

A Perceptual Model for Eccentricity-dependent Spatio-temporal Flicker Fusion and its Applications to Foveated Graphics

BROOKE KRAJANCICH, Stanford University
 PETR KELLNHOFER, Stanford University and Raxium
 GORDON WETZSTEIN, Stanford University

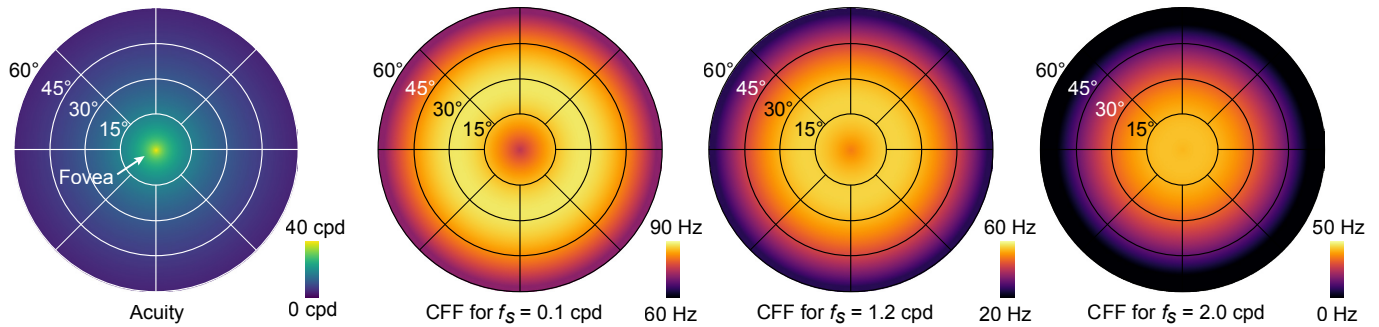


Fig. 1. Foveated graphics techniques rely on eccentricity-dependent models of human vision. While such models are well understood for spatial acuity (left, [Geisler and Perry 1998]), our work is the first to experimentally derive a more comprehensive model for the spatio-temporal aspects over the retina under conditions close to virtual reality applications. As seen in the three plots on the right, we model critical flicker fusion (CFF) thresholds in an eccentricity-dependent manner. The CFF varies with spatial frequency f_s and luminance and it exhibits an anti-foveated effect, with the highest thresholds observed in the near-mid periphery of the visual field. Our perceptual model and its experimental validation could provide the foundation of future spatio-temporally foveated graphics systems.

Virtual and augmented reality (VR/AR) displays strive to provide a resolution, framerate and field of view that matches the perceptual capabilities of the human visual system, all while constrained by limited compute budgets and transmission bandwidths of wearable computing systems. Foveated graphics techniques have emerged that could achieve these goals by exploiting the falloff of spatial acuity in the periphery of the visual field. However, considerably less attention has been given to temporal aspects of human vision, which also vary across the retina. This is in part due to limitations of current eccentricity-dependent models of the visual system. We introduce a new model, experimentally measuring and computationally fitting eccentricity-dependent critical flicker fusion thresholds jointly for both space and time. In this way, our model is unique in enabling the prediction of temporal information that is imperceptible for a certain spatial frequency, eccentricity, and range of luminance levels. We validate our model with an image quality user study, and use it to predict potential bandwidth savings 7× higher than those afforded by current spatial-only foveated models. As such, this work forms the enabling foundation for new temporally foveated graphics techniques.

CCS Concepts: • **Hardware** → **Displays and imagers**; • **Computing methodologies** → **Computer graphics**; **Mixed / augmented reality**.

Additional Key Words and Phrases: applied perception, flicker fusion, foveated rendering, foveated compression, virtual reality, augmented reality

1 INTRODUCTION

Emerging virtual and augmented reality (VR/AR) systems have unlocked new user experiences by offering unprecedented levels of immersion. With the goal of showing digital content that is indistinguishable from the real world, VR/AR displays strive to match the perceptual limits of human vision. That is, a resolution, framerate, and field of view (FOV) that matches what the human eye can perceive. However, the required bandwidths for graphics processing units to render such high-resolution and high-framerate content in interactive applications, for networked systems to stream, and for displays and their interfaces to transmit and present this data is far from achievable with current hardware or standards.

Foveated graphics techniques have emerged as one of the most promising solutions to overcome these challenges. These approaches use the fact that the acuity of human vision is eccentricity dependent, i.e., acuity is highest in the fovea and drops quickly towards the periphery of the visual field. In a VR/AR system, this can be exploited using gaze-contingent rendering, shading, compression, or display (see Sec. 2). While all of these methods build on the insight that acuity, or spatial resolution, varies across the retina, common wisdom also suggests that so does temporal resolution. In fact it may be anti-foveated in that it is higher in the periphery of the visual field than in the fovea. This suggests that further bandwidth savings, well beyond those offered by today’s foveated graphics approaches, could be enabled by exploiting such perceptual limitations with novel gaze-contingent hardware or software solutions. To the best

Authors’ addresses: Brooke Krajancich, Stanford University, brookek@stanford.edu; Petr Kellnhöfer, Stanford University and Raxium, pkellnho@stanford.edu; Gordon Wetzstein, Stanford University, gordon.wetzstein@stanford.edu.

*Project website:
<https://www.computationalimaging.org/publications/cff/>

of our knowledge, however, no perceptual model for eccentricity-dependent spatio-temporal aspects of human vision exists today, hampering the progress of such foveated graphics techniques.

The primary goal of our work is to experimentally measure user data and computationally fit models that adequately describe the eccentricity-dependent spatio-temporal aspects of human vision. Specifically, we design and conduct user studies with a custom, high-speed VR display that allow us to measure data to model critical flicker fusion (CFF) in a spatially-modulated manner. In this paper, we operationally define CFF as a measure of spatio-temporal flicker fusion thresholds. As such, and unlike current models of foveated vision, our model is unique in predicting what temporal information may be imperceptible for a certain eccentricity, spatial frequency, and luminance.

Using our model, we predict potential bandwidth savings of factors up to 3,500 $\times$ over unprocessed visual information and 7 $\times$ over existing foveated models that do not account for the temporal characteristics of human vision.

Specifically we make the following contributions:

- We design and conduct user studies to measure and validate eccentricity-dependent spatio-temporal flicker fusion thresholds with a custom display.
- We fit several variants of an analytic model to these data and also extrapolate the model beyond the space of our measurements using data provided in the literature, including an extension for varying luminance.
- We analyze bandwidth considerations and demonstrate that our model may afford significant bandwidth savings for foveated graphics.

Overview of Limitations. The primary goals of this work are to develop a perceptual model and to demonstrate its potential benefits to foveated graphics. However, we are not proposing new foveated rendering algorithms or specific compression schemes that directly use this model.

2 RELATED WORK

2.1 Perceptual Models

The human visual system (HVS) is limited in its ability to sense variations in light intensity over space and time. Furthermore, these abilities vary across the visual field. The drop in spatial sensitivity with eccentricity has been well studied, attributed primarily to the drop in retinal ganglion cell densities [Curcio and Allen 1990] but also lens aberrations [Shen et al. 2010; Thibos et al. 1987]. This is often described by the spatial contrast sensitivity function (sCSF), defined as the inverse of the smallest contrast of sinusoidal grating that can be detected at each spatial frequency [Robson 1993]. Spatial resolution, or acuity, of the visual system is then defined as the highest spatial frequency that can be seen, when contrast is at its maximum value of 1 or the log of the sCSF is zero.

On the contrary, temporal sensitivity has been observed to peak in the periphery, somewhere between 20 – 50° eccentricity [Hartmann et al. 1979; Rovamo and Raninen 1984; Tyler 1987], attributed to faster cone cell responses in the mid visual field by Sinha et al. [2017].

Table 1. Existing models of contrast sensitivity (CSF) and flicker fusion (CFF). Note that CSF models perception continuously across a range of conditions while CFF only models the limit of temporal perception at high temporal frequencies. Unlike ours, none of these models accounts for spatial and temporal variation as well as eccentricity and luminance.

Model	Spat.	Temp.	Eccentr.	Lum.	Prop.
Tyler [1990]	✗	✓	✗	✓	CFF
Kelly [1979]	✓	✓	✗	✗	stCSF
Watson [2016]	✓	✓	✗	✓	stCSF
Watson [2018]	✓	✗	✓	✓	sCSF
Ours	✓	✓	✓	✓	CFF

Analogous to spatial sensitivity, temporal sensitivity is often described by the temporal contrast sensitivity function (tCSF) and detection at each temporal frequency. CFF thresholds then define the temporal resolution of the visual system. While Davis et al. [2015] recently showed that retinal image decomposition due to saccades in specific viewing conditions can elicit a flicker sensation beyond this limit, in our work we model only fixated sensitivities and highlight that further modeling would be needed to incorporate gaze-related effects (see Sec. 7).

