudinal sections, as has been observed in previous studies (Chiang et al. 2016; d'Eon et al. 2011). Moreover, we find that the medulla is not purely scattering, and can also absorb light (Carrlee and Horelick 2011). With a more general model taking these factors into consideration, our error for fitting measured data is lower than in Yan et al. (2015) in most cases, even with the simplification of IORs discussed above, as shown in Fig. 11.

We also observed that the precomputed data for medulla scattering provided by Yan et al. (2015), essentially a set of 4D tables, is low rank. We exploit tensor decomposition, the high dimensional analogue to 2D singular value decomposition (SVD), to compress it to as low as 150 KB, compared to more than 600 MB originally.

Analytic near/far field solution: Our model starts with analytic near-field solutions, as opposed to the implicit ray tracing required previously. Then we integrate our near-field model over the azimuthal section, by partitioning the range of integration into a few (< 5) segments. Finally, we analytically integrate for each segment using piecewise linear approximation. Moreover, we show how our analytic integration transitions between near and far field fur rendering, enabling multi-scale rendering for the first time (Fig. 1). This is especially useful when a pixel covers a small range over the azimuthal section. Our multi-scale rendering benefits hair rendering as well, as shown in Fig. 18.

Significant speed up: Due to the simplicity of our reflectance model and the efficiency of our analytic integration, compared with Yan et al. (2015), we achieve a $6 - 8\times$ speed up in generating equal quality results. We show more results and comparisons in Sec. 7 as well as in the accompanying video.

2 RELATED WORK

There is a large body of work on reflectance models for hair. Here we only discuss physically-based methods. Then we analyze the double cylinder fur reflectance model, and its limitations.

Hair reflectance models: Marschner et al. (2003) proposed the initial physically-based hair reflectance model. They use a longitudinal-azimuthal parameterization of hair fibers (Fig. 2), approximating them as rough dielectric cylinders (Fig. 3 (a)). Their model has three lobes: *R*, *TT* and *TRT*, where *R* and *TT* stand for reflection and transmission, respectively. The Marschner model has an analytic solution for evaluating the azimuthal lobes, but is based on the assumption that the azimuthal sections of hair fibers are perfectly smooth, which is not physically correct. d'Eon et al. (2011) extended the Marschner model to account for azimuthal roughness. However, their computational cost of evaluation is significant, using Gaussian quadrature and Taylor expansion. Chiang et al. (2016) adopted a near-field formulation by considering accurate incident positions azimuthally, leaving expensive integration across the entire azimuthal section to the renderer. In summary, apart from the original Marschner model that is not physically correct, there are no efficient closed-form solutions for azimuthal integration. In contrast, our method supports azimuthal roughness, and is analytic and efficient for multi-scale hair and fur rendering.

Double cylinder fur reflectance model: Yan et al. (2015) proposed a double cylinder model for physically accurate fur reflectance based on anatomical literature and measurements, in which the cuticle, the cortex and the medulla are accurately modeled (Fig. 3 (b)). The cuticle is layered so that its reflectance can be adjusted. The cortex only absorbs light, while the medulla only scatters. Both the cortex and the medulla have their own indices of refraction. However, this model is complicated, considering all kinds of light interactions with the inner cylinder such as *TrT* and *TrRrT* (capitalized and uncanceled letters represent interactions with the outer and inner cylinder, respectively). To incorporate the medulla's complex scattering effects, Yan et al. precompute for all possible incident positions or directions of medulla scattering of different kinds of fur fibers. The precomputed data puts a heavy burden on both storage and memory, limiting the model's practicality. Instead, we develop a simplified fur reflectance model that allows fast analytic integrations, but with only 5 lobes to consider, and requiring negligible precomputed data.

Near field scattering and far field approximation: Far-field approximation (d'Eon et al. 2011; Kajiya and Kay 1989; Marschner et al. 2003) assumes that hair fibers are very thin, usually thinner than a pixel. Hence, the accurate incident position on the azimuthal section is not important, compared to the integral over it. Thus, they always assume collimated incident light covering the width of a hair fiber. Since the positional information is lost, far field approximation produces a flat appearance when viewed close up so that a hair fiber covers more than one pixel (Fig. 5). Zinke et al. (2007) introduced near-field scattering by considering accurate azimuthal incident positions, and mathematically revealed the relationship between near and far field scattering. Yan et al. (2015) is also based on near-field scattering. To perform azimuthal integrations when far-field approximation is needed, they refer to the same numerical integration technique that is used in (d'Eon et al. 2011). In contrast, we propose an accurate and analytic method for both near-field and far-field rendering, benefiting both hair and fur models.

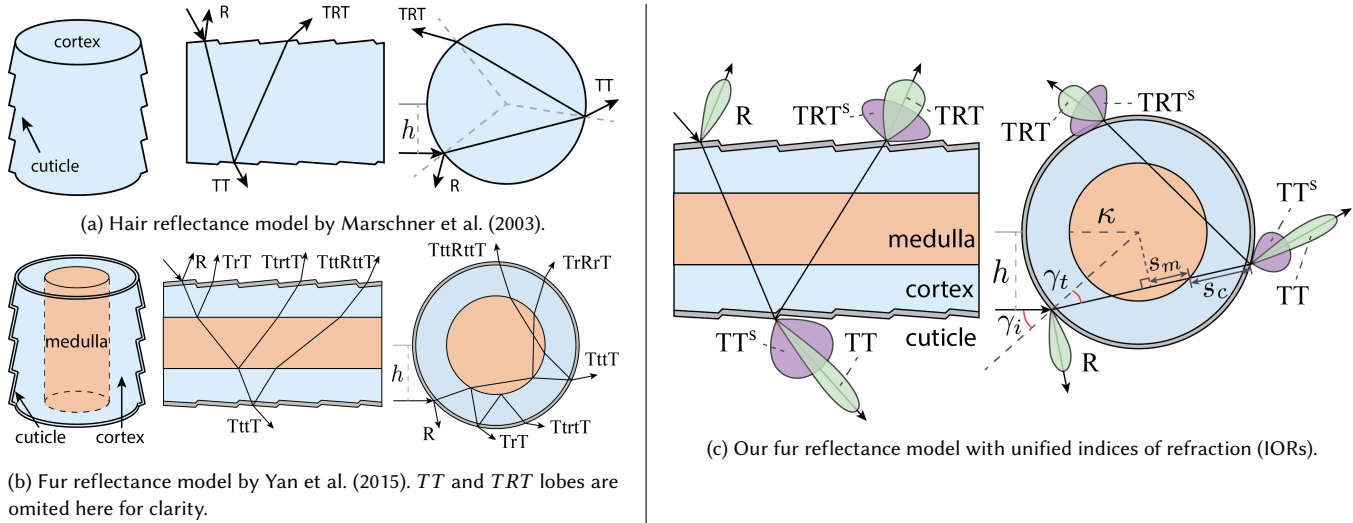


Fig. 3. Illustration of various reflectance models for hair and fur. All the models are illustrated using the factored representation longitudinally and azimuthally from left to right, and the differences between hair and fur structures can be found in the cylinders in the leftmost column. Note that light paths in our model don't refract as they go through the medulla. This is the same with (a), but much simpler than (b).

3 BACKGROUND

Most hair/fur reflectance models represent individual hair/fur fibers as cylinders. Similar to BRDFs, they use BCSDFs (Bidirectional Curve Scattering Distribution Function) to characterize how hair/fur fibers scatter light,

$$L_r(\omega_r) = \int L_i(\omega_i) S(\omega_i, \omega_r) \cos \theta_i d\omega_i, \quad (1)$$

where ω_i and ω_r are the incident and outgoing directions, and L_i and L_r are the incoming and outgoing radiance. S is the BCSDF.

We use the longitudinal-azimuthal parameterization by Marschner et al. (2003) as shown in Fig. 2,

$$L_r(\theta_r, \phi_r) = \int_{-\pi}^{\pi} \int_{-\frac{\pi}{2}}^{\frac{\pi}{2}} L_i(\theta_i, \phi_i) S(\theta_i, \theta_r, \phi_i, \phi_r) \cos^2 \theta_i d\theta_i d\phi_i. \quad (2)$$

Hair models. The Marschner hair model factors the BCSDF S into a product of M and N profiles, representing longitudinal and azimuthal events separately (Fig. 3 (a)). It takes three types of light paths $p \in R, TT, TRT$ into consideration, where R and T stand for reflection and transmission respectively. Its BCSDF is

$$\begin{aligned} S(\theta_i, \theta_r, \phi_i, \phi_r) &= \sum_p S_p(\theta_i, \theta_r, \phi_i, \phi_r) / \cos^2 \theta_d \\ &= \sum_p M_p(\theta_h) \cdot N_p(\phi; \eta') / \cos^2 \theta_d, \end{aligned} \quad (3)$$

where $\theta_h = (\theta_r + \theta_i)/2$ and $\theta_d = (\theta_r - \theta_i)/2$ are the longitudinal half angle and difference angle, $\phi = \phi_r - \phi_i$ is the relative exiting azimuth, and $\eta' = \sqrt{\eta^2 - \sin^2 \theta_d} / \cos \theta_d$ is the cortex's virtual index of refraction (IOR), treating elliptical azimuthal sections due to inclined longitudinal incident directions as if still being circular but with changed IOR.

