

Uniwersytet Wrocławski
Wydział Matematyki i Informatyki
Instytut Informatyki

Mateusz Hobgarski

Przegląd metod mapowania tonów obrazów wysokokontrastowych

Praca Magisterska

Praca wykonana pod kierunkiem
dr Andrzeja Łukaszewskiego

Wrocław 2005

Spis treści

Wstęp	3
1. Podstawy teoretyczne	5
1.1. Światło	5
1.2. Radiometria	5
1.3. Fotometria	7
1.4. Kolorymetria	8
1.5. Działanie oka ludzkiego	11
1.5.1. Czułość widzenia	13
1.5.2. Zmiana ostrości	14
1.5.3. Odblaski (<i>glare</i>)	14
1.5.4. Zmiany w postrzeganiu kolorów	15
1.5.5. Przebieg adaptacji w czasie	16
1.5.6. Jasność	16
1.6. Korekcja gamma	17
2. Metody mapowania tonów	19
2.1. Podział metod	19
2.2. Oznaczenia	20
2.3. Metoda C. Schlicka z 1994 roku	20
2.3.1. Postępowanie z kolorami	22
2.4. Metoda G. Warda et al. z 1997 roku	23
2.4.1. Poprawianie histogramu	23
2.4.2. Symulacja ludzkiego widzenia	26
2.4.3. Całą metoda	30
2.5. Metoda J. Tumblina i G. Turka z 1999 roku	30
2.5.1. Dyfuzja anizotropowa	31
2.5.2. Filtr LCIS	32
2.5.3. Implementacja LCIS	34
2.5.4. Redukcja kontrastu	36
2.6. Metoda E. Reinharda et al. z 2002 roku	37
2.6.1. System strefowy	37
2.6.2. Algorytm	38
2.6.3. Szczegóły implementacyjne	41
2.7. Metoda F. Duranda i J. Dorsey z 2002 roku	41
2.7.1. Filtr bilateralny	42
2.7.2. Przyspieszanie filtrowania	44
2.7.3. Redukcja kontrastu	45
2.8. Metoda M. Ashikhmina z 2002 roku	46

2.8.1. Wyznaczanie poziomu lokalnej adaptacji	48
2.8.2. Funkcja mapująca TM	49
2.8.3. Pełna metoda	51
3. Opis programu	53
3.1. Sesja pracy z programem	54
3.2. Dokładny opis opcji programu	54
4. Porównanie działania metod.....	59
4.1. Parametry domyślne metod.....	59
4.2. Szybkości działania metod	59
4.3. Otrzymane obrazy	63
4.3.1. Mapowanie z parametrami domyślnymi.....	64
4.3.2. Mapowanie z dopasowanymi parametrami	68
4.4. Podsumowanie	73
Bibliografia	75

Wstęp

Zakres jasności, jaki dociera do nas na codzień jest bardzo duży. Ludzki układ widzenia jest w stanie postrzegać obrazy, których kontrast wynosi pięć rzędów wielkości, może też przyzwyczajać się do jasności z przedziału około dziesięciu rzędów wielkości. Przeniesienie obrazów o charakterystykach odpowiadających rzeczywistości na komputery nie stanowi już problemu. Stosowane w grafice komputerowej, zaawansowane metody renderingu, takie jak radiosity czy metody Monte Carlo, pozwalają na tworzenie obrazów zgodnych jasnościami i kontrastem ich rzeczywistym odpowiednikom. Obrazy takie nazywane są obrazami wysokokontrastowymi lub po prostu obrazami HDR (od *ang. High Dynamic Range*). Można je tworzyć również za pomocą specjalnych kamer, oraz poprzez łączenie kilku zdjęć tego samego ujęcia zrobionych z różnymi parametrami ekspozycji [8]. Problem polega na tym, że zakres jasności, który możemy oddać na ekranie lub wydruku, obejmuje w najlepszym wypadku dwa rzędy wielkości. Tu właśnie pojawia się pytanie: w jaki sposób odwzorować jasności obrazów HDR na monitorach?

W niniejszej pracy przedstawimy kilka odpowiedzi na to pytanie, czyli kilka metod mapowania tonów (nazywanych też często operatorami mapowania tonów). Częścią pracy będzie program implementujący opisane metody.

Najpierw jednak, w rozdziale pierwszym, przedstawimy podstawy teoretyczne przydatne do rozumienia specyfiki problemu i jego rozwiązań.

W rozdziale drugim przedstawimy siedem metod mapowania tonów. Będą one reprezentować różne podejścia do problemu. Znajdą się wśród nich zarówno metody starsze, jedne z pierwszych propozycji rozwiązania problemu mapowania tonów, oraz metody z ostatnich lat.

Rozdział trzeci zawierać będzie opis programu implementującego metody.

W rozdziale czwartym porównamy przedstawione operatory. Zwrócimy uwagę na efekty ich działania, czasy wykonywanych obliczeń, oraz łatwość korzystania z metod.

Rozdział 1

Podstawy teoretyczne

1.1. Światło

Światło to promieniowanie elektromagnetyczne. Światło widzialne, czyli promieniowanie odbierane przez siatkówkę oka ludzkiego, stanowi jedynie mały fragment całego spektrum promieniowania elektromagnetycznego, rozciągającego się od bardzo długich fal radiowych, przez mikro fale, podczerwień, światło widzialne, ultrafiolet, promieniowanie rentgenowskie, do promieniowanie gamma. Precyzyjne ustalenie zakresu długości fal elektromagnetycznych nie jest możliwe do ustalenia, gdyż wzrok każdego człowieka charakteryzuje się inną wrażliwością, stąd za wartości graniczne przyjmuje się maksymalnie 380 – 780 nm, choć często podaje się mniejsze zakresy (szczególnie od strony fal najdłuższych).



Rysunek 1.1: Widmo światła widzialnego.

Pojęcie światła jest szersze, gdyż w praktyce zalicza się do niego nie tylko fale widzialne, ale i sąsiednie zakresy, czyli bliski ultrafiolet i bliską podczerwień można obserwować i mierzyć korzystając z podobnego zestawu przyrządów, a jednocześnie wyniki tych badań można opracowywać korzystając z tych samych praw fizyki.

W naukach ścisłych przyjęto, że światłem nazywa się fale elektromagnetyczne o długości fali od 10nm do 1mm, które podzielono na trzy zakresy - podczerwień, światło widzialne oraz ultrafiolet. Zgodnie z dualizmem korpuskularno-falowym, postrzega się światło jednocześnie jako falę poprzeczną oraz jako strumień cząstek nazywanych fotonami.[26]

1.2. Radiometria

Radiometria to nauka zajmująca się mierzaniem energii światła w dowolnej części spektrum elektromagnetycznego. W praktyce ogranicza się do pomiarów fal o długości od

10 nm do 1 mm w obrębie ultrafioletu, światła widzialnego i podczerwieni, przy użyciu urządzeń optycznych.

Światło to **energia promienista** (*ang. radiant energy*). Jest ona przenoszona w przestrzeni przez promieniowanie elektromagnetyczne. Oznacza się ją jako Q i mierzy w **Joulach** (J).

Źródła szerokopasmowe, takie jak Słońce, emitują promieniowanie elektromagnetyczne w większości spektrum, od fal radiowych do promieniowania gamma, jednak większość jego energii promienistej skupiona jest w widzialnej części spektrum. Z drugiej strony, wąskoemisyjny laser emituje promieniowanie o jednej, konkretnej długości fali.

W związku z powyższym definiuje się **spektralną energię promienistą** (*ang. spectral radiant energy*), która określa ilość energii promienistej na jednostkowy przedział długości fali dla długości fali λ . Opisuje ją wzór:

$$Q_\lambda = \frac{dQ}{d\lambda}.$$

Spektralną energię promienistą mierzy się w J/nm .

Strumień promienisty (*ang. radiant flux*), lub inaczej **moc promienista** (*ang. radiant power*), oznaczany jako Φ , to energia promienista przepływająca, bądź emitowana w jednostce czasu.

$$\Phi = \frac{dQ}{dt}$$

Mierzy się ją w **Wattach** ($W = J/s$).

Gęstość strumienia promienistego (*ang. radiant flux density*) to strumień promienisty na jednostkę powierzchni w punkcie powierzchni. Jeśli strumień promienisty skierowany jest do powierzchni, gęstość strumienia promienistego nazywa się **irradiancją** i oznacza jako E :

$$E = \frac{d\Phi}{dA}$$

gdzie A to powierzchnia do której dochodzi strumień Φ . Strumień może napływać do punktu z dowolnego kierunku na półsferze osaczającej powierzchnię. Jeśli strumień promienisty skierowany jest od powierzchni, gęstość strumienia promienistego nazywa się **promienistością** lub **natężeniem wypromieniowanym** (*ang. radiosity*) i oznacza jako B :

$$B = \frac{d\Phi}{dA}$$

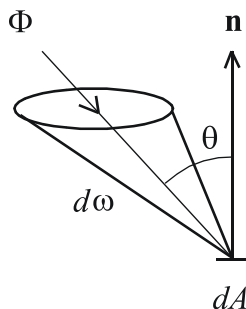
Obydwie jednostki mierzone są w W/m^2 .

Intensywność promieniowania (*ang. radiant intensity*), oznaczana jako I , to strumień promienisty na jednostkowy kąt bryłowy:

$$I = \frac{d\Phi}{d\omega}$$

gdzie ω to kąt bryłowy. Intensywność promieniowania mierzy się w Wattach na steradian (W/sr).

Radiancja (inaczej luminancja energetyczna, *ang. radiance*) L jest najczęściej stosowaną jednostką radiometryczną w grafice komputerowej. Określa stosunek mocy promienistej na kąt bryłowy i jednostkę obszaru rzutu powierzchni rys. 1.2:



Rysunek 1.2: Radiancja dochodząca do powierzchni dA . n to wektor normalny do dA .

$$L = \frac{d^2\Phi}{dA d\omega \cos \theta}$$

mierzy się ją w $W/m^2 sr$.

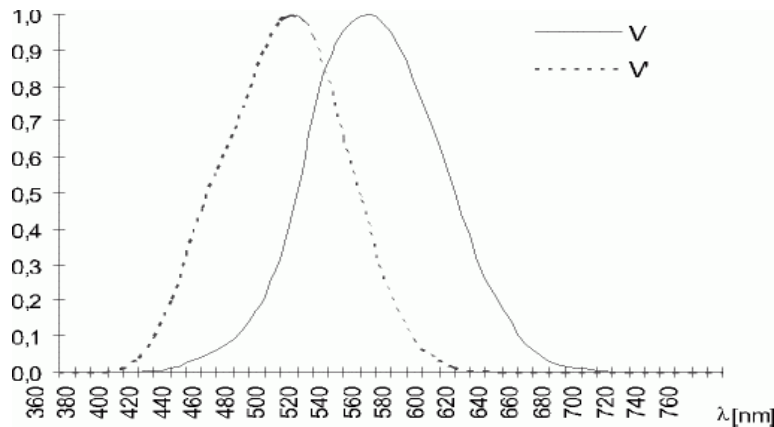
W przeciwieństwie do gęstości strumienia definicja radiancji nie odróżnia strumienia dochodzącego od strumienia wychodzącego. Wynika to z niezmienności radiancji wzdłuż prostych ścieżek, tzn. radiancja opuszczająca punkt x w kierunku punktu y jest równa radiancji padającej z punktu y na punkt x . Własność ta nie jest spełniona, jeśli między punktami x i y znajduje się ośrodek absorbujący lub rozpraszający energię.

1.3. Fotometria

Fotometria to nauka zajmująca się pomiarami światła widzialnego przy użyciu jednostek wagowanych według wrażliwości oka ludzkiego na konkretną długość fali. Opiera się ona na statystycznym modelu ludzkiej percepcji światła. Oko ludzkie to bardzo skomplikowany, nieliniowy, detektor promieniowania elektromagnetycznego. Reaguje na fale długości w przybliżeniu od 380 do 780 nm.

Czułość oka ludzkiego na światło zmienia się w zależności od długości fali. W 1924 roku CIE (*fr. Commission Internationale d'Eclairage* – Międzynarodowa Komisja Oświetlenia), organizacja zajmująca się wprowadzaniem standardów oświetleniowych, przeprowadziła badania dużej grupy ludzi. Badani mieli za zadanie w kontrolowanych warunkach oświetleniowych, dopasować "jasność" światła do długości fali monochromatycznego źródła światła (czyli emitującego fale jednej długości). Rezultatem badań było powstanie krzywej względnej skuteczności oka ludzkiego. Z krzywej tej odczytać można, że źródło światła o radiancji $1 \text{ watt}/m^2 sr$ będzie odbierane jako jaśniejsze jeśli emituje fale o długości 550 nm, od tego źródła o tej samej mocy emitującego fale długości 440 nm.

Fotometria nie określa w jaki sposób postrzegany jest kolor. Mierzone światło może być monochromatyczne (czyli składać się z fal o jednej i tej samej długości) lub może składać się z fal o różnych długościach. Reakcja oka określana jest na podstawie krzywej względnej skuteczności oka. Jedyną różnicą między radiometrią i fotometrią są jednostki miary używane w obu gałęziach nauki.



Rysunek 1.3: Krzywa względnej skuteczności oka ludzkiego, V określa krzywą dla widzenia dziennego (fotopowego), V' dla nocnego (skotopowego). Rysunek z [3]

Dla punktowego źródła światła (tzn. źródła, którego rozmiar jest zanedbywalnie mały w porównaniu do odległości do niego) używa się pojęcia **natężenia światła** (ang. *luminous intensity*). Jest to odpowiednik radiometrycznej intensywności promieniowania. Definiuje się je jako strumień promieniowania źródła, przypadający na jednostkowy kąt bryłowy. Jednostką natężenia światła jest **kandela** (cd). Jest to natężenie światła, w określonym kierunku, źródła emitującego promieniowanie monochromatyczne o długości fali 555 nm , równe $1/683\text{ W/sr}$. 555 nm to długość odpowiadająca maksymalnej względnej skuteczności oka. Często maksymalną względną skuteczność określa się za pomocą częstotliwości, która wynosi ok $5,4 \cdot 10^{14}\text{ Hz}$.

Krzywa względnej skuteczności razem z definicją kandelii umożliwia konwersję między danymi radiometrycznymi i fotometrycznymi. Rozpatrzmy, dla przykładu, monochromatyczne, punktowe źródło światła emitujące fale długości 510 nm i intensywności promieniowania $1/683\text{ W/sr}$. Fotopowa skuteczność dla 510 nm wynosi $0,503$, źródło ma w związku z tym natężenie światła równe $0,503$ kandelii.

Odpowiednikiem radiometrycznego strumienia promienistego jest **strumień świetlny** (ang. *luminous flux* lub *luminous power*). Wyraża się go w lumenach $lm = cd \cdot sr$.

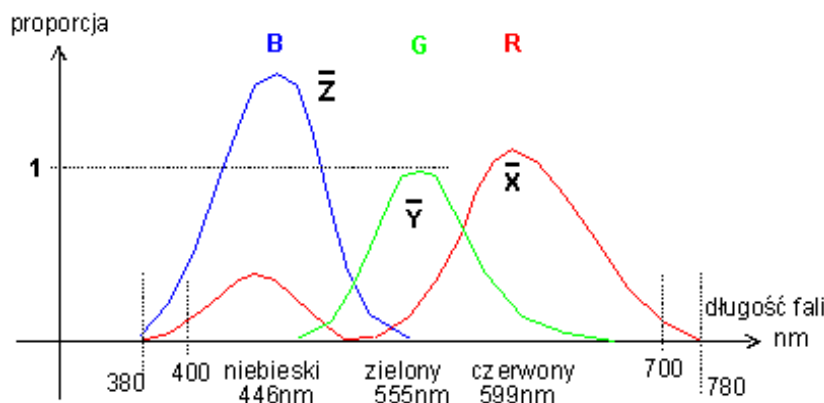
Energia świetlna (ang. *luminous energy*) to odpowiednik energii promienistej. Mierzy się ją w lumenach razy sekundę.

Jednostką, która najczęściej pojawiać się będzie w dalszym ciągu pracy jest **luminancja** (ang. *luminance*). Jest to odpowiednik radiancji. W kontekście ludzkiej percepcji określa ona jak jasna wydaje się nam dana powierzchnia, kiedy patrzymy na nią z określonego kierunku. Luminancję mierzy się w cd/m^2 .

1.4. Kolorymetria

Kolorymetria to zespół metod pomiaru i opisu ilościowego barw. Każdy człowiek postrzega kolory trochę inaczej, CIE wprowadziło pojęcie obserwatora standardowego,

czyli hipotetycznego obserwatora reprezentującego "normalne" postrzeganie kolorów. Dodatkowo określono zestaw standardowych warunków, w których powinno dokonywać się pomiarów. CIE wykonało zestaw badań z udziałem grupy ludzi, polegający na dopasowaniu długości fal trzech źródeł światła oświetlających ten sam obszar tak, aby otrzymana barwa zgadzała się z inną z góry zadaną. Rezultatem badań było powstanie w 1931 roku krzywej kolorymetrycznej układu CIE XYZ 1931 (rys. 1.4).



Rysunek 1.4: Krzywe kolorymetryczne $\bar{x}$, $\bar{y}$, $\bar{z}$ układu CIE XYZ 1931

Znając rozkład widmowy promieniowania $P(\lambda)$, można obliczyć współczynniki XYZ w następujący sposób:

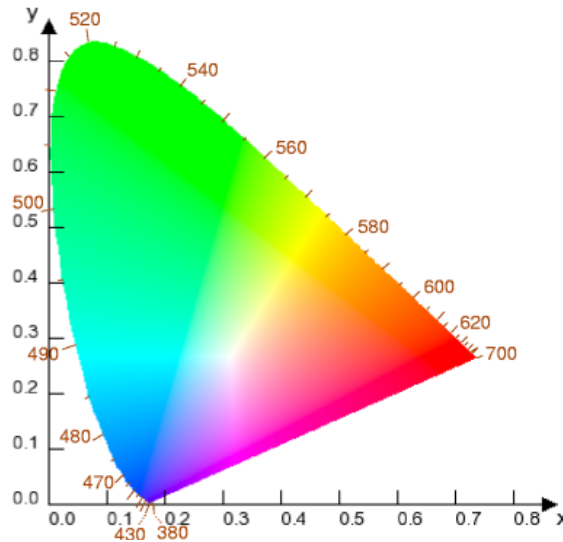
$$\begin{aligned}
 X &= k \int_{380}^{780} P(\lambda) \bar{x} d\lambda \\
 Y &= k \int_{380}^{780} P(\lambda) \bar{y} d\lambda \\
 Z &= k \int_{380}^{780} P(\lambda) \bar{z} d\lambda \\
 k &= 683 \text{ lumen/Watt}
 \end{aligned}$$

Wartości X , Y i Z definiują kolor w przestrzeni kolorów CIE XYZ. Jest to trójwymiarowa, liniowa przestrzeń kolorów. W powyższej postaci jest dosyć niewygodna w użyciu, dlatego często stosuje się jej rzut na płaszczyznę $X + Y + Z = 1$. Rezultatem jest dwuwymiarowa płaszczyzna nazywana diagramem chromatycznym CIE (*ang. chrominance diagram*) przedstawiona na rys. 1.5. Współrzędne na diagramie nazy-

wane są zwykle x , y i powstają z XYZ w następujący sposób:

$$\begin{aligned}x &= \frac{X}{X + Y + Z} \\y &= \frac{Y}{X + Y + Z} \\z &= \frac{Z}{X + Y + Z} = 1 - x - y\end{aligned}$$

Jako że z nie niesie żadnej dodatkowej informacji, zwykle jest pomijany. Przestrzeń xy jest tylko projekcją trójwymiarowej przestrzeni XYZ , w związku z tym jednemu punktowi w xy odpowiada wiele punktów z XYZ . Brakującą informacją jest Y odpowiadający luminancji. Ostatecznie kolor opisywany jest jako trójka Yxy , gdzie Y określa luminancję, a x i y punkt na diagramie chromatycznym.



Rysunek 1.5: Diagram chromatyczny CIE. Obrazek z [26].

Kształt przedstawiony na rys. 1.5 zawiera całe spektrum. Linia prosta między niebieskim o najmniejszej długości fali i czerwonym o długości największej nie należy do spektrum i nazywana jest "linią purpurową". Punkt, w którym znajduje się kolor biały nazywa się punktem bieli. Usytuowany jest w okolicy środka diagramu. Natomiast dokładne położenie środka zależy od iluminantu, czyli charakterystyki źródła światła (iluminant D65, określający światło dzienne, ma kolor biały w punkcie $x = 0.312727$, $y = 0.329024$). Jeśli poprowadzimy linię l przez wybrany kolor określony punktem (x, y) i punkt bieli, to stosunek odległości od punktu bieli do punktu (x, y) i odległości od punktu bieli do brzegu diagramu (linii spektralnej) określa nasycenie koloru. Jeśli kolor jest blisko linii spektralnej, jego nasycenie jest duże. Przecięcie linii l z linią spektralną określa dominującą długość fali, czyli barwę koloru. Wszystkie możliwe mieszanki kolorów między (x_1, y_1) i (x_2, y_2) znajdują się na linii łączącej te dwa punkty. Wszystkie możliwe kolory złożone z trzech kolorów (x_1, y_1) , (x_2, y_2) , (x_3, y_3) znajdują

się wewnątrz trójkąta wyznaczonego przez te punkty. Widać stąd, że gamut określony za pomocą trzech współrzędnych, czyli również gamut zwykłego monitora CRT, zawiera jedynie część wszystkich widzialnych kolorów.

Istnieje wiele modyfikacji i usprawnień modelu CIE XYZ (np. CIE LUV czy CIE LAB), jednak dla potrzeb tej pracy używane używane będą tylko modele CIE XYZ i CIE Yxy.

1.5. Działanie oka ludzkiego

Zakres luminancji, z którym wszyscy mamy do czynienia w życiu codziennym jest bardzo szeroki. Światło słońca w południe może być nawet 10 milionów razy mocniejsze od światła księżyca. Rys.1.6 przedstawia zakres luminancji obserwowanych na codzień. Ludzki układ widzenia (HVS - *ang. human visual system*) radzi sobie z tak dużym zakresem przez adaptowanie się do przeważających luminancji. Dzięki adaptacji układ widzenia może działać na przedziale prawie 14 jednostek logarytmicznych. Adaptacja jest złożona. Składają się na nią skoordynowane procesy mechaniczne, fotochemiczne i nerwowe.



Rysunek 1.6: Zakres luminancji obserwowany w życiu codziennym. Użyto logarytmu dziesiętnego. Obrazek z pracy [11].