While related, CFF and tCSF are used by vision scientists to quantify slightly different aspects of human vision. The CFF is considered to be a measure of conscious visual detection dependent on the temporal resolution of visual neurons, since at the CFF threshold, an identical flickering stimulus varies in percept from flickering to stable [Carmel et al. 2006]. While contrast sensitivity is considered to be more of a measure of visual discrimination, as evidenced by the sinusoidal grating used in its measure requiring orientation recognition [Pelli and Bex 2013].

Several studies measure datapoints for these functions independently [de Lange Dzn 1958; Eisen-Enosh et al. 2017; Van Nes and Bouman 1967], across eccentricity [Hartmann et al. 1979; Koenderink et al. 1978a], as a function of luminance [Koenderink et al. 1978b] and color [Davis et al. 2015], however few present quantitative models, possibly due to sparse sampling of the measured data. Additionally, data is often captured at luminances very different from typical VR displays. The Ferry-Porter [Tyler and Hamer 1990] and Granit-Harper [Rovamo and Raninen 1988] laws are exceptions to this, describing CFF as increasing linearly with log retinal illuminance and log stimulus area, respectively. Subsequent work by Tyler et al. [1993] showed that the Ferry-Porter law also extends to higher eccentricities. Koenderink et al. [1978] also derived a complex model for sCSF across eccentricity for arbitrary spatial patterns.

However, spatial and temporal sensitivity are not independent. This relationship can be captured by varying both spatial and temporal frequency to obtain the spatio-temporal contrast sensitivity function (stCSF) [Robson 1966]. While Kelly [1979] was the first to fit an analytical function to stCSF data, this was limited to the fovea and a single luminance. Similarly, Watson et al. [2016] devised the pyramid of visibility, a simplified model that can be used if only higher frequencies are relevant. While also modeling luminance dependence, this model was again not applicable to higher eccentricities. Subsequent work saw the refitting of this model for higher

eccentricity data but the original equation was simplified to only consider stationary content and the prediction of sCSF [Watson 2018].

To the best of our knowledge, however, there is no unified model parameterized by all axes of interest, as illustrated in Table 1. With this work, we address this gap, measuring and fitting the first model that describes CFF in terms of spatial frequency, eccentricity and luminance. In particular, we measure CFF rather than sCSF to be more conservative in modeling the upper limit of human perception, for use in foveated graphics applications.

2.2 Foveated Graphics

Foveated graphics encompasses a plethora of techniques that is enabled by eye tracking (see e.g. [Duchowski et al. 2004; Koulieris et al. 2019] for surveys on this topic). Knowing where the user fixates allows for manipulation of the visual data to improve perceptual realism [Krajancich et al. 2020; Mauderer et al. 2014; Padmanaban et al. 2017], or save bandwidth [Albert et al. 2017; Arabadzhyska et al. 2017], by exploiting the drastically varying acuity of human vision over the retina. The most prominent approach that does the latter is perhaps foveated rendering, where images and videos are rendered, transmitted, or displayed with spatially varying resolutions without affecting the perceived image quality [Friston et al. 2019; Geisler and Perry 1998; Guenter et al. 2012; Kaplanyan et al. 2019; Patney et al. 2016; Sun et al. 2017]. These methods primarily exploit spatial aspects of eccentricity-dependent human vision, but Tursun et al. [2019] recently expanded this concept by considering local luminance contrast across the retina, Xiao et al. [2020] additionally consider temporal coherence, and Kim et al. [2019] showed hardware-enabled solutions for foveated displays. Related approaches also use gaze-contingent techniques to reduce bit-depth [McCarthy et al. 2004] and shading or level-of-detail [Luebke and Hallen 2001; Murphy and Duchowski 2001; Ohshima et al. 1996] in the periphery. All of these foveated graphics techniques primarily aim at saving bandwidth of the graphics system by reducing the number of vertices or fragments a graphics processing unit has to sample, raytrace, shade, or transmit to the display. To the best of our knowledge, none of these methods exploit the eccentricity-dependent temporal characteristics of human vision, perhaps because no existing model of human perception accounts for these in a principled manner (see Tab. 1). The goal of this paper is to develop such a model for eccentricity-dependent spatio-temporal flicker fusion that could enable significant bandwidth savings for all of the aforementioned techniques well beyond those enabled by existing eccentricity-dependent acuity models.

Many approaches in graphics, such as frame interpolation, temporal upsampling [Chen et al. 2005; Denes et al. 2019; Didyk et al. 2010; Scherzer et al. 2012] and multi-frame rate rendering and display [Springer et al. 2007], do account for the limited temporal resolution of human vision. However, spatial and temporal sensitivities are not independent. Spatio-temporal video manipulation is common in foveated compression (e.g., [Ho et al. 2005; Wang et al. 2003]) while other studies have investigated how trading one for the other affects task performance, such as in first-person shooters [Claypool and Claypool 2007, 2009]. Denes et al. [2020] recently

studied spatio-temporal resolution tradeoffs with perceptual models and applications to VR. However, none of these approaches rely on eccentricity-dependent models of spatio-temporal vision, which could further enhance their performance.

3 ESTIMATING FLICKER FUSION THRESHOLDS

To develop an eccentricity-dependent model of flicker fusion, we need a display that is capable of showing stimuli at a high framerate and over a wide FOV. In this section, we first describe a custom high-speed VR display that we built to support these requirements. We then proceed with a detailed discussion of the user study we conducted and the resulting values for eccentricity and spatial frequency-dependent CFF we estimated.

3.1 Display Prototype

Our prototype display is designed in a near-eye display form factor to support a wide FOV. As shown in Figure 2(a), we removed the back panel of a View-Master Deluxe VR Viewer and mounted a semi-transparent optical diffuser (Edmund Optics #47-679) instead of a display panel, which serves as a projection screen. This View-Master was fixed to an SR Research headrest to allow users to comfortably view stimuli for extended periods of time. To support a sufficiently high framerate, we opted for a Digital Light Projector (DLP) unit (Texas Instruments DLP3010EVM-LC Evaluation Board) that rear-projects images onto the diffuser towards the viewer. A neutral density (ND16) filter was placed in this light path to reduce the brightness to an eye safe level, measured to be 380 cd/m² at peak.

The DLP has a resolution of 1280 × 720, and a maximum frame rate of 1.5 kHz for 1-bit video, 360 Hz for 8-bit monochromatic video, or 120 Hz for 24-bit RGB video. We positioned the projector such that the image matched the size of the conventional View-Master display. Considering the magnification of the lenses, this display provides a pixel pitch of 0.1' (arc minutes) and a monocular FOV of 80° horizontally and 87° vertically.

To display stimuli for our user study, we used the graphical user interface provided by Texas Instruments to program the DLP to the 360 Hz 8-bit grayscale mode, deemed sufficient for our measurements of CFF, which unlike CSF, does not require precise contrast tuning. The DLP was unable to support the inbuilt red, green and blue light-emitting diodes (LEDs) being on simultaneously, so we chose to use a single LED to minimize the possibility of artifacts from temporally multiplexing colors. Furthermore, since the HVS is most sensitive to mid-range wavelengths, the green LED (OSRAM: LE CG Q8WP) of peak wavelength 520 nm and a 100 nm full width at half maximum, was chosen so as to most conservatively measure the CFF thresholds. We used Python's PsychoPy toolbox [Peirce 2007] and a custom shader to stream frames to the display by encoding them into the required 24-bit RGB format that is sent to the DLP via HDMI.

Stimuli. The flicker fusion model we wish to acquire could be parameterized by spatial frequency, rotation angle, eccentricity (i.e., distance from the fovea), direction from the fovea (i.e., temporal, nasal, etc.) and other parameters. A naive approach may sample across all of these dimensions, but due to the fact that each datapoint needs to be recorded for multiple subjects and for many temporal

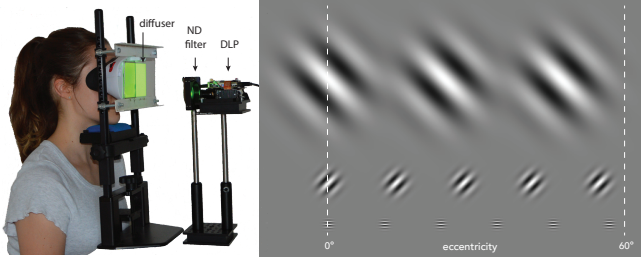


Fig. 2. Our set-up for measuring user CFF thresholds. A photograph of a user in our custom VR display prototype is shown on the left. An ND filter is used to reduce the brightness of a DLP which projects onto a semi-transparent diffuser. On the right we show an illustration of the last 3 orders of test Gabor wavelets shown at a contrast of 1. Eccentricities were chosen to cover a FOV of 60° at each scale. Orientation was chosen randomly from 45°, 90°, 135° and 180° at the point of beginning a QUEST staircase.

frequencies per subject to determine the respective CFFs, this seems infeasible. Therefore, similar to several previous studies [Allen and Hess 1992; Eisen-Enosh et al. 2017; Hartmann et al. 1979], we make the following assumptions to make data acquisition tractable:

- (1) The left and right eyes exhibit the same sensitivities and monocular and binocular viewing conditions are equivalent. Thus, we display the stimuli monocularly to the right eye by blocking the left side of the display.
- (2) Sensitivity is rotationally symmetric around the fovea, i.e. independent of nasal, temporal, superior, and inferior direction, thus being only a function of absolute distance from the fovea. It is therefore sufficient to measure stimuli only along the temporal direction starting from the fovea.
- (3) Sensitivity is orientation independent. Thus, the rotation angle of the test pattern is not significant.