The Marschner model assumes that the azimuthal section is perfectly smooth, thus leading to an analytic solution. However, azimuthal roughness has been demonstrated to be important (Chiang

et al. 2016; d'Eon et al. 2011) and should be taken into account for physical accuracy. Unfortunately, with the additional complexity, these models with rough azimuthal sections are no longer analytical for far field approximation.

Fur model. Yan et al. (2015) treat each fur fiber as two concentric cylinders. The outer cylinder is similar to the Marschner model, except that its reflectance can vary with a cuticle layers parameter l , adjusting the ratio between the reflected and refracted light. The inner cylinder represents the medulla. It doesn't absorb light, but scatters light when light travels inside. Between these two cylinders is the cortex, which simply absorbs light. The IORs of the two cylinders need not be the same. Figure 3 (b) illustrates the double cylinder model.

With the double cylinder model, types of paths such as TrT , $TrRrT$, $TtrT$, $TtrtT$ and $TttRttT$ are introduced. Furthermore, due to the scattering property of the medulla, light paths that go through the medulla generate both unscattered and scattered lobes when they exit it. The double cylinder BCSDF model can be written as:

$$\begin{aligned} S(\theta_i, \theta_r, \phi_i, \phi_r, h) &= \frac{\sum_p M_p^u(\theta_i, \theta_r) N_p^u(h, \phi)}{\cos^2 \theta_i} \\ &\quad + M^s(\theta_i, \theta_r, \phi) \frac{\sum_p N_p^s(h, \phi)}{\cos^2 \theta_i}, \end{aligned} \quad (4)$$

where M and N still represent longitudinal and azimuthal lobes, but with superscripts u and s specifying whether they are unscattered or scattered. p includes these complex types of paths mentioned above, as well as classic R, TT and TRT paths from the Marschner model. h is the azimuthal offset specifying the incident position in the azimuthal plane, as will be described in more detail at the end of the section.

The unscattered lobes are represented using Gaussian distributions both longitudinally and azimuthally. The longitudinal lobes M_p are Gaussians around the reflected longitudinal angle:

$$M_p(\theta_i, \theta_r) = G(\theta_r; -\theta_i + \alpha_p, \beta_p), \quad (5)$$

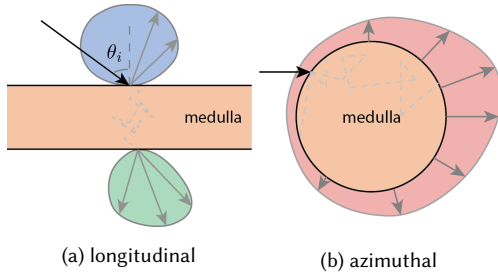


Fig. 4. Illustration of pre-computation. (a) Longitudinal. (b) Azimuthal.

where α_p is the lobe shift due to the cuticle tilt, and β_p is the roughness of each lobe. Both are empirically given.

The azimuthal unscattered lobes N_p have two parts: the attenuation term A_p and the distribution term D_p . A_p is the attenuation along path p azimuthally, with the incident position specified using the azimuthal offset h . It considers the Fresnel reflection/refraction at each intersection, and the attenuation along segments within the cortex. $D_p = G(\Phi_p(h) - \phi; \beta_p)$ is a Gaussian that describes how the attenuated energy distributes. It is centered at the exiting direction $\Phi_p(h)$ along path p , and it accumulates its variance using the surface roughness at each intersection. Thus, the azimuthal unscattered lobes can be written as

$$N_p(h, \phi) = A_p(h) \cdot D_p(h, \phi), \quad (6)$$

$$D_p = G(\Phi_p(h) - \phi; \beta_p). \quad (7)$$

The scattered lobes are queried from pre-computed scattering profiles. As shown in Fig. 4, Yan et al. (2015) pre-compute the medulla's scattering profiles by enumerating every longitudinal incident angle θ_i and every azimuthal offset h , simulating the medulla's scattering for different scattering coefficients σ_s and anisotropy g , and recording the scattered energy for every outgoing direction. This results in two 4D tables C^M and C^N , with storage of more than 600 MB.

Since the longitudinal scattering profile records both backward and forward lobes, i.e. longitudinal scattered lobes at $\phi = 0$ and $\phi = \pi$ (blue and green in Fig. 4), any longitudinal scattered lobe can be linearly interpolated between them, independent of path p :

$$M^S(\theta_i, \theta_r, \phi) = \mu F_t \cdot \text{lerp}_\phi(C_{\phi=0}^M(\theta_i', \theta_r'), C_{\phi=\pi}^M(\theta_i', \theta_r')), \quad (8)$$

where F_t are Fresnel transmissions through the cuticle of the incident and outgoing light, θ_i' and θ_r' are angles that enter the medulla and μ is a normalization factor.

The azimuthal scattered lobes are queried similar to the unscattered lobes as

$$N_p^S(h, \phi) = A_p^S(h) \cdot D_p^S(h, \phi). \quad (9)$$

The attenuation term A_p^S records the path p with Fresnel terms and absorptions until it goes into the medulla with direction $\Phi_p^S(h)$. Then the distribution of the scattered lobe is queried as

$$D_p^S(h, \phi) = C^N(\Phi_p^S(h) - \phi). \quad (10)$$

The scattered lobe is then attenuated again as it exits through the cortex. When it reaches the inner side of the cuticle, it either refracts through, and is thus attenuated by another Fresnel transmission, or reflects back, producing a diffuse lobe.

The biggest limitation of the reflectance model by Yan et al. (2015) is the complexity. First, it has too many lobes, 11 in total. Second,

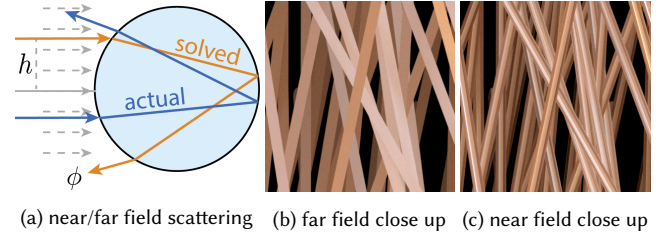


Fig. 5. (a) Illustration of near/far field scattering. Far field scattering always assumes collimated incident light covering the azimuthal section. Assuming the azimuthal section is perfectly smooth (Marschner et al. 2003), it is possible to solve for certain azimuthal offset h , given the relative exiting azimuth ϕ . Near field scattering considers actual incident h and calculates how much it contributes to ϕ . (b) A close-up view of fur with far field scattering, note the flat ribbon-like appearance. (c) A close-up view of fur with near field scattering rendered using our model in Sec. 4, note the cylindrical appearance and the clearly visible medulla inside.

it needs to store significant pre-computed data. Third, interactions between the scattered lobes and the cuticle are complicated. And perhaps most critically, a path p needs to be traced from the starting offset h until it exits the double cylinder, to find all the intersections and segments along p . Although tracing rays against circles is implicit, i.e. intersections can be solved, it is always required to compute the entire path p explicitly. So, there is no analytic BCSDf.

Note that the azimuthal scattering profile from the double cylinder model depends on specific azimuthal offset h , which means that it is limited to near field. We will describe near and far field models next.

Near/far field scattering. When an incident ray hits a hair/fur fiber, near-field scattering considers the actual incident position, thus an actual azimuthal offset $h \in [-1, 1]$ is specified in the azimuthal plane. In contrast, far-field approximation always assumes that the incident light is a collimated beam covering the entire azimuthal section. Figure 5 illustrates the difference.

Far field approximation produces ribbon-like flat appearance when viewed close, but becomes accurate when viewed far away. In this case, the width of a hair/fur fiber is so thin that distinguishing individual h is unnecessary. In fact, the azimuthal scattering function N_p from far field approximation can be found by performing multiple near field evaluations, i.e. integrating near field results over all possible offsets h ,

$$N_p(\phi) = \frac{1}{2} \int_{-1}^1 N_p(h, \phi) dh. \quad (11)$$

Far field approximation is useful, especially for fur that is usually near an order of magnitude thinner than human hair. Besides, far field approximation is expected to perform faster than near field scattering. This is because near field scattering essentially leaves the integration problem to the renderer, leading to an inefficient point-sampled far field model. However, as mentioned earlier, fast analytic integration of Eqn. 11 is already difficult for human hair, let alone for fur fibers with even more complex types of paths.

4 SIMPLIFIED NEAR FIELD BCSDf MODEL

In this section, we propose our near field BCSDf model for fur reflectance, focusing on its simplicity and generality. We validate

Parameter	Definition
η	refractive index of cortex and medulla
κ	medullary index (rel. radius length)
α	scale tilt for cuticle
β_m	longitudinal roughness of cuticle (stdev.)
β_n	azimuthal roughness of cuticle (stdev.)
$\sigma_{c,a}$	absorption coefficient in cortex
$\sigma_{m,s}$	scattering coefficient in medulla
$\sigma_{m,a}$	absorption coefficient in medulla
g	anisotropy factor of scattering in medulla
l	layers of cuticle

Table 1. Parameters used in our BCSDf model.

(a) Yan et al. (2015)
 $\eta_c = 1.45, \eta_m = 1.0$ (b) ours
 $\eta = 1.45$

(c) Microscopic photo

Fig. 6. Setting the medulla's index of refraction different from the cortex results in dark edges at the interface between them (a) for a fur fiber of polar bear lit from behind with directional lighting (1.45 for the cortex and 1.0 for medulla). This deviates from the photometric ground truth, even though polar bears' medullas are filled with air inside complex structures (c). Our local illumination model with unified IOR (1.45) does not have this problem (b). Microscopic photo courtesy of Carrlee et al. (2011).