To że możemy obserwować sceny o bardzo różnej jasności nie znaczy jednak, że widzimy tak samo w każdych warunkach. Dla przykładu w nocy, w bardzo słabym oświetleniu, wzrok jest bardzo czuły na wszelkie zmiany w luminancji. Jednak detale obrazu nie są zbyt ostre a do tego rozróżnianie kolorów jest mocno zaburzone. Z kolei w ciągu dnia, przy mocnym świetle, obraz jest ostry, dobrze rozpoznajemy kolory, jednak absolutna czułość na zmiany jasności jest mała. Różnice w luminancji muszą być bardzo duże żebyśmy je dostrzegli.

Ponadto adaptacja nie następuje natychmiastowo. Przyzwyczajanie się wzroku przy przejściu z jasnego pomieszczenia do ciemnego może trwać do kilkunastu minut. W sytuacji odwrotnej widzenie powinno wrócić do normy po około minucie.

Na proces adaptacji składają cztery opisane niżej mechanizmy.

Żrenica

To najbardziej oczywisty z dostępnych mechanizmów regulowania ilości światła dochodzącego do układu widzenia. W obrębie 10 jednostek logarytmicznych luminancji źrenica zmienia swoją średnicę od ok. 7 do 2 mm. Taka zmiana średnicy powoduje ograniczenie dochodzącej do siatkówki luminancji o niewiele ponad jedną jednostkę logarytmiczną, około 16 razy. Nie jest więc w stanie samodzielnie przeprowadzić adap-

tacji. Uważa się, że rolą źrenicy jest w większym stopniu łagodzenie aberracji soczewki oka, niż jej udział w procesie adaptacji. W jasnym otoczeniu, kiedy do oka dochodzi duża ilość światła źrenica zmniejsza się, zmniejszając tym samym aberrację. W otoczeniu ciemnym, kiedy światła brakuje, źrenica rozszerza się tak aby do siatkówki dochodziło możliwie dużo światła.

Pręciki i czopki

Wewnętrzna część gałki ocznej pokryta jest warstwą komórek nerwowych, którą nazywa się siatkówką. Siatkówka zaopatrzona jest w fotoreceptory: pręciki i czopki, wyspecjalizowane neurony posiadające światłoczułe barwniki wzrokowe. Są one odpowiedzialne za wstępną fazę przetwarzania światła na impulsy nerwone. Zwykle na siatkówce znajduje się 75 do 150 milionów pręcików i ok. 6 – 7 milionów czopków. Pręciki są niezwykle czułe na światło i zapewniają widzenie achromatyczne (czarno-białe) dla zakresu luminancji od 10^{-6} do 10 cd/m^2 , jest to widzenie skotopowe. Czopki odpowiedzialne są za tzw. widzenie fotopowe. Są mniej wrażliwe od pręcików, za to umożliwiają widzenie kolorowe. Reagują na luminancję od 0.01 do 10^8 cd/m^2 . W przedziale od 0.01 do 10 cd/m^2 aktywne są jednocześnie czopki i pręciki. Nazywa się to widzeniem mezopowym lub zmierzchowym.

Fotoreceptory nie są rozmieszczone równomiernie. Na powierzchni siatkówki wyróżnia się miejsce nazywane plamką żółtą. Jest to miejsce o średnicy ok. 1.5 mm, odpowiadające środkowi naszego pola widzenia. W obszarze tym występuje maksymalne zagęszczenie czopków (w sumie około dwa miliony), natomiast zupełnie brak jest pręcików. Gęstość rozmieszczenia czopków rośnie wraz ze zbliżaniem się do samego środka plamki żółtej, nazywanego dołkiem środkowym (*ang. fovea*). Jest to miejsce, w którym widziany obraz jest najbardziej ostry. W dołku gęstość rozmieszczenia czopków wynosi ok. $190,000/\text{mm}^2$, 500 μm od dołka gęstość spada o 50% (ok. $100,000/\text{mm}^2$), a w odległości 4 mm (ok 20 stopni) od środka spada do mniej niż 5% (ok $10,000/\text{mm}^2$). Dalej czopki występują z bardzo małą gęstością [6]. W dołku środkowym nie ma wogóle pręcików, poza tym obszarem występują one na całej powierzchni siatkówki. Maksymalna gęstość rozmieszczenia pręcików przypada na ok. 20 stopni od dołka środkowego. Ogólnie gęstość rozmieszczenia pręcików zmienia się dużo wolniej niż czopków.

Fotoreceptory są czułe na światło o długości fali od ok. 380 do 700 nm. Pręciki są najbardziej czułe na światło długości ok. 505 nm, czopki ok. 555 nm. Pręciki i czopki nie są tak samo czułe na światło we wszystkich długościach fal. Zależność między długością fali a intensywnością wywołwanego przez nią impulsu opisuje, wspomnianą już wcześniej przy okazji fotometrii, krzywa względnej skuteczności oka (rys. 1.3). A dokładniej dwie krzywe. Pierwsza, oznaczana zwykle jako V , określająca czułość widzenia fotopowego i druga, V' określająca widzenie skotopowe.

Zmiana koncentracji barwników wzrokowych

Barwniki wzrokowe zawarte w fotoreceptorach rozkładają się pod wpływem pochłanianego światła. Produkty ich rozpadu są chemicznym bodźcem dla fotoreceptorów. Przy dużym natężeniu światła barwniki rozkładają się szybciej niż mogą być produkowane. Przez to fotoreceptory są mniej czułe.

Procesy nerwowe

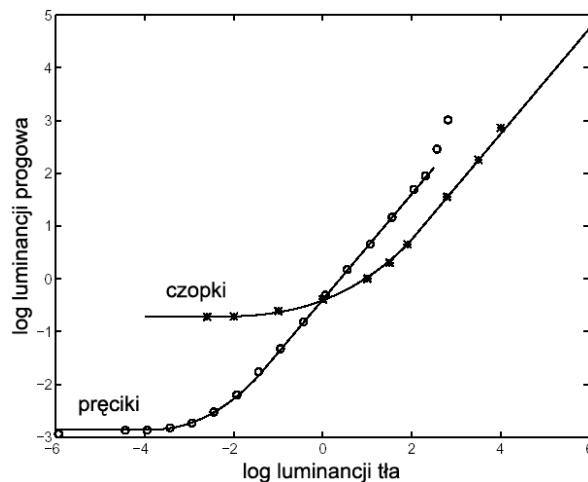
Bodźce wzrokowe tworzone przez fotoreceptory zależą od procesów chemicznych zachodzących w komórkach fotoreceptorów. Jeśli procesy te zachodzą z intensywnością bliską maksymalnej i ilość dochodzącego do receptora światła jeszcze się zwiększy, wtedy komórka nie będzie mogła w pełni zasygnalizować wzrostu ilości światła. Sytuację taką nazywa się nasyceniem. Aby jej zapobiec, powyżej pewnego poziomu luminancji, bodziec jest kompresowany, tzn. kolejne zwiększenie ilości światła powodować będzie coraz mniejsze zmiany w uzyskiwanym bodźcu.

Fizjologiczne mechanizmy przedstawione powyżej są podstawą procesu adaptacji. Ich działanie uwidacznia się w zmianach w czułości i ostrości widzenia, oraz postrzeganiu kolorów na przedziale obserwowanych luminancji.

1.5.1. Czułość widzenia

Czułość widzenia mierzy się za pomocą progów detekcji (*ang. detection threshold*). Progi te wyznacza się przez wykonanie następującego eksperymentu: obserwator siedzi naprzeciw pustego ekranu emitującego określoną luminancję (nazywaną luminancją tła lub luminancją adaptacyjną). Przed rozpoczęciem eksperymentu obserwator unieruchamia wzrok na środku ekranu. Po zaadaptowaniu się wzroku obserwatora do luminancji tła, w kolejnych próbach, rozświetla się na kilkaset milisekund obszar w kształcie dysku na środku ekranu. Obserwator zgłasza czy widział dysk czy nie. Jeśli nie widział, w następnej próbie zwiększa się luminancję dysku. Jeśli widział, to luminancję dysku się zmniejsza. W ten sposób wyznacza się najmniejszą rozróżnialną luminancję dla danej luminancji tła, nazywa się ją progiem detekcji lub luminancją progową.

Zależność między luminancją tła a luminancją progową nazywa się **ledwo dostrzegalną różnicą** (*ang. just noticeable difference*) i oznacza jako **JND** lub **TVI** (*ang. threshold versus intensity*) przedstawia ją rys. 1.7.



Rysunek 1.7: Wykres ledwo dostrzegalnej różnicy. Obrazek z pracy [11].

1.5.2. Zmiana ostrości

Ostrość wzroku to zdolność rozróżniania dwóch punktów leżących blisko siebie. Do badania ostrości służą tablice Snellena. Są to namalowane na białym matowym tle, czarne matowe znaki nazywane optotypami, stopniowo zmniejszające się ku dołowi (rys. 1.8). Optotypami najczęściej są litery. Znaki te mają określoną grubość i oddzielone są

Odległość w m		Ostrość wzroku
50	E	0.1
25	F H	0.2
16.5	E N T	0.3
12.5	T N H L	0.4
10	L E F N H	0.5
8.3	Z L P O H F	0.6
7.1	H L A E Z T P T	0.7
6.2	N Z E P L P F N	0.8
5.5	F P Z H T L E Z	0.9
5.0	P N P Z H T L E Z	1.0

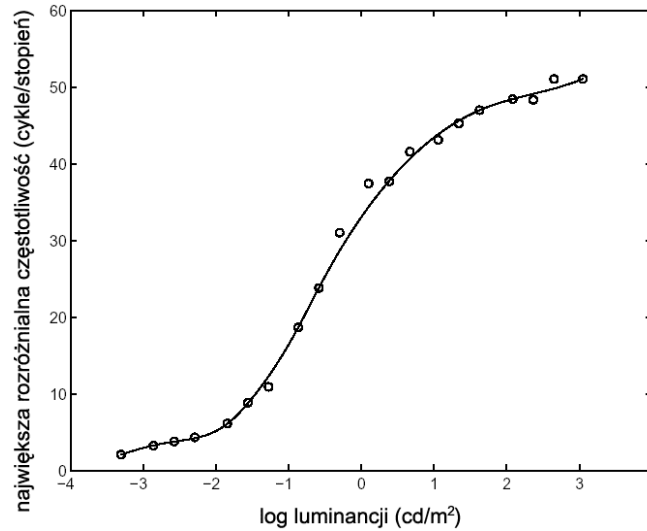
Rysunek 1.8: Typowa tablica Snellena. Obrazek pobrany z <http://www.mediweb.pl>.

przerwami o tej samej odległości, tak, że oko zdrowe widzi cały znak z ostatniego rzędu pod kątem 5 minut, a jego szczegóły pod kątem 1 minuty, czyli kątem pod którym dwa najbliższe położone siebie punkty są rozróżnialne. Siłę ostrości wzroku osoby badanej wyraża się stosunkiem odległości z której optotyp jest rozpoznawany (d), do odległości, z której jego elementy są widziane pod kątem 1 stopnia (D). Jeśli obserwator jest w stanie z odległości $d = 5\text{ m}$ odczytać rząd przeznaczony do odczytania z odległości $D = 5$, to jego ostrość wzroku wyniesie $5/5$ (pełna ostrość wzroku) lub w postaci ułamka dziesiętnego jako 1.0. Natomiast jeśli z odległości $d = 5\text{ m}$ odczyta znaki przeznaczone do odczytania z odległości $D = 50$, to ostrość wzroku równa się $5/50$ lub 0.1 ([14]).

Ostrość widzenia w świetle skotopowym jest mniejsza niż w świetle fotopowym. Krzywa z rys. 1.9 przedstawia zależność ostrości widzenia od luminancji tła. Krzywa powstała z wyników pomiarów rozróżniania fal kwadratowych o różnych częstotliwościach w różnym oświetleniu. Z jej wykresu odczytać można, że rozróżnialna częstotliwość zmienia się od ok. 50 *cykli/stopień* przy $3\log\text{ cd/m}^2$ do 2 *cykli/stopień* przy $-3.3\log\text{ cd/m}^2$. Odpowiada to zmianie ostrości widzenia od ok. $5/2.5$ przy świetle dziennym do $5/75$ w nocy, przy minimalnym oświetleniu. Dzięki tej krzywej jesteśmy w stanie przewidywać widoczność detali przy różnym oświetleniu.

1.5.3. Odblaski (*glare*)

Mocne źródła światła znajdujące się na obrzeżach pola widzenia mogą powodować zmniejszanie kontrastu widzenia. Objawia się to często występowaniem wokół źródeł



Rysunek 1.9: Ostrość widzenia w zależności od luminancji tła. Obrazek zapożyczony z pracy [11].

światła m.in. mgiełek (*ang. disability glare* lub *veiling luminance*) czy rozchodzących się gwiazdziście od źródła promieni światła (*ang. ciliary corona*). Powodem występowania tego zjawiska jest rozpraszanie światła przez soczewkę oka, rogówkę oraz siatkówkę. Dokładniejszy opis tego zjawiska można znaleźć w pracy [22].

1.5.4. Zmiany w postrzeganiu kolorów

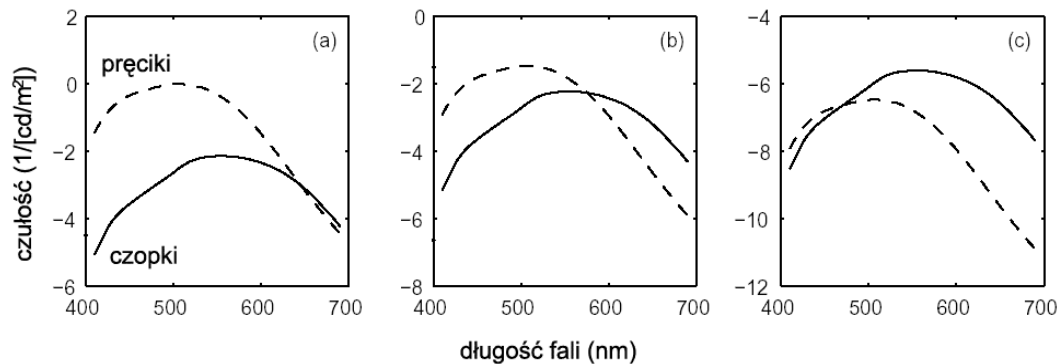
Spektralną wrażliwość pręcików i czopków określają, omawiane już przy radiometrii, skotopowe i fotopowe krzywe względnej skuteczności oka ludzkiego. Zwykle, tak jak na rys. 1.3 ich wykresy są znormalizowane, przez co nie widać, że pręciki i czopki zdecydowanie różnią się czułością i działają na różnych zakresach luminancji.

Na rys. 1.10 przedstawione są nieznormalizowane krzywe względnej skuteczności dla trzech poziomów jasności obserwowanej sceny.

Na rys. 1.10 (a) prezentowana jest wrażliwość spektralna dla poziomu skotopowego (dokładniej dla logarytmu luminancji tła równego $-4 \log \text{cd/m}^2$). Na tym poziomie detekcja zdominowana jest przez pręciki. Absolutna czułość jest wysoka, a z powodu achromatyczności pręcików, kolory nie są rozróżniane.

Na rys. 1.10 (b) przedstawiona jest czułość spektralna dla poziomu mezopowego (luminancja tła wynosi $0 \log \text{cd/m}^2$). Absolutna czułość pręcików i czopków jest tu bardzo podobna. Światło o danej długości fali będzie wykrywane przez bardziej czułe fotoreceptory. Z rysunku odczytać można, że pręciki wykrywać będą długości fal krótszych od ok. 575 nm , a czopki fale dłuższe. Przy przejściu z widzenia skotopowego do mezopowego stopniowo aktywować się będą czopki, czego rezultatem będzie stopniowe rozpoznawanie kolorów zaczynając od długofalowych czerwieni a następnie zieleni, której odpowiadają fale średniej długości.

Na rys. 1.10 (c) przedstawiona jest czułość spektralna dla poziomu fotopowego (lu-



Rysunek 1.10: Zmiany w czułości spektralnej dla luminancji tła równej (a) $-4 \log \text{cd/m}^2$ (widzenie skotopowe), (b) $0 \log \text{cd/m}^2$ (widzenie mezopowe), (c) $4 \log \text{cd/m}^2$ (widzenie fotopowe). Na wszystkich rysunkach linia ciągła oznacza czułość czopków, a linia przerywana pręcików. Obrazek zapożyczony z pracy [11].

minancja tła wynosi $4 \log \text{cd/m}^2$). Absolutna czułość jest znacznie mniejsza niż przy widzeniu skotopowym. Na tym poziomie dominują czopki, dzięki czemu widoczne są kolory. Jednak dla najkrótszych długości fal (odpowiadających kolorowi niebieskiemu) bardziej czułe są jeszcze pręciki. Przez co trzeba relatywnie wysokiej (wyższej od $4 \log \text{cd/m}^2$) luminancji aby kolor niebieski był widoczny.

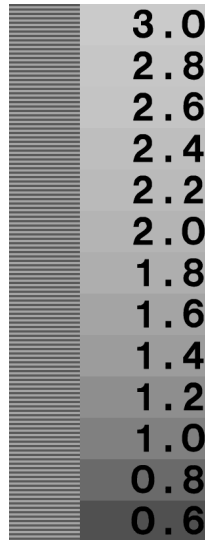
1.5.5. Przebieg adaptacji w czasie

Adaptacja nie jest procesem natychmiastowym. Czas jej trwania zależy od luminancji, do której przyzwyczajony był wzrok przed rozpoczęciem adaptacji, od luminancji, do której wzrok musi się dostosować, oraz od kierunku zmiany luminancji. Ogólnie przystosowanie ze światła ciemniejszego do jaśniejszego jest dużo krótsze niż ze światła jaśniejszego do ciemniejszego. Z grubsza można stwierdzić, że adaptacja do światła jaśniejszego dla pręcików trwa kilka sekund, a 75% procesu odbywa się w czasie krótszym niż jedna sekunda. Przystosowanie czopków trwa dłużej, nawet do kilkunastu minut, jednak większość procesu dokonuje się w przeciągu pierwszych dwóch minut. Przystosowanie do światła ciemniejszego jest dłuższe, może trwać nawet do godziny. Różnica między szybkością zachodzenia procesu adaptacji w czopkach i pręcikach jest niewielka.

1.5.6. Jasność

Przydatne w dalszej części pracy będzie pojęcie jasności (*ang. brightness*). Jest to subiektywna wielkość wrażenia wywołana przez światło widzialne. Bywa ona utożsamiana z luminancją. Nie jest to jednak luminancja, między innymi dlatego, że dwie różne wartości luminancji mogą być odbierane jako taka sama jasność. Powodem tego jest, opisywany wyżej, mechanizm adaptacji oraz ograniczenia widzenia ludzkiego. Człowiek bardzo słabo rozróżnia absolutną luminancję, lepiej radzi sobie z luminancją relatywną.

Jako przybliżoną wartość jasności używa się często logarytmu luminancji, lub pierwiastka (kwadratowego bądź sześciennego) luminancji.



Rysunek 1.11: Obrazek służący do wyznaczania przybliżonej wartości gammy monitora. Pobrany z <http://radsite.lbl.gov/radiance/refer/Notes/gamma.html>

1.6. Korekcja gamma

Zależność między wartością opisującą kolor w komputerze a odpowiadającą tej wartości jasnością na monitorze nie jest liniowa. Generowana luminancja jest proporcjonalna do sygnału wejściowego podniesionego do pewnej potęgi *gamma* (γ), czyli dla danej wartości wejściowej L_{in} luminancja na monitorze wynosi L_{in}^γ . Gamma zwykle mieści się między 2.0 a 2.5.

Korekcja gamma to sposób na wyeliminowanie nieliniowości. Polega ona na zastosowaniu do koloru, przed jego wyświetleniem na monitorze, funkcji potęgowej o wykładniku $\frac{1}{\gamma}$, czyli:

$$L_{out} = L_{in}^{\frac{1}{\gamma}}$$

gdzie L_{out} to otrzymana na monitorze, a L_{in} to luminancja koloru wejściowego. Dla koloru w postaci RGB takiej korekcji podlega każdy kanał.

Korekcja gamma powinna być stosowana zawsze kiedy ważne jest wyświetlenie obrazu zgodnie z jego rzeczywistym wyglądem. Aby wyznaczyć dokładną wartość gammy monitora potrzebny jest fotometr. Można jednak znaleźć jej przybliżoną wartość. W tym celu należy najpierw ustawić kontrast i jasność monitora, a następnie skorzystać z obrazka podobnego do tego z rys. 1.11. Korzystanie z obrazka polega na dobraniu do średniej jasności czarno białych linii z lewej strony odpowiadającego im jasnością prostokąta z prawej strony. Liczba na prostokącie określa wartość gammy monitora.

Rozdział 2

Metody mapowania tonów

2.1. Podział metod

W pracy [9] Devlin dzieli metody mapowania tonów na dwie grupy: metody globalne (z angielskiego nazywane również *spatially uniform*) oraz metody lokalne (ang. *spatially varying*). Operatory globalne (operator mapowania tonów to metoda mapowania tonów, obydwie te pojęcia będą stosowane zamiennie) do każdego piksela używają takiej samej transformacji. Działanie operatora zależy jedynie od obrazu jako całości. Wartość na jaką odwzorowany zostanie piksel zależy tylko od jego luminancji i ogólnych cech obrazu, jak na przykład maksymalnej i minimalnej luminancji. Operatory lokalne mogą działać inaczej w różnych obszarach tego samego obrazu. Wartość na jaką odwzorowany zostanie piksel zależy, poza jego luminancji i ogólnymi cechami obrazu, również od luminancji pikseli z pewnego jego otoczenia.

Reinhard w pracy [7] przedstawia inny podział. Wyróżnia on dwie główne kategorie. Pierwsza to metody oparte na działaniu ludzkiego układu widzenia (ang. *Human Visual System* – HVS). Druga kategoria to metody opierające się na podziale obrazu na warstwy, oddzielające w ten sposób detale od tła. Korzystają one z filtrów selektywnie wygładzających, bardziej niż z zasad działania HVS. Metody z pierwszej kategorii Reinhard dzieli podobnie jak Devlin na lokalne i globalne.

Przedstawione w tym rozdziale metody mapowania tonów należą do często stosowanych i ważnych "historycznie". Znajdują się wśród nich reprezentanci wszystkich kategorii. Operatory Schlicka, Warda et al., Reinharda et al. i Ashikhmina zaliczyć można do metod opartych na działaniu ludzkiego układu widzenia, dwie pierwsze z nich są globalne, dwie ostatnie są lokalne. Kluczowy dla metod Tumblina i Turka oraz Duranda i Dorsey jest podział obrazu na warstwy.