These assumptions allow us to reduce the sample space to only two dimensions: eccentricity e and spatial frequency f_s . Later we also analyze retinal illuminance I as an additional factor.

Another fact to consider is that an eccentricity-dependent model that varies with spatial frequency must adhere to the uncertainty principle. That is, low spatial frequencies cannot be well localized in eccentricity. For example, the lowest spatial frequency of 0 cpd is a stimulus that is constant across the entire retina whereas very high spatial frequencies can be well localized in eccentricity. This behavior is appropriately modeled by wavelets. As such, we select our stimuli to be a set of 2D Gabor wavelets. As a complex sinusoid modulated by a Gaussian envelope, these wavelets are defined by the form:

$$g(\mathbf{x}, \mathbf{x}_0, \theta, \sigma, f_s) = \exp\left(\frac{-\|\mathbf{x} - \mathbf{x}_0\|^2}{2\sigma^2}\right) \cos(2\pi f_s \mathbf{x} \cdot [\cos \theta, \sin \theta]), \quad (1)$$

where $\mathbf{x}$ denotes the spatial location on the display, $\mathbf{x}_0$ is the center of the wavelet, σ is the standard deviation of the Gaussian in visual degrees, and f_s and θ are the spatial frequency in cpd and angular orientation in degrees for the sinusoidal grating function.

Table 2. Parameters of 18 test Gabor wavelets. We define 6 orders by spatial frequency (and radius) of stimuli. The number and eccentricity locations per order were chosen based on radius to uniformly sample the available eccentricity range. f_s : spatial frequency, σ : wavelet standard deviation, e : eccentricity. (*) Note that in practice the extent was limited by our display FOV and $f_s = 0.0055$ cpd is used for the analysis in Sec. 4.1.

Order	f_s (cpd)	σ (°)	e (°)
0	0.000(*)	inf(*)	0.0
1	0.011	63.0	0.0
2	0.041	17.2	0.0, 19.2
3	0.154	4.6	0.0, 24.5, 48.2
4	0.571	1.2	0.0, 14.8, 29.2, 42.7, 55.0
5	2.000	0.5	0.0, 12.3, 24.4, 35.9, 46.8, 56.8

We provide details of the conversion between pixels and eccentricity e in the Supplement. Of note however, is that we use a scaled and shifted version of this form to be suitable for display, namely $0.5 + 0.5 g(\mathbf{x}, \mathbf{x}_0, \theta, \sigma, f_s)$, such that the pattern modulates between 0 and 1, with an average gray level of 0.5. The resulting stimulus exhibits three clearly visible peaks and smoothly blends into the uniform gray field which covers the entire field of view of our display and ends sharply at its edge with a dark background.

This choice of Gabor wavelet was motivated by many previous works in vision science, including the standard measurement procedure for the CSF [Pelli and Bex 2013] (see Sec. 2), and image processing [Kamarainen et al. 2006; Weldon et al. 1996], where Gabor functions are frequently used for their resemblance to neural activations in human vision. For example, it has been shown that 2D Gabor functions are appropriate models of the transfer function of simple cells in the visual cortex of mammalian brains and thus mimicking the early layers of human visual perception [Daugman 1985; Marčelja 1980].

As a tradeoff between sampling the parameter space as densely as possible while keeping our user studies to a reasonable length, we converged on using 18 unique test stimuli, as listed in Table 2. We sampled eccentricities ranging from 0° to almost 60°, moving the fixation point into the nasal direction with increasing eccentricity, such that the target stimulus is affected as little as possible by the lens distortion. We chose not to utilize the full 80° horizontal FOV of our display due to lens distortion becoming too severe in the last 10°. We chose to test 6 different spatial frequencies, with the highest being limited to 2 cpd due to the lack of a commercial display with both high enough spatial and temporal resolution. However it should be noted that the acuity of human vision is considerably higher; 60 cpd based on peak cone density [Deering 1998] and 40–50 cpd based on empirical data [Guenter et al. 2012; Robson and Graham 1981; Thibos et al. 1987]. The Gaussian windows limiting the extent of the Gabor wavelets are scaled according to spatial frequency, i.e., $\sigma = 0.7/f_s$ such that each stimulus exhibits the same number of periods, defining the 6 wavelet orders. Finally, eccentricity values were chosen based on the radius of the wavelet order to uniformly sample the available eccentricity range, as illustrated in Figure 2(b).

The wavelets were temporally modulated by sinusoidally varying the contrast from $[-1, 1]$ and added to the background gray level of 0.5. In this way, at high temporal frequencies the Gabor wavelet would appear to fade into the background. The control stimulus was modulated at 180 Hz, which is far above the CFF observed for all of the conditions.

3.2 User Study

Participants. Nine adults participated (age range 18–53, 4 female). Due to the demanding nature of our psychophysical experiment, only a few subjects were recruited, which is common for similar low-level psychophysics (see e.g. Patney et al. [2016]). Furthermore, previous work measuring the CFF of 103 subjects found a low variance [Romero-Gómez et al. 2007], suggesting sufficiency of a small sample size. All subjects in this and the subsequent experiment had normal or corrected-to-normal vision, no history of visual deficiency, and no color blindness, but were not tested for peripheral-specific abnormalities. All subjects gave informed consent. The research protocol was approved by the Institutional Review Board at the host institution.

Procedure. To start the session, each subject was instructed to position their chin on the headrest such that several concentric circles centered on the right side of the display were minimally distorted (due to the lenses). The threshold for each Gabor wavelet was then estimated in a random order with a two-alternative forced-choice (2AFC) adaptive staircase designed using QUEST [Watson and Pelli 1983]. The orientation of each Gabor patch was chosen randomly at the beginning of each staircase from 0° , 45° , 90° and 135° . At each step, the subject was shown a small (1°) white cross for 1.5 s to indicate where they should fixate, followed by the test and control stimuli in a random order, each for 1 s. For stimuli at 0° eccentricity, the fixation cross was removed after the initial display so as not to interfere with the pattern. The screen was momentarily blanked to a slightly darker gray level than the gray background to indicate stimuli switching. The subject was then asked to use a keyboard to indicate which of the two randomly ordered patterns (1 or 2) exhibited more flicker. The ability to replay any trial was also given via key press and the subjects were encouraged to take breaks at their convenience. Each of the 18 stimuli were tested once per user, taking approximately 90 minutes to complete.

Results. Mean CFF thresholds across subjects along with the standard error (vertical bars) and extent of the corresponding stimulus (horizontal bars) are shown in Figure 3. The table of measured values is available in the Supplemental Material.

The measured CFF values have a maximum above 90 Hz. This relatively large magnitude can be explained by the Ferry–Porter law [Tyler and Hamer 1990] and the high adaptation luminance of our display. Similarly large values have previously been observed in corresponding conditions [Tyler and Hamer 1993].

As expected, the CFF reaches its maximum for the lowest f_s values. This trend follows the Granit–Harper law predicting a linear increase of CFF with stimuli area [Hartmann et al. 1979].

For higher f_s stimuli, we observe an increase of the CFF from the fovea towards a peak between 10° and 30° of eccentricity. Similar

trends have been observed by Hartmann et al. [1979], including the apparent shift of the peak position towards fovea with increasing f_s and decreasing stimuli size.

Finally, our subjects had difficulty to detect flicker for the two largest eccentricity levels for the maximum $f_s = 2$ cpd. This is predictable as acuity drops close to or below this value for such extreme retinal displacements [Geisler and Perry 1998].

4 AN ECCENTRICITY-DEPENDENT PERCEPTUAL MODEL FOR SPATIO-TEMPORAL FLICKER FUSION

The measured CFFs establish an envelope of spatio-temporal flicker fusion thresholds at discretely sampled points within the resolution afforded by our display prototype. Practical applications, however, require these thresholds to be predicted continuously for arbitrary spatial frequencies and eccentricities. To this end, we develop a continuous eccentricity-dependent model for spatio-temporal flicker fusion that is fitted to our data. Moreover, we extrapolate this model to include spatial frequencies that are higher than those supported by our display by incorporating existing visual acuity data and we account for variable luminance adaptation levels by adapting the Ferry–Porter law [Tyler and Hamer 1993].

4.1 Model Fitting

Each of our measured data points is parametrized by its spatial frequency f_s , eccentricity e , and CFF value averaged over all subjects. Furthermore, it is associated with a localization uncertainty determined by the radius of its stimulus u . In our design, u is a function of f_s and, for 13.5% peak contrast cut-off, we define $u = 2\sigma$ where $\sigma = 0.7/f_s$ to be the standard deviation of our Gabor patches.