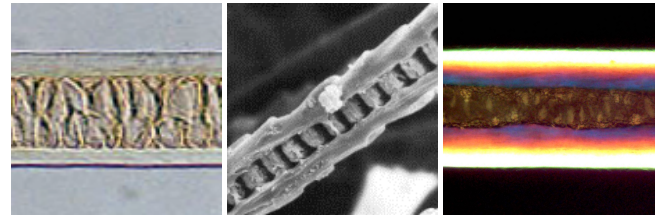
our simplified model against measured data in Fig. 11 and Table 4. Section 5 discusses an efficient piecewise analytic BCSDf model that unifies near and far field rendering.

4.1 Overview

Our first observation is that, different indices of refraction (IORs) of the cortex and the medulla are mostly responsible for complex paths and lobes. However, in most of the fitting results presented in Yan et al. (2015), the IORs between the cortex and the medulla are close. The similarity indicates that, complex types of paths such as TrT and $TtrtT$ are often too weak to be observed. Furthermore, Fig. 6 shows a rare and extreme case where the IORs are clearly different. In this case, the model by Yan et al. (2015) with different IORs of the cortex and the medulla is supposed to be more accurate, but still fails to match the ground truth.

Based on this observation, we unify the IORs for the cortex and the medulla. Despite the assumption, this leads to a much simpler model, an analytic solution, and comparably accurate results as Yan et al. (2015). With unified IORs, the light path no longer reflects or changes direction at the interface between the cortex and the medulla. Our model now shares the same light paths from hair models¹ — R , TT and TRT . The only difference is that, TT and TRT paths can be scattered passing through the medulla, forming scattered lobes TT^s and TRT^s . Thus, we write our near field BCSDf

¹We also need to consider attenuation by absorption in the medulla, in the formula for the R , TT and TRT lobes, as given in Table 3.



(a) Mounted medulla

(b) Thick medulla

(c) Pigmented medulla

Fig. 7. The medulla absorbs light due to both its complex structure and the pigments inside. (a) Photograph of a fur fiber with the medulla filled with a mounting medium to minimize medulla's scattering. We can still see the structure of the medulla. (b) Solid ladder-like medulla which is comparably thick as the cortex. (c) Pigments are found filling brown bear's medulla, as reported by Carrlee et al. (2011).



(a) All

(b) R (c) TT (d) TRT (e) TT^s (f) TRT^s

Fig. 8. (a) A rendering of a lock of hair with medulla. (b-f) With the unification of IORs, our reflectance model has only 5 lobes. From left to right, results rendered using lobe R , TT , TRT , TT^s and TRT^s , with path traced global illumination (Sec. 6).

as:

$$S(\theta_i, \theta_r, \phi_i, \phi_r, h) = \frac{S_R + S_{TT} + S_{TRT} + S_{TT}^s + S_{TRT}^s}{\cos^2 \theta_i}, \quad (12)$$

where $S_p = M_p(\theta_i, \theta_r)N_p(h, \phi)$ represent unscattered lobes and $S_p^s = M^s(\theta_i, \theta_r)N_p^s(h, \phi)$ represent scattered lobes. p is from the set of reduced types of paths $\{R = 0, TT = 1, TRT = 2\}$. S_R^s is always zero.

Apart from unifying the IORs, we improve Yan et al. (2015) by introducing the medulla's absorption and different longitudinal/azimuthal roughness. Figure 7 demonstrates that the medulla can absorb light because of its complex internal structure, and that there are cases where pigments are found within the medulla (Carrlee and Horelick 2011). We also take different longitudinal/azimuthal roughness into account, since these differences are often observed (GalatÅnjk et al. 2011) and used in practice (Chiang et al. 2016; d'Eon et al. 2011).

Table 1 lists all the parameters used in our model, Fig. 3 (c) illustrates all the lobes in our BCSDf model, and Fig. 8 shows decomposed renderings using each lobe. Except for the R lobe that does not pass through the cortex, all other lobes produce colored appearance. The medulla blocks most TT light paths from previous hair models, thus producing a dark TT lobe and bright TT^s lobe. Though much simpler with only 5 lobes, our model is slightly more

Lobe p	Shift α_p	Roughness β_p
R	α	β_m
TT	$-\alpha/2$	$\beta_m/2$
TRT	$-3\alpha/2$	$3\beta_m/2$
Lobe p	Distribution M^s	
TT^s	$\text{lerp}_\phi(C_{\phi=0}^M, C_{\phi=\pi}^M)$ (Eqn. 19)	
TRT^s	$\text{lerp}_\phi(C_{\phi=0}^M, C_{\phi=\pi}^M)$ (Eqn. 19)	

Table 2. Shift and roughness of longitudinal lobes.

accurate than Yan et al. (2015), as validated in Table 4. Furthermore, individual lobes in our model are more efficient to evaluate, replacing more costly implicit ray tracing with closed-form expressions (Sec. 4.2, Fig. 14). The pre-computed data for scattering lobes can be accurately compressed (Sec. 4.3) to make storage negligible (much less than a megabyte in all).

4.2 Unscattered lobes (R , TT , TRT)

Since the light paths in our model no longer deviate from those in hair models, our model unifies hair and fur rendering (Figs. 8 and 18). Thus, the unscattered lobes are very similar to the Marschner model.

Longitudinal unscattered lobes. As in previous work, we approximate the longitudinal scattering profiles for each lobe using an empirical Gaussian distribution $M_p(\theta_i, \theta_r) = G(\theta_r; -\theta_i + \alpha_p, \beta_p)$ (Eqn. 5). Their centers and variances are listed in Table 2.

Similarly, the azimuthal unscattered lobes are evaluated using $N_p(h, \phi) = A_p(h) \cdot D_p(h, \phi)$ (Eqn. 6, Table 3). However, with the unification of IORs, we are able to derive closed-form expressions for both the attenuation term and the distribution term, rather than needing to use implicit ray tracing.

Azimuthal attenuation. As shown in Fig. 3 (c), the R lobe ($p = 0$) will be reflected by the cuticle directly, so it is attenuated by the Fresnel reflection $F(\eta', \gamma_i)$. The TT ($p = 1$) and TRT ($p = 2$) lobes both refract through the cuticle twice, thus attenuated by two Fresnel transmissions, i.e. $(1 - F)^2$. And the TRT lobe has an additional internal Fresnel reflection. Besides, both TT and TRT are attenuated traveling through the cortex and possibly the medulla, which are exponential falloffs with the distances $2s_c$ and $2s_m$ traveled within the cortex and the medulla, respectively. This is how colors are introduced. We list the attenuation terms for individual unscattered lobes R , TT and TRT in Table 3, and give a general representation² as

$$A_0(h) = F, \quad (13)$$

$$A_p(h) = (1 - F)^2 F^{p-1} T_c T_m \quad p \geq 1, \quad (14)$$

where $F = F(\eta', \gamma_i, l)$ is the Fresnel term with respect to cuticle layers l as in previous work, $T_c = \exp(-2ps_c \sigma_{c,a} / \cos \theta_d)$ and $T_m = \exp(-2ps_m (\sigma_{m,a} + \sigma_{m,s}) / \cos \theta_d)$ are the attenuations within the cortex and medulla, respectively. The division by $\cos \theta_d$ is to account for elongated azimuthal paths when viewed from an oblique

²These general representations for all our attenuation and distribution computations also hold for arbitrary higher-ordered lobes, such as $TRRT$ and $TRRT^s$. However, we ignore them in our renderings for simplicity.

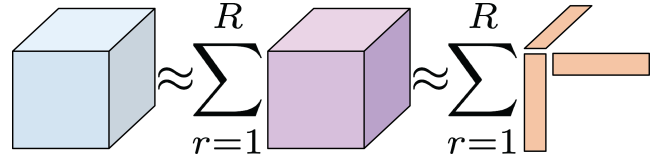


Fig. 9. Illustration of tensor decomposition.

longitudinal angle. Here, different azimuthal offsets h decide different Fresnel terms F , as well as distances s_c and s_m a path travels, and thus the attenuation terms T_c and T_m .

Azimuthal distribution. For the distribution term D_p of unscattered lobes, since they are Gaussians $G(\Phi_p(h) - \phi; \beta_p)$ (Eqn. 7), what we need are their centers Φ_p and variance β_p . To find their centers, we follow the corresponding light paths, performing mirror reflections or refractions at intersections until the path leaves the double cylinder, as shown in Fig. 3. In this way, the exiting azimuth Φ_p can be calculated, such that each refraction makes the path deviate its direction by $\gamma_t - \gamma_i$, and each internal reflection introduces a deviation of $\pi + 2\gamma_t$. For the distribution's variance, accumulating the squared roughness β_n^2 at each cuticle intersection is a simple multiplication with the number of intersections $p + 1$. Similar to the attenuation terms, we also list the centers and variances of the distribution terms for lobes R , TT and TRT in Table 3, and give a generalized representation as

$$\Phi_p(h) = 2p\gamma_t - 2\gamma_i + p\pi, \quad (15)$$

$$\beta_p^2 = (p + 1)\beta_n^2. \quad (16)$$

4.3 Scattered lobes (TT^s , TRT^s)

As introduced in Sec. 3, the scattered lobes involve two large pre-computed lookup tables C^M and C^N . In this subsection, we first describe how the pre-computed data can be compressed. Then we minimize interactions between the scattered lobes and the cuticle, so that querying the scattered lobe is simplified. Finally, we derive closed-form expressions for their attenuation and distribution terms (Table 3), so that the previous implicit ray tracing is no longer required.