2.2. Oznaczenia

$\log(x)$	logarytm naturalny z x
$\log_{10}(x)$	logarytm dziesiętny z x
$L_w(x, y)$	(<i>ang. world luminance</i>) - luminancja sceny w punkcie (x, y) lub inaczey luminancja wejściowa
$L_{w \min}$	(<i>ang. minimal world luminance</i>) - minimalna luminancja sceny
$L_{w \max}$	(<i>ang. maximal world luminance</i>) - maksymalna luminancja sceny
$B_w(x, y)$	(<i>ang. world brightness</i>) - jasność sceny w punkcie (x, y)
$L_d(x, y)$	(<i>ang. display luminance</i>) - luminancja wyświetlana w punkcie (x, y) lub inaczey luminancja wyjściowa
$L_a(x, y)$	(<i>ang. adaptation luminance</i>) - luminancja adaptacyjna w punkcie (x, y) czyli luminancja do której przyzwyczajają się oko patrząc w kierunku punktu (x, y)
$L_{wa}(x, y)$	(<i>ang. world adaptation luminance</i>) - luminancja adaptacyjna sceny w punkcie (x, y)
$L_{da}(x, y)$	(<i>ang. display adaptation luminance</i>) - luminancja adaptacyjna obrazu wyjściowego w punkcie (x, y)
$(f \otimes g)(x, y)$	(<i>ang. convolution</i>) - spłot f i g (dyskretny) w punkcie (x, y) , czyli $\sum_{i=-s/2}^{s/2} \sum_{j=-s/2}^{s/2} f(x-i, y-j)g(i, j)$ - gdzie s to rozmiar g

2.3. Metoda C. Schlicka z 1994 roku

Istnieje wiele prostych metod mapowania tonów. Jedną z częściej stosowanych jest mapowanie liniowe z korekcją gamma (*ang. gamma-corrected linear mapping*). Polega ono na przemnożeniu obrazu (każdego z kanałów RGB lub tylko luminancji) przez stałą a następnie poddaniu korekcji gamma opisanej w rozdziale pierwszym. Metodę tę zapisać można następująco:

$$L_d = (aL_w)^{\frac{1}{q}} \quad (2.1)$$

gdzie a to współczynnik mapowania liniowego, równy na przykład $a = \frac{1}{L_{w \max}}$, a q to parametr korekcji gamma. Jego dokładna wartość zależy głównie od charakterystyk i ustawień monitora. Może ona należeć do przedziału $[1, 3]$ i zwykle wynosi około 2.2.

Często stosowaną modyfikacją metody z równania 2.1 jest obcinanie z korekcją gamma (*ang. gamma corrected clamping*), zdefiniowane jako:

$$L_d = \begin{cases} \left(\frac{L_w}{p}\right)^{\frac{1}{q}} & \Leftrightarrow L_w < p \\ 1 & \Leftrightarrow L_w \geq p \end{cases} \quad (2.2)$$

gdzie q tak jak wcześniej należy do przedziału $[1, 3]$, a p to pewna luminancja progowa. Zwykle mapowanie to daje najlepsze rezultaty gdy p jest równe luminancji najjaśniejszego obszaru nie będącego źródłem światła.

Kolejną modyfikacją mapowania liniowego z korekcją gamma jest mapowanie eksponencjalne (*ang. exponentiation mapping*):

$$L_d = (aL_w)^{\frac{p}{q}} \quad (2.3)$$

gdzie $q \in [1, 3]$, a $p \in [0, 1]$. Mapowanie to jest równoważne zwykłemu mapowaniu liniowemu z korekcją gamma z wartością gammy większej od zwykle przyjmowanej.

Innym ważnym operatorem jest mapowanie logarytmiczne (*ang. logarithmic mapping*):

$$L_d = \left[\frac{\log(L_w + 1)}{\log(L_{w \max} + 1)} \right]^{\frac{1}{q}}. \quad (2.4)$$

Motywacją do stworzenia tego operatora był fakt, że funkcja logarytmiczna jest dobrą aproksymacją jasności, czyli wrażeniu jakie wywołuje w ludzkim układzie widzenia luminancja.

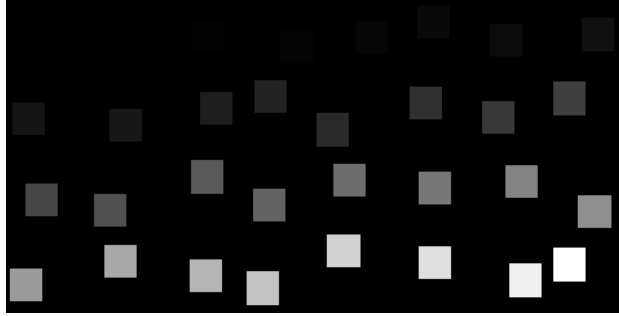
Metoda przedstawiona przez Schlicka w pracy [21] z założenia zachowuje się podobnie do mapowania logarytmicznego, przy czym jest od niego dużo mniej kosztowna. Korzysta ona z następującej funkcji mapującej:

$$L_d = \frac{pL_w}{pL_w - L_w + L_{w \max}} \quad (2.5)$$

gdzie $p \in [1, \infty)$ a jego dokładną wartość określa wzór:

$$p = \frac{ML_{w \max} - ML_{w \min}}{NL_{w \min} - ML_{w \min}} \approx \frac{ML_{w \max}}{NL_{w \min}} \quad (2.6)$$

w którym N to liczba różnych jasności obrazu wyjściowego (zwykle przyjmuje się 256), a M to najciemniejszy odcień szarości postrzegany na monitorze jako różny od czerni.



Rysunek 2.12: Obrazek służący do wyznaczania najciemniejszego, odróżnialnego od czerni, odcienia szarości. Obrazek z pracy [21].

Sposób wyznaczania parametru p z równania 2.6 sprawia, że najmniejsza niezerowa luminancja sceny wejściowej $L_{w \min}$ odwzorowana jest na najciemniejszy nie czarny odcień szarości M . Wartość M zależy od ustawień i charakterystyk konkretnego monitora. Można ją wyznaczyć w następujący interaktywny sposób: użytkownikowi przedstawiona jest czarna prostokątna plansza, a na niej kolejne, coraz ciemniejsze, szare kwadraty. Użytkownik wybiera najciemniejszy widoczny kwadrat a jego jasność przypisana zostaje parametrowi M . Regularne rozmieszczenie kwadratów mogłoby powodować, że wybierany byłby kwadrat czarny, dlatego są one rozmieszczone losowo. Przykładową planszę przedstawia rys. 2.12.

Dzięki opisanym parametrom p i M , niezależnie od warunków, w mapowanych obrazach zachowywany jest ten sam poziom detali.

Przedstawiony operator można nazwać monotonicznym. Luminancje L_{w1} i L_{w2} takie, że $L_{w1} \leq L_{w2}$, zostaną odwzorowane przez operator na wartości L_{d1} i L_{d2} ,

spełniające $L_{d1} \leq L_{d2}$. Wzrok ludzki nie jest monotoniczny. Jasność obserwowanej sceny w dużym stopniu zależy od otoczenia. W celu odwzorowania monotoniczności wzroku autor, dla każdego piksela obrazu, modyfikuje parametr p . Korzysta przy tym z przybliżonej wartości luminancji adaptacyjnej piksela. Modyfikacja Schlicka działa w następujący sposób: gdy luminancja adaptacyjna jest mała, wtedy piksel powinien być postrzegany jako jaśniejszy, co oznacza, że p' (zmodyfikowana wartość p) powinna być mniejsza. Dla wysokich wartości luminancji adaptacyjnej ma miejsce sytuacja odwrotna. Do oceny jasności piksela używa się luminancji średniej:

$$L_{mid} = \sqrt{L_{w \min} L_{w \max}}. \quad (2.7)$$

Jeśli L_a/L_{mid} jest mniejsze od 1, wtedy piksel uważa się za ciemny, w przeciwnym wypadku piksel uważa się za jasny. Parametr p' definiujemy jako:

$$p' = p \left(1 - k + k \frac{L_a}{L_{mid}} \right) \quad (2.8)$$

gdzie $k \in [0, 1]$, określa siłę z jaką modyfikujemy p . Jeśli $k = 0$, wtedy $p' = p$. Schlick używał parametru $k = 0.5$. Wartość luminancji adaptacyjnej aproksymowana została przez luminancję w pikselu, czyli $L_a(k) = L_w(k)$ dla piksela k .

2.3.1. Postępowanie z kolorami

Większość operatorów mapowania tonów, tak jak operator Schlicka, działa jedynie na luminancji, czyli na odcieniach szarości. Aby rozszerzyć działanie operatorów na obrazy w formacie RGB wykonuje się następujące czynności: najpierw dla wszystkich pikseli obrazu obliczamy luminancję L_w , następnie mapujemy luminancję. Otrzymujemy w ten sposób L_d . Ostatni krok polega na mnożeniu każdego z kanałów RGB, wszystkich pikseli, przez współczynnik L_d/L_w . Tego sposobu używa Schlick, jest on również używany przez część metod przedstawionych w dalszym ciągu pracy. Pozostałe metody obchodzą się z kolorem inaczej. Transformują kolor z formatu RGB do XYZ lub Yxy (Y jest równe luminancji L_w). Następnie mapują L_w , nie zmieniając przy tym pozostałych składowych (XZ lub xy), otrzymują L_d , poczym wykonują transformację odwrotną z XYZ (lub Yxy) do RGB.

Transformację koloru z przestrzeni RGB do XYZ przedstawia poniższe równanie:

$$\begin{bmatrix} X \\ Y \\ Z \end{bmatrix} = M \begin{bmatrix} R \\ G \\ B \end{bmatrix} \quad (2.9)$$

wartości współczynników macierzy M mogą się zmieniać w zależności od przyjętego luminantu i luminoforu (substancji wytwarzającej światło w monitorach). W tej pracy używać będziemy, za Reinhardem, następujących wartości współczynników M :

$$M = \begin{bmatrix} 0.5141364 & 0.3238786 & 0.16036376 \\ 0.265068 & 0.6702343 & 0.06409157 \\ 0.0241188 & 0.1228178 & 0.84442666 \end{bmatrix}. \quad (2.10)$$

Transformacja odwrotna, z XYZ do RGB, używa macierzy M^{-1} o współczynnikach:

$$M^{-1} = \begin{bmatrix} 2.5651 & -1.1665 & -0.3986 \\ -1.0217 & 1.9777 & 0.0439 \\ 0.0753 & -0.2543 & 1.1892 \end{bmatrix}. \quad (2.11)$$

2.4. Metoda G. Warda et al. z 1997 roku

W obrębie przetwarzania obrazu powstało wiele metod służących do polepszenia kontrastu obrazu. Jedną z najbardziej znanych jest metoda wyrównywania histogramu. W metodzie tej rozkład poziomów szarości obrazu jest wyrównywany na całym możliwym zakresie, czego skutkiem jest lepsze wykorzystanie zakresu i zwiększenie kontrastu obrazu. Metoda mapowania tonów zaprezentowana przez Warda et al. w [25] składa się z dwóch kroków. Pierwszym jest zmniejszanie kontrastu oparte na wyrównywaniu histogramu, nazwane przez autorów poprawianiem histogramu (*ang. histogram adjustment*). Drugim jest dodanie efektów symulujących specyficzne cechy ludzkiego widzenia, takie jak zmiana ostrości widzenia, zmiana postrzegania kolorów, odbłaski.

2.4.1. Poprawianie histogramu

Pierwszym krokiem poprawianiem histogramu jest oczywiście stworzenie samego histogramu. Ward et al. oblicza histogram z przybliżonych wartości luminancji adaptacyjnych obrazu wejściowego, a dokładniej z ich logarytmów, czyli jasności.

W celu wyznaczenia luminancji adaptacyjnych autorzy korzystają z faktu, że największy wpływ na wartość luminancji adaptacyjnej ma w poziomie fotopowym (czyli również tym emitowanym przez monitory) obraz odbierany przez dołek środkowy (*ang. fovea*) obejmujący ok. 1 stopień kąta widzenia. Obliczenie luminancji adaptacyjnych sprowadza się do odpowiedniego przeskalowania obrazu wejściowego, tak aby rozmiar każdego piksela odpowiadał w przybliżeniu jednemu stopniowi kąta widzenia. Postępowanie takie wymaga znajomości pionowego i poziomego kąta widzenia. Dokładne ich wartości często nie są znane, mogą się one mocno zmieniać w zależności od obrazu wejściowego, dlatego, jeśli nie są znane trzeba przyjąć wartości orientacyjne. Dla danych wymiarów obrazu wejściowego i kątów widzenia autorzy używali następującej formuły do obliczania wymiarów obrazu luminancji adaptacyjnych:

$$S = \frac{2 \tan(\Theta/2)}{0.01745} \quad (2.12)$$

gdzie S to wysokość lub szerokość w pikselach, Θ to pionowy lub poziomy kąt widzenia a 0.01745 to liczba radianów w jednym stopniu ($\frac{\pi}{180}$). Korzystając z powyższej formuły otrzymujemy obraz, w którym piksele blisko środka obrazu będą odpowiadać rozmiarowi dokładnie jednego stopnia kąta widzenia, a im bliżej brzegów obrazu piksele będą odpowiadać coraz mniejszemu wycinkowi pola widzenia. Dodatkowo, przez nieliniowość formuły 2.12 obraz będzie miał inne proporcje od wejściowego, co stanowić będzie pewne utrudnienie w dalszych etapach metody. W związku z powyższymi wadami przekształcenia 2.12 w implementacji metody użyłem innego, prostszego podejścia. Wysokość i szerokość obrazu adaptacyjnego równa jest odpowiednio pionowemu i

poziomemu kątowemu widzenia. Sposób ten również nie jest dokładny jednak ułatwi dalsze korzystanie z obrazu. Przy skalowaniu używany jest filtr prostokątny (*ang. box filter*), w którym każdy piksel ma jednakową wagę.

Jak już wcześniej zaznaczyłem histogram obliczany będzie nie z wartości luminancji adaptacyjnych L_a ale z odpowiadających im jasności, czyli logarytmów luminancji adaptacyjnej $B_a = \log(L_a)$. Umożliwi to lepsze (bardziej jednorodne) odwzorowanie rozkładu luminancji na szerokim zakresie. Wymaga to ustalenia minimalnej i maksymalnej wartości luminancji $L_{w \min}$ i $L_{w \max}$. Jako minimum ustalamy większą z wartości 10^{-4} cd/m^2 (dolny próg ludzkiego widzenia) lub najmniejszą wartość z obliczonych luminancji adaptacyjnych. Za maksimum przyjmujemy największą wartość luminancji adaptacyjnej.

Histogram obliczany jest na przedziale od $L_{w \min}$ do $L_{w \max}$. Przedział ten dzielony jest na N równej długości podprzedziałów (gęstość histogramu) $b_1, b_2, \dots, b_N$ tak, że $b_1 = [B_{w \min}, B_{w \min} + \Delta b]$, a $b_N = [B_{w \max} - \Delta b, B_{w \max}]$, gdzie

$$\Delta b = \frac{B_{w \max} - B_{w \min}}{N} \quad (2.13)$$

czyli jest to rozmiar długości każdego z podprzedziałów histogramu w skali logarytmicznej, $B_{w \max} = \log(L_{w \max})$ oraz $B_{w \min} = \log(L_{w \min})$. N jest zwykle równe 100. Histogram określa funkcja $f(b_i)$ zwracająca liczbę pikseli obrazu adaptacyjnego, których luminancja mieści się w przedziale b_i .

Mając już histogram obliczamy dystrybucję kumulacyjną jako:

$$P(b_i) = \frac{\sum_{b_j < b_i} f(b_j)}{T} \quad (2.14)$$

gdzie

$$T = \sum_{b_i} f(b_i)$$

czyli T to liczba wszystkich próbek (pikseli obrazu adaptacyjnego). W późniejszej fazie poprawiania histogramu przydatna będzie pochodna dystrybucji kumulacyjnej. Autorzy przybliżają ją następującym wzorem:

$$\frac{dP(b_i)}{db} = \frac{f(b_i)}{T \Delta b} \quad (2.15)$$

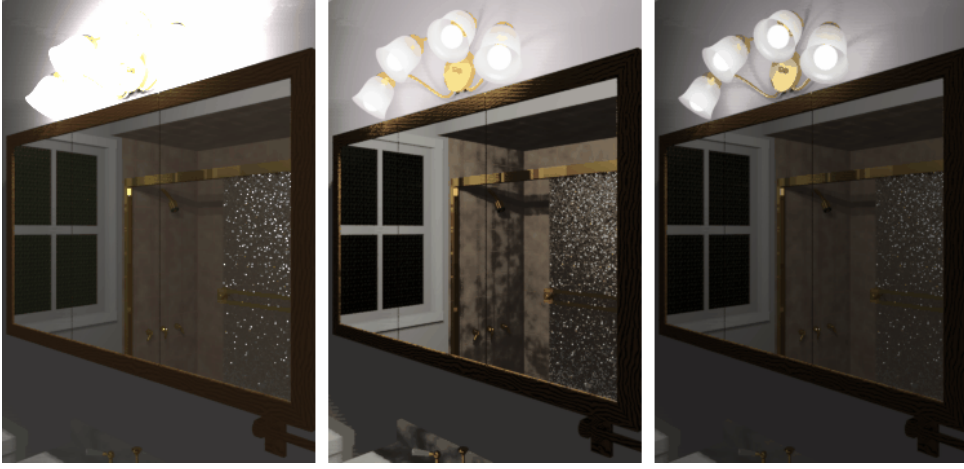
gdzie Δb to określony wzorem 2.13 rozmiar pojedynczego podprzedziału.

Zwykle wyrównywanie histogramu korzystając z dystrybucji kumulacyjnej zdefiniować można jako:

$$B_d = B_{d \min} + [B_{d \max} - B_{d \min}] P(B_w) \quad (2.16)$$

gdzie $B_{d \max} = \log(L_{d \max})$ i $B_{d \min} = \log(L_{d \min})$ czyli jest to odpowiednio największa i najmniejsza wyświetlana jasność.

Metoda ta zmniejsza przedział zakresu monitora przeznaczony do odwzorowywania jasności rzadko występujących, zmniejszając tym samym ich kontrast. Jednocześnie zwiększa ona przedział zakresu przeznaczony dla jasności częściej występujących. Może to prowadzić do zbyt dużego zwiększenia kontrastu niektórych rejonów obrazu i nie naturalnego ich wyglądu. Ilustruje to rys. 2.13. Z lewej strony rysunku widzimy obrazek



Rysunek 2.13: Porównanie rezultatów mapowania liniowego (lewa strona), zwykłego wyrównywania histogramu (środek), i mapowania histogramowego z liniowym ograniczeniem (prawa strona). Wszystkie obrazki z pracy [25].

odwzorowany operatorem liniowym. Nie można na nim rozróżnić szczegółów przy lampach, jednak cała reszta obrazka wygląda zadowalająco. Obrazek środkowy został zmapowany za pomocą wyrównywania histogramu przedstawionego wyżej. Lampy i ich otoczenie wyglądają dobrze, jednak na kafelkach wewnątrz kabiny pojawiają się nienaturalnie wyglądające plamy. Są to podobne jasności, relatywnie licznie reprezentowane w całym obrazie przez co odwzorowane na dużym przedziale jasności monitora.

Aby wyeliminować opisany problem autorzy wprowadzili ograniczenie: dla żadnego przedziału jasności kontrast nie powinien być większy od tego, który daje mapowanie liniowe. Będziemy je dalej nazywać ograniczeniem liniowym. Można je zapisać jako:

$$\frac{dL_d}{dL_w} = \frac{L_d}{L_w}. \quad (2.17)$$

Czyli pochodna z luminancji wyjściowej po luminancji wejściowej nie może być większa od luminancji wyjściowej podzielonej przez luminancję wejściową. Równanie 2.16 jest funkcją luminancji wyjściowej od luminancji wejściowej. Obliczając jej pochodną i korzystając z równania 2.15 przekształcamy liniowe ograniczenie z 2.17 do stałego ograniczenia postaci:

$$f(b_i) \leq \frac{T\Delta b}{B_{d\max} - B_{d\min}}. \quad (2.18)$$

Jeśli warunek ten będzie spełniony dla wszystkich przedziałów b_i , otrzymany histogram nie będzie nadmiernie zwiększać kontrastu.

Aby ze zwykłego histogramu otrzymać histogram liniowo ograniczony, wszystkie $f(b_i)$, które nie spełniają warunku 2.18 zostają zmniejszone tak, aby były równe ograniczeniu. Powoduje to zmianę całkowitej liczby próbek T a co za tym idzie wartości ograniczenia. Dlatego postępowanie jest iteracyjnie powtarzane z nową wartością ograniczenia do momentu kiedy nie więcej niż 2.5% początkowej liczby próbek przekracza ograniczenie. Liczba 2.5 to arbitralnie wybrana mała wartość. Jej zmiana nie

powoduje nie powoduje istotnych różnic w działaniu algorytmu. Pseudokod procedury wygląda następująco:

```

1.      void OgraniczenieHistogramu()
2.          tolerancja = 2.5% * T
3.          repeat
4.              reszta = 0
5.              oblicz nową wartość T
6.              if(T < tolerancja)
7.                  return
8.              for(i = 1..N)
9.                  oblicz ograniczenie
10.                 if(f(bi) > ograniczenie)
11.                     reszta += f(bi) - ograniczenie
12.                     f(bi) = ograniczenie
13.             until(reszta <= tolerancja)
14.         return

```

Procedura ta nie będzie zbieżna i zakończy się wyjściem w linii 6 tylko gdy zakres obrazu wejściowego jest mniejszy od zakresu monitora. Dlatego przed uruchomieniem procedury należy to sprawdzić i ewentualnie użyć mapowania liniowego.

Cała metoda tworzenia histogramu razem z jego modyfikacją nazywana jest przez autorów poprawianiem histogramu. Rezultat jej działania przedstawiony jest z prawej strony rys. 2.13.

2.4.2. Symulacja ludzkiego widzenia

Pierwszą cechą charakterystyczną symulowaną w metodzie Warda et al. będzie zmiana czułości widzenia wraz ze zmianą jasności sceny, następnie dodana zostanie symulacja odbłasków, wrażliwości na kolory i ostrości widzenia.