We formulate our model as

$$\begin{aligned} \Psi(e, f_s) &= \max(0, p_0 \tau(f_s)^2 + p_1 \tau(f_s) + p_2 \\ &\quad + (p_3 \tau(f_s)^2 + p_4 \tau(f_s) + p_5) \cdot \zeta(f_s) e \\ &\quad + (p_6 \tau(f_s)^2 + p_7 \tau(f_s) + p_8) \cdot \zeta(f_s) e^2) \\ \zeta(f_s) &= \exp(p_9 \tau(f_s)) - 1 \\ \tau(f_s) &= \max(\log_{10} f_s - \log_{10} f_{s_0}, 0), \end{aligned} \quad (2)$$

where $\mathbf{p} = [p_0, \dots, p_9] \in \mathbb{R}^{10}$ are the model parameters (see Table 3), $\zeta(f_s)$ restricts eccentricity effects for small f_s and $\tau(f_s)$ offsets logarithmic f_s relative to our constant function cut-off.

We build on three domain-specific observations to find a continuous CFF model $\Psi(e, f_s) : \mathbb{R}^2 \rightarrow \mathbb{R}$ that fits our measurements.

First, both our measurements and prior work indicate that the peak CFF is located in periphery, typically between 20° and 50° of eccentricity [Hartmann et al. 1979; Rovamo and Raninen 1984;

Table 3. Parameters $p_0 \dots p_9$ for our model fitted for conservative (cons.) and relaxed (rel.) assumptions as well as the full modeled extended using acuity data. The degrees-of-freedom adjusted R^2 shows the fit quality.

Model	Parameters										R^2
	p_0	p_1	p_2	p_3	p_4	p_5	p_6	p_7	p_8	p_9	
Cons.	94.4	-10.1	-4.08	0.431	-0.280	0.0484	-0.00912	0.00679	-0.00140	1.56	0.889
Rel.	94.3	-10.1	-4.06	0.430	-0.282	0.0464	-0.00929	0.00672	-0.00129	1.58	0.938
Full	94.3	-10.1	-4.06	0.435	-0.281	0.0440	-0.00877	0.00613	-0.00111	1.58	0.938

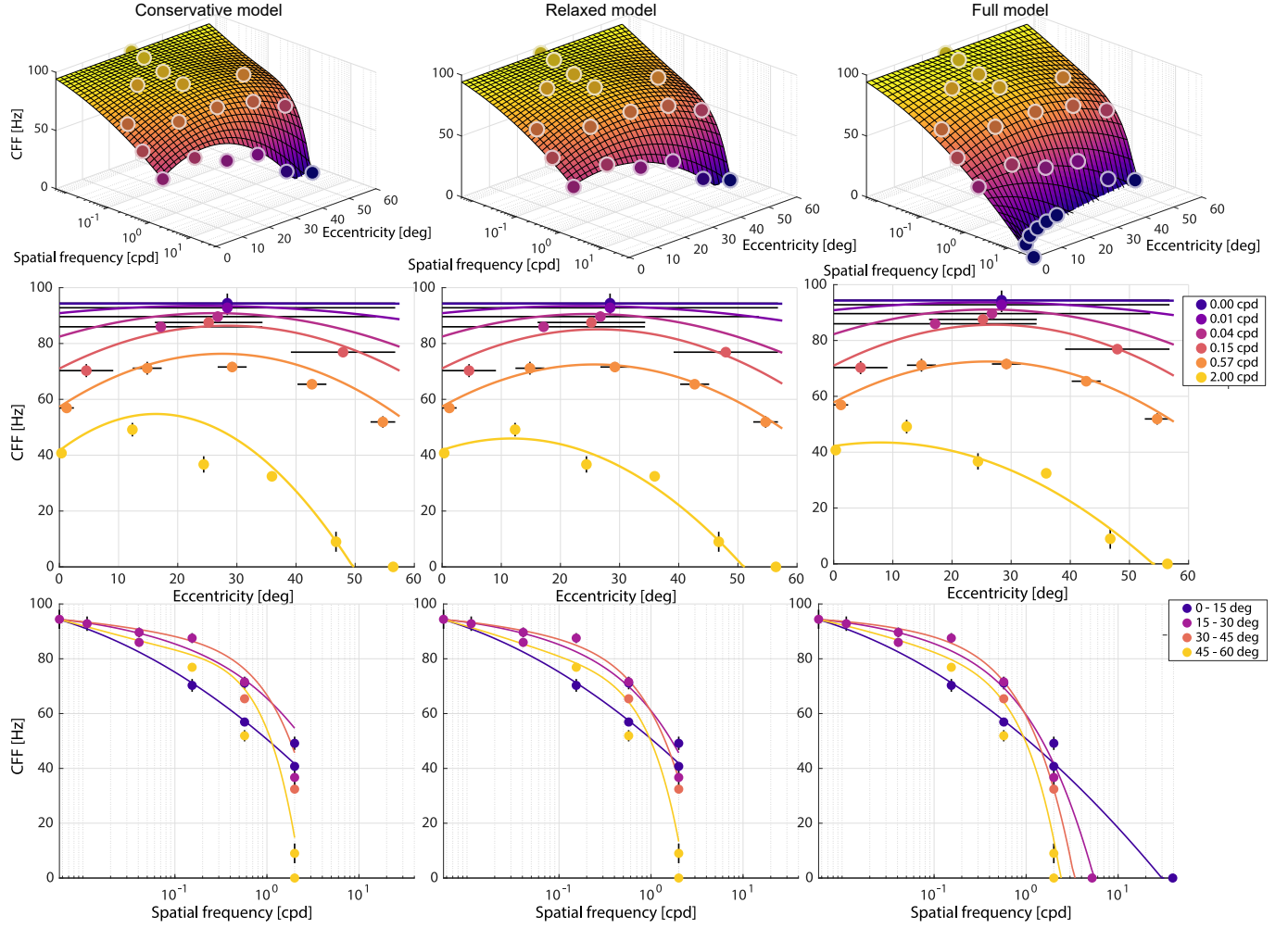


Fig. 3. Three different parameter fits for our flicker fusion model $\Psi(e, f_s)$ (columns). Orthographic views of the eccentricity–CFF (second row) and spatial frequency–CFF (third row) planes show subject-averaged measured points along with their variances (standard error of means, vertical bars). The horizontal bars in the eccentricity dimension denote spatial extents of respective stimuli. The curves represent corresponding cross sections of the fitted models.

Tyler 1987]. For both fovea and far periphery the CFF drops again forming a convex shape which we model as a quadratic function of e .

Second, because the stimuli with very low f_s are not spatially localized, their CFF does not vary with e . Consequently, we enforce the dependency on e to converge to a constant function for any f_s below $f_{s_0} = 0.0055$ cpd. This corresponds to half reciprocal of the full-screen stimuli visual field coverage given our display dimensions.

Finally, following common practices in modeling the effect of spatial frequencies on visual effects, such as contrast [Koenderink et al. 1978c] or disparity sensitivities [Bradshaw and Rogers 1999], we fit the model for logarithmic f_s .

Before parameter optimization, we need to consider the effect of eccentricity uncertainty. The subjects in our study detected flicker regardless of its location within the stimuli extent $\mathbf{m} = [e \pm u]$ deg.

Therefore, Ψ achieves its maximum within $\mathbf{m}$ and is upper-bounded by our measured flicker frequency $f_t \in \mathbb{R}$. At the same time, nothing can be claimed about the variation of Ψ within $\mathbf{m}$ and therefore, in absence of further evidence, a conservative model has to assume that Ψ is not lower than f_t within $\mathbf{m}$. These two considerations delimit a piece-wise constant Ψ . In practice, based on previous work [Hartmann et al. 1979; Tyler 1987] it is reasonable to assume that Ψ follows a smooth trend over the retina and its value is lower than f_t value at almost all eccentricities within $\mathbf{m}$.

Consequently, in Table 3 we provide two different fits for the parameters. Our conservative model strictly follows the restrictions from the measurement and tends to overestimate the range of visible flicker frequencies which prevents discarding potentially visible signal. Alternatively, our relaxed model follows the smoothness assumption and applies the measured values as upper bound.

To fit the parameters, we used the Adam solver in PyTorch initialized by the Levenberg–Marquardt algorithm and we minimized the mean-square prediction error over all extents $\mathbf{m}$. The additional constraints were implemented as soft linear penalties. To leverage data points with immeasurable CFF values, we additionally force $\Psi(e, f_s) = 0$ at these points. This encodes imperceptibility of their flicker at any temporal frequency.

Figure 3 shows that the fitted Ψ represents the expected effects well. The eccentricity curves (row 2) flatten for low f_s and their peaks shift to lower e for large f_s . The conservative fit generally yields larger CFF predictions though it does not strictly adhere to the stimuli extents due to other constraints.

4.2 Extension for High Spatial Frequencies

Due to technical constraints, the highest f_s measured was 2 cpd. At the same time, the acuity of human vision has an upper limit of 60 cpd based on peak cone density [Deering 1998] and 40–50 cpd based on empirical data [Guenter et al. 2012; Robson and Graham 1981; Thibos et al. 1987]. To minimize this gap and to generalize our model to other display designs we extrapolate the CFF at higher spatial frequencies using existing models of spatial acuity.