Compression. We treat these precomputed longitudinal and azimuthal scattering profiles C^M and C^N as 4D tensors, and refer to tensor decomposition techniques to compress them.

Tensors are high dimensional analogues to vectors or matrices, and can be regarded as multidimensional arrays. Tensor decomposition is a generalization of matrix singular value decomposition (SVD). In tensor decomposition, a d -dimensional tensor $\mathcal{A}$ is represented as a linear combination of R "simpler" tensors, each represented as the *tensor product* of d vectors:

$$\mathcal{A} = \sum_{r=1}^R \lambda_r \mathcal{A}^{(r)} = \sum_{r=1}^R \lambda_r \cdot \mathbf{a}^{(r,1)} \otimes \mathbf{a}^{(r,2)} \otimes \dots \otimes \mathbf{a}^{(r,d)}, \quad (17)$$

where the $(i_1, i_2, \dots, i_d)$ -th element of $\mathcal{A}^{(r)}$ is the product of the i_1 -th element of $\mathbf{a}^{(r,1)}$, the i_2 -th element of $\mathbf{a}^{(r,2)}$, $\dots$, and the i_d -th element of $\mathbf{a}^{(r,d)}$ (Fig. 9).

Here, R is the *rank* of tensor $\mathcal{A}$, and each $\mathcal{A}^{(r)}$ is rank 1. Similar to SVD for matrices, by assuming that the coefficients λ_r are sorted in decreasing order and keeping only the largest k coefficients, we

Lobe p	Attenuation A_p or A_p^s	Center Φ_p	Variance β_p^2 or Distribution D_p^s
R	F	$-2\gamma_i$	β_n^2
TT	$(1-F)^2 \exp\left(-\frac{2s_c\sigma_{c,a}+2s_m(\sigma_{m,a}+\sigma_{m,s})}{\cos\theta_d}\right)$	$2(\gamma_t - \gamma_i) + \pi$	$2\beta_n^2$
TRT	$(1-F)^2 F \exp\left(-\frac{4s_c\sigma_{c,a}+4s_m(\sigma_{m,a}+\sigma_{m,s})}{\cos\theta_d}\right)$	$(\pi + 2\gamma_t) + 2(\gamma_t - \gamma_i) + \pi$	$3\beta_n^2$
TT^s	$F \exp\left(-\frac{(s_c+1-\kappa)\sigma_{c,a}+\kappa\sigma_{m,a}}{\cos\theta_d}\right)$	$\gamma_t - \gamma_i$	$C^N(\Phi_p^s - \phi)$ (Eqn. 10)
TRT^s	$(1-F)F \exp\left(-\frac{(3s_c+1-\kappa)\sigma_{c,a}+(2s_m+\kappa)\sigma_{m,a}+2s_m\sigma_{m,s}}{\cos\theta_d}\right)$	$3\gamma_t - \gamma_i + \pi$	$C^N(\Phi_p^s - \phi)$ (Eqn. 10)

Table 3. Attenuation and distribution of each azimuthal lobe.

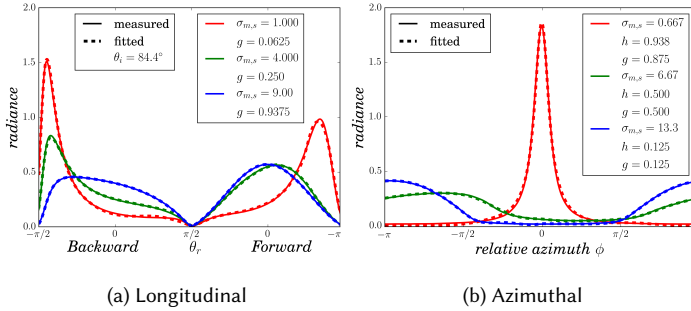


Fig. 10. Comparison of measured (solid lines) and compressed (dotted lines) scattering profiles. Using our tensor decomposition scheme, the compressed profiles have good matches with the measured data over a wide range of parameters.

are able to reconstruct tensor $\mathcal{A}$ approximately as

$$\mathcal{A} \approx \sum_{r=1}^k \lambda_r \cdot \mathbf{a}^{(r,1)} \otimes \mathbf{a}^{(r,2)} \otimes \dots \otimes \mathbf{a}^{(r,d)}. \quad (18)$$

If a relatively small k is needed to approximate tensor $\mathcal{A}$ accurately enough, we say that $\mathcal{A}$ is low rank. In this case, we only need to store $k \times d$ vectors to accurately reconstruct it, which is significantly less than the storage for $\mathcal{A}$ itself.

We apply this tensor decomposition scheme to compress the precomputed scattering profiles C^M and C^N , both with resolution $24 \times 16 \times 16 \times 720$, storing radiance at 24 values of $\sigma_{m,s}$, 16 values of g and 16 values of h or θ_i towards 720 outgoing directions. We use scikit-tensor, a Python module for multilinear algebra and tensor factorizations, to perform tensor decomposition using the alternating least squares (CP-ALS) algorithm. Our experiments show that, it usually takes less than one minute (single-threaded) to decompose either longitudinal or azimuthal precomputed data, with a maximum of 500 iterations. After the decomposition, we find that using up to rank 16 is good enough to accurately capture the complex shapes of all precomputed data.

The resulting coefficients take only 150 KB in storage, which is negligible compared to the 600 MB raw data in Yan et al. (2015). Figure 10 verifies the accuracy of our compression. Animated comparisons can be found in the accompanying video.

Note that previous tensor decomposition techniques for visual data (Tsai and Shih 2006; Vasilescu and Terzopoulos 2004; Wang et al. 2005) usually perform more complicated factorizations, using N-mode SVD with a core tensor. However, in our application, we find using a combination of rank 1 tensors suffices.

Longitudinal scattered lobes. The longitudinal scattered lobes are still interpolated between the precomputed lobes at $\phi = 0$ and $\phi = \pi$. However, we find that in Eqn. 8, the normalization factor μ is costly to compute but still approximate. Besides, the Fresnel transmittance cancels the normalization in most practical cases. Thus, we simplify the queries to use the incident and outgoing longitudinal angles directly as

$$M^s(\theta_i, \theta_r, \phi) = \text{lerp}_{\phi}(C_{\phi=0}^M(\theta_i, \theta_r), C_{\phi=\pi}^M(\theta_i, \theta_r)). \quad (19)$$

Azimuthal attenuation. Our azimuthal attenuation term A_p consists of two parts. The first part is the attenuation from the beginning to the point where the p -th segment starts intersecting the medulla, i.e. before the medulla scatters. Specifically, for TT^s ($p = 1$), we consider the first segment, and for TRT^s ($p = 2$), we consider the second rather than the first two, since the first segment's contribution is already accounted for in TT^s . In analogy to the unscattered lobes, we write the first part of the attenuation as $T_a = \exp(-[(2p-1)s_c\sigma_{c,a} + 2(p-1)s_m(\sigma_{m,a} + \sigma_{m,s})]/\cos\theta_d)$. The second part is the attenuation after the medulla's scattering, attenuated by the medulla of distance κ and by the cortex of distance $1 - \kappa$. We assume no reflection/refraction events happen when exiting the cuticle. Thus, the second part of the attenuation becomes $T_b = \exp(-[\kappa\sigma_{m,a} + (1-\kappa)\sigma_{c,a}]/\cos\theta_d)$. The overall attenuation term is thus

$$A_p^s(h) = (1-F)F^{p-1}T_aT_b. \quad (20)$$

Azimuthal distribution. For the distribution term $D_p^s(h, \phi) = C^N(\Phi_p^s(h) - \phi)$ (Eqn. 10), we derive the direction Φ_p^s of the p -th segment, and use the difference angle $\Phi_p^s - \phi$ to query from the precomputed azimuthal scattering profile C^N . Similar to the unscattered lobes, the direction of the p -th segment that enters the inner cylinder is given by

$$\Phi_p^s(h) = (\gamma_t - \gamma_i) + (p-1)(\pi + 2\gamma_t). \quad (21)$$

Intuitively, Eqn. 21 is the result of one refraction into the outer cylinder for the TT^s lobe ($p = 1$), or one refraction plus one internal reflection for the TRT^s lobe ($p = 2$).

The final azimuthal scattered lobes are written as $N_p^s(h, \phi) = A_p^s(h) \cdot D_p^s(h, \phi)$ (Eqn. 9).