Poprawianie histogramu z uwzględnieniem czułości widzenia

Czułość widzenia określa, opisana już w rozdziale pierwszym, krzywa *JND*. Przedstawia ona najmniejszą rozróżnialną luminancję dla danej luminancji adaptacyjnej. Dokładniej mówiąc są to dwie krzywe, pierwsza określa czułość pręcików, druga czopków, przecinają się one dla logarytmu luminancji adaptacyjnej równego $10^{-0.0184} \text{ cd/m}^2$. Dla ułatwienia korzystania z krzywych autorzy łączą je w jedną krzywą, oznaczoną jako ΔL_t , która dla danej luminancji adaptacyjnej przyjmuje wartości bardziej czulej z krzywych. Dodatkowo połączona krzywa przybliżana jest łamaną składającą się z pięciu segmentów. Aproksymacja ΔL_t wygląda następująco:

$$\begin{aligned}
 & \log_{10}(\Delta L_t(L_a)) \\
 = & \begin{cases} -2.86 & \Leftrightarrow \log_{10}(L_a) < -3.94 \\ (0.405 \log_{10}(L_a) + 1.6)^{2.18} - 2.86 & \Leftrightarrow -3.94 \leq \log_{10}(L_a) < -1.44 \\ \log_{10}(L_a) - 0.395 & \Leftrightarrow -1.44 \leq \log_{10}(L_a) < -0.0184 \\ (0.249 \log_{10}(L_a) + 0.65)^{2.7} - 0.72 & \Leftrightarrow -0.0184 \leq \log_{10}(L_a) < 1.9 \\ \log_{10}(L_a) - 1.255 & \Leftrightarrow \log_{10}(L_a) \geq 1.9 \end{cases} \quad (2.19)
 \end{aligned}$$

Aby zagwarantować, że kontrast w obrazie wyjściowym nie będzie większy od postrze-

ganago w rzeczywistej scenie autorzy wprowadzili nowe ograniczenie postaci:

$$\frac{dL_d}{dL_w} \leq \frac{\Delta L_t(L_d)}{\Delta L_t(L_w)}. \quad (2.20)$$

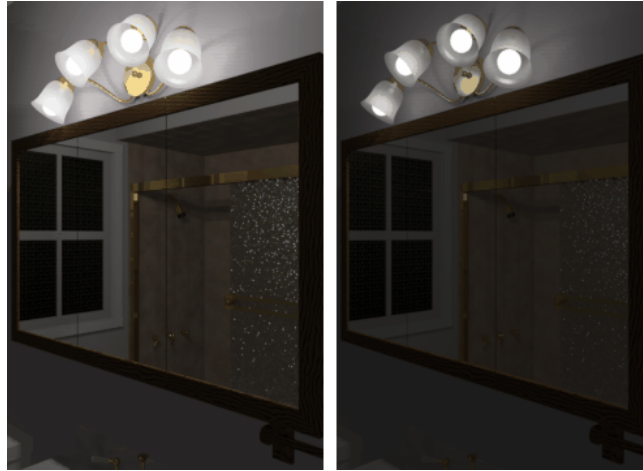
Przekształcając wzór 2.20 otrzymujemy ograniczenie:

$$f(b_i) \leq \frac{\Delta L_t(L_d)}{\Delta L_t(L_w)} \cdot \frac{T \Delta b L_w}{(B_{d\max} - B_{d\min}) L_d} \quad (2.21)$$

gdzie $L_d = \exp(B_d)$, a B_d określa równanie 2.16.

Tak jak w przypadku ograniczenia liniowego musimy iteracyjnie zmniejszać liczbę próbek w tych przedziałach histogramu, które nie spełniają nierówności 2.21. Możemy w tym celu ponownie wykorzystać procedurę `OgraniczenieHistogramu()`, zamieniając tylko ograniczenie liniowe z właśnie przedstawionym.

Działanie poprawiania z uwzględnieniem czułości jest najlepiej widoczne gdy luminancje adaptacyjne obrazu wejściowego różnią się mocno od luminancji adaptacyjnych na monitorze. Rysunek rys. 2.14 przedstawia porównanie rezultatów mapowania sceny,



Rysunek 2.14: Porównanie działania ograniczenia liniowego (lewa strona) i ograniczenia uwzględniającego czułość wzroku (prawa strona), dla sceny z rys. 2.13 o luminancji podzielonej przez 100. Obrazki z pracy [25].

której luminancja została 100-krotnie zmniejszona. Lewa strona została zmapowana z ograniczeniem liniowym (nierówność 2.18), po prawej znajduje się wynik mapowania z ograniczeniem z uwzględnieniem czułości widzenia (nierówność 2.21). Jak widać duża zmiana bezwzględnej wartości luminancji wogóle nie zmienia rezultatu działania mapowania z ograniczeniem liniowym. Wadę tę eliminuje mapowanie z uwzględnieniem czułości, które zmniejsza kontrast w większej części sceny, przez co sprawia wrażenie ciemniejszej.

Symulacja odbłasków

Mocne źródła światła w obrzeżach pola widzenia, jak wspomniano w rozdziale pierwszym, mogą prowadzić do zmniejszenia kontrastu w obserwowanej scenie.

W celu odwzorowania powyższego zjawiska autorzy skorzystali ze zmodyfikowanej definicji luminancji adaptacyjnej, uwzględniającej źródła odbłasków. Wygląda ona następująco:

$$L_a = 0.913L_f + \frac{K}{\pi} \int_{\Theta > \Theta_f} \int \frac{L(\Theta, \phi)}{\Theta^2} \cos(\Theta) \sin(\Theta) d\Theta d\phi \quad (2.22)$$

gdzie L_a to ostateczna luminancja adaptacyjna, L_f to średnia luminancja odbierana przez plamkę żółtą, czyli to co wcześniej nazywaliśmy luminancją adaptacyjną, $L(\Theta, \phi)$ to luminancja w kierunku (Θ, ϕ) , Θ_f to połowa kąta plamki żółtej (ok. 0.5 stopnia, czyli 0.00873 radiana), K to stała zmierzona przez Holladay'a równa 0.0096. Luminancji spoza środka pola widzenia stanowi mniej niż 9% całkowitej luminancji adaptacyjnej, stąd stała 0.913 przy L_f . Całkę z równania 2.22 autorzy przybliżyli w następujący sposób:

$$L_{vi} = 0.087 \frac{\sum_{j \neq i} \frac{L_j \cos(\Theta_{i,j})}{\Theta_{i,j}^2}}{\sum_{j \neq i} \frac{\cos(\Theta_{i,j})}{\Theta_{i,j}^2}} \quad (2.23)$$

gdzie L_{vi} to luminancja spoza punktu i , L_j to podstawowa luminancja adaptacyjna, czyli wartości j -tego piksela z obrazu luminancji adaptacyjnych, a Θ_{ij} to kąt między pikselem i a j w radianach.

Symulacja odbłasków polega na obliczeniu L_{vi} zgodnie z wzorem 2.23, tworząc tym samym obraz odbłasków o rozdzielczości takiej samej jak obraz luminancji adaptacyjnych. Następnie dodajemy do oryginalnych luminancji luminancję odbłasków:

$$L_{pvk} = 0.913L_{pk}(k) + L_v(k) \quad (2.24)$$

gdzie L_{pvk} to luminancja wejściowa z uwzględnieniem odbłasków w punkcie k , $L_{pk}(k)$ to luminancja świata w punkcie k a $L_v(k)$ to luminancja odbłasków interpolowana dwuliniowo z czterech najbliższych punktów obrazu odbłasków. Dodatkowo zmodyfikowany powinien zostać obraz luminancji adaptacyjnych w ten sam sposób co obraz wejściowy ale bez interpolacji. Tak zmodyfikowanych obrazów należy użyć do stworzenia histogramu i wykonania dalszych faz metody.

Lewa strona rys. 2.15 przedstawia rezultat działania metody z dodanym symulowaniem odbłasków. W otoczeniu lamp kontrast jest zmniejszony przez otaczającą je poświatę, ale poświata szybko się zmniejsza i reszta sceny oddalona od lamp wygląda tak jak wcześniej.

Wrażliwość na kolory

Symulowanie utraty rozróżniania kolorów w ciemnym otoczeniu odbywa się przez zmianę barwy piksela z obrazu wejściowego przez odpowiedni odcień szarości jeśli luminancja piksela mieści się w obszarze skotopowym, oraz przez interpolację między barwą skotopową (odcieniem szarości) i barwą fotopową w przedziale luminancji mezopowych. Dla luminancji mniejszych od 0.0056 cd/m^2 , gdy aktywne są tylko pręciki, kolor z obrazu wejściowego zastępowany jest odcieniem szarości określonej przez luminancję skotopową, którą przybliżyć można w następujący sposób:

$$L_{scot} = Y \left[1.33 \left(1 + \frac{Y + Z}{X} \right) - 1.68 \right] \quad (2.25)$$



Rysunek 2.15: Rezultaty symulacji odbłasków (lewa strona), zmiany postrzegania kolorów razem z odbłaskami (środek), oraz zmiany ostrości widzenia razem ze zmianą postrzegania kolorów i odbłaskami (prawa strona).

gdzie L_{scot} to luminancja skotopowa a X , Y , Z to kolor piksela w modelu CIE XYZ (Y to luminancja). Dla luminancji większych od 5.6 cd/m^2 , gdy aktywne są tylko czopki, używamy zwykłej luminancji i koloru, czyli piksele obrazu wejściowego pozostają bez zmian. Między 0.0056 cd/m^2 a 5.6 cd/m^2 , w przedziale mezopowym, aktywne są zarówno czopki jak i pręciki. Wartość pikseli w tym przedziale jest liniowo interpolowana, w zależności od luminancji, między kolorem z przedziału skotopowego i mezopowego.

Środkowa część rys. 2.15 przedstawia rezultat działania metody z symulowaniem odbłasków oraz wrażliwości na kolory dla obrazu testowego o stukrotnie zmniejszonej luminancji. Jak widać w jasnych obszarach obrazka, przy lampach kolory są dobrze widoczne i stopniowo ich widzialność słabnie w coraz ciemniejszych obszarach sceny.

Ostrość widzenia

Poza zmianami w postrzeganiu kontrastu i koloru, wzrok ludzki, wraz ze zmniejszaniem się jasności otoczenia, traci ostrość widzenia. Zależność ostrości widzenia od luminancji adaptacyjnej przedstawia rys. 1.9 z rozdziału pierwszego. Ward et al., w celu zasy-mulowania zmiany ostrości widzenia użyli następującej aproksymacji krzywej z rys. 1.9:

$$R(L_a) \approx 17.25 \arctan(1.4 \log_{10}(L_a) + 0.35) + 25.72 \quad (2.26)$$

gdzie $R(L_a)$ to ostrość widzenia w cyklach na stopień.

Symulacja zmiany widzenia przebiega następująco: dla każdego piksela obrazu obliczamy odpowiadającą mu luminancję adaptacyjną przez interpolację czterech najbliższych pikseli z obrazu luminancji adaptacyjnych. Następnie, dla obliczonej luminancji adaptacyjnej obliczamy wartość $R(L_a)$. W zależności od wartości $R(L_a)$ oryginalne piksele zastępowane są przez piksele odpowiednio rozmyte ("zblurowane"). Wartości pikseli rozmytych obliczane są przy użyciu piramidy obrazu (*ang. image pyramid*, opisywane m. in. w [5]). Każdy z poziomów piramidy jest obrazem o coraz większym rozmy-

ciu, czyli o coraz mniejszej ostrości. Interpolując wartości pikseli między poszczególnymi poziomami piramidy możemy otrzymać piksele o dowolnej ostrości (oczywiście nie większej od ostrości oryginalnego obrazu). Powyższy proces jest bardzo podobny do często stosowanego obecnie filtrowania trójliniowego (*ang. trilinear filtering*) i mip map.

Rezultatem działania powyższej procedury jest obraz, którego ostrość zmniejsza się w najciemniejszych obszarach scen, i pozostaje niezmienną w obszarach o normalnej jasności. Ilustruje to prawa część rysunku rys. 2.15.

2.4.3. Całą metoda

W celu uzyskania odpowiedniego rezultatu opisane wcześniej części składowe metody muszą być wykonane w odpowiedniej kolejności. Kolejność ta przedstawiona jest w poniższej procedurze:

1. void Mapuj()
2. oblicz obraz luminancji adaptacyjnych
3. oblicz obraz odbłasków
4. dodaj odblaski do obrazu luminancji adaptacyjnych
5. dodaj odblaski do obrazu wejściowego
6. symuluj zmianę ostrości
7. symuluj wrażliwość na kolory
8. oblicz histogram
9. popraw histogram z uwzględnieniem czułości widzenia
10. zastosuj histogram

Jest to kolejność dla pełnej metody. Jeśli nie zależy nam na symulowaniu ludzkiego widzenia możemy pominąć linie 3, 4, 5, 6, 7, lub tylko część z nich jeśli jesteśmy zainteresowani tylko częścią efektów.

2.5. Metoda J. Tumblina i G. Turka z 1999 roku

Metoda przedstawiona przez J. Tumblina i G. Turka w pracy [24] naśladuje technikę malarską. W celu zachowania perceptualnej zgodności malowanej sceny i obrazu artysta zwykle maluje od najogólniejszych do najdokładniejszych elementów. Tworzy w ten sposób hierarchię warstw składających się z krawędzi i cieniowań sceny. Zaczyna od szkicu, zawierającego tylko ostro zarysowane krawędzie otaczające duże gładko cieniowane obszary, przedstawia najbardziej kontrastowe i najważniejsze elementy. Następnie stopniowo dodaje kolejne krawędzie i odcienie, wypełniając wizualnie puste obszary, jednocześnie oddając coraz subtelniejsze detale.

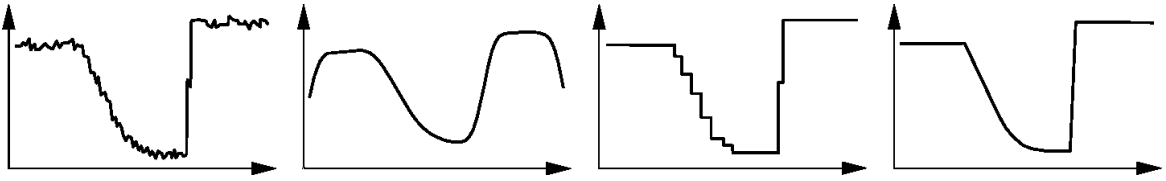
Powyższa metoda działa szczególnie dobrze dla scen wysokokontrastowych. Pozwala na oddzielną regulację kontrastu dla każdej warstwy detalu. Malarz mocno zmniejsza kontrast najogólniejszej warstwy. Następnie dodaje detale, których kontrast zmniejsza nieznacznie lub wcale, przez co zapewnia ich dobrą widoczność w końcowym obrazie.

Metoda Tumblina i Turka naśladuje schemat postępowania malarza tworząc podobną hierarchię stopniowo upraszczanych obrazów. Niskokontrastowy obraz wyjściowy konstruowany jest z hierarchii obrazów przez kompresowanie warstw zawierających naj-

ogólniejsze elementy, oraz dodanie do nich nieskompresowanych warstw zawierających mniejsze detale. Do upraszczania obrazów użyty został filtr nazwany przez autorów "low curvature image simplifier", w skrócie *LCIS*. Filtr ten redukuje obraz do gładkich obszarów ograniczonych ostrymi nieciągłościami gradientów.

2.5.1. Dyfuzja anizotropowa

Dyfuzja anizotropowa, przedstawiona przez Perone i Malika w [18], jest bazą dla filtra *LCIS*. Metoda ta odniosła duży sukces wśród filtrów selektywnie wygładzających (*ang. selective smoothing filter* lub *edge preserving filter*). Jest ona stopniowym, zależnym od czasu procesem upraszczania do obrazu o luminancji przedziałami stałej (rys. 2.16). Zmiana luminancji obrazu w czasie określona jest za pomocą cząstkowych



Rysunek 2.16: Od lewej: oryginalny obraz, dyfuzja izotropiczna, dyfuzja anizotropiczna, LCIS.

równań różniczkowych. Równania te należą do klasy równań opisujących rozchodzenie się ciepła w ciałach stałych. Jeśli potraktujemy luminancję $L(x, y)$ jako temperaturę w punkcie dużej, płaskiej płyty zbudowanej z jednakowego materiału o takiej samej grubości w każdym punkcie, oraz potraktujemy temperaturę jako miarę przepływu ciepła (*ang. heat flow*) na jednostkę powierzchni, wtedy zmiana L w czasie określona jest jako:

$$L_t = \nabla \cdot (-\Phi) = \nabla \cdot (C \nabla L) \quad (2.27)$$

Dolne indeksy oznaczają pochodne częściowe, czyli L_t to $(\partial/\partial t)L(x, y, t)$, L_x to $(\partial/\partial x)L(x, y, t)$, L_{xx} to $(\partial^2/\partial x^2)L(x, y, t)$, itd. C to skalar określający przewodnictwo ciepła materiału, Φ to strumień ciepła, czyli wektor prędkości przyływu ciepła.

W tym klasycznym równaniu ciepła, przepływa ono z obszarów cieplejszych do chłodniejszych, czyli od większych L do mniejszych. Ruchy ciepła określone są przez wektor strumienia Φ , przeciwny do wektora gradientu i przeskalowany przez przewodnictwo materiału, czyli $\Phi = -C \nabla L = (-CL_x, -CL_y)$. Jeśli przewodnictwo jest stałe i równe C_0 , wtedy równanie 2.27 redukuje się do $L_t = C_0(L_{xx} + L_{yy})$. Wraz z upływem czasu obraz zmienia się tak jakby był rozmywany coraz większym filtrem gausowskim, zmierzając do obrazu o takiej samej wartości w każdym pikselu. Szybkość wygładzania zależy od parametru C . Powyższy proces ze stałym C nazywany jest dyfuzją izotropową.

W dyfuzji anizotropowej przewodnictwo zależy od obrazu i zmienia się wraz z upływem czasu. Określa ono lokalną szybkość wygładzania obrazu. Według Perony i Malika [18] w celu znalezienia, utrzymania i wyostrenia krawędzi powinno zmieniać się w stosunku odwrotnym do długości gradientu w punkcie, ponieważ długość gra-

dientu określa prawdopodobieństwo, że piksel znajduje się blisko krawędzi. Niskie przewodnictwo w pobliżu nieciągłości gradientów oraz wysokie w pozostałych rejonach, zachowuje krawędzie i ostre detale jednocześnie szybko wygładza drobniejsze detale i faktury między nimi.

W związku z powyższym dyfuzję anizotropową opisuje równanie:

$$L_t = \nabla \cdot (C(x, y, t) \nabla L) \quad (2.28)$$

$$= \frac{\partial}{\partial x} (C(x, y, t) L_x) + \frac{\partial}{\partial y} (C(x, y, t) L_y) \quad (2.29)$$

gdzie $C(x, y, t) = g(\|\nabla L\|) = g(\sqrt{L_x^2 + L_y^2})$, i (wybrane z [18]):

$$g(m) = \frac{1}{1 + \left(\frac{m}{K}\right)^2} \quad (2.30)$$

gdzie K to przewodnictwo progowe dla m .

Dyfuzja anizotropowa jest szczególnie interesująca ponieważ obydwie jej cechy charakterystyczne: zachowywanie krawędzi oraz wygładzanie, same się wzmacniają. Przedstawia to rysunek rys. 2.17. Obszary małego gradientu mają wysokie przewodnictwo, które umożliwia tym samym dalszą redukcję gradientu. Obszary dużego gradientu mają niskie przewodnictwo, co hamuje przepływ i tym samym tworzy stałą barierę. Wysokie przewodnictwo otaczające bariery powoduje utratę ciepła z dołu krawędzi i napłynięcie ciepła z góry krawędzi, przez co bariera staje się coraz większa i bardziej stroma. Obszar taki ewoluuje do skokowych nieciągłości z nieskończonym gradientem i zerowym przewodnictwem, nazywany "wstrząsem" (*ang. shock*). W rezultacie, dyfuzja anizotropowa przekształca obraz w przedziałami stały, ze skokowymi nieciągłościami w rejonach zawierających krawędzie.



Rysunek 2.17: Samowzmacniające się działanie dyfuzji anizotropowej

Mimo tego, że metoda ta jest stabilna numerycznie i zbieżna dla $t \rightarrow \infty$ do rozwiązania przedziałami stałego, Nitzberg i Shiota w pracy [15] pokazali, że dyfuzja anizotropowa jest źle uwarunkowana. Niewielka zmiana obrazu wejściowego może powodować bardzo duże zmiany w obrazie wyjściowym. Wstrząsy zwykle formują się w lokalnych maksimach gradientu i odpowiadają dokładnie krawędziom obrazu. Jednak tak nie musi być zawsze. Obszary o wysokim, w przybliżeniu jednolitym, gradientie mogą powodować powstawanie wstrząsów w dowolnym miejscu tych obszarów. Zamiast jednej dużej krawędzi, dyfuzja anizotropowa może powodować powstanie kilku rozmieszczonych pozornie w losowy sposób.

2.5.2. Filtr LCIS

Na podstawie dyfuzji anizotropowej Tumblin i Turk stworzyli podobne cząstkowe równanie różniczkowe. Jego działanie jest również samowzmacniające się. Oparte jest na pochodnych wyższego rzędu niż dyfuzja anizotropowa. Dyfuzja anizotropowa prowadzi do powstania obrazu przedziałami stałego poprzez stopniową zmianę gradientów do zera lub nieskończoności. Rezultatem działania równania Tumblina i Turka jest obraz przedziałami liniowy powstający przez stopniową zmianę krzywizny do zera lub nieskończoności. Zwykle warunki brzegowe zapobiegają rozwiązaniom o zerowych krzywiznach, stąd angielska nazwa metody to *"low curvature image simplifier"* (LCIS).

Tak jak dyfuzja anizotropowa, LCIS można traktować jako opis rozkładu ciepła w ciele stałym. Jednak zarówno "siły powodujące" przemieszczanie się ciepła, jak i przewodnictwo obliczane są inaczej. Siłą powodującą w dyfuzji anizotropowej jest ujemny gradient $-\nabla L$. W celu wyznaczenia "siły powodującej" dla LCIS w punkcie (x, y) obrazu autorzy definiują linię l przez ten punkt, skierowaną w kierunku θ i parametr α określający odległość na linii l . Obliczamy wartość luminancji L_α obrazu wzdłuż linii l , poczym znajdujemy jej trzecią pochodną, czyli zmianę krzywizny w kierunku θ :

$$A = x_0 + \alpha \cos \theta \quad (2.31)$$

$$B = y_0 + \alpha \sin \theta \quad (2.32)$$

$$L_\alpha(A, B) = L_x(A, B) \cos \theta + L_y(A, B) \sin \theta \quad (2.33)$$

$$= \left(\left(\frac{\partial}{\partial x} \right) \cdot \cos \theta + \left(\frac{\partial}{\partial y} \right) \cdot \sin \theta \right) L \quad (2.34)$$

$$L_{\alpha\alpha\alpha} = \left(\left(\frac{\partial}{\partial x} \right) \cdot \cos \theta + \left(\frac{\partial}{\partial y} \right) \cdot \sin \theta \right)^3 L \quad (2.35)$$

gdzie (A, B) to para współrzędnych (x, y) na linii l , L_α to pochodna L wzdłuż linii l a $L_{\alpha\alpha\alpha}$ to trzecia pochodna wzdłuż linii l .