For this purpose, we utilize the acuity model of Geisler and Perry [1998]. It predicts acuity limit A for e as

$$A(e) = \ln(64) \frac{2.3}{0.106 \cdot (e + 2.3)}, \quad (3)$$

with parameters fitted to measurements of Robson and Graham [1981]. Their study of pattern detection rather than resolution is well aligned with our own study design and conservative visual performance assessment. Similarly, their bright adaptation luminance of 500 cd/m^2 is also close to our display.

$A(e)$ predicts limit of spatial perception. We reason that at this absolute limit flicker is not detectable and, therefore, the CFF is not defined. We represent this situation by zero CFF values in the same way as for imperceptible stimuli in our study and force our model to satisfy $\Psi(e, A(e)) = 0$.

In combination with our relaxed constraints we obtain our final full model as shown in Figure 3. It follows the same trends as our original model within the bounds of our measurement space and intersects the zero plane at the projection of $A(e)$.

4.3 Adaptation luminance

Our experiments were conducted at half of our display peak luminance $L = 380 \text{ cd}/\text{m}^2$. This is relatively bright compared to the 50–200 cd/m^2 luminance setting of common VR systems [Mehrfard et al. 2019]. Consequently, our estimates of the CFF are conservative because the Ferry–Porter law predicts the CFF to increase linearly with logarithmic levels of retinal illuminance [Tyler and Hamer 1993]. While the linear relationship is known, the actual slope and intercept varies with retinal eccentricity [Tyler and Hamer 1993]. For this reason, we measured selected points from our main experiment for two other display luminance levels.

Four of the subjects from experiment 1 performed the same procedure for a subset of conditions with a display modified first with one and then two additional ND8 filters yielding effective luminance

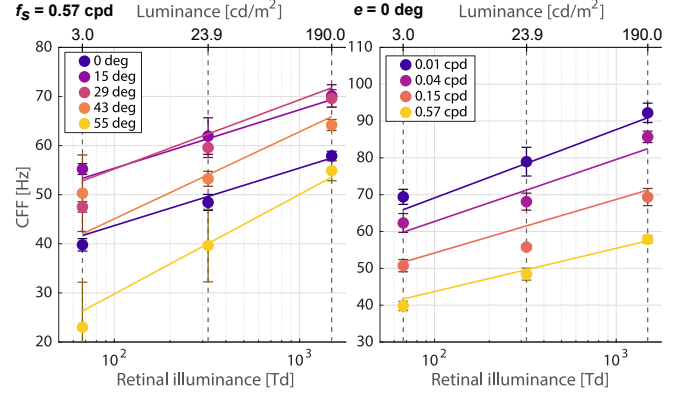


Fig. 4. Luminance scaling of our full model fitted for $l_0 = 1488 \text{ Td}$. The points represent mean measured CFF values and the lines the prediction of the transformed model. Vertical bars are standard errors. The left plot shows varying slopes across e (with $f_s = 0.57 \text{ cpd}$). The right plot shows varying intercepts across f_s (with $e = 0 \text{ deg}$).

values of 23.9 and 3.0 cd/m^2 . We then applied a formula of Stanley and Davies [1995] to compute the pupil diameter under these conditions as

$$d(L) = 7.75 - 5.75 \left(\frac{(La/846)^{0.41}}{(La/846)^{0.41} + 2} \right), \quad (4)$$

where $a = 80 \times 87 = 6960 \text{ deg}^2$ is the adapting area of our display. This allows us to derive corresponding retinal illuminance levels for our experiments using $l(L) = \pi d(L)^2 / 4 \cdot L$ as 67.3, 321 and 1488 Td and obtain a linear transformation of our original model $\Psi(e, f_s)$ to account for L with an eccentricity-dependent slope as

$$\hat{\Psi}(e, f_s, L) = (s(e, f_s) \cdot (\log_{10}(l(L)/l_0)) + 1) \Psi(e, f_s) \quad (5)$$

$$s(e, f_s) = \zeta(f_s)(q_0 e^2 + q_1 e) + q_2 \quad (6)$$

where $l_0 = 1488 \text{ Td}$ is our reference retinal illuminance, $\zeta(f_s)$ encodes localization uncertainty for low f_s as in Equation 3 and $\mathbf{q} = [5.71 \cdot 10^{-6}, -1.78 \cdot 10^{-4}, 0.204]$ are parameters obtained by a fit with our full model.

Figure 4 shows that this eccentricity-driven model of Ferry–Porter luminance scaling models not only the slope variation over the retina but also the sensitivity difference over range of f_s well (degree-of-freedom-adjusted $R^2 = 0.950$).

5 EXPERIMENTAL VALIDATION

Our eccentricity-dependent spatio-temporal model is unique in that it allows us to predict what temporal information may be imperceptible for a certain eccentricity and spatial frequency. One possible application for such a model is in the development of new perceptual video quality assessment metrics (VQMs). Used to guide the development of different video codecs, encoders, encoding settings, or transmission variants, such metrics aim to predict subjective video quality as experienced by users. While it is commonly known that many existing metrics, such as peak signal-to-noise ratio (PSNR) and structural similarity index metric (SSIM), do not capture many eccentricity-dependent spatio-temporal aspects of human vision

well, in this section we discuss a user study we conducted which shows that our model could help better differentiate perceivable and non-perceivable spatio-temporal artifacts. Furthermore, we also use the study to test the fit derived in the previous section, showing that our model makes valid predictions for different users and points beyond those measured in our first user study.

The study was conducted using our custom high-speed VR display with 18 participants, 13 of which did not participate in the previous user study. Users were presented with videos, made up of a single image frame perturbed by Gabor wavelet(s) such that when modulated at a frequency above the CFF they become indistinguishable from the background image. Twenty-two unique videos were tested, where the user was asked to rank the quality of the video from 1 (“bad”) to 5 (“excellent”). We then calculate the difference mean opinion score (DMOS) per stimulus, as described in detail by Seshadrinathan et al [2010], and 3 VQMs, namely PSNR, SSIM and one of the most influential metrics used today for traditional contents: the Video Multimethod Assessment Fusion (VMAF) metric developed by Netflix [Li et al. 2016]. The results of which are shown in Figure 5. We acknowledge that neither of these metrics has been designed to specifically capture the effects we measure. This often results in uniform scores across the stimuli range (see e.g. Figure 5 III). We include them into our comparison to illustrate these limits and the need for future research in this direction.

The last row shows a comparison of all measured DMOS values with each of the standard VQMs and our own flicker quality predictor. This is a simple binary metric that assigns videos outside of the CFF volume where flicker is invisible a good label while it assigns a bad label to others. Our full $\Psi(e, f_s)$ model was used for this purpose. We computed Pearson Linear Correlation Coefficients (PLCC) between DMOS and each of the metrics while taking into account their semantics. Only our method exhibits statistically significant correlation with the user scores ($r(20) = 0.988, p < 0.001$ from a two-tail t-test).

Further, we broke our stimuli into 5 groups to test the ability of our model to predict more granular effects.

- (1) *Unseen stimulus.* We select $e = 30^\circ$ and $f_s = 1.00$ cpd, a configuration not used to fit the model, and show that users rate the video to be of poorer quality (lower DMOS) when it is modulated at a frequency lower than the predicted CFF (see panel I in Fig. 5).
- (2) *f_s dependence.* We test whether our model captures spatio-temporal dependence. Lowering f_s with fixed radius makes the flicker visible, despite not changing f_t (panel II).
- (3) *e change.* By moving the fixation point towards the outer edges of the screen, thereby increasing e , we show that our model captures and predicts the higher temporal sensitivity of the mid-periphery. Furthermore, by moving the fixation point in both nasal and temporal directions, we validate the assumed symmetry of the CFF. As a result, users were able to notice flicker that was previously imperceptible in the fovea (panel III).
- (4) *Direction independence.* We confirm our assumption that sensitivity is approximately rotationally symmetric around the

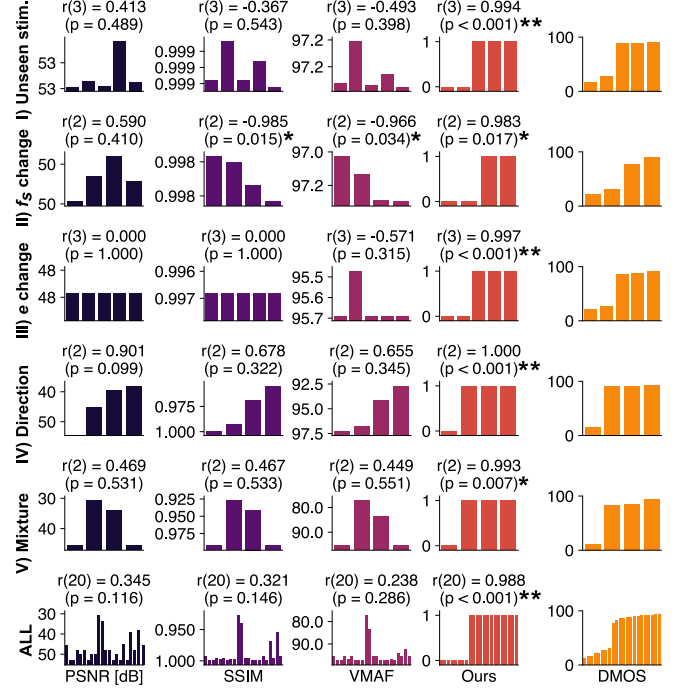


Fig. 5. Results of our validation study. Each row represents one of the 5 tested effects along with a cumulative plot in the last row. Results of different metrics (vertical columns) are compared to DMOS (right column). Each bar represents the score for one stimulus rescaled so that a higher bar indicates a better quality. PLCC of each metric and DMOS are shown along the p-values (t-test). (*) and (**) mark statistical significance at $p = 0.05$ and $p = 0.001$ levels respectively.

fovea. A single visible stimulus from Group 1 was also added for the analysis as a control (panel IV).