In summary, for longitudinal scattered lobes, we simplify the interactions of the scattered lobes with the cuticle by ignoring the cuticle refraction. In this way, complex normalization is avoided along with Fresnel transmittance. For azimuthal scattered lobes, we assume that they originate from the center of the double cylinder and are attenuated evenly for all directions by the medulla and the cortex successively. So, compared to Yan et al. (2015), our model does

Parameter	Unit	Bobcat	Cat	Deer	Dog	Mouse	Rabbit	Raccoon	Red fox	Springbok	Human
κ	unitless	0.88	0.87	0.91	0.68	0.66	0.79	0.65	0.86	0.82	0.36
η	unitless	1.69	1.36	1.60	1.58	1.35	1.47	1.19	1.49	1.48	1.20
α	degree	5.48	3.65	3.52	2.94	0.55	3.14	1.81	2.64	4.61	0.70
β_m	degree	11.64	5.66	7.00	5.77	8.39	11.91	7.44	9.45	8.02	2.05
β_n	degree	7.49	1.34	4.53	18.94	2.80	10.52	6.88	17.63	11.46	3.75
$\sigma_{c,a}$	diameter ⁻¹	0.64	0.06	1.39	0.01	0.04	0.24	0.25	0.39	0.32	0.41
$\sigma_{m,s}$	diameter ⁻¹	1.69	2.47	2.51	2.44	1.34	0.78	2.30	3.15	2.45	3.49
$\sigma_{m,a}$	diameter ⁻¹	0.17	0.12	0.09	0.00	0.06	0.10	0.14	0.21	0.31	0.00
g	unitless	0.44	0.60	0.46	0.26	0.36	0.12	0.08	0.79	0.19	0.28
l	unitless	0.47	0.44	0.45	0.60	2.36	1.03	2.00	0.68	0.46	1.79
NRMSE ((Yan et al. 2015))		7.2%	5.3%	7.9%	9.1%	8.5%	8.4%	10.1%	6.3%	7.0%	19.3%
NRMSE (ours)		6.8%	6.4%	7.1%	7.3%	4.7%	6.0%	9.7%	6.2%	8.1%	16.1%

Table 4. (Top) Optimized parameters fit from our measured data using our far field model. All length-related parameters are calculated assuming the azimuthal section of every fiber is a unit circle. All angle-related parameters are in degrees. (Bottom) Normalized RMS error of Yan et al. (2015) and our model.

not have the additional diffusive lobe, and it accounts for absorption from the medulla. The simplicity of our model naturally leads to the ease of implementation. We provide implementation details in Sec. 6.

5 PIECEWISE ANALYTIC BCSDF MODEL

So far, we've derived a near field BCSDF. Now we show how to make a far field BCSDF approximation, which is especially efficient in reducing variance where hair or fur fibers are much thinner than a pixel. We then describe how to transition between near and far field models to make a multi-scale BCSDF that integrates per pixel. Our multi-scale BCSDF is the first model that is able to produce near field appearance when viewed close up, and requires no sampling when viewed far away.

5.1 Far field BCSDF model

To enable far field approximation, we need to integrate the near field azimuthal scattering profiles N_p and N_p^s over the azimuthal offset h . Both unscattered and scattered lobes share the same representation

$$N(\phi) = \frac{1}{2} \int_{-1}^1 A(h) \cdot D(h, \phi) dh, \quad (22)$$

by replacing N with $A \cdot D$ from Eqn. 6. For simplicity, we focus on the range $h \in [0, 1]$ since it is always equivalent for symmetric queries (h, ϕ) and $(-h, -\phi)$. Thus, Eqn. 22 becomes

$$\begin{aligned} N(\phi) &= \frac{1}{2} \int_0^1 A(h) \cdot D(h, \phi) dh + \frac{1}{2} \int_0^1 A(h) \cdot D(h, -\phi) dh \\ &\triangleq N^{(+)}(\phi) + N^{(-)}(\phi). \end{aligned} \quad (23)$$

Unscattered lobes. We first look at the attenuation term. Our observation is that, the attenuation term A_p from Eqn. 14 can be treated as the product of four components: $(1 - F)^2$, F^{p-1} , T_c and T_m . These four components are all functions of h — for the former two components with Fresnel terms, h defines different incident angles γ_i , and for the latter two absorption components, h decides the distances light travels within cortex and medulla s_c and s_m . Another observation is that, these components are either monotonic or smooth when h changes. So we start with partitioning h into a few segments (Fig. 12 (a)). Then, for maximum accuracy, we linearize the entire

first two components $(1 - F)^2$ and F^{p-1} and the distances s_c and s_m in the latter two components. Thus, the attenuation term $A(h)$ becomes the product of two linear functions and two exponentials of linear functions.

Then we analyze the Gaussian distribution term $D_p(h, \phi) = G(\Phi_p(h) - \phi)$. Here we're not interested in its variance which is constant with h . Instead, we focus on its center that varies with h . As shown in Eqn. 15, when h changes, γ_t and γ_i change with it. So, similar to the attenuation term, with segmented h , we are able to represent γ_t and γ_i with linear functions. Since a Gaussian is an exponential of squared variables, the distribution term ends up with an exponential of a quadratic polynomial.

With the piecewise polynomial representation inside both the attenuation term A_p and the distribution term D_p , we're able to represent the azimuthal scattering profile N_p for unscattered lobes in a simple form. Note that the product of two linear functions from A_p makes a quadratic polynomial, and the product of two exponentials of linear functions from A_p and the exponential of a quadratic polynomial from D_p together makes another exponential of a quadratic polynomial. So, we have the following form for unscattered lobes:

$$N_p^{(+|-)}(\phi) = \frac{1}{2} \sum_{i=1}^n \int_{h_{i-1}}^{h_i} Q_1(h) \cdot \exp(Q_2(h)) dh, \quad (24)$$

where Q_1 and Q_2 are quadratic polynomials. This can be easily solved analytically, as will be described with details in the Appendix.

Scattered lobes: The scattered lobes are similarly handled. For the attenuation term A_p^s in Eqn. 20, we linearize the Fresnel terms $1 - F$, F^{p-1} as well as the distances s_c and s_m , resulting in the product of two linear functions and two exponentials of linear functions.

Since the distribution term D_p is queried rather than computed as a Gaussian, it is even simpler so that we directly linearize it. So the entire distribution term is a linear function.

Thus, the final result of the azimuthal scattering profile N_p^s for scattered lobes has the form

$$N_p^{s(+|-)}(\phi) = \frac{1}{2} \sum_{i=1}^{n-1} \int_{h_{i-1}}^{h_i} C(h) \cdot \exp(L(h)) dh, \quad (25)$$

where C is a cubic polynomial that is the product of linearized components $(1 - F)$, F^{p-1} and D_p^s , and L is a linear function of the

product from T_a and T_b . This integral can also be solved analytically with even simpler results than unscattered lobes.

Segmentation. We observed that when h is large, i.e. the incident position is away from the center, the linear terms change more rapidly. So we partition $h \in [0, 1]$ quadratically, i.e. $h_i = \sqrt{i/n}$. For unscattered lobes, we find that using 5 segments is enough in most practical renderings, while using 8 segments generates indistinguishable scattering profiles. For scattered lobes, since they are even smoother, we find 4 segments good enough throughout all computations.

Acceleration. In practice, we usually don't have to integrate all n segments. Given a specific h , we immediately know its relative exiting azimuth $\Phi_p(h)$. For unscattered lobes, since the distribution term is a Gaussian around $\Phi_p(h)$, it is safe to assume that a path that is incident from h contributes only within this outgoing range $[\Phi_p(h) - 3\beta_p, \Phi_p(h) + 3\beta_p]$. Based on this observation, for a segment $h \in [h_1, h_2]$, we can limit its contribution within the range $[\min\{\Phi_p(h_1), \Phi_p(h_2)\} - 3\beta_p, \max\{\Phi_p(h_1), \Phi_p(h_2)\} + 3\beta_p]$. So we simply throw away all the queries with ϕ that are not within this range. In this way, each query for unscattered lobes now requires an average of only 2–4 integrations in practice. However, for scattered lobes, since the distribution term is precomputed for every direction, the acceleration scheme does not apply.

Validation. To verify the accuracy of all these simplifications / improvements, we re-fitted all the measured fur reflectance profiles from Yan et al. (2015) and compared the NRMSE³ with previous fitting results.

We use Ceres Solver (Agarwal et al. 2010), a nonlinear least squares minimizer, to fit the measured profiles. The fittings are performed in logarithmic space, with the cost function defined as the sum of all per-pixel differences of the log-measured and log-fitted values, divided by the range between the minimum and maximum log-measured values. The initial values of all parameters are manually set. The fitted profiles are generated using our far field model, with 8 segments for unscattered lobes and 5 segments for scattered lobes. Fitting each profile takes 2 ~ 3 minutes in our test platform.

Figure 11 shows 3-way comparison of the measured data, fitted profiles in Yan et al. (2015) and our fitted profiles. From the fitted profiles, we can see that even with only 5 lobes, our method is still able to produce similar forward scattered lobes (e.g. red fox) and backward scattered lobes (e.g. raccoon). Also, with the introduction of medulla absorption, and the unification of IORs thus reducing the complexity of light scattering, our model has much “cleaner” forward scattered regions near $(\theta = 10^\circ, \phi = 180^\circ)$ (e.g. mouse and rabbit), which cannot be handled previously. The introduced azimuthal roughness intuitively smoothes the fitted profiles azimuthally. It is especially obvious for the R and TT lobes, so that the high-intensity regions fit better (e.g. raccoon, rabbit and dog). Note that our method may not be consistently better than Yan et al. (2015) over these regions. This is expected, since the fitting procedure is aimed at global optimization.

One limitation of our method would be that, our results are slightly more blurred longitudinally, indicating larger longitudinal

³Normalized RMS Error, or precisely, RMS error of the fitting result divided by the range of measured data.