Jeśli $L_{\alpha\alpha\alpha} > 0$, to krzywa $L_{\alpha\alpha}$ rośnie w punkcie (x_0, y_0) linii l , przepływ w kierunku θ linii l pomógłby wyrównać krzywiznę po obu stronach (x_0, y_0) i zmniejszyć $L_{\alpha\alpha\alpha}$. W związku z powyższym siłę powodującą wzdłuż linii l określać będzie $L_{\alpha\alpha\alpha}$, a zsumowane $L_{\alpha\alpha\alpha}$ ze wszystkich kierunków dadzą składowe x i y wektora siły powodującej. Aby składowe nie myliły się z indeksami pochodnych częściowych będziemy oznaczali je jako E (*ang. East* czyli Wschód) i N (*ang. North* czyli Północ) i tak $F = (f_E, f_N)$ definiujemy jako:

$$f_E = \frac{1}{\pi} \int_0^{2\pi} L_{\alpha\alpha\alpha} \cos \theta d\theta = \frac{3}{4} (L_{xxx} + L_{yyx}), \quad (2.36)$$

$$f_N = \frac{1}{\pi} \int_0^{2\pi} L_{\alpha\alpha\alpha} \sin \theta d\theta = \frac{3}{4} (L_{xxy} + L_{yyy}). \quad (2.37)$$

Przewodnictwo w LCIS określone jest tak jak w równaniu 2.30 z parametrem m spełniającym równanie:

$$m^2 = 0.5 (L_{xx}^2 + L_{yy}^2) + L_{xy}^2 \quad (2.38)$$

Ogólnie równanie LCIS to:

$$L_t(x, y, t) = \nabla \cdot (C(x, y, t)F(x, y, t)) \quad (2.39)$$

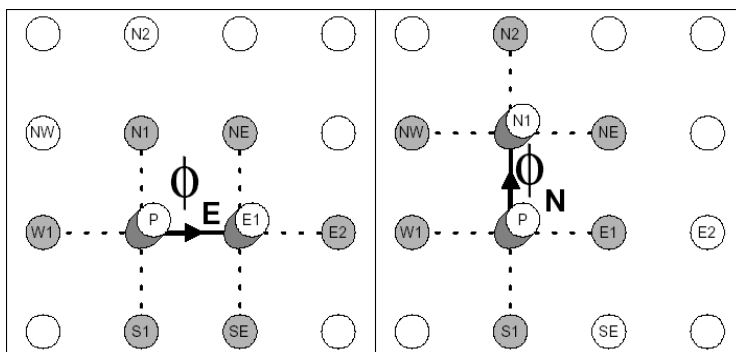
gdzie $F(x, y, t) = (f_E, f_N)$, a $C(x, y, t) = g(m)$ z m obliczonym z równania 2.38.

Podobnie jak duże obszary wysokich gradientów mogą powodować powstawanie wielu nieprzewidywalnych skoków w przypadku dyfuzji anizotropowej, tak LCIS może powodować powstawanie skoków w dużych rejonach jednakowej i dużej krzywizny. Jednak obszary dużej krzywizny są zwykle mniejsze, ponieważ wymagają szerokich przedziałów luminancji. Poza tym skoki te są mniej widoczne od skoków tworzonych przez dyfuzję anizotropową.

Obydwie metody zdefiniowane są dla ciągłych zmiennych x i y . Aby je zaimplementować trzeba użyć ich dyskretnej aproksymacji.

2.5.3. Implementacja LCIS

Do implementacji LCIS autorzy użyli jawnego schematu z ustalonym i stałym krokiem Δt . Dla każdego kolejnego t zwiększanego każdorazowo o Δt , obliczany jest nowy obraz $L(t)$. W każdym przedziale czasowym Δt przewodnictwo i przepływ są stałe. Każdy kolejny obraz $L(t + \Delta t)$ zależy tylko od poprzedniego obrazu $L(t)$. Każdy piksel obliczany jest z ustalonego zbioru sąsiednich pikseli poprzedniego obrazu.



Rysunek 2.18: Połączenia między pikselami (strzałki), EW po lewej, NS po prawej. W każdej iteracji strumień Φ_E obliczany jest z wartości szarych pikseli po lewej stronie obrazka a Φ_N z szarych pikseli po prawej stronie obrazka.

Kontynuując analogię z przepływem, wyobraźmy sobie obraz jako siatkę zbiorników zawierających ciecz. Każdy piksel to osobny zbiornik, zawierający ilość cieczy odpowiadającą luminancji piksela. Piksel jest połączony ze swoimi czterema sąsiadami. Wymiana cieczy (bądź ciepła) ma miejsce tylko przez połączenia między pikselami. Jak widać na rys. 2.18 piksel P połączony jest z pikselami $E1$, $N1$, $W1$ i $S1$. W ciągu przedziału czasowego Δt przez połączenia przepływa strumień Φ , zmieniający luminancję źródła i zwiększający luminancję celu. W każdym przedziale czasowym oblicza się wartość strumienia dla każdego połączenia, następnie reguluje się luminancję obrazu zgodnie z obliczonymi wartościami strumieni dla połączeń, tworząc w ten sposób nowy obraz.

Rozróżniamy dwa rodzaje połączeń: poziome, oznaczone EW (*ang. East West*, czyli Wschód Zachód) i pionowe, oznaczone NS (*ang. North South*, czyli Północ Południe). Lewa strona rys. 2.18 przedstawia połączenie EW łączące piksele P i $E1$. Strumień Φ_E przepływający przez to połączenie obliczany jest z wartości ośmiu sąsiednich pikseli (piksele oznaczone kolorem szarym i połączone przerywaną linią po lewej stronie rys. 2.18). Kierunek strumienia określa znak Φ , jeśli $\Phi_E > 0$ to strumień przepływa z P do $E1$, zmniejszając P i zwiększając o taką samą wartość $E1$. Analogicznie połączenie NS łączy piksele P i $N1$ a $\Phi_N > 0$ oznacza przepływ z P do $N1$.

Siła powodująca w każdym połączeniu obliczana jest za pomocą oszacowań trzecich pochodnych częściowych luminancji obrazu w środkach połączeń. Dla połączenia EW między P i $E1$ jest:

$$L_{xxx} = (E2 - W1) + 3(P - E1), \quad (2.40)$$

$$L_{yyx} = (NE - N1) + (SE - S1) + 2(P - E1), \quad (2.41)$$

a dla połączenia NS między P i $N1$:

$$L_{yyy} = (N2 - S1) + 3(P - N1), \quad (2.42)$$

$$L_{xxy} = (NE - E1) + (NW - W1) + 2(P - N1). \quad (2.43)$$

Przewodnictwo połączeń obliczane jest z przybliżeń drugich pochodnych częściowych. Definiujemy:

$$\begin{aligned} P_{xx} &= E1 + W1 - 2P, \\ P_{yy} &= N1 + S1 - 2P, \\ E_{xx} &= E2 + P - 2E1, \\ E_{yy} &= NE + SE - 2E1, \\ N_{xx} &= NE + NW - 2N1, \\ N_{yy} &= N2 + P - 2N1, \\ N_{xy} &= (NE - E1) - (N1 - P), \\ S_{xy} &= (E1 - SE) - (P - S1), \\ W_{xy} &= (N1 - P) - (NW - W1). \end{aligned}$$

Kwadrat parametru m z równania 2.30 wynosi dla połączenia EW :

$$m_{EW}^2 = (P_{xx}^2 + P_{yy}^2 + E_{xx}^2 + E_{yy}^2)/4 + (N_{xy}^2 + S_{xy}^2)/2 \quad (2.44)$$

i dla połączeń NS :

$$m_{NS}^2 = (P_{xx}^2 + P_{yy}^2 + N_{xx}^2 + N_{yy}^2)/4 + (W_{xy}^2 + N_{xy}^2)/2. \quad (2.45)$$

Strumień przepływający przez każde połączenie określony jest jako $\Phi_E = \Delta t f_E C_E$ i $\Phi_N = \Delta t f_N C_N$. Autorzy zalecają użycie $\Delta t \leq \frac{1}{32}$ w celu zachowania stabilności. f_E i f_N to siły powodujące zdefiniowane w równaniach 2.36 i 2.37. Przewodnictwo połączeń NS i EW to:

$$C_E = \frac{1}{1 + \left(\frac{m_{EW} D_E}{K}\right)^2} \quad (2.46)$$

$$C_N = \frac{1}{1 + \left(\frac{m_{NS} D_N}{K}\right)^2} \quad (2.47)$$

gdzie m_{EW} i m_{NS} określają równania 2.44 i 2.45, D_E i D_N to parametry opisane poniżej, początkowo równe 1.0.

Przybliżanie pochodnych obrazu różnicami wartości pikseli prowadzi do problemu nazwanego przez autorów wyciekaniem (*ang. leakage*). Wstrząsy w ciągłym obrazie są doskonale nieprzepuszczalne, hamując jakikolwiek przepływ. Problem wyciekania polega na tym, że wstrząsy, zarówno gradientów jak i krzywizn, w dyskretnym obrazie, w związku ze skończeniem małymi odległościami między pikselami, nie są nieskończone. Powoduje to tym samym niewielki przepływ przez krawędzie. Drobne wycieki nawarstwiają się wraz z upływem czasu, może to prowadzić nawet do całkowitego zaniku niektórych krawędzi.

Proponowanym przez autorów rozwiązaniem problemu jest dodanie samodostrajających się parametrów D_E i D_N do równań 2.46 i 2.47. Obydwa parametry zapamiętywane są dla każdego połączenia EW i NS odpowiednio. Działanie parametrów D_E i D_N opiera się na obserwacji, że tworzenie się wstrząsów powoduje szybkie oddalanie się m od wartości przewodnictwa progowego K . Na krawędziach m zwiększa się szybko do nieskończoności, w pozostałych miejscach spada w okolice zera. W celu identyfikacji i zapobieżeniu wyciekom, w każdej iteracji metody porównujemy m do K , aby znaleźć połączenia przecinające krawędzie obrazu. Dla tych połączeń dostrajamy odpowiednio D_E lub D_N tak aby zwiększyć parametr m i zmniejszyć przewodnictwo do zera. Parametry D_E i D_N rosną wykładniczo w połączeniach o m stale większych od K i ustalają się w pobliżu 1.0 jeśli m będzie mniejsze od K . Początkowo $D_E = D_N = 1.0$ i dalej, w każdej iteracji zmienia się zgodnie z następującą definicją:

$$D_E = \begin{cases} D_E(1.0 + m_{EW}) & \iff m_{EW} > K \\ 0.9D_E + 0.1 & \iff m_{EW} \leq K \end{cases} \quad (2.48)$$

$$D_N = \begin{cases} D_N(1.0 + m_{NS}) & \iff m_{NS} > K \\ 0.9D_N + 0.1 & \iff m_{NS} \leq K \end{cases} \quad (2.49)$$

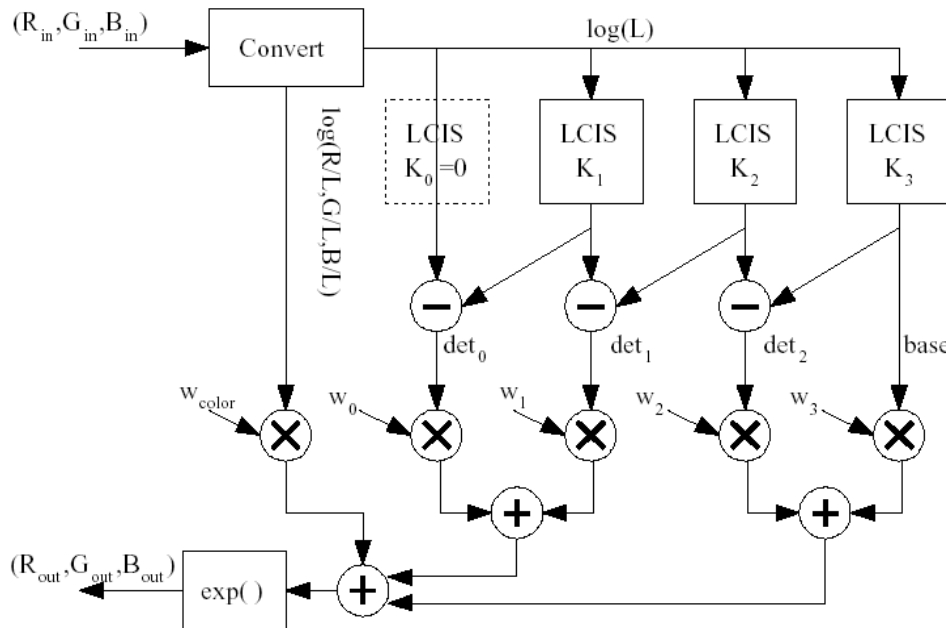
Każde połączenie o D_E lub D_N większym od 10 zaznaczane jest jako graniczne. Dodatkowo piksele między nimi również zaznaczone zostają jako graniczne i wstrzymuje się wszelki przepływ z ich udziałem.

Powyższe postępowanie eliminuje pojawianie się widocznych artefaktów powodowanych przez problem wyciekania.

2.5.4. Redukcja kontrastu

Filtr LCIS naśladuje w odwrotnej kolejności malowanie obrazu. Stopniowo usuwa detale z obrazu, zostawiając tylko gładkie obszary oddzielone ostrymi krawędziami. Usunięte detale można odzyskać przez odjęcie od oryginału obrazu wygładzonego LCISem. Następnie kontynuując schemat postępowania malarza, w celu redukcji kontrastu, mocno kompresuje się obraz uproszczony i dodaje się do niego lekko lub wcale nie skompresowany detal.

Cała metoda przedstawiona przez Tumblina i Turka wykorzystuje LCIS kilkakrotnie, zwykle 4, za każdym razem upraszczając obraz coraz bardziej. Tworzy w ten sposób ich hierarchię, którą przedstawia rys. 2.19. Zaczynamy od zlogarytmowania obrazu (logarytmem dziesiętnym) tak aby różnice pikseli odpowiadały bezpośrednio kontrastowi. LCIS stosowany jest tylko do luminancji obrazu, kolor rekonstruowany jest na końcu



Rysunek 2.19: Redukcja kontrastu przy użyci hierarchii obrazów LCIS.

w sposób przedstawiony przez Schlicka w [21]. Następnie tworzymy coraz bardziej uproszczone obrazy, stosując filtr LCIS o coraz większych wartościach K . Zaczynamy od $K_0 = 0$ (filtr nie zmienia obrazu), kolejne wartości $K_{1,2,3}$ mieszczą się zwykle między 0.02 a 0.32 (mogą być równe np. $K_1 = 0.06$, $K_2 = 0.10$, $K_3 = 0.16$). Filtr z największym K tworzy tzw. warstwę bazową. Obliczamy obrazy zawierające coraz większe szczegóły (det_0 , det_1 , ..., od *ang. detail* czyli szczegół) przez odjęcie obrazów K_i od oryginału, z najmniejszymi szczegółami w det_0 . Następnie kompresujemy ich kontrast, mnożąc przez wagi $w_{0,1,2,3}$ (o wartościach np. 1.0, 0.8, 0.4, 0.16), $w_3 = w_{color}$, dodajemy do siebie i potęgujemy w celu uzyskania skompresowanych luminancji obrazu.

2.6. Metoda E. Reinharda et al. z 2002 roku

Zadanie stawiane przed metodami mapowania tonów, czyli mapowanie potencjalnie wysoko kontrastowych luminancji świata rzeczywistego na zakres nisko kontrastowy, jest jednym z zadań fotografii. Dlatego też metoda przedstawiona przez Reinharda et al. [20] bazuje na schemacie postępowania używanym od dawna w fotografii zwanym systemem strefowym [1].

2.6.1. System strefowy

System strefowy pomaga przewidzieć fotografowi w jaki sposób luminancja kadrowanej sceny przeniesiona zostanie na negatyw fotograficzny. Strefę definiujemy jako rzymską liczbę z przypisaną jej luminancją rzeczywistej sceny oraz odpowiadającą jej reflek-

tancją wydruku. W systemie strefowym wyróżnia się jedenaście stref od 0 do X. Strefa 0 odpowiada czystej czerni, strefa X czystej bieli, każda kolejna strefa podwaja intensywność poprzedniej. Schemat postępowania fotografa jest następujący: fotograf mierzy luminancję powierzchni, którą ocenia jako średnią jasność w kadrowanej scenie, jasność tą nazywa się średnią szarością (*ang. middle – gray*), używa się też określenia szara kartka Kodaka, zwykle mapuje się ją na V strefę, która odpowiada 18% refleksyjności zdjęcia. Następnie fotograf wykonuje pomiary luminancji obszarów, które na zdjęciu miałyby być najjaśniejsze i najciemniejsze. Mając pomiary jasnych i ciemnych obszarów, fotograf może określić zakres dynamiczny kadrowanej sceny. Jeśli zakres ten nie przekracza 9 stref, to przy odpowiednim wyborze średniej szarości na zdjęciu da się odwzorować scenę bez straty w widoczności szczegółów poszczególnych powierzchni. Jeśli zakres będzie większy niż 9 stref, wtedy część jaśniejszych bądź ciemniejszych (lub jednych i drugich) zostanie odwzorowana odpowiednio na czystą biel, lub na czystą czerń. W tych obszarach informacja o fakturze zostanie stracona. Czasami taka utrata szczegółów jest pożądana, aby bardzo jasne obiekty, źródła światła lub rozbłyśki zwierciadlane zostały odwzorowane na czystą biel. W obszarach, na których tak nie jest, fotograf może zastosować technikę nazywaną wysłanianiem i doświetlaniem (*ang. dodging and burning*). Polega to na niedoświetleniu (wysłonieniu), bądź prześwietleniu (doświetleniu) części zdjęcia podczas naświetlania. Można przy jej użyciu ściemnić bądź rozjaśnić pewne obszary zdjęcia tak aby nie były one mapowane na brzegowe strefy.

2.6.2. Algorytm

Powyższy schemat postępowania stanowi bazę dla operatora Reinharda et al. Autorzy nie starali się dokładnie odwzorować procesu przedstawionego powyżej, zamiast tego zapożyczyli jego warstwę konceptualną. Na początku skalowali luminancję obrazu co odpowiadało ustawieniu parametrów ekspozycji zdjęcia (wielkość i czas otwarcia przysłony). Następnie, jeśli było to konieczne używali zautomatyzowanego wysłaniania i doświetlania.

Proces skalowania określa równanie:

$$L(x, y) = \frac{a}{\bar{L}_w} L_w(x, y) \quad (2.50)$$

gdzie $L(x, y)$ to przeskalowana luminancja. $\bar{L}_w$ to przybliżenie średniej jasności sceny, zdefiniowane wzorem:

$$\bar{L}_w = \frac{1}{N} \exp \left(\sum_{x, y \in \Omega} \log(\delta + L_w(x, y)) \right) \quad (2.51)$$

gdzie Ω to zbiór wszystkich pikseli, N to liczba wszystkich pikseli obrazu, δ to mała wartość, zapobiegająca nieoznaczoności logarytmu, gdy $L_w(x, y)$ jest równe zero. Parametr a określa średnią szarość obrazu (*middle – gray*). W normalnych warunkach chcielibyśmy odwzorować średnią jasność na typową średnią szarość, czyli 0.18 w skali od 0 do 1, dlatego zwykle $a = 0.18$. Jednak w scenach jaśniejszych, bądź ciemniejszych niż przeciętne, średnią jasność obrazu należałoby odwzorować na inną wartość a . Typowe wartości a to 0.18, 0.36, 0.72 oraz 0.09, 0.045. W dalszym ciągu parametr a będzie nazywany wartością kluczową. Operator 2.50 nie poradzi sobie z często spotykanymi

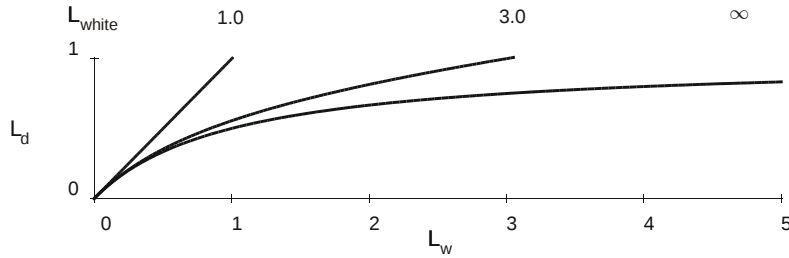
scenami, których zakres jest standardowy za wyjątkiem kilku bardzo jasnych obszarów. Tradycyjna fotografia radzi sobie z tego rodzaju problemami poprzez kompresję wysokich oraz niskich luminancji. Filmy fotograficzne mają krzywą odwzorowania tonów kształtu "S". Autorzy proponują prosty operator kompresujący tylko wysokie luminancje, definiuje go następujący wzór:

$$L_d(x, y) = \frac{L(x, y)}{1 + L(x, y)} \quad (2.52)$$

luminancja L zaaplikowana do równania 2.52 zostanie przeskalowana w przybliżeniu o $\frac{1}{L}$, jeśli L jest duże, oraz w przybliżeniu o L w przypadku małego L . Wzór 2.52 gwarantuje utrzymanie każdej luminancji L w zakresie wyświetlacza (czyli $[0..1]$), takie zachowanie nie zawsze jest pożądane. Aby umożliwić sterowanie od jakiej wartości luminancja będzie mapowana na czystą biel równanie 2.52 modyfikuje się do następującej postaci:

$$L_d(x, y) = \frac{L(x, y)(1 + \frac{L(x, y)}{L_{white}^2})}{1 + L(x, y)} \quad (2.53)$$

gdzie L_{white} to najmniejsza luminancja sceny mapowana na czystą biel. Mapowanie dla różnych wartości L_{white} przedstawia rys. 2.20. Dla $L_{white} = \infty$ równanie 2.53 jest



Rysunek 2.20: Funkcja z równania 2.53 dla różnych wartości parametru L_{white} .

tożsame z równaniem 2.52. Dla $L_{white} = L_{w \max}$, $L_{w \max}$ mapowane jest na $L_{d \max}$, czyli wykorzystany zostaje cały zakres wyświetlacza. Domyślną wartością L_{white} jest $L_{w \max}$.

Dla wielu obrazów powyższa technika daje zadowalające rezultaty. Zachowuje detale w obszarach o niskim kontraście, kompresuje luminancję sceny do zakresu wyświetlacza. Jednak nie jest ona w stanie zachować szczegółów w obszarach o wysokim kontraście. Mapowanie takich obrazów wymaga zastosowania metody lokalnej, implementującej wysłanianie i doświetlanie (*ang. dodging and burning*).