- (5) *Mixed stimuli.* We show that the model predictions hold up even with concurrent viewing of several stimuli of different sizes, eccentricities and spatial frequencies (panel V).

The perturbation applied to the videos in Groups 1–3 only varied in modulation frequency or eccentricity with respect to a moving gaze position. In response, the dependency of other metrics on f_t is on the level of noise, since changes related to variation of f_t or f_s of a fixed-size wavelet average out over several frames, and moving the gaze position does not change the frames at all. Our metric is able to capture the perceptual effect, maintaining significant correlation with DMOS ($r(3) = 0.994, p = 0.001, r(2) = 0.983, p < 0.05$ and $r(3) = 0.977, p < 0.001$ respectively).

Groups 4 and 5 were designed to be accumulating sets of Gabor wavelets, but all with parameters that our model predicts as imperceptible. We observe that VQMs drop with increasing total distortion area, while the DMOS scores stay relatively constant and significantly correlated with our metric ($r(2) = 1.000, p < 0.001$ and $r(3) = 0.993, p < 0.01$ respectively). These results indicate that existing compression schemes may be able to utilize more degrees of freedom to achieve higher compression, without seeing a drop in the typical metrics used to track perceived visual quality. The

exact list of chosen parameters, along with the CFFs predicted by our model are listed in the supplemental material.

In conclusion, the study shows that our metric predicts visibility of temporal flicker for independently varying spatial frequencies and retinal eccentricities. While existing image and video quality metrics are important for predicting other visual artifacts, our method is a novel addition in this space that significantly improves perceptual validity of quality predictions for temporally varying content. In this way, we also show that our model may enable existing compression schemes to utilize more degrees of freedom to achieve higher compression, without compromising several conventional VQMs.

6 ANALYZING BANDWIDTH CONSIDERATIONS

Our eccentricity-dependent spatio-temporal CFF model defines the gamut of visual signals perceivable to human vision. Hence any signal outside of this gamut can be removed to save bandwidth in computation or data transmission. In this section we provide a theoretical analysis of the compression gain factors this model may potentially enable for foveated graphics applications when used to allocate resources such as bandwidth. In practice, developing foveated rendering and compression algorithms holds other nuanced challenges and thus the reported gains represent an upper bound.

To efficiently apply our model, a video signal should be described by a decomposition into spatial frequencies, retinal eccentricities, and temporal frequencies. We consider discrete wavelet decompositions (DWT) as a suitable candidate because these naturally decompose a signal into eccentricity-localized spatio-temporal frequency bands. Additionally, the multitude of filter scales in the hierarchical decomposition closely resembles scale orders in our study stimuli.

For the purpose of this thought experiment, we use a biorthogonal Haar wavelet, which, for a signal of length N , results in $\log_2(N)$ hierarchical planes of recursively halving lengths. The total number of resulting coefficients is the same as the input size. From there our baseline is retaining the entire original set of coefficients yielding compression gain of 1.

For traditional spatial-only foveation, we process each frame independently and after a 2D DWT we remove coefficients outside of the acuity limit. We compute eccentricity assuming gaze in the center of the screen and reject coefficients for which $\Psi(e, f_s) = 0$. From our model definition, such signals cannot be perceived regardless of f_t as they lie outside of the vision acuity gamut. The resulting compression gain compared to the baseline can be seen for various display configurations in Figure 6.

Next, for our model we follow the same procedure but we additionally decompose each coefficient of the spatial DWT using 1D temporal DWT. This yields an additional set of temporal coefficients f_t and we discard all values with $\Psi(e, f_s) < f_t$.

For this experiment, we assume a screen with the same peak luminance of 380 cd/m^2 as our display prototype and a pixel density of 60 pixel per degree (peak $f_s = 30 \text{ cpd}$) for approximately retinal display or a fifth of that for a display closer to existing VR systems. The tested aspect ratio of 165:135 approximates the human binocular visual field [Ruch and Fulton 1960]. We consider displays with a maximum framerate of 100 Hz (peak reproducible $f_t = 50 \text{ Hz}$) and

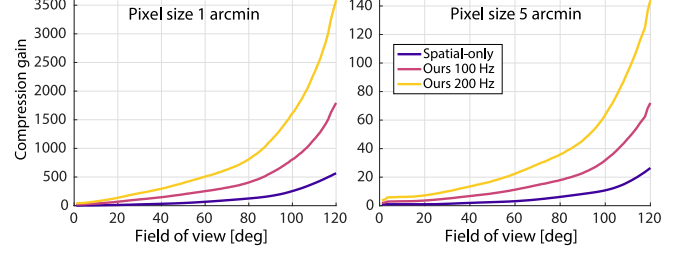


Fig. 6. Compression gain analysis as a fraction of original and retained decomposition coefficients depending on the screen pixel density (left and right plots) and the covered visual field (horizontal axis). Note the difference in the gain axes scales.

200 Hz, capable of reaching the maximum CFF in our model. The conventional spatial-only foveation algorithm is not affected by this choice.

Figure 6 shows that this significantly improves compression for both small and large field-of-view displays regardless of pixel density. This is because the spatial-only compression needs to retain all temporal coefficients in order to prevent the worst-scenario flicker at e and f_s with highest CFF. Therefore, it only discards signal components in the horizontal $e \times f_s$ plane. On the other hand, our scheme allows to also discard signal in the third temporal dimension closely copying the shape of the Ψ which allows to reduce retained f_t for both fovea and far periphery in particular for high f_s coefficients which require the largest bandwidth.

7 DISCUSSION

The experimental data we measure and the models we fit to them further our understanding of human perception and lay the foundation of future spatio-temporal foveated graphics techniques. Yet, several important questions remain to be discussed.

Limitations and Future Work. First, our data and model work with CFF, not the CSF. CFF models the temporal thresholds at which a visual stimulus is predicted to become (in)visible. Unlike the CSF, however, this binary value of visible/invisible does not model relative sensitivity. This makes it not straightforward to apply the CFF directly as an error metric, for example to optimize foveated rendering or compression algorithms. Obtaining an eccentricity-dependent model for the spatio-temporal CSF is an important goal for future work, yet it requires a dense sampling of the full three-dimensional space spanned by eccentricity, spatial frequency, and time with user studies. The CFF, on the other hand only requires us to determine a single threshold for the two-dimensional space of eccentricity and spatial frequency. Considering that obtaining all 18 sampling points in our 2D space for a single user already takes about 90 minutes, motivates the development of a more scalable approach to obtaining the required data in the future.

Although we validate the linear trends of luminance-dependent behavior of the CFF predicted by the Ferry-Porter law, it would be desirable to record more users for larger luminance ranges and more densely spaced points in the 2D sampling space. Again, this would require a significantly larger amount of user studies, which

seems outside the scope of this work. Similarly, we extrapolate our model for f_s above the 2 cpd limit of our display. Future work may utilize advancements in display technology or additional optics to confirm the model extension. We also map the spatio-temporal thresholds only for monocular viewing conditions. Some studies have suggested that such visual constraint lowers the measured CFF in comparison to binocular viewing [Ali and Amir 1991; Eisen-Enosh et al. 2020]. Further work is required to confirm the significance of such an effect in the context of displays.

It should also be noted that our model assumes a specific fixation point. We did not account for the dynamics of ocular motion, which may be interesting for future explorations. By simplifying temporal perception to be equivalent to the temporal resolution of the visual system we also do not account for effects caused by our pattern sensitive system, such as our ability to detect spatio-temporal changes like local movements or deformations. Similarly, we do not model the effects of crowding in the periphery, which has been shown to dominate spatial resolution in pattern recognition [Rosenholtz 2016]. Finally, all of our data is measured for the green color channel of our display and our model assumes rotational invariance as well as orientation independence of the stimulus. A more nuanced model that studies variations in these dimensions may be valuable future work. All these considerations should be taken into account when developing practical compression algorithms, which is a task beyond the scope of this paper.

Finally, we validate that existing error metrics, such as PSNR, SSIM, and VMAF, do not adequately model the temporal aspects of human vision. The correlation between our CFF data and the DMOS of human observers is significantly stronger, motivating error metrics that better model these temporal aspects. Yet, we do not develop such an error metric nor do we propose practical foveated compression or rendering schemes that directly exploit our CFF data. These investigations are directions of future research.