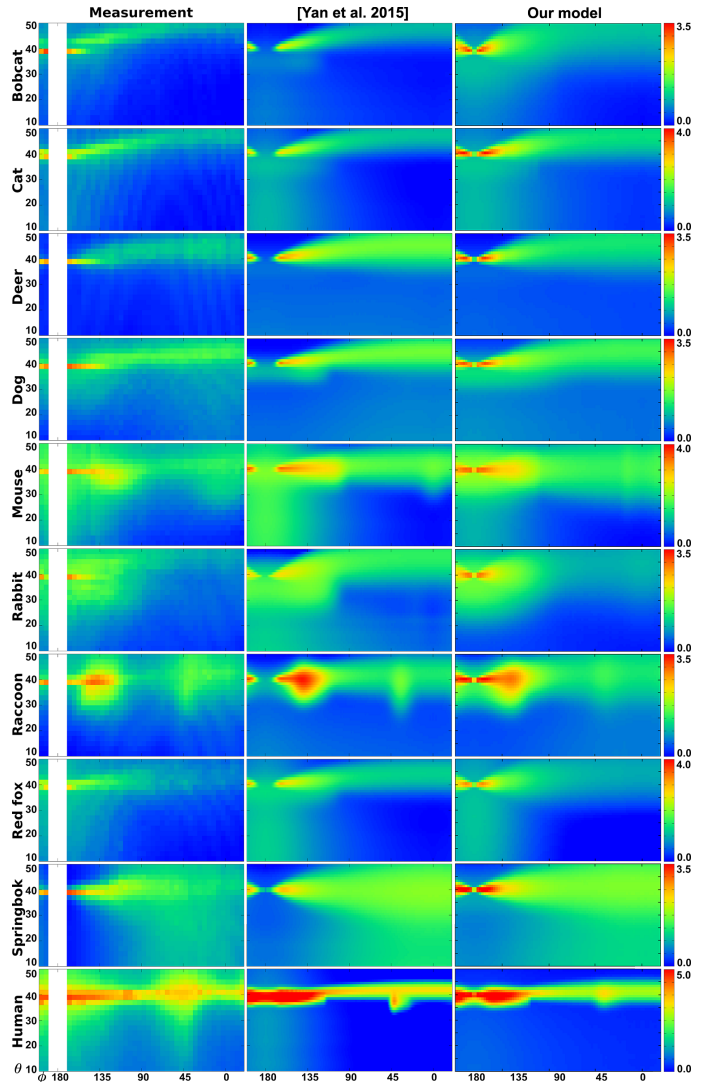


Fig. 11. (Left) Reflectance profiles measured from different animals' fur fibers in Yan et al. (2015). (Middle) Synthesized profiles using the factored rendering model in Yan et al. (2015). (Right) Synthesized profiles using our analytic far field BCSDf model in Sec. 5. All profiles are scaled and displayed in logarithmic space for perceptual brightness.

roughness fitted. This may relate to a cuticle scattering phenomenon observed as stripes in almost all measured profiles, and we leave it as future work. Also, there are cases (e.g. human) where neither our method nor Yan et al. (2015) produce good fits, indicating that the double cylinder model can be further improved.

Table 4 lists all the fitted parameters and NRMSE values. As analyzed above, though our near field model makes several approximations for scattered lobes, and our far field model builds upon it with further piecewise linear approximations, our model still has better results in most cases. Note that our model is much simpler and more general, and the new parameters in our model (β_n and $\sigma_{c,a}$) provide better flexibility for artist control.

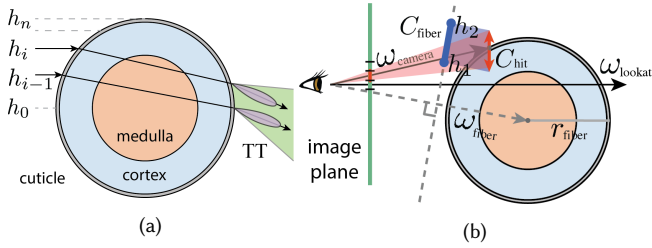


Fig. 12. (a) Illustration of far field integration. The longitudinal section is partitioned into n segments. For each segment, we compute its contribution to the queried relative exiting azimuth. (b) Illustration of calculating a pixel's coverage for multi-scale rendering. A pixel (marked red on the image plane) is first projected to the hit point (red segment at the hit point), then projected again towards the fiber (blue segment).

5.2 Multi-scale BCSDf model

Far field approximation is accurate when hair or fur fibers are thinner than a pixel. However, when viewed close up, a fiber's width may cover several pixels, i.e. each pixel actually covers a small range over the azimuthal section. In these cases, far field approximation will produce ribbon-like appearance (Fig. 5 (b)). To deal with these cases, we propose our multi-scale BCSDf model. We use each pixel's coverage on hair or fur fibers to decide the range of azimuthal offset h it covers, and integrate per pixel instead of per fiber. In this way, we integrate similar to far field approximation, but keep the accurate near field appearance.

Pixel-wise integration. Suppose we know that a pixel covers a range of azimuthal offset $h \in [h_1, h_2]$. We extend Eqn. 22 to integrate only within this range as

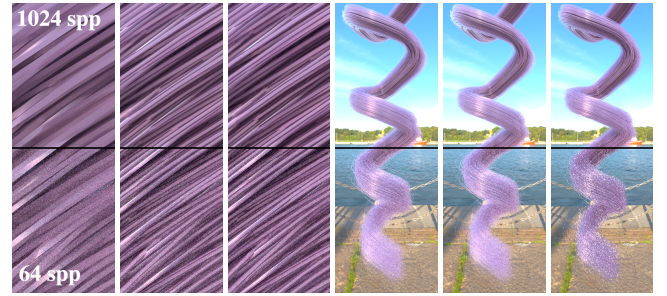
$$N(\phi) = \frac{1}{h_2 - h_1} \int_{h_1}^{h_2} A(h) \cdot D(h, \phi) dh, \quad (26)$$

where the term $1/(h_2 - h_1)$ guarantees energy conservation. When a pixel fully covers the entire azimuthal section, Eqn. 26 degenerates to the far field case Eqn. 22. And in the limit case where h_1 and h_2 are infinitesimally close, it becomes the near field scattering representation $N(h, \phi) = A(h) \cdot D(h, \phi)$. These two cases indicate that our multi-scale BCSDf model bridges both near and far field scattering, and is consistent when scaling between them.

Calculating a pixel's coverage. Now that we have a multi-scale BCSDf model, what remains is to find a pixel's coverage $[h_1, h_2]$. Figure 12 illustrates the way to calculate it. Assuming that pixels are round rather than square on the image plane, we can tell how large a pixel's coverage is at the hit point in world coordinates using its diameter, denoted as C_{hit} . This is very similar to ray differentials (Igehy 1999). Then, the projected pixel at the hit point becomes a disk, facing along the camera's look-at direction ω_{lookat} . We project the disk again towards the hit fiber, i.e. onto the direction ω_{fiber} . Finally, we compare it with the fiber's radius r_{fiber} in world coordinates to get the pixel's coverage. So, we have

$$C_{fiber} \approx (\omega_{fiber} \cdot \omega_{lookat}) C_{hit} / r_{fiber} \quad (27)$$

as the pixel's coverage in the azimuthal section of the fiber, with the same unit as $h \in [-1, 1]$. Here ω_{fiber} is the direction from the camera to the fiber's center within the same azimuthal section with the hit point. The fiber's center can be calculated when performing



Far field Multi-scale Near field Far field Multi-scale Near field
14.0min/51s 11.8min/45s 9.3min/34s 3.5min/14s 3.7min/14s 2.4min/9s

Fig. 13. Validation of far field and multi-scale rendering. We render the same scene viewed close up (left three columns) and viewed from far (right three columns), using 1024 spp (top half) and 64 spp (bottom half). Timings are listed for 1024 spp and 64 spp, respectively. When zoomed out, our far field and multi-scale models perform around $1.5\times$ slower than near field. When zoomed in, our multi-scale model performs closer to near field, since the range to integrate becomes smaller for each pixel.

ray-cylinder intersections. However, unless viewed from extremely close so that a fiber covers a very large area in the image plane, ω_{fiber} is always close to the camera ray's direction ω_{camera} . So, in practice, we replace ω_{fiber} with ω_{camera} in Eqn. 27 for simplicity.

After getting a pixel's coverage, the azimuthal range to integrate can be written as $[h - C_{fiber}/2, h + C_{fiber}/2]$. This range is then clamped to be within $[-1, 1]$, in case a pixel is much larger than a fiber's width, or it covers the boundary of a fiber. To integrate, we use the same segmentation scheme for far field approximation, clipping segments to be within this range.

Validation. We render the same insets using our near field, far field and multi-scale BCSDf models. As shown in Fig. 13, when viewed far away, the differences between our three methods are barely visible, but the far field and multi-scale models produce significantly less noise. When viewed close up, our multi-scale model still generates the same results as compared to the near field model, but the far field results look flat. A similar effect is seen in Fig. 1, where the multi-scale model produces the same appearance as the near field, but is much less noisy. In the accompanying video, we include a zooming-in sequence to verify that our multi-scale scattering model scales smoothly and does not produce flickering appearance between frames.

6 RENDERING

In this section, we provide a brief summary of how to implement our reflectance model within global illumination renderers. We discuss two relevant aspects: evaluation and sampling.

BCSDf evaluation. Since our reflectance model unifies hair and fur rendering with only two additional scattered lobes, our BCSDf evaluation easily fits in a hair rendering system. We provide a dependency tree in Fig. 14 of variables to compute, separating the classic R , TT and TRT lobes for hair and new scattered lobes TT^s , TRT^s .