W tradycyjnym przysłanianiu i doświetlaniu dla każdego obszaru zdjęcia czas naświetlania może być inny. Wydłużając czas naświetlania jasnych rejonów można je przyciemnić (doświetlanie), na odwrót dany obszar można rozjaśnić skracając jego czas naświetlania (przysłanianie). Zautomatyzowanie procesu umożliwi kompresję tonów obrazów HDR nawet o bardzo dużym zakresie dynamicznym. Proces ten można wyobrazić sobie jako dobieranie wartości kluczowej (czyli a w równaniu 2.50) dla każdego piksela obrazu.

Przysłanianie i doświetlanie zwykle stosuje się na całej powierzchni ograniczonej



Rysunek 2.21: Zjawisko halo (odwrotne gradienty) wokół źródeł światła. Z lewej poprawny obraz, w środku lekkie halo, z prawej mocne halo. Oryginalny obrazek to `cornellbox.hdr` pobrany ze strony E. Reinharda (<http://www.cs.utah.edu/~reinhard/cdrom/hdr/>)

wysokim kontrastem, na przykład drzewo na tle nieba. Do wyznaczenia powierzchni Reinhard używa funkcji V zwracającej kontrast w pewnym otoczeniu punktu (x, y) , funkcja ta używa filtrów gausowskich postaci:

$$R_i(x, y, s) = \frac{1}{\pi(\alpha_i s)^2} \exp\left(-\frac{x^2 + y^2}{(\alpha_i s)^2}\right) \quad (2.54)$$

gdzie x, y to współrzędne piksela, s określa bazową wielkość filtra a α_i stosunek wielkości kolejnych filtrów. Za pomocą R_i obliczany jest spłot V_i :

$$V_i(x, y, s) = L(x, y) \otimes R_i(x, y, s) \quad (2.55)$$

Funkcja V ocenia kontrast w punkcie jako różnicę dwóch spłotów obrazu z różnej z różnej wielkości filtrami R_i . Ostatecznie V wygląda następująco:

$$V(x, y, s) = \frac{V_1(x, y, s) - V_2(x, y, s)}{2^\phi a / s^2 + V_1(x, y, s)} \quad (2.56)$$

gdzie V_1 i V_2 to funkcje ze wzoru 2.55, a to wartość kluczowa, ϕ to parametr odpowiedzialny za ostrość. V_1 w mianowniku ma na celu uniezależnienie wartości V od absolutnej luminancji, natomiast $2^\phi a / s^2$ zapobiega przyjmowaniu przez V zbyt dużych wartości gdy V_1 jest bliskie zeru.

Obszar niskiego kontrastu to największe otoczenie piksela, w którym nie ma żadnych dużych zmian kontrastu. Aby obliczyć rozmiar tego obszaru obliczana jest wartość V dla różnych wartości parametru s . V_1 określa średnią luminancję w otoczeniu piksela (x, y) , V_2 również, lecz w większym otoczeniu. Wartości V_1 i V_2 powinny być bardzo zbliżone w obszarach o niskim kontraście i różne w obszarach o wysokim kontraście. Aby obliczyć rozmiar największego obszaru niskiego kontrastu szukamy, zaczynając od najmniejszego rozmiaru, takiego s aby

$$|V(x, y, s)| < \varepsilon \quad (2.57)$$

było fałszywe, gdzie ε to ustalona wartość progowa.

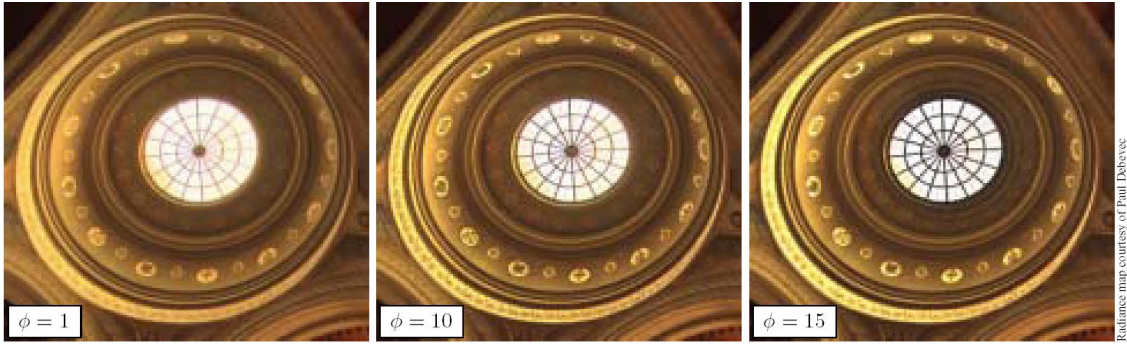
Mając już wyznaczony rozmiar s dla piksela (x, y) , $V_1(x, y, s)$ można użyć jako średniej luminancji otoczenia. W końcu z operatora globalnego z równania 2.52 otrzymu-

jemy operator lokalny postaci:

$$L_d(x, y) = \frac{L(x, y)}{1 + V_1(x, y, s(x, y))} \quad (2.58)$$

Dla $L < V_1$ w relatywnie jasnym otoczeniu L_d będzie mniejsze niż ta otrzymana z równania 2.52, sytuacja taka odpowiada doświelaniu. Podobnie luminancja piksela jaśniejszego od swojego relatywnie ciemnego otoczenia będzie mniej skompresowana, czyli piksel zostanie przysłonięty. W obydwu przypadkach lokalny kontrast zwiększy się.

Bardzo ważne jest prawidłowe wyznaczenie s . Jeśli jest ono zbyt małe V_1 jest zbyt bliskie L i w rezultacie operator lokalny zachowuje się podobnie do operatora globalnego. Jeśli z kolei s jest za duże na obrazie powstają czarne obwódki (*ang. halo artefacts*) wokół jasnych obszarów rys. 2.21. Generalnie użycie większego rozmiaru s zwiększa kontrast oraz wyostża krawędzie. Za pomocą ε z równania 2.57, oraz ϕ z równania 2.56 można sterować ostrością krawędzi (rys. 2.22). Zmniejszenie ε , tak samo jak zwiększenie ϕ , powoduje wybieranie większych s .



Rysunek 2.22: Różne wartości parametru ϕ . Obrazek z pracy [20].

2.6.3. Szczegóły implementacyjne

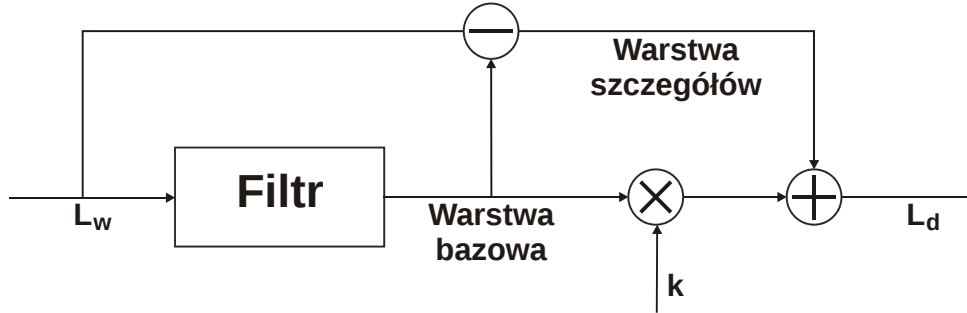
Autorzy używali następujących ustawień domyślnych :

- ε wynosiło 0.005,
- ϕ zwykle 8,
- używali filtrów gausowskich R_i w 8 rozmiarach, najmniejszy miał szerokość jednego piksela, największy 43 piksele, odpowiednie wielkości filtrów otrzymali stosując $\alpha_1 = 1/2\sqrt{2} \approx 0.35$, $\alpha_i = \alpha_1(1.6)^{i-1}$ ($i = 2, \dots, 8$)
- wartość kluczowa (parametr a z równania 2.50) była ustalana osobno dla każdego obrazka

2.7. Metoda F. Duranda i J. Dorsey z 2002 roku

Metoda przedstawiona przez Duranda i Dorsey [10] opiera się, tak jak metoda

Tumblina i Turka [24], na dekompozycji obrazu na warstwy. W odróżnieniu od metody Tumblina i Turka, obraz rozkładany jest tylko na dwie warstwy: warstwę bazową zawierającą "ogólne" informacje o obrazie, tzn. poziom jasności i mocno widoczne kontury, oraz na warstwę zawierającą szczegóły. Następnie zmniejszeniu ulega tylko kontrast warstwy bazowej, po czym warstwy są z powrotem łączone. Dzięki zmniejszeniu kontrastu jedynie warstwy bazowej zachowana zostaje widoczność szczegółów. Schemat działania metody przedstawia rys. 2.23.



Rysunek 2.23: Schemat działania operatora dwu-warstwowego. k to współczynnik kompresji warstwy bazowej.

2.7.1. Filtr bilateralny

W celu uzyskania warstwy bazowej, należy zastosować filtr selektywnie wygładzający (*ang. selective smoothing* lub inaczej *edge preserving filter*). Ma on na celu wygładzenie małych wahań jasności obrazu, takich jak szum i faktura powierzchni, jednocześnie zachowując duże wahania, czyli krawędzie. Działanie filtra na zwykłe obrazki przedstawia rys. 2.24.

Durand i Dorsey użyli filtra bilateralnego przedstawionego przez Tomasiego i Manduchiego w 1998 roku [23]. Filtrowanie bilateralne jest alternatywą dla znanej w grafice komputerowej metody dyfuzji anizotropowej [18], która stanowi podstawę dla metody Tumblina i Turka przedstawionej w rozdziale 2.5. Sporowadza się ona do iteracyjnego rozwiązania częściowego równania różniczkowego. Atutami filtra bilateralnego jest nieiteracyjność, prostota i możliwość optymalizacji.

Działanie filtra polega na zastosowaniu do piksela (x, y) zwykłego filtra gausowskiego f , którego wagi zależą dodatkowo od określonej na luminancji obrazu, funkcji g , zmniejszającej wagi pikseli o dużej różnicy luminancji. Wynikiem działania filtra na piksel $s = (x, y)$ jest:

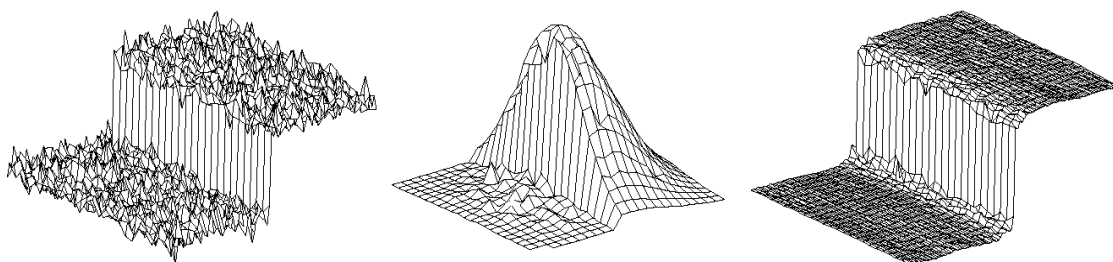
$$J(s) = \frac{1}{k(s)} \sum_{p \in \Omega} f(p - s) g(L(p) - L(s)) L(p) \quad (2.59)$$

gdzie s i p to piksele o współrzędnych (x_s, y_s) i (x_p, y_p) , k to funkcja normalizująca:

$$k(s) = \sum_{p \in \Omega} f(p - s) g(L(p) - L(s)) \quad (2.60)$$



Rysunek 2.24: Działanie filtra bilateralnego na zwykły obraz. Z lewej oryginalny obraz, w środku obraz rozmyty zwykłym filtrem gausowskim, z prawej obraz przefiltrowany bilateralnie. Lewy i prawy obrazek zapożyczony z pracy [4].



Rysunek 2.25: Filtrowanie bilateralne, z lewej sygnał niefiltrowany, w środku filtr dla sygnału z lewej, po prawej przefiltrowany sygnał. Obrazki z pracy [23].

W gładkim obszarze piksele w bliskim otoczeniu mają podobną luminancję, a filtr zachowuje się jak zwykle wygładzanie, uśredniając drobne szumy. W przypadku ostrej krawędzi między jasnym i ciemnym obszarem (rys. 2.25 lewo), jeśli rozważamy piksel po jasnej stronie, funkcja g będzie zwracała wartości bliskie 1 dla pikseli leżących po jasnej stronie, oraz wartości bliskie 0 dla pikseli po ciemnej stronie (rys. 2.25 środek). W rezultacie filtr zamieni jasny piksel przy krawędzi przez średnią jasnych pikseli z jego otoczenia, ignorując ciemne piksele (rys. 2.25 prawo) [23].

Zwykle g jest, tak jak f , postaci gausowskiej, więc:

$$g(t) = \exp\left(-\frac{t^2}{\sigma_r^2}\right), \quad (2.61)$$

$$f(x, y) = \exp\left(-\frac{x^2 + y^2}{\sigma_s^2}\right) \quad (2.62)$$

Parametry σ_r (r od angielskiego *range*), σ_s (s od angielskiego *spacial*) określają rozmiary filtrów.

2.7.2. Przyspieszanie filtrowania

Implementacja filtrowania bilateralnego bezpośrednio ze wzoru 2.59 wymagałaby $O(nm)$ czasu, gdzie n to liczba pikseli obrazu, a m to liczba pikseli filtra f . Autorzy opracowali dwie strategie poprawy szybkości filtrowania: przedziałami liniowe aproksymowanie filtra i wykonywanie części obliczeń na obrazie o zmniejszonej rozdzielczości.

Przedziałami liniowy filtr bilateralny

Obliczanie splotów może być znacznie przyspieszone przy użyciu szybkiej transformaty Fouriera (*FFT* – *ang. Fast Fourier Transform*). Zamiast w $O(nm)$ (m to rozmiar filtra) obliczenia można wykonać w czasie $O(n \log n)$, co przy dużym m daje spory zysk. Niestety filtr bilateralny nie jest splotem, ponieważ zależy od luminancji obrazu, na którym działa. Powodem tego jest funkcja g pobierająca dla punktu s jako argument różnicę luminancji $L(p) - L(s)$. Jednak, jeśli rozważymy równanie 2.59 dla ustalonego piksela s , staje się ono splotem funkcji $H^{L(s)}$ postaci:

$$H^{L(s)} : p \rightarrow g(L(p) - L(s))L(p)$$

z filtrem f . Analogicznie czynnik normalizujący k jest splotem funkcji $G^{L(s)}$ postaci:

$$G^{L(s)} : p \rightarrow g(L(p) - L(s))$$

z f .

W związku z powyższym cały zakres luminancji sceny dzielony zostaje na $NB_SEGMENTS$ (liczba segmentów) wartości $\{l^j\}$. Dla każdej wartości l^j oblicza się:

$$\begin{aligned} J^j(s) &= \frac{1}{k^j(s)} \sum_{p \in \Omega} f(p - s) g(L(p) - L(s)) L(p) \\ &= \frac{1}{k^j(s)} \sum_{p \in \Omega} f(p - s) H^{l^j}(p) \end{aligned} \quad (2.63)$$

oraz

$$\begin{aligned} k^j(s) &= \sum_{p \in \Omega} f(p - s) g(L(p) - L(s)) \\ &= \sum_{p \in \Omega} f(p - s) G^{l^j}(p) \end{aligned} \quad (2.64)$$

Ostatecznie wynikiem działania filtra jest dla piksela s wartość liniowo interpolowana między dwiema wartościami l^j najbliższymi $L(s)$. Rysunek rys. 2.26 przedstawia pseudokod metody. Wykres rys. 2.27 przedstawia wyniki działania metody zwykłej i przedziałami liniowej.

Zmniejszenie rozdzielczości

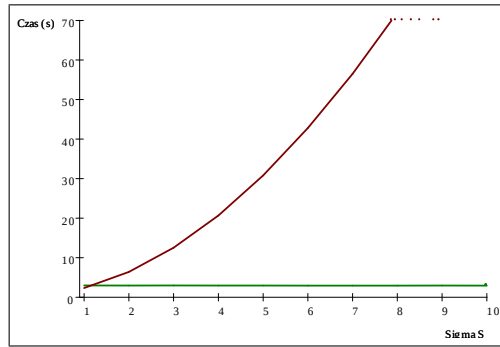
W celu dalszego przyspieszenia należy zauważyć, że wszystkie operacje, poza ostatnią interpolacją, działają na niskie częstotliwości obrazu. Dlatego jedynie z małą stratą jakości, można użyć zmniejszonej wersji obrazu do wykonania wszystkich operacji,

```

PrzedziałamiLiniowyFiltrBilateralny
    (obraz  $L$ , filtr  $f_{\sigma_s}$ , filtr  $g_{\sigma_r}$ )
     $J = 0$  // zerujemy wszystkie piksele  $J$ 
    for  $j = 0..NB\_SEGMENTS$ 
         $l^j = \min(L) + j \times (\max(L) - \min(L)) / NB\_SEGMENTS$ 
         $G^j = g_{\sigma_r}(L - l^j)$  // obliczamy  $g_{\sigma_r}$  w każdym piksle
         $H^j = G^j \times L$  // obliczamy  $H$  w każdym piksle
         $K^j = G^j \otimes f_{\sigma_s}$  // obliczamy czynnik normalizujący
         $H^j = H^j \otimes f_{\sigma_s}$ 
         $J^j = H^j / K^j$ 
     $J = J + J^j \times \text{InterpolowanaWaga}(L, l^j)$ 

```

Rysunek 2.26: Pseudokod przedziałami liniowego filtra bilateralnego.



Rysunek 2.27: Porównanie szybkości działania zwykłego filtra bilateralnego z filtrem przedziałami liniowym. Linia zielona reprezentuje wersję przedziałami liniową ($NB_SEGMENTS = 15$) a czerwona wersję zwykłą. Pomiary przeprowadzone zostały na obrazie lamp.hdr o rozdzielczości 400x300.

poza interpolacją. Interpolacja musi być wykonana w pełnej rozdzielczości, ponieważ wykonanie jej w zmniejszonej rozdzielczości powodowałoby występowanie widocznych artefaktów. Na rys. 2.28 znajduje się pseudokod metody z obydwoma usprawnieniami. Autorzy zmniejszali rozdzielczość metodą "najbliższego sąsiada". Uzyskane przyspieszenie przedstawia rys. 2.29.

2.7.3. Redukcja kontrastu

Metoda kompresji tonów Duranda i Dorsey opiera się na dwuwarstwowym schemacie dekompozycyjnym (rys. 2.23), czyli na rozkładzie obrazu na dwie warstwy: warstwę bazową oraz warstwę szczegółów. Warstwę bazową otrzymujemy aplikując do obrazu szybki filtr bilateralny, warstwa szczegółu powstaje w wyniku podzielenia luminancji sceny przez warstwę bazową. Następnie kompresujemy kontrast warstwy bazowej

```

SzybkiFiltrBilateralny
    (obraz L, filtr  $f_{\sigma_s}$ , filtr  $g_{\sigma_r}$ , wsp. zmiany rozmiaru z)
    J = 0    // zerujemy wszystkie piksele J
    L' = ZmniejszRozdzielczość(L, z)
     $f'_{\sigma_s/z}$  = ZmniejszRozdzielczość( $f_{\sigma_s}$ , z)
    for j = 0..NB_SEGMENTS
         $l^j$  = min(L) + j × (max(L) – min(L))/NB_SEGMENTS
         $G^j$  =  $g_{\sigma_r}(L' - l^j)$     // obliczamy  $g_{\sigma_r}$  w każdym pikselu
         $H^j$  =  $G^j \times L'$     // obliczamy H w każdym pikselu
         $K^j$  =  $G^j \otimes f_{\sigma_s/z}$     // obliczamy czynnik normalizujący
         $H^j$  =  $H^j \otimes f_{\sigma_s/z}$ 
         $J^j$  =  $H^j / K^j$ 
         $J^j$  = ZwiększRozdzielczość( $J^j$ , z)
    J = J +  $J^j \times$  InterpolowanaWaga(L,  $l^j$ )
    
```

Rysunek 2.28: Szybki filtr bilateralny, czyli filtr z obydwoma usprawnieniami.

mnożąc ją przez współczynnik k , równy:

$$k = \text{kontrastBazowy} / (L_{base\ max} - L_{base\ min}) \quad (2.65)$$

gdzie $L_{base\ max}$, to maksymalna luminancja w warstwie bazowej, $L_{base\ min}$ to minimalna luminancja warstwy bazowej a kontrastBazowy to kontrolowany przez użytkownika parametr, domyślnie równy 5.

Kolor traktowany jest podobnie jak w [21], tzn. do redukcji kontrastu wykorzystywana jest jedynie luminancja, kolor jest rekonstruowany po redukcji.

Wszystkie obliczenia wykonywane są na logarytmach luminancji pikseli, powodem tego jest fakt, że różnice logarytmów luminancji odpowiadają bezpośrednio kontrastowi, poza tym umożliwia to jednorodne traktowanie całego zakresu luminancji.

Autorzy używali jako wartości domyślnych σ_s równego 2% rozmiaru obrazu i $\sigma_r = 0.4$

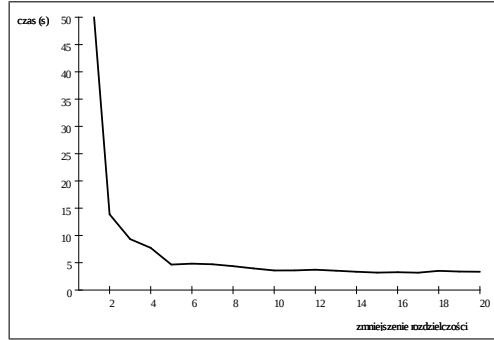
2.8. Metoda M. Ashikhmina z 2002 roku

Metoda Ashikhmina przedstawiona w [2] symuluje dwie ważne funkcje ludzkiego układu wzrocznego (*ang.* *HVS – Human Visual System*): sygnalizowanie absolutnej jasności oraz lokalnego kontrastu.

Lokalny kontrast c w punkcie (x, y) to:

$$c(x, y) = \frac{L(x, y)}{L_a(x, y)} - 1 \quad (2.66)$$

gdzie L_a , czyli luminancja adaptacyjna, to średnia luminancja pewnego otoczenia piksela (x, y) .



Rysunek 2.29: Przyspieszenie uzyskane dzięki zmniejszaniu rozdzielczości. Oś pozioma reprezentuje ile razy zmniejszona została rozdzielczość obrazu. Obliczenia przeprowadzono na obrazie bigfogmap.hdr o rozdzielczości 751x1130.