Conclusion. At the convergence of applied vision science, computer graphics, and wearable computing systems design, foveated graphics techniques will play an increasingly important role in emerging augmented and virtual reality systems. With our work, we hope to contribute a valuable foundation to this field that helps spur follow-on work exploiting the particular characteristics of human vision we model.

ACKNOWLEDGMENTS

B.K. was supported by a Stanford Knight-Hennessy Fellowship. G.W. was supported by an Okawa Research Grant, a Sloan Fellowship, and a PECASE by the ARO. Other funding for the project was provided by NSF (award numbers 1553333 and 1839974). The authors would also like to thank Brian Wandell, Anthony Norcia, and Joyce Farrell for advising on temporal mechanisms of the human visual system, Keith Winstein for expertise used in the development of the validation study, Darryl Krajancich for constructing the apparatus for our custom VR display, and Yifan (Evan) Peng for assisting with the measurement of the display luminance.

REFERENCES

- Rachel Albert, Anjul Patney, David Luebke, and Joohwan Kim. 2017. Latency requirements for foveated rendering in virtual reality. *ACM Transactions on Applied Perception (TAP)* 14, 4 (2017), 1–13.
- MR Ali and T Amir. 1991. Critical flicker frequency under monocular and binocular conditions. *Perceptual and motor skills* 72, 2 (1991), 383–386.
- Dale Allen and Robert F Hess. 1992. Is the visual field temporally homogeneous? *Vision research* 32, 6 (1992), 1075–1084.
- Elena Arabadzhiyska, Okan Tarhan Tursun, Karol Myszkowski, Hans-Peter Seidel, and Piotr Didyk. 2017. Saccade landing position prediction for gaze-contingent rendering. *ACM Transactions on Graphics (TOG)* 36, 4 (2017), 1–12.
- Mark F Bradshaw and Brian J Rogers. 1999. Sensitivity to horizontal and vertical corrugations defined by binocular disparity. *Vision research* 39, 18 (1999), 3049–3056.
- David Carmel, Nilli Lavie, and Geraint Rees. 2006. Conscious awareness of flicker in humans involves frontal and parietal cortex. *Current biology* 16, 9 (2006), 907–911.
- Hanfeng Chen, Sung-Soo Kim, Sung-Hee Lee, Oh-Jae Kwon, and Jun-Ho Sung. 2005. Nonlinearity compensated smooth frame insertion for motion-blur reduction in LCD. In *2005 IEEE 7th Workshop on Multimedia Signal Processing*. IEEE, 1–4.
- Kajal T. Claypool and Mark Claypool. 2007. On frame rate and player performance in first person shooter games. *Multimedia Systems* 13, 1 (Sept. 2007), 3–17.
- Mark Claypool and Kajal Claypool. 2009. Perspectives, Frame Rates and Resolutions: It's All in the Game. In *Proc. Int. Conference on Foundations of Digital Games*. Association for Computing Machinery, New York, NY, USA, 42–49.
- Christine A Curcio and Kimberly A Allen. 1990. Topography of ganglion cells in human retina. *Journal of comparative Neurology* 300, 1 (1990), 5–25.
- John G Daugman. 1985. Uncertainty relation for resolution in space, spatial frequency, and orientation optimized by two-dimensional visual cortical filters. *JOSA A* 2, 7 (1985), 1160–1169.
- James Davis, Yi-Hsuan Hsieh, and Hung-Chi Lee. 2015. Humans perceive flicker artifacts at 500 Hz. *Scientific Reports* 5, 1 (Feb. 2015), 7861.
- H de Lange Dzn. 1958. Research into the dynamic nature of the human fovea→cortex systems with intermittent and modulated light. I. Attenuation characteristics with white and colored light. *Josa* 48, 11 (1958), 777–784.
- Michael F Deering. 1998. The limits of human vision. In *2nd International Immersive Projection Technology Workshop*, Vol. 2. 1.
- Gyorgy Denes, Akshay Jindal, Aliaksei Mikhailiuk, and Rafal K. Mantiuk. 2020. A perceptual model of motion quality for rendering with adaptive refresh-rate and resolution. *ACM Transactions on Graphics* 39, 4 (July 2020), 133:133:1–133:133:17.
- Gyorgy Denes, Kuba Maruszczyk, George Ash, and Rafal K. Mantiuk. 2019. Temporal Resolution Multiplexing: Exploiting the limitations of spatio-temporal vision for more efficient VR rendering. *IEEE Transactions on Visualization and Computer Graphics* 25, 5 (May 2019), 2072–2082.
- Piotr Didyk, Elmar Eisemann, Tobias Ritschel, Karol Myszkowski, and Hans-Peter Seidel. 2010. Perceptually-motivated Real-time Temporal Upsampling of 3D Content for High-refresh-rate Displays. *Computer Graphics Forum (Eurographics)* 29, 2 (2010), 713–722.
- Andrew T. Duchowski, Nathan Cournia, and Hunter A. Murphy. 2004. Gaze-Contingent Displays: A Review. *Cyberpsychology & behavior* 7 (2004), 621–634.
- Auria Eisen-Enosh, Nairouz Farah, Zvia Burgansky-Eliash, Idit Maharshak, Uri Polat, and Yossi Mandel. 2020. A dichoptic presentation device and a method for measuring binocular temporal function in the visual system. *Experimental Eye Research* 201 (2020), 108290.
- Auria Eisen-Enosh, Nairouz Farah, Zvia Burgansky-Eliash, Uri Polat, and Yossi Mandel. 2017. Evaluation of critical flicker-fusion frequency measurement methods for the investigation of visual temporal resolution. *Scientific reports* 7, 1 (2017), 1–9.
- Sebastian Friston, Tobias Ritschel, and Anthony Steed. 2019. Perceptual Rasterization for Head-mounted Display Image Synthesis. *ACM Trans. Graph. (Proc. SIGGRAPH 2019)* 38, 4 (2019).
- Wilson S. Geisler and Jeffrey S. Perry. 1998. Real-time foveated multiresolution system for low-bandwidth video communication. In *Human Vision and Electronic Imaging III*, Vol. 3299. International Society for Optics and Photonics, 294–305.
- Brian Guenter, Mark Finch, Steven Drucker, Desney Tan, and John Snyder. 2012. Foveated 3D Graphics. *ACM Trans. Graph. (SIGGRAPH Asia)* (2012).
- E Hartmann, B Lachenmayr, and H Brettel. 1979. The peripheral critical flicker frequency. *Vision Research* 19, 9 (1979), 1019–1023.
- Chia-Chiang Ho, Ja-Ling Wu, and Wen-Huang Cheng. 2005. A practical foveation-based rate-shaping mechanism for MPEG videos. *IEEE transactions on circuits and systems for video technology* 15, 11 (2005), 1365–1372.
- J-K Kamarainen, Ville Kyrki, and Heikki Kalviainen. 2006. Invariance properties of Gabor filter-based features-overview and applications. *IEEE Transactions on image processing* 15, 5 (2006), 1088–1099.
- Anton S. Kaplanyan, Anton Sochenov, Thomas Leimkühler, Mikhail Okunev, Todd Goodall, and Gizem Rufo. 2019. DeepFovea: Neural Reconstruction for Foveated Rendering and Video Compression Using Learned Statistics of Natural Videos. *ACM Trans. Graph.* 38, 6, Article 212 (Nov. 2019), 13 pages.