When near field rendering is required, the azimuthal lobes are queried at specific offset h . On the other hand, for far field approximation or multi-scale rendering, queries happen on the ends of all

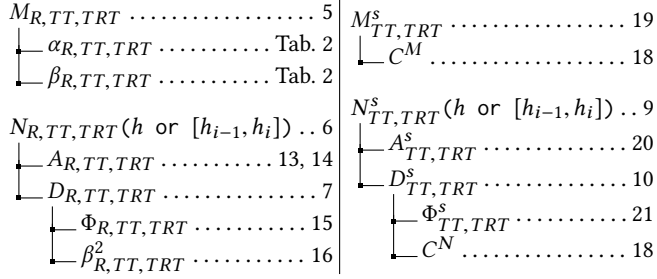


Fig. 14. A dependency tree of variables for BCSDf evaluation. Equation numbers are marked. All the variables required for both unscattered and scattered lobes are also listed in Tables 2 and 3.

segments. Note again that every step in the tree of Fig. 14 is either analytic or queried, thanks to the unification of IORs.

Importance sampling. Similar to d'Eon et al. (2013), our importance sampling works in four steps as follows:

- Choosing azimuthal offset h . For near field scattering, since it is fixed, there's no need to choose. For far field approximation and multi-scale rendering, we randomly pick an h from the corresponding range, i.e. $[-1, 1]$ for far field and $[h_1, h_2]$ for multi-scale.
- Choosing a lobe p to sample. The lobes are weighted according to the energy they carry. Since the longitudinal integral of M_p is always 1, and the azimuthal integral of the distribution term D_p is also 1 for any selected h , the energy that lobe p carries depends on its attenuation term A_p . So, we calculate the attenuation term for all 5 lobes, and choose one with the probabilities in proportion to their values.
- Sampling azimuthally. Once the azimuthal offset h and the lobe p have been selected, for unscattered lobes, we immediately know how they distribute with center Φ_p and standard variance β_p . We perform a Gaussian sampling according to this distribution to get the relative exiting azimuth ϕ . Then the actual outgoing azimuth $\phi_r = \phi + \phi_i$ can be computed. For scattered lobes, since they're smooth, we sample them uniformly over all azimuthal angles.
- Sampling Longitudinally. This is very similar to the azimuthal case. Since we know which lobe is to be sampled, for unscattered lobes, we just need to sample according to the Gaussian M_p . For scattered lobes, we use cosine sampling.
- Calculating PDF and sampling weight. The final probability density function (PDF) is the product of PDFs sampling the selected lobe p azimuthally and longitudinally, followed by a conversion from (θ, ϕ) -measure to solid angle measure. The final sampling weight is the BCSDf value of the selected lobe p over the final PDF, then divided by the probability of selecting it. Since the unscattered lobes are importance sampled and the scattered lobes are usually smooth, the sampling weight is usually smaller than 2 in practice.

With the importance sampling scheme, we perform standard path tracing to determine accurate global illumination.

	Fig.	#Strands	#Segs	#Samples	Time	Method
Hair Lock	8	1K	210	1024	3.7min	N/M/F
Raccoon	1	260K	22	1024	14.1min	M
Hamster	16	580K	15	1024	36.9min	N
Cat	17	267K	9	256	3.8min	M
Hair	18	53K	64	1024	17.3min	M

Table 5. Statistics for our scenes, all rendered in 720p, using different rendering methods (N for near field, M for multi-scale, F for far field). Each of the (# Strands) fur fibers is represented using (# Segs) line segments. # Samples is the number of samples per pixel.

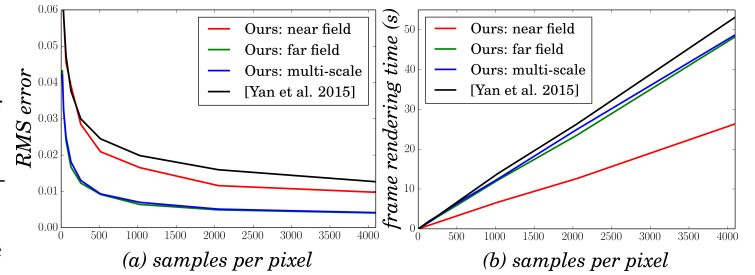


Fig. 15. We compare our method with Yan et al. (2015) in (a) convergence and (b) frame rendering time. The experiments are conducted on the central 32×32 patch of the hair lock scene. Our near field model converges twice as fast as Yan et al. (2015), and our multi-scale model converges an order of magnitude faster. Also, our near field model performs $2\times$ faster than Yan et al. (2015). Moreover, since our far field and multi-scale models require solving integrals rather than querying the integrand as near field models do, a slight performance drop is expected. However, they still outperform Yan et al. (2015) even though it is near field, indicating that our analytic integration is efficient. Also note that since the scene is viewed from far away, our far field and multi-scale models converge almost as fast in this case.

While our importance sampling method is similar to Yan et al. (2015), their model has 11 lobes, most of which make a very small contribution to the final image. When a lobe with low energy is selected with low probability, the sampling weight will be large, and the result will be noisy (have high variance). In contrast, our reflectance model has only 5 lobes, each of which carries a significant amount of energy. Thus, even when comparing only our near field results, our method still has less noise (Figs. 1, 15 and 16).

7 RESULTS

In this section, we show rendering results generated using our practical reflectance model, and compare them with previous work. We implement our model in the Mitsuba renderer (Jakob 2010). Scene configurations, including number of hair or fur fibers and samples per pixel, are listed in Table 5. Parameters for raccoon, cat and human hair/fur models are taken from our best fit results in Table 4. The hamster model uses parameters from mouse. Since the measured data from Yan et al. (2015) is grayscale only, to introduce color, we convert colored textures at different positions to absorption coefficients $\sigma_{c,a}$ in the cortex, similar to Yan et al. (2015). All scenes are rendered using path tracing on an Amazon EC2 c4.8xlarge instance

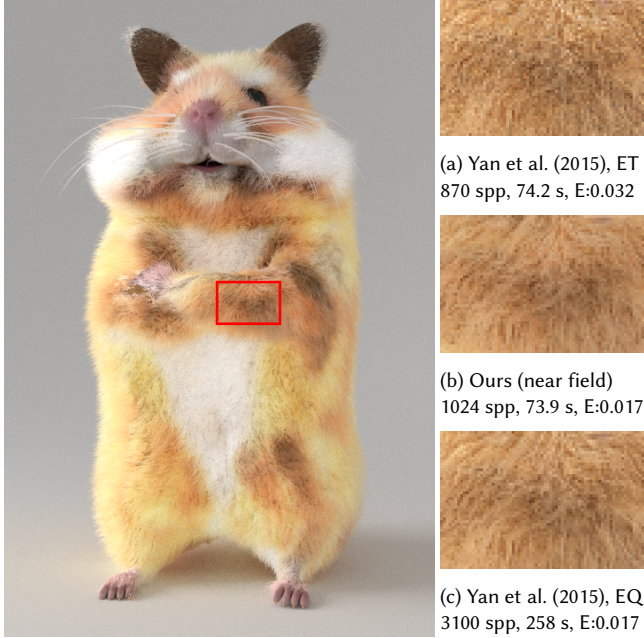


Fig. 16. A *Hamster* model rendered under studio lighting with a diffuse backdrop, using our near field model. Insets compare (b) our method with Yan et al. (2015) for (a) equal time (ET) and (c) equal quality (EQ) (E in the sub-captions stands for RMS error). The two models have different parameter spaces, thus small differences can be found in the EQ comparison. In the ET comparison, noise can be clearly seen when zoomed in. For equal quality, our method performs 3.5× faster.

with 36 vCPUs. The source code and compressed pre-computed data are available on <http://viscomp.ucsd.edu/projects/fur2>.

We measure and compare the entire frame rendering time, including BVH traversal and ray-cylinder intersections. Even so, our near field reflectance model still performs around 3.5× faster than Yan et al. (2015) in terms of equal quality comparison, and our multi-scale rendering scheme performs even better with up to a 8× speed-up. Note that since the parameter space in our reflectance model is different from Yan et al. (2015), slight differences can be observed in their rendering results.⁴

Hair lock. Figure 8 shows decomposed renderings from each of the 5 lobes in our model, using the fitted parameters of human hair in Table 4. We can clearly see different lobes' contribution, indicating that our model is concise and effective. Figure 15 compares the convergence curves and rendering time using different methods. Yan et al. (2015) is near field and does not require integration. However, our multi-scale model not only evaluates slightly faster, but also converges fastest among all the models. Our near field model evaluates fastest, but still converges 2× faster than Yan et al. (2015). In the accompanying video, we show a zooming in and out sequence, showing that our multi-scale rendering scheme is accurate when viewed close, and is efficient when viewed far away (Fig. 13).

Raccoon. The raccoon scene is rendered with an HDR environment map. We use a ground projection scheme similar to Autodesk

⁴Hence, errors are computed with respect to the converged result for each reflectance model separately.