Jedną z dwu najważniejszych cech przedstawianej metody jest zachowanie lokalnego kontrastu świata w obrazie wyjściowym, czyli zachowanie

$$c_d(x, y) = c_w(x, y) \quad (2.67)$$

lub inaczej

$$\frac{L_d(x, y)}{L_{da}(x, y)} = \frac{L_w(x, y)}{L_{wa}(x, y)} \quad (2.68)$$

Najprostszą metodą zachowania kontrastu jest mapowanie liniowe. Niestety, jedno mapowanie liniowe użyte do całego obrazu nie jest dobrym rozwiązaniem. Może ono dać dobre rezultaty w przypadku obrazu o małym kontraście dynamicznym, jednak mapowanie obrazu wysoko kontrastowego spowoduje obcięcie najmniejszych bądź największych (lub obu) wartości luminancji. Skutkiem tego będzie utracenie całej wiadomości o kontraście w najciemniejszych i najjaśniejszych rejonach obrazu. Dodatkowo, jeśli użyty współczynnik mapujący nie będzie zależeć od luminancji obrazu, wtedy utracone zostanie ogólne wrażenie jasności sceny.

Zamiast zwykłego mapowania liniowego, w metodzie Ashikhmina stosuje się mapowanie lokalnie liniowe, ze współczynnikami zależnym od otoczenia konkretnego piksela, tak aby zachować detale w obrębie całego obrazu. Mając obliczone L_{wa} dla całego obrazu wejściowego, oraz stosując do L_{wa} funkcję mapującą TM , otrzymujemy obraz $TM(L_a)$. Funkcja TM kompresuje luminancję świata do zakresu monitora, jednocześnie zachowując wrażenie perceptualnej jasności obrazu.

Korzystając z obrazów L_{wa} i $TM(L_{wa})$ oraz z równania 2.68 uzyskujemy ostateczne mapowanie:

$$L_d(x, y) = \frac{L_w(x, y)TM(L_{wa}(x, y))}{L_{wa}(x, y)} \quad (2.69)$$

gdzie $TM(L_{wa}(x, y))$ odpowiada $L_{da}(x, y)$ z równania 2.68.

Mapowanie określone równaniem 2.69 jest lokalnie liniowe, a co za tym idzie zachowuje lokalny kontrast, tylko wtedy, gdy L_{wa} (i w konsekwencji $TM(L_{wa})$) jest wolno zmieniającą się funkcją na większości powierzchni obrazu. Żeby tak było L_{wa} nie może być wyznaczone jako średnia luminancja otoczenia stałego rozmiaru. Stały rozmiar

otoczenia powodowałby pojawianie się zjawiska *halo* (odwrotnych gradientów) w obszarach o wysokim kontraście rys. 2.21. Rozmiar otoczenia powinien być duży w rejonach o niskim kontraście oraz zmniejszać się wraz ze wzrostem kontrastu, aż do wielkości jednego piksela w rejonie bardzo wysokiego kontrastu. Metoda wyznaczania L_{wa} zostanie przedstawiona poniżej.

2.8.1. Wyznaczanie poziomu lokalnej adaptacji

W celu wyznaczenia poziomu lokalnej adaptacji stosowane jest podejście balansujące dwa przeciwstawne wymogi stawiane przed *HVS*. Są to zachowanie lokalnego kontrastu w rozsądnych granicach i jednocześnie utrzymanie wystarczającej informacji o detalach obrazu. Na tym etapie ważny jest jedynie pomiar kontrastu. Zakładamy, że przekłamania mogące wynikać z powyższego będą redukowane przez opisaną wcześniej funkcję *TM*.

Aby wyznaczyć lokalną adaptację w pikselu należy uśrednić wartości luminancji w otoczeniu piksela. Rozmiar otoczenia powinien być możliwie największy, jednocześnie piksele tego obszaru muszą mieć podobne luminancje, tak aby kontrast w obszarze był mały. Ważne jest również aby poziom lokalnej adaptacji zmieniał się płynnie w obszarach o niedużym kontraście.

Powyższe wymagania spełnia często wykorzystywana strategia opierająca się na porównywaniu ze sobą pikseli z dwóch obrazów wygładzonych filtrami (najczęściej gausowskimi) o różnych wielkościach jądra ([17], [20]). Dokładniej, obliczymy lokalny kontrast lc (*ang. local contrast*) w pikselu jako:

$$lc(s, x, y) = \frac{G_s(x, y) - G_{2s}(x, y)}{G_s(x, y)} \quad (2.70)$$

gdzie

$$G_s(x, y) = L(x, y) \otimes R_s(x, y) \quad (2.71)$$

to piksel z obrazu wygładzonego filtrem gausowskim o rozmiarze s , czyli

$$R_s(x, y) = \frac{1}{\pi s^2} \exp\left(-\frac{x^2 + y^2}{s^2}\right) \quad (2.72)$$

Wybór G_{2s} zamiast G_s w mianowniku równania 2.70 jest arbitralny. Parametr s jest równy rozmiarowi otoczenia, z którego oblicza się lokalny kontrast w pikselu.

Poniżej opisany został proces wyznaczania optymalnego s dla równania 2.70. Zaczynamy od najmniejszego otoczenia, czyli $s = 1$ piksel i obliczamy dla niego lc . Jeśli wartość lc jest zbyt duża, należy uznać, że luminancja adaptacyjna jest równa luminancji w punkcie. Sytuacja taka ma miejsce gdy dany piksel stanowi wysokokontrastowy detal w obrazie. Jeśli lc nie jest zbyt duże, powtarzamy postępowanie ze zwiększaniem za każdym razem s o 1 dopóki obliczony kontrast nie będzie większy od określonej wartości brzegowej lub dojdziemy do największego dopuszczalnego rozmiaru otoczenia. Rozmiar ten, co wynika z pomiarów fizjologicznych, powinien wynosić około jeden stopień pola widzenia. Jednak z powodu braku informacji o rozmiarze pola widzenia obrazów, jako maksymalny rozmiar otoczenia używana jest stała wartość, niezależna od obrazu wejściowego. Odpowiednia dla większości okazała się $s_{\max} = 10$.

Kryterium ograniczające kontrast ma postać:

$$|lc(s, x, y)| < ac \quad (2.73)$$

gdzie ac (*ang. allowed contrast*) to parametr określający dopuszczalny kontrast obszaru, zwykle wynosi on 0.5 ale możliwe są też inne wartości. rys. 2.30 przedstawia rezultaty użycia innych wartości ac . Opisany proces wyznacza największe



Rysunek 2.30: Wyniki mapowania z różnymi wartościami parametru ac . Po lewej $ac = 0$, w środku $ac = 0.5$, z prawej $ac = 0.1$. Obrazki z pracy [2].

$s \in \{1, 2, \dots, s_{\max}\}$ spełniające warunek 2.73, po czym jako $L_{wa}(x, y)$ przyjmuje się $G'_s(x, y)$. $G'_s(x, y)$ to wartość interpolowana między $G_s(x, y)$ i $G_{s+1}(x, y)$, odpowiadająca $|lc(s)| = ac$.

2.8.2. Funkcja mapująca TM

Funkcja TM ma za zadanie mapowanie luminancji sceny L_w do luminancji wyjściowej L_d . Nie bierze ona pod uwagę widoczności detali, czyli lokalnego kontrastu. Dzięki temu funkcja zależy tylko od luminancji w pikselu, nie ma znaczenia położenie piksela w obrazie.

Do skonstruowania funkcji TM potrzebne będzie pojęcie "postrzeganej pojemności" (*ang. perceptual capacity*) przedziału luminancji. Postrzegana pojemność danego, niewielkiego przedziału luminancji jest długością tego przedziału przeskalowaną przez wartość funkcję TVI (ledwo dostrzegalna różnica opisana w rozdziale pierwszym). Pojęcie to umożliwia porównywanie postrzeganej ważności różnych przedziałów luminancji.

Parametrem funkcji TVI powinna być luminancja adaptacyjna L_{wa} , ponieważ wartość dostrzegalnej różnicy zależy od otoczenia. Niestety użycie jako parametru TVI , wartości L_{wa} obliczonej metodą przedstawioną poniżej łamałoby założenie o globalności TM . W związku z tym autor wybrał jako aproksymację $L_{wa}(x, y)$ wartość $L_w(x, y)$. Alternatywna aproksymacja mogła by polegać na wybraniu jednej wartości L_a dla całego obrazu, dałaby ona dobre rezultaty dla obrazów o niskim kontraście dynamicznym. Dla obrazów o wysokim kontraście, w których luminancja adaptacyjna mocno się waha, przybliżenie za pomocą $L_w(x, y)$ jest lepszym rozwiązaniem.

Definiujemy pomocniczą funkcję pojemności C jako:

$$C(L) = \int_0^L \frac{dl}{TVI(l)} \quad (2.74)$$

Określa ona postrzeganą pojemność przedziału $[0, L]$. Przy takiej definicji pojemność przedziału $[L_1, L_2]$ to $C(L_2) - C(L_1)$.

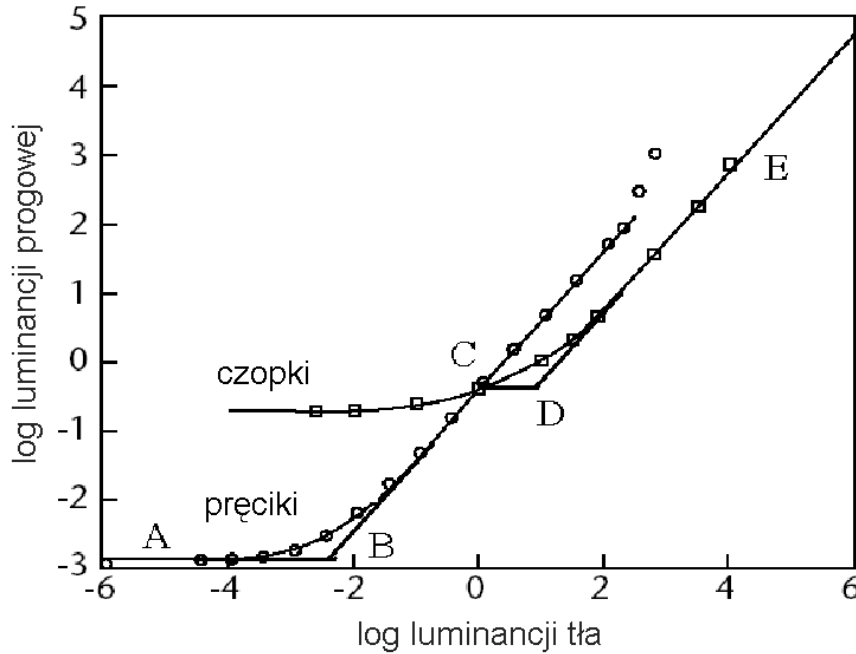
Tak jak Ward et al. w pracy [25], Ashikhmin używa aproksymacji funkcji TVI łamaną. rys. 2.31 przedstawia funkcję TVI oraz jej aproksymującą. Korzystając z powyższego funkcja $C(L)$ przybiera następującą postać:

$$C(L) = \begin{cases} L/0.0014 & L < 0.0034 \\ 2.4483 + \log(L/0.0034)/0.4027 & 0.0034 \leq L < 1 \\ 16.5630 + (L - 1)/0.4027 & 1 \leq L < 7.2444 \\ 32.0693 + \log(L/7.2444)/0.0556 & L \geq 7.2444 \end{cases} \quad (2.75)$$

Definiujemy również postrzeganą pojemność dla obrazu wyjściowego jako:

$$C_d(L) = \frac{L}{TVI(L_{ad})} \quad (2.76)$$

gdzie L_{ad} to pojedyncza wartość aproksymująca luminancję adaptacyjną całego obrazu.



Rysunek 2.31: Funkcja TVI oraz jej aproksymacja łamaną $ABCD$, luminancja brzegowa i tła w cd/m^2 . Rysunek z pracy [2].

Główną zasadą działania funkcji TM jest zachowanie proporcji między postrzeganą pojemnością przedziałów luminancji wyjściowej i luminancji świata. Zasadę tę przedstawić można jako:

$$\frac{C_d(L_d) - C_d(L_{d\min})}{C_d(L_{d\max}) - C_d(L_{d\min})} = \frac{C(L_w) - C(L_{w\min})}{C(L_{w\max}) - C(L_{w\min})} \quad (2.77)$$

W celu wykorzystania całego przedziału luminancji monitora, należałoby użyć $L_{d\min} = 0$ oraz $L_{d\max} = LDMAX$, czyli maksymalnej luminancji monitora. Jako $L_{w\min}$ i $L_{w\max}$ dobrze jest wziąć najmniejszą i największą luminancję obrazu wygładzonego. Uniezależni to działanie operatora od występowania w obrazie nielicznych, bardzo jasnych lub bardzo ciemnych pikseli.

Równanie 2.75 jednoznacznie określa L_d . Korzystając z równania 2.76 oraz opisanych wyżej ustawień $L_{d\min}$ i $L_{d\max}$, L_d przybiera postać:

$$L_d = TM(L_w) = LDMAX \frac{C(L) - C(L_{w\min})}{C(L_{w\max}) - C(L_{w\min})} \quad (2.78)$$

2.8.3. Pełna metoda

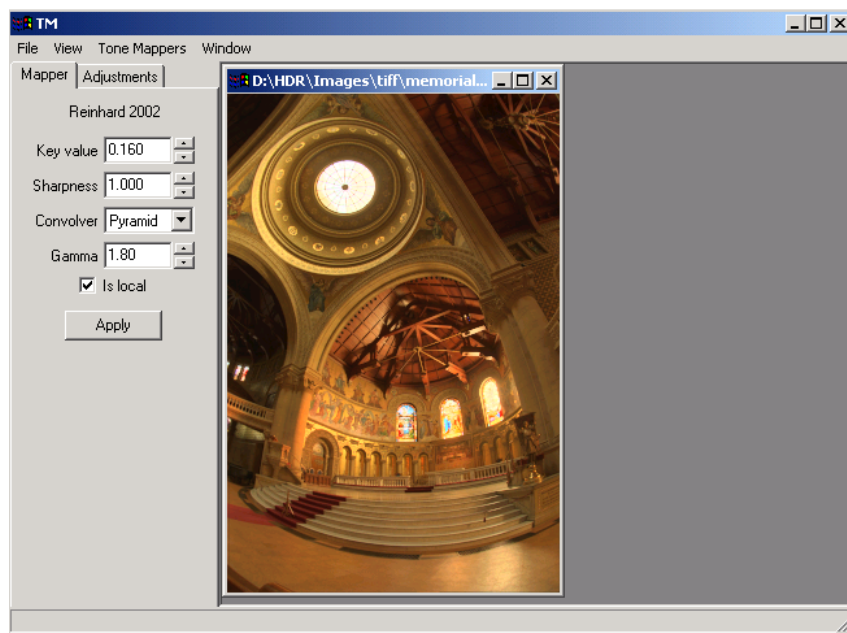
Dysponując obrazem luminancji adaptacyjnej L_a , należy obliczyć obraz $TM(L_a)$ a następnie korzystając z równania 2.69 obliczyć ostateczny odwzorowany obraz L_d . Rezultat to czarno-biały obraz luminancji. Kolory traktowane są jak w metodzie Duranda i Dorsey [10] czy Schlicka [21], tzn. każdy z kanałów RGB oryginalnego obrazu mnożony jest przez $L_d(x, y)/L_w(x, y)$ w każdym pikselu (x, y) .

Nie wykonuje się żadnych dodatkowych skalowań rezultatu. Jeśli wartości obrazu wyjściowego nie mieszczą się w zakresie są przycinane do zakresu. Sytuacja taka, jeśli wogóle zaistnieje, zwykle zdarza się w bardzo jasnym obszarze obrazu o bardzo dużym kontraście lokalnym. W sytuacji tego typu obcięcie luminancji nie powoduje pogorszenia jakości odwzorowanego obrazu.

Rozdział 3

Opis programu

W celu porównania działania metod mapowania tonów, napisany został specjalny program komputerowy o nazwie **TM** (rys. 3.32). Program ten zawiera moją implementację metod z rozdziału drugiego. Umożliwia zastosowanie ich do obrazów HDR oraz oglądanie otrzymanych nisko kontrastowych rezultatów mapowania. Możliwa jest także praca nad kilkoma obrazami jednocześnie oraz łatwa zmiana parametrów używanych metod. Zawiera również proste narzędzia do korekcji obrazów nisko kontrastowych.



Rysunek 3.32: Program **TM**. Z lewej strony okna programu znajduje się panel operatora mapowania tonów. Z prawej strony znajduje się okno zawierające obraz HDR poddany działaniu operatora.

Program został napisany przy użyciu środowiska programistycznego Microsoft Visual C++ 6.0, korzysta z bibliotek wxWindows (wersja 2.4.2) oraz fftw (wersja 3.0.1).

3.1. Sesja pracy z programem

Poniżej przedstawiam przykładową sesję pracy z programem.

Zaczynamy od otworzenia pliku HDR, wybierając z menu **File** opcję **Open**, następnie znajdujemy odpowiedni plik HDR, oraz wciskamy przycisk **Open**. Z prawej strony okna programu pojawia się okno zawierające wybrany plik HDR do którego zastosowany został operator liniowy z korekcją gamma. Każdy nowo otwarty plik HDR przed pokazaniem go jest mapowany przy użyciu operatora liniowego z korekcją gamma. Z lewej strony okna pojawia się panel operatora. Każdy panel operatora zawiera wartości wszystkich możliwych do zmiany parametrów operatora, oraz przycisk **Apply**, po wciśnięciu którego zaczyna się mapowanie obrazu, z bieżącymi wartościami parametrów. Program nie ogranicza liczby otwartych okien z obrazkami. Jeden obrazek może być również otwarty wielokrotnie.

Po otworzeniu obrazka wybieramy z menu **Tone Mappers**, jeden z operatorów. Zmieniamy w dowolny sposób parametry, po czym wciskamy przycisk **Apply**. Zaawansowanie procesu mapowania obrazuje pojawiające się okno z paskiem postępu. Obrazek w aktualnie aktywnym oknie zostanie zamieniony przez rezultat przeprowadzonego właśnie mapowania.

Obrazek nisko kontrastowy, otrzymany w wyniku mapowania, można dodatkowo modyfikować używając standardowych metod obróbki obrazu. Można w ten sposób zmieniać jego jasność, nasycenie lub kontrast. Aby tak zrobić należy na panelu, z lewej strony okna programu, wybrać zakładkę **Adjustments**, a następnie wcisnąć przycisk **Color**, dla zmiany jasności i nasycenia kolorów, lub **Brightness/Contrast** aby zmodyfikować jasność i kontrast. Wszystkie modyfikacje dokonane z zakładki **Adjustments** można anulować wciskając przycisk **Reset**.

Rezultaty mapowania można zapisać na dysku jako zwykły, nisko kontrastowy, plik graficzny wybierając z menu **File** opcję **Save As** i wpisując nazwę pliku wyjściowego.

3.2. Dokładny opis opcji programu

Menu **File**, zawiera podstawowe opcje manipulacji obrazami:

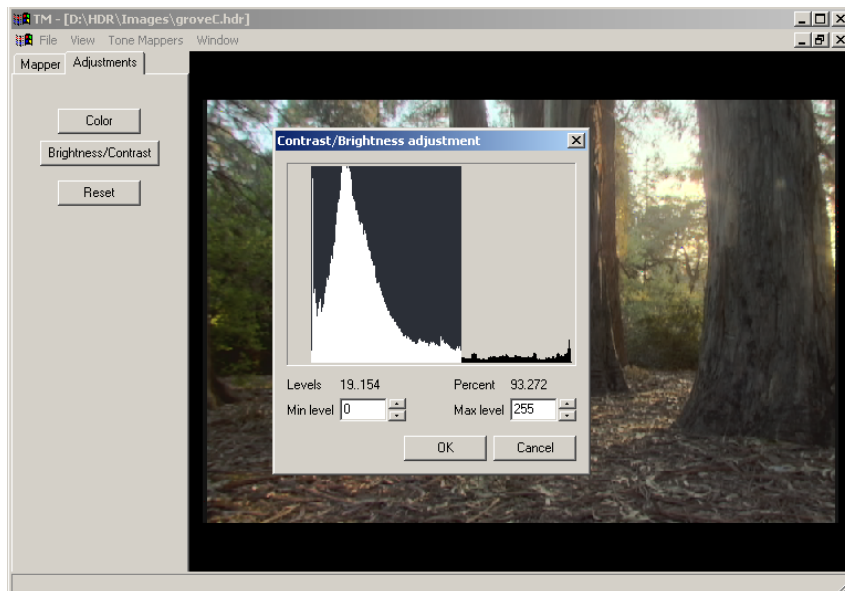
- **Open** (klawisz Ctrl+O) otwiera obraz HDR znajdujący się w określonym przez użytkownika miejscu na dysku. Program **TM** obsługuje pliki formatu: **hdr**, **pic**, czyli pliki generowane przez program **Radiance** [19], pliki **exr** opracowane przez **Industrial Light & Magic** [16], oraz zmodyfikowaną wersję plików **tiff** zawierającą obrazy HDR [12].
- **Save As** (klawisz Ctrl+S) zapisuje w podanym przez użytkownika miejscu na dysku nisko kontrastowy obraz. Możliwe jest zapisanie obrazu w trzech formatach: **ppm**, **png** i **jpg**,
- **About** – otwiera okno zawierające krótką informację o programie,
- **Exit** (klawisz Ctrl+X) - wyjście z programu.

Menu **View**, zawiera opcje ułatwiające obróbkę i oglądanie obrazów. Są to:

- **Zoom In** (klawisz =) – zwiększa aktywny obraz o 10%, nie zmieniając rozmiaru okna w którym prezentowany jest obraz. Aby przesunąć aktualnie prezentowany

obszar obrazu należy przeciągnąć myszą obraz, lub skorzystać z pasków przewijania.

- **Zoom Out** (klawisz **-**) – zmniejsza aktywny obraz o 10%.
- **Zoom Actual Size** (klawisz **Ctrl+A**) – ustawia oryginalny rozmiar obrazu.
- **Fit Window** (klawisz **Ctrl+F**) – zmienia rozmiar obrazu tak, aby mieścił się on w oknie oraz zajmował możliwie dużą jego powierzchnię.
- **Image Info** (klawisz **Ctrl+I**) – pokazuje okno zawierające informacje o aktywnym obrazie. Znajdują się tam kolejno: ścieżka pod którą znajduje się plik z obrazem HDR (**Path**), szerokość obrazu (**Width**), wysokość obrazu (**Height**), aktualne powiększenie (**Current Zoom**), nazwa operatora użytego do mapowania obrazu oraz wszystkie parametry operatora (**Mapped with**), parametry modyfikacji koloru (**Color**), parametry modyfikacji jasności i kontrastu (**Brightness/Contrast**), oraz czas ostatniego mapowania (**Mapping Time**).



Rysunek 3.33: Oknem **Brightness/Contrast**. W górnej części histogram jasności obrazu, z zaznaczonym przedziałem jasności.

Menu **Tone Mappers** służy do wybierania operatora mapowania tonów. Wybranie operatora powoduje wyświetlenie odpowiedniego panelu z lewej strony okna aplikacji, w zakładce **Mapper**. Panel ten pozwala na zmianę parametrów operatora.