- D. H. Kelly. 1979. Motion and vision. II. Stabilized spatio-temporal threshold surface. *JOSA* 69, 10 (Oct. 1979), 1340–1349.
- Jonghyun Kim, Youngmo Jeong, Michael Stengel, Kaan Aksit, Rachel Albert, Ben Boudaoud, Trey Greer, Joohwan Kim, Ward Lopes, Zander Majercik, et al. 2019. Foveated AR: dynamically-foveated augmented reality display. *ACM Transactions on Graphics (TOG)* 38, 4 (2019), 1–15.
- JJ Koenderink and AJ Van Doorn. 1978. Visual detection of spatial contrast; influence of location in the visual field, target extent and illuminance level. *Biological Cybernetics* 30, 3 (1978), 157–167.
- Jan J Koenderink, Maarten A Bouman, Albert E Bueno de Mesquita, and Sybe Slappendel. 1978c. Perimetry of contrast detection thresholds of moving spatial sine wave patterns. II. The far peripheral visual field (eccentricity 0° – 50°). *JOSA* 68, 6 (1978), 850–854.
- Jan J Koenderink, Maarten A Bouman, Albert E Bueno de Mesquita, and Sybe Slappendel. 1978b. Perimetry of contrast detection thresholds of moving spatial sine wave patterns. III. The target extent as a sensitivity controlling parameter. *JOSA* 68, 6 (1978), 854–860.
- Jan J Koenderink, Maarten A Bouman, Albert E Bueno de Mesquita, and Sybe Slappendel. 1978a. Perimetry of contrast detection thresholds of moving spatial sine wave patterns. IV. The influence of the mean retinal illuminance. *JOSA* 68, 6 (1978), 860–865.
- George Alex Koulouris, Kaan Aksit, Michael Stengel, Rafal K Mantiuk, Katerina Mania, and Christian Richardt. 2019. Near-eye display and tracking technologies for virtual and augmented reality. In *Computer Graphics Forum*, Vol. 38. Wiley Online Library, 493–519.
- Brooke Krajancich, Petr Kellnhofer, and Gordon Wetzstein. 2020. Optimizing depth perception in virtual and augmented reality through gaze-contingent stereo rendering. *ACM Transactions on Graphics (TOG)* 39, 6 (2020), 1–10.
- Zhi Li, Anne Aaron, Ioannis Katsavounidis, Anush Moorthy, and Megha Manohara. 2016. Toward a practical perceptual video quality metric. *The Netflix Tech Blog* 6, 2 (2016).
- David Luebke and Benjamin Hallen. 2001. Perceptually driven simplification for interactive rendering. In *Proceedings of the 12th Eurographics conference on Rendering (EGWR'01)*. Eurographics Association, Goslar, DEU, 223–234.
- S Marčelja. 1980. Mathematical description of the responses of simple cortical cells. *JOSA* 70, 11 (1980), 1297–1300.
- Michael Mauderer, Simone Conte, Miguel A Nacenta, and Dhanraj Vishwanath. 2014. Depth perception with gaze-contingent depth of field. In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*. 217–226.
- John D. McCarthy, M. Angela Sasse, and Dimitrios Miras. 2004. Sharp or smooth? comparing the effects of quantization vs. frame rate for streamed video. In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems (CHI '04)*. Association for Computing Machinery, New York, NY, USA, 535–542.
- Arian Mehrfard, Javad Fotouhi, Giacomo Taylor, Tess Forster, Nassir Navab, and Bernhard Fuerst. 2019. A comparative analysis of virtual reality head-mounted display systems. *arXiv preprint arXiv:1912.02913* (2019).
- Hunter Murphy and Andrew T. Duchowski. 2001. Gaze-Contingent Level Of Detail Rendering. (2001).
- Toshikazu Ohshima, Hiroyuki Yamamoto, and Hideyuki Tamura. 1996. Gaze-directed adaptive rendering for interacting with virtual space. In *Proceedings of the IEEE 1996 Virtual Reality Annual International Symposium*. IEEE, 103–110.
- Nitish Padmanaban, Robert Konrad, Tal Stramer, Emily A Cooper, and Gordon Wetzstein. 2017. Optimizing virtual reality for all users through gaze-contingent and adaptive focus displays. *Proceedings of the National Academy of Sciences* 114, 9 (2017), 2183–2188.
- Anjul Patney, Marco Salvi, Joohwan Kim, Anton Kaplanyan, Chris Wyman, Nir Bentley, David Luebke, and Aaron Lefohn. 2016. Towards foveated rendering for gaze-tracked virtual reality. *ACM Transactions on Graphics (TOG)* 35, 6 (Nov. 2016), 179:1–179:12.
- Jonathan W Peirce. 2007. PsychoPy—psychophysics software in Python. *Journal of neuroscience methods* 162, 1–2 (2007), 8–13.
- Denis G Pelli and Peter Bex. 2013. Measuring contrast sensitivity. *Vision research* 90 (2013), 10–14.
- JG Robson. 1993. Contrast sensitivity: One hundred years of clinical measurement in Proceedings of the Retina Research Foundation Symposia, Vol. 5 Eds R. Shapley, D. Man-Kit Lam.
- John G Robson. 1966. Spatial and temporal contrast-sensitivity functions of the visual system. *Josa* 56, 8 (1966), 1141–1142.
- J. Gordon Robson and Norma Graham. 1981. Probability summation and regional variation in contrast sensitivity across the visual field. *Vision research* 21, 3 (1981), 409–418.
- Manuel Romero-Gómez, Juan Córdoba, Rodrigo Jover, Juan A Del Olmo, Marta Ramírez, Ramón Rey, Enrique De Madaria, Carmina Montoliu, David Nuñez, Montse Flavia, et al. 2007. Value of the critical flicker frequency in patients with minimal hepatic encephalopathy. *Hepatology* 45, 4 (2007), 879–885.
- Ruth Rosenholtz. 2016. Capabilities and limitations of peripheral vision. *Annual Review of Vision Science* 2 (2016), 437–457.
- Jyrki Rovamo and Antti Raninen. 1984. Critical flicker frequency and M-scaling of stimulus size and retinal illuminance. *Vision research* 24, 10 (1984), 1127–1131.
- Jyrki Rovamo and Antti Raninen. 1988. Critical flicker frequency as a function of stimulus area and luminance at various eccentricities in human cone vision: a revision of Granit-Harper and Ferry-Porter laws. *Vision research* 28, 7 (1988), 785–790.
- Theodore C. Ruch and John F. Fulton. 1960. Medical Physiology and Biophysics. *Journal of Medical Education* 35 (1960), 1067. Issue 11.
- Daniel Scherzer, Lei Yang, Oliver Mattausch, Diego Nehab, Pedro V. Sander, Michael Wimmer, and Elmar Eisemann. 2012. Temporal Coherence Methods in Real-Time Rendering. *Computer Graphics Forum* 31, 8 (2012), 2378–2408.
- Kalpna Seshadrinathan, Rajiv Soundararajan, Alan Conrad Bovik, and Lawrence K Cormack. 2010. Study of subjective and objective quality assessment of video. *IEEE transactions on Image Processing* 19, 6 (2010), 1427–1441.
- Jie Shen, Christopher A Clark, P Sarita Soni, and Larry N Thibos. 2010. Peripheral refraction with and without contact lens correction. *Optometry and vision science: official publication of the American Academy of Optometry* 87, 9 (2010), 642.
- Raunak Sinha, Mrinalini Hoon, Jacob Baudin, Haruhisa Okawa, Rachel OL Wong, and Fred Rieke. 2017. Cellular and circuit mechanisms shaping the perceptual properties of the primate fovea. *Cell* 168, 3 (2017), 413–426.
- Jan P Springer, Stephan Beck, Felix Weiszig, Dirk Reiners, and Bernd Froehlich. 2007. Multi-frame rate rendering and display. In *2007 IEEE Virtual Reality Conference*. IEEE, 195–202.
- Philip A Stanley and A Kelvin Davies. 1995. The effect of field of view size on steady-state pupil diameter. *Ophthalmic and Physiological Optics* 15, 6 (1995), 601–603.
- Qi Sun, Fu-Chung Huang, Joohwan Kim, Li-Yi Wei, David Luebke, and Arie Kaufman. 2017. Perceptually-guided foveation for light field displays. *ACM Transactions on Graphics* 36, 6 (Nov. 2017), 192:1–192:13.
- LN Thibos, DJ Walsh, and FE Cheney. 1987. Vision beyond the resolution limit: aliasing in the periphery. *Vision Research* 27, 12 (1987), 2193–2197.
- Okan Tarhan Tursun, Elena Arabadzhiyska-Koleva, Marek Wernikowski, Radosław Mantiuk, Hans-Peter Seidel, Karol Myszkowski, and Piotr Didyk. 2019. Luminance-contrast-aware foveated rendering. *ACM Transactions on Graphics (TOG)* 38, 4 (2019), 1–14.
- Christopher W Tyler. 1987. Analysis of visual modulation sensitivity. III. Meridional variations in peripheral flicker sensitivity. *JOSA A* 4, 8 (1987), 1612–1619.
- Christopher W Tyler and Russell D Hamer. 1990. Analysis of visual modulation sensitivity. IV. Validity of the Ferry–Porter law. *JOSA A* 7, 4 (1990), 743–758.
- Christopher W Tyler and Russell D Hamer. 1993. Eccentricity and the Ferry–Porter law. *JOSA A* 10, 9 (1993), 2084–2087.
- Floris L Van Nes and Maarten A Bouman. 1967. Spatial modulation transfer in the human eye. *JOSA* 57, 3 (1967), 401–406.
- Zhou Wang, Ligang Lu, and Alan C Bovik. 2003. Foveation scalable video coding with automatic fixation selection. *IEEE Transactions on Image Processing* 12, 2 (2003), 243–254.
- Andrew B Watson. 2018. The Field of View, the Field of Resolution, and the Field of Contrast Sensitivity. *Journal of Perceptual Imaging* 1, 1 (2018), 10505–1.
- Andrew B Watson and Albert J Ahumada. 2016. The pyramid of visibility. *Electronic Imaging* 2016, 16 (2016), 1–6.
- Andrew B Watson and Denis G Pelli. 1983. QUEST: A Bayesian adaptive psychometric method. *Perception & psychophysics* 33, 2 (1983), 113–120.
- Thomas P Weldon, William E Higgins, and Dennis F Dunn. 1996. Efficient Gabor filter design for texture segmentation. *Pattern recognition* 29, 12 (1996), 2005–2015.
- Lei Xiao, Salah Nouri, Matt Chapman, Alexander Fix, Douglas Lanman, and Anton Kaplanyan. 2020. Neural Supersampling for Real-Time Rendering. *ACM Trans. Graph.* 39, 4, Article 142 (July 2020), 12 pages.