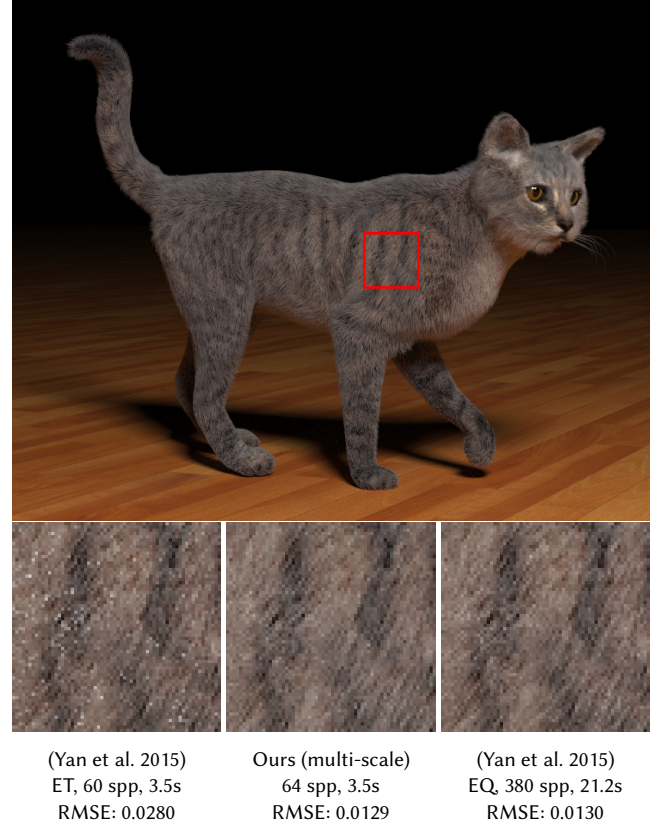


Fig. 17. A *Cat* model rendered multi-scale with an area light in front. We compare with Yan et al. (2015) for equal time and equal quality (same RMSE). For the same quality, our multi-scale BCSDf achieves a 6.0× speed-up.

Fusion 360's implementation, so that the raccoon stands on an actual ground rather than floating. As shown in Fig. 1, our near field model is capable of generating very similar diffusive appearance as compared to Yan et al. (2015), but is much simpler and has less noise. Furthermore, our multi-scale rendering scheme converges significantly faster with minimum overhead. This is because the fur fibers are so thin that near field sampling is very inefficient, while multi-scale rendering integrates efficiently, successfully removing the high frequency noise.

Hamster. This scene (Fig. 16) shows a hamster model, rendered under studio lighting with several area lights on top. The hamster model is located inside a capsule-like diffuse backdrop, encompassing the top, bottom and back sides. Since everything is diffuse, near field reflectance is efficient enough. We compare our near field model with Yan et al. (2015) which is also near field, showing that our model has better convergence because of its simplicity.

Cat. The cat model (Fig. 17) is rendered using an area light in front. We compare our multi-scale model with Yan et al. (2015). Our scene is rendered noise-free with only 256 samples per pixel. In the video, we move the light back and forth, so the noise can be observed more clearly. Note that since the cat fur fibers are very thin, our multi-scale and far field models generate exactly the same results in approximately the same time.



(a) Without medulla. ($\kappa = 0$) (b) With medulla. ($\kappa = 0.15$)

Fig. 18. A hair model rendered with and without medulla using our multi-scale model. Our model unifies hair and fur rendering with the same light paths, regardless of the medulla's size. The difference between these two renderings is clearly visible, indicating the importance of the medulla as well as our scattered lobes TT^s and TRT^s , even with a small κ .

Hair. Our reflectance model and multi-scale integration also work on human hair. As shown in Fig. 18, even a small medulla ($\kappa = 0.15$) makes a difference in hair's overall appearance. Intuitively, this is because the light that goes through the medulla is spread more, making each hair more diffusive. The result indicates the importance of the medulla even for human hair, showing that the diffusive appearance comes not only from global illumination between hair fibers, but also within hair fibers. Note that the same light paths are computed through hair or fur fibers regardless of the medulla. Moreover, our multi-scale integration benefits both hair and fur rendering.

8 CONCLUSION AND FUTURE WORK

We present a practical reflectance model for efficient fur rendering. By unifying the IORs of cortex and medulla, our model is capable of representing complex scattering within hair and fur fibers with only 5 lobes. By introducing medulla's absorption and different longitudinal and azimuthal roughness, and using tensor approximation to minimize the storage overhead, our model achieves both accuracy and practicality. Along with the simplified model, we propose an analytic integration scheme for efficient far field approximation, and extend it to handle multi-scale rendering for the first time.

In the future, it is straightforward to implement our BCSDf model for real-time rasterization based applications. It is also worth thinking about ways to further accelerate the far field approximation. An artist-friendly perspective for our models can benefit the industry. Introducing explicit eccentricity or irregular shaped azimuthal sections would also help. In summary, we believe that our model is an important step in practical physically-based fur rendering, which unifies hair and fur reflectance models, as well as near and far field rendering schemes.

ACKNOWLEDGMENTS

This work was supported in part by NSF grant 1451828, an NVIDIA fellowship, and the UC San Diego Center for Visual Computing.

REFERENCES

- Sameer Agarwal, Keir Mierle, and Others. 2010. Ceres Solver. <http://ceres-solver.org>. (2010).
- Ellen Carrlee and Lauren Horelick. 2011. The Alaska Fur ID Project: A virtual resource for material identification. In *Objects Specialty Group Postprints*, Vol. 18. American Institute for Conservation of Historic and Artistic Works, 149–171.
- Matt Jen-Yuan Chiang, Benedikt Bitterli, Chuck Tappan, and Brent Burley. 2016. A Practical and Controllable Hair and Fur Model for Production Path Tracing. *Computer Graphics Forum* 35, 2 (2016), 275–283. DOI: <https://doi.org/10.1111/cgf.12830>
- Eugene d'Eon, Guillaume Francois, Martin Hill, Joe Letteri, and Jean-Marie Aubry. 2011. An Energy-conserving Hair Reflectance Model. In *ACM Transactions on Graphics (TOG)*. 1181–1187. DOI: <https://doi.org/10.1111/j.1467-8659.2011.01976.x>
- Eugene d'Eon, Steve Marschner, and Johannes Hanika. 2013. Importance Sampling for Physically-based Hair Fiber Models. In *SIGGRAPH Asia 2013 Technical Briefs*. Article 25, 4 pages. DOI: <https://doi.org/10.1145/2542355.2542386>
- AntonĀņn GalatĀņk, Jan GalatĀņk, Zdislav Krul, and AntonĀņn GalatĀņk Jr. 2011. Furskin Identification. <http://www.furskin.cz>. (2011).
- Homan Igchey. 1999. Tracing Ray Differentials. In *SIGGRAPH* 99. 179–186.
- Wenzel Jakob. 2010. Mitsuba renderer. <http://www.mitsuba-renderer.org>. (2010).
- James Kajiya and Timothy Kay. 1989. Rendering Fur with Three Dimensional Textures. In *SIGGRAPH* 89. 271–280.
- Stephen R. Marschner, Henrik Wann Jensen, Mike Cammarano, Steve Worley, and Pat Hanrahan. 2003. Light Scattering from Human Hair Fibers. *ACM Transactions on Graphics (TOG)* 22, 3 (2003), 780–791. DOI: <https://doi.org/10.1145/882262.882345>
- Yuting Tsai and Zenchung Shih. 2006. All-frequency precomputed radiance transfer using spherical radial basis functions and clustered tensor approximation. *ACM Transactions on Graphics (TOG)* 25, 3 (2006), 967–976.
- M Alex O Vasilescu and Demetri Terzopoulos. 2004. TensorTextures: multilinear image-based rendering. *ACM Transactions on Graphics (TOG)* 23, 3 (2004), 336–342.
- Hongcheng Wang, Qing Wu, Lin Shi, Yizhou Yu, and Narendra Ahuja. 2005. Out-of-core tensor approximation of multi-dimensional matrices of visual data. *ACM Transactions on Graphics (TOG)* 24, 3 (2005), 527–535.
- Ling-Qi Yan, Chi-Wei Tseng, Henrik Wann Jensen, and Ravi Ramamoorthi. 2015. Physically-accurate fur reflectance: modeling, measurement and rendering. *ACM Transactions on Graphics (TOG)* 34, 6 (2015), 185.
- Arno Zinke and Andreas Weber. 2007. Light Scattering from Filaments. *IEEE Transactions on Visualization and Computer Graphics* 13, 2 (2007), 342–356. DOI: <https://doi.org/10.1109/TVCG.2007.43>

APPENDIX: SOLVING EQNS. 24 AND 25

Equations. 24 and 25 are piecewise integrations of a polynomial and the exponential of a polynomial. The key to solving them analytically is to integrate the forms $Q_1 \cdot \exp(Q_2)$ and $C \cdot \exp(L)$, where Q_1 and Q_2 are quadratic polynomials, C is cubic, and L is linear. For simplicity, here we present the analytic result of both forms as indefinite integrations.

Integrating $Q_1 \cdot \exp(Q_2)$:

$$\int (dx^2 + ex + f) \cdot \exp(-ax^2 + bx + c) dx$$

$$= \frac{\sqrt{\pi} \exp\left(\frac{b^2}{4a} + c\right) \operatorname{erf}\left(\frac{2ax-b}{2\sqrt{a}}\right) (4a^2f + 2abe + 2ad + b^2d)}{8a^{5/2}} - \exp(-ax^2 + bx + c)(2adx + 2ae + bd)/(4a^2) + K$$

where erf is the error function, which can be approximated with high precision using polynomials. K is the integration constant in the indefinite integral.

Integrating $C \cdot \exp(L)$:

$$\int (cx^3 + dx^2 + ex + f) \cdot \exp(-ax + b) dx$$

$$= -\exp(-ax + b) \left(a^3cx^3 + (a^3d + 3a^2c)x^2 + (a^3e + 2a^2d + 6ac)x + a^3f + a^2e + 2ad + 6c \right) / a^4 + K$$

where K is a constant.