Początkowe wartości parametrów w panelach są w większości przypadków wartościami domyślnymi, których używali autorzy prac. W niektórych przypadkach są to wartości neutralne, które powinny dawać dobre rezultaty dla obrazów o niedużym kontraście i średniej jasności. Do parametrów tych zalicza się parametr a z mapowania liniowego z korekcją gamma, parametr a określający wartość kluczową z mapowania Reinharda.

Każdy z paneli umożliwia zmianę parametru **Gamma** (γ) odpowiedzialnego za korekcję gamma opisaną w rozdziale 1.6. Poniżej przedstawione są dostępne operatory wraz z opisem ich parametrów.

- **Gamma-corrected Linear Mapping** – realizuje mapowanie liniowe z korekcją gamma przedstawione w rozdziale 2.3 (równanie 2.1). Panel tego operatora umożliwia zmianę parametru **Exposure**, odpowiadającego parametrowi a z opisu metody.
- **Schlick 1984** – to operator przedstawiony w rozdziale 2.3 (równanie 2.5). W jego panelu znajdują się następujące elementy:
 1. przycisk **Choose DNBG**, po wciśnięciu którego pojawia się okno wyznaczania najciemniejszego, odróżnialnego od czerni, odcienia szarości (rys. 2.12),
 2. pole **DNBG**, określające najciemniejszy, odróżnialny od czerni, odcień szarości
 3. pole **k** z równania 2.8,
 4. pole kontrolne **Non Uniform**, określające czy użyć modyfikacji operatora ze wzoru 2.8.
- **Ward 1997** – to operator przedstawiony w rozdziale 2.4. W jego panelu znajdują się następujące elementy:
 1. pole **Num Bins**, określające gęstość histogramu (oznaczana w rozdziale 2.4 jako N),
 2. pole **Sight Angle**, określające kąt widzenia dłuższego z wymiarów, pionowego lub poziomego,
 3. pole wyboru **Adjustment**, określające rodzaj użytego ograniczenia histogramu, wybór **No** powoduje nie stosowanie żadnego ograniczenia, wybór **Linear** powoduje użycie ograniczenia liniowego (wzór 2.17), a **Jnd** użycie ograniczenia uwzględniającego ludzką czułość widzenia (wzór 2.20),
 4. pole kontrolne **Color Sensitivity**, określające czy symulować wrażliwość na kolory,
 5. pole kontrolne **Glare**, określające czy symulować odbłaski,
 6. pole kontrolne **Acuity**, określające czy symulować ostrość widzenia.
- **Tumblin 1999** – to operator przedstawiony w rozdziale 2.5. W jego panelu znajdują się następujące elementy:
 1. pole **Num Steps** określa liczbę kroków wykonywanych podczas obliczania uproszczonych obrazów LCIS
 2. pole **Step Length** określa długość jednego kroku obliczania obrazów LCIS,
 3. pole **w color** zawiera wartość współczynnika kompresji koloru,
 4. pole **LCIS Images** określa liczbę obrazów LCIS użytych przez operator, poniżej pola **LCIS Images** znajdują się, dla każdego obrazu LCIS, dwa pola określające parametry k (przewodnictwo progowe), oraz współczynnik kompresji tego obrazu.

- **Reinhard 2002** – to operator przedstawiony w rozdziale 2.6. W jego panelu znajdują się następujące elementy:
 1. pole **Key Value** określa wartość kluczową obrazu (oznaczane w rozdziale 2.6 jako a),
 2. pole **Sharpness** określa wartość parametru ϕ ,
 3. pole wyboru **Convolver** określa sposób obliczania splotów; możliwe wybory to:
 - **Simple** – obliczanie splotów bez optymalizacji,
 - **FFT** – obliczanie splotów za pomocą szybkiej transformaty Fouriera,
 - **Pyramid** – obliczanie splotów za pomocą piramidy obrazów,
 4. zaznaczone pole kontrolne **Is Local** oznacza użycie wersji lokalnej operatora (równanie 2.58), odznaczenie pola oznacza użycie wersji globalnej (równanie 2.53).
- **Ashikhmin 2002** – to operator przedstawiony w rozdziale 2.8. W jego panelu znajdują się następujące elementy:
 1. pole **Threshold** określa dopuszczalny kontrast z równania 2.78,
 2. pole **Convolver**, tak jak w metodzie Reinharda, określa sposób obliczania splotów,
 3. działanie pola kontrolnego **Is Local** jest takie same jak w panelu metody Reinharda, odznaczenie pola powoduje mapowanie ze wzoru 2.73, bez obliczania lokalnego kontrastu.
- **Durand 2002** – to operator przedstawiony w rozdziale 2.7. W jego panelu znajdują się następujące elementy:
 1. pole **Contrast** określa wartość kontrastu (wzór 2.65),
 2. pole **Sigma S** to parametr filtra bilateralnego (wzory 2.59–2.62),
 3. pole **Sigma R** to drugi parametr filtra bilateralnego (wzory 2.59–2.62),
 4. pole wyboru **Method** określa implementację filtra bilateralnego, wartość **Simple** oznacza implementację prosto ze wzorów 2.59–2.62, **Piecewise** odpowiada implementacji przedziałami liniowego filtra bilateralnego (rys. 2.26), a **Fast** to implementacja szybka (rys. 2.28),
 5. pole **Segments** określa liczbę segmentów używanych w szybkiej lub przedziałami liniowej implementacji filtra,
 6. pole **Downsample** określa ile razy zmniejszyć rozdzielczość oryginalnego obrazu przy szybkiej implementacji filtra bilateralnego.

Menu **Window** jest standardowe dla systemu operacyjnego Windows. Ułatwia korzystanie z kilku okien jednocześnie. Poza spisem wszystkich otwartych obrazów zawiera ono następujące opcje:

- **Cascade** – rozmieszcza wszystkie otwarte obrazy kaskadowo,
- **Tile Horizontally** – rozmieszcza wszystkie otwarte obrazy jeden pod drugim, rozciągając je na całą szerokość roboczej części okna aplikacji
- **Tile Vertically** – rozmieszcza wszystkie otwarte obrazy jeden przy drugim, rozciągając je na całą wysokość roboczej części okna aplikacji,
- **Arrange Icons** – porządkuje rozmieszczenie zminimalizowanych obrazów,
- **Next** – ustawia następne ze spisu otwartych obrazów jako aktywne,
- **Previous** – ustawia jako aktywne poprzednie z obrazów ze spisu.

Zakładka **Adjustments** zawiera proste narzędzia do manipulowania obrazem nisko kontrastowym. Znajdują się na niej trzy przyciski **Color**, **Brightness/Contrast** oraz **Reset**.

- Przycisk **Color** otwiera okno, w którym znajdują się dwa suwaki: **Saturation**, zmieniający nasycenie obrazu, i **Lightness**, zmieniający jego jasność. Obraz modyfikowany jest zgodnie z wartościami ustawionymi na suwakach po wciśnięciu przycisku **OK**.
- Przycisk **Brightness/Contrast** otwiera okno (rys. 3.33), w którego górnej części znajduje się histogram luminancji aktualnego obrazu. Po zaznaczeniu części histogramu myszą pod histogramem można odczytać jaki dokładnie obszar jasności został zaznaczony oraz ile procent stanowią piksele o jasności z zaznaczonego przedziału. Poniżej informacji o zaznaczeniu znajdują się pola **Min Level** i **Max Level**. Określają one wielkość przedziału, do którego zostaną przetransformowane wszystkie jasności obrazu. Odpowiednio ustawiając wartości **Min Level** i **Max Level** można zwiększyć kontrast obrazu, jednocześnie zwiększając bądź zmniejszając jego jasność.
- Przycisk **Reset** anuluje zmiany spowodowane przez modyfikacje **Color** i **Brightness/Contrast**.

Kliknięcie prawym przyciskiem myszy na obrazie otwiera menu kontekstowe. Znajdują się w nim opcje z menu **View** oraz opcja **Save As** z menu **File**.

Kliknięcie lewym przyciskiem na obrazie wypisuje w pasku stanu głównego okna współrzędne klikniętego piksela a następnie składowe RGB oraz luminancję z oryginalnego obrazu HDR, oraz składowe RGB i luminancję z obrazu nisko kontrastowego.

Rozdział 4

Porównanie działania metod

W tym rozdziale porównamy metody mapowania tonów przedstawione wcześniej. Analizie poddamy szybkość ich działania, łatwość korzystania z metod oraz jakość obrazów otrzymanych w wyniku mapowania. Najpierw jednak przedstawimy domyślne parametry metod co jest niezbędne do przeprowadzenia porównań.

4.1. Parametry domyślne metod

Parametry domyślne metod są następujące:

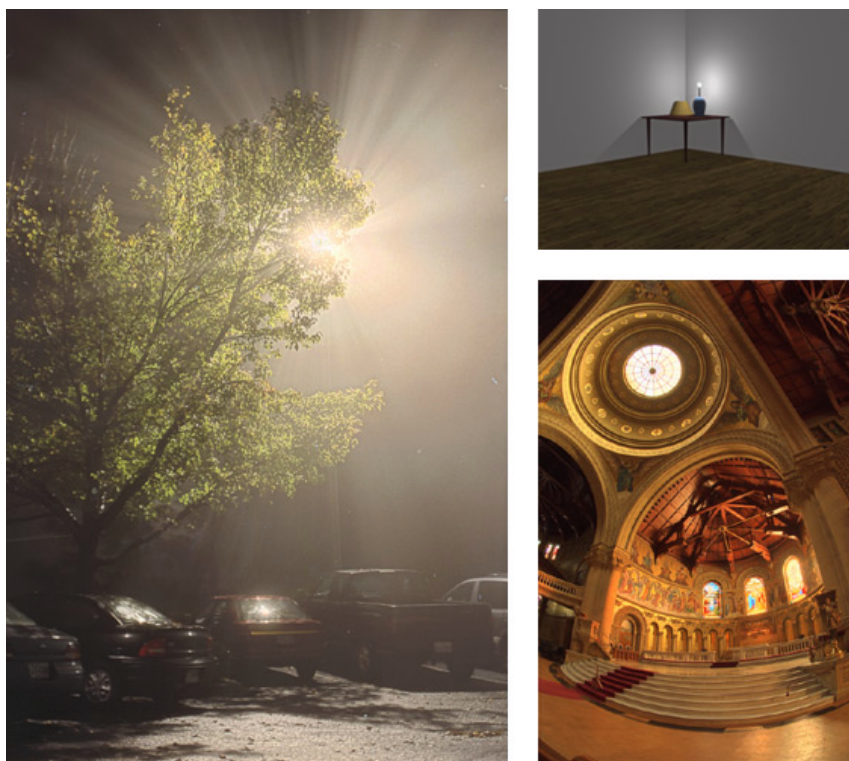
- mapowanie liniowe z korekcją gamma: Exposure = 1.0,
- metoda Schlicka z 1994 roku: DNBG = 1, $k = 0.5$, opcja Non Uniform włączona,
- metoda Warda et al. z 1997 roku: Num Bins = 100, Sight Angle = 70, Adjustment = jnd, opcje Color Sensitivity, Glare i Acuity wyłączone,
- metoda Tumblina i Turka z 1999 roku: Num Steps = 100, Step Length = 0.031, w Color = 0.6, LCIS Images = 4, $w_0 = 1.0$, $w_1 = 0.7$, $w_2 = 0.5$, $w_3 = 0.3$, $k_0 = 0.0$, $k_1 = 0.06$, $k_2 = 0.1$, $k_3 = 0.16$,
- metoda Reinharda et al. z 2002 roku: Key Value = 0.18, Sharpness = 1.0, Convolver = Pyramid, opcja Is Local włączona,
- metoda Ashikhmina z 2002 roku: Threshold = 0.5, Convolver = Pyramid, opcja Is Local włączona,
- metoda Duranda i Dorsey z 2002 roku: Contrast = 5.0, Sigma S = 10.0, Sigma R = 0.4, Method = Fast, Segments = 12, Downsample = 9.0.

Wszystkie metody używają domyślnie parametru Gamma = 2.2.

4.2. Szybkości działania metod

Testy szybkości polegały na użyciu wszystkich metod do mapowania trzech wybranych obrazów o coraz większych rozmiarach. Wszystkie testy zostały przeprowadzone za pomocą programu TM opisanego w rozdziale trzecim, na komputerze z procesorem AMD Athlon 900 MHz i 256 MB pamięci RAM działającym pod systemem operacyjnym Microsoft Windows 2000. Zamieszczone poniżej czasy obejmują, poza długością obliczeń wykonywanych przez metody, również dodatkowy czas obejmujący operacje wymagane przez program TM, m. in. uaktualnianie paska postępu. Dodatkowy czas nie utrudnia porównywania metod ponieważ jest zwykle dużo mniejszy od czasu potrzebnego na

wykonanie mapowania, poza tym jest on bardzo podobny we wszystkich metodach.



Rysunek 4.34: Obrazy wykorzystane do testowania szybkości operatorów. Z lewej strony znajduje się obraz bigfogmap.hdr. Z prawej strony u góry obraz lamp.hdr, a z prawej strony u dołu memorial.hdr. Zamieszczone obrazy są rezultatem mapowania metodą Reinharda z 2002 roku.

Poniżej znajdują się czasy mapowania wszystkimi metodami trzech obrazów.

Operator	Czasy obliczeń		
Mapowanie liniowe z korekcją gamma	0.241 sek.	0.717 sek.	1.415 sek.
Schlick 1984	0.290 sek.	0.915 sek.	1.966 sek.
Ward 1997	0.448 sek.	1.336 sek.	2.734 sek.
Tumblin 1999	12.154 sek.	44.429 sek.	98.887 sek.
Reinhard 2002	0.571 sek.	1.399 sek.	3.338 sek.
Ashikhmin 2002	0.912 sek.	2.231 sek.	5.688 sek.
Durand 2002	0.487 sek.	1.662 sek.	3.835 sek.

W kolumnach tabeli znajdują się wyniki mapowania kolejnych obrazów: lamp.hdr (obraz z prawego górnego rogu rys. 4.34) po lewej, memorial.hdr (prawy dolny róg na rys. 4.34) w środkowej kolumnie wyników, bigfogmap.hdr (lewa strona rys. 4.34) z prawej strony.

Rozdzielczości obrazów testowych wynoszą kolejno: 751x1130 dla bigfogmap.hdr, 400x300 dla lamp.hdr, 512x768 dla memorial.hdr.

Wyniki w tabeli są czasami mapowania przy standardowych ustawieniach para-

metrów. Dla niektórych metod wartości części parametrów mogą istotnie wpłynąć na czas mapowania. Poniżej znajdują się tabele z wynikami otrzymanymi z różnych wartości parametrów dla metod, których czasy wykonania zmieniają się istotnie w zależności od wartości parametrów. Wszystkie tabele zawierają, tak jak tabela wyżej, w kolumnie drugiej czasy mapowania obrazu lamp.hdr, w kolumnie trzeciej obrazu memorial.hdr i w czwartej obrazu bigfogmap.hdr.

Metoda Warda et al. z 1997 roku.

Parametry	Czasy obliczeń		
tylko poprawianie histogramu	0.448 sek.	1.336 sek.	2.734 sek.
poprawianie + wrażliwość na kolory	0.528 sek.	1.605 sek.	3.202 sek.
poprawianie + odbłaski	1.078 sek.	1.909 sek.	3.412 sek.
poprawianie + ostrość widzenia	0.788 sek.	2.227 sek.	4.476 sek.
poprawianie + wszystkie efekty	1.469 sek.	3.027 sek.	5.632 sek.

Czasy wykonywania nie zmieniają się istotnie w zależności od rodzaju ograniczenia histogramu. Wyniki z tabeli zostały uzyskane stosując ograniczenie uwzględniające ludzką czułość widzenia. Mały wpływ na czas obliczeń ma również gęstość histogramu.

Jak widać w metodzie Warda najbardziej czasochłonne jest obliczanie odbłasków oraz ostrości widzenia. Istotny wpływ na czas obliczania odbłasków ma parametr Sight Angle, określający kąt widzenia dłuższego z wymiarów obrazu. Jego zwiększenie ze standardowej wartości – 70 do wartości 140 wydłuża czas działania ponad dziesięciokrotnie, jednak jego zmiana powoduje jedynie niewielkie zmiany w otrzymywanych obrazach. W przypadku nie obliczania odbłasków parametr Sight Angle minimalnie wpływa na czas obliczeń metody.

Metoda Tumblina i Turka z 1999 roku

Parametry	Czasy obliczeń		
3xLCIS, 100 kroków	8.188 sek.	29.813 sek.	66.806 sek.
3xLCIS, 250 kroków	19.863 sek.	72.159 sek.	162.708 sek.
4xLCIS, 100 kroków	12.154 sek.	44.429 sek.	98.887 sek.
4xLCIS, 250 kroków	29.533 sek.	108.191 sek.	238.903 sek.
5xLCIS, 100 kroków	16.068 sek.	58.95 sek.	132.034 sek.
5xLCIS, 250 kroków	39.186 sek.	144.392 sek.	318.623 sek.

W tej metodzie na czas obliczeń najbardziej wpływa liczba wykorzystywanych obrazów LCIS oraz liczba kroków wykonywanych przy obliczaniu każdego z obrazów LCIS. Parametry określające długość każdego kroku oraz k i w nie zmieniają czasu działania metody, poza przypadkiem gdy $k = 0$. Wtedy obraz LCIS to po prostu luminancja oryginalnego obrazu. W takiej sytuacji obliczanie obrazu LCIS sprowadza się do skopiowania luminancji obrazu, co trwa nieporównywalnie krócej od standardowych obliczeń (należy pamiętać, że wartość domyślna dla zerowego obrazu LCIS wynosi właśnie 0). Potwierdzeniem powyższego są wyniki zamieszczone w tabeli, w których widać dosyć wyraźnie liniową zmianę czasów w zależności od liczby kroków oraz liczby obrazów LCIS.

Metoda Reinharda et al. z 2002 roku

Parametry	Czas obliczeń		
Sharpness = 1, Convolver = Pyramid	0.538 sek.	1.389 sek.	3.352 sek.
Sharpness = 16, Convolver = Pyramid	0.598 sek.	1.706 sek.	3.836 sek.
Sharpness = 1, Convolver = FFT	1.119 sek.	5.561 sek.	14.08 sek.
Sharpness = 16, Convolver = FFT	1.178 sek.	5.761 sek.	14.34 sek.
Sharpness = 1, Convolver = Simple	5.034 sek.	7.793 sek.	31.167 sek.
Sharpness = 16, Convolver = Simple	6.169 sek.	20.68 sek.	44.261 sek.
globalna wersja operatora	0.338 sek.	1.008 sek.	2.173 sek.

Jak można się było spodziewać, z racji nie obliczania splotów, najszybsza jest wersja globalna operatora. Operator lokalny obliczający sploty za pomocą piramidy obrazów jest niecałe dwa razy wolniejszy od wersji globalnej oraz od 5 do 10 razy szybszy od metody obliczającej sploty z definicji (przy Sharpness = 1, przy większych wartościach tego parametru różnica jest jeszcze większa na korzyść piramidy obrazów). Zdecydowanie najwolniejsza jest wersja lokalna, obliczająca sploty bezpośrednio z definicji. Obliczanie splotów za pomocą szybkiej transformaty Fouriera jest szybsze od obliczania z definicji oraz wolniejsze od obliczania za pomocą piramidy obrazów.

Na czas wykonywania operatora lokalnego w wersji Simple i Pyramid wyraźny wpływ ma wartość parametru Sharpness. Jest to szczególnie widoczne przy nieprzyspieszonym obliczaniu splotów. Szybkość wersji korzystającej z szybkiej transformaty Fouriera prawie nie zależy od wartości parametru Sharpness.

Metoda korzystająca z piramidy obrazów ma niestety również wady. Rezultaty jej działania nie są dokładnie takie jak przy wersji Simple lub FFT. Przy dużych wartościach parametru Sharpness (> 12) w wielu scenach metoda korzystająca z piramidy obrazów może tworzyć mocno widoczne artefakty w pobliżu najbardziej kontrastowych obszarów obrazu.

Metoda Ashikhmina z 2002 roku

Parametry	Czas obliczeń		
Threshold = 0.1, Convolver = Pyramid	0.853 sek.	2.2 sek.	5.718 sek.
Threshold = 0.5, Convolver = Pyramid	0.936 sek.	2.861 sek.	6.543 sek.
Threshold = 1.0, Convolver = Pyramid	0.951 sek.	2.991 sek.	6.669 sek.
Threshold = 0.1, Convolver = FFT	2.284 sek.	12.141 sek.	30.764 sek.
Threshold = 0.5, Convolver = FFT	2.291 sek.	12.524 sek.	32.453 sek.
Threshold = 1.0, Convolver = FFT	2.295 sek.	12.608 sek.	32.011 sek.
Threshold = 0.1, Convolver = Simple	13.649 sek.	17.749 sek.	79.344 sek.
Threshold = 0.5, Convolver = Simple	15.626 sek.	45.956 sek.	113.138 sek.
Threshold = 1.0, Convolver = Simple	15.866 sek.	50.839 sek.	114.519 sek.
globalna wersja operatora	0.31 sek.	0.985 sek.	2.06 sek.

Wyniki działania metody Ashikhmina są analogiczne do wyników metody Reinharda. Najszybsza jest globalna wersja operatora i jest ona ponad dwa razy szybsza do wersji lokalnej z przyspieszonymi piramidami obliczaniem splotów. Najwolniejsza jest metoda lokalna nieprzyspieszająca obliczania splotów. Wersja korzystająca z

szybkiej transformaty Fouriera daje, tak jak wcześniej, wyniki pośrednie. Warto zauważyć również, że globalna wersja operatora Ashikhmina jest minimalnie szybsza od globalnej wersji operatora Reinharda, a pozostałe warianty są o ok. jedną trzecią wolniejsze.

Parametrowi Sharpness z metody Reinharda odpowiada Threshold, zwiększenie jego wartości powoduje wydłużenie czasu mapowania w wariantach Pyramid i Simple. Czas obliczeń wersji FFT praktycznie nie zależy od wartości parametru Threshold.

Wykorzystanie pyramid do obliczania splotów daje w rezultacie tylko trochę inny obraz niż wersje Simple i FFT. Nie powoduje powstawania artefaktów, jednak uzyskiwany obraz jest mniej wyraźny, tak jakby był otrzymany z mniejszej wartości parametru Threshold.

Metoda Duranda i Dorsey z 2002 roku

Filtr bilateralny	Czas obliczeń		
prosty	118.035 sek.	348.71 sek.	880.84 sek.
przedziałami liniowy	2.25 sek.	13.166 sek.	42.558 sek.
szybki	0.487 sek.	1.662 sek.	3.835 sek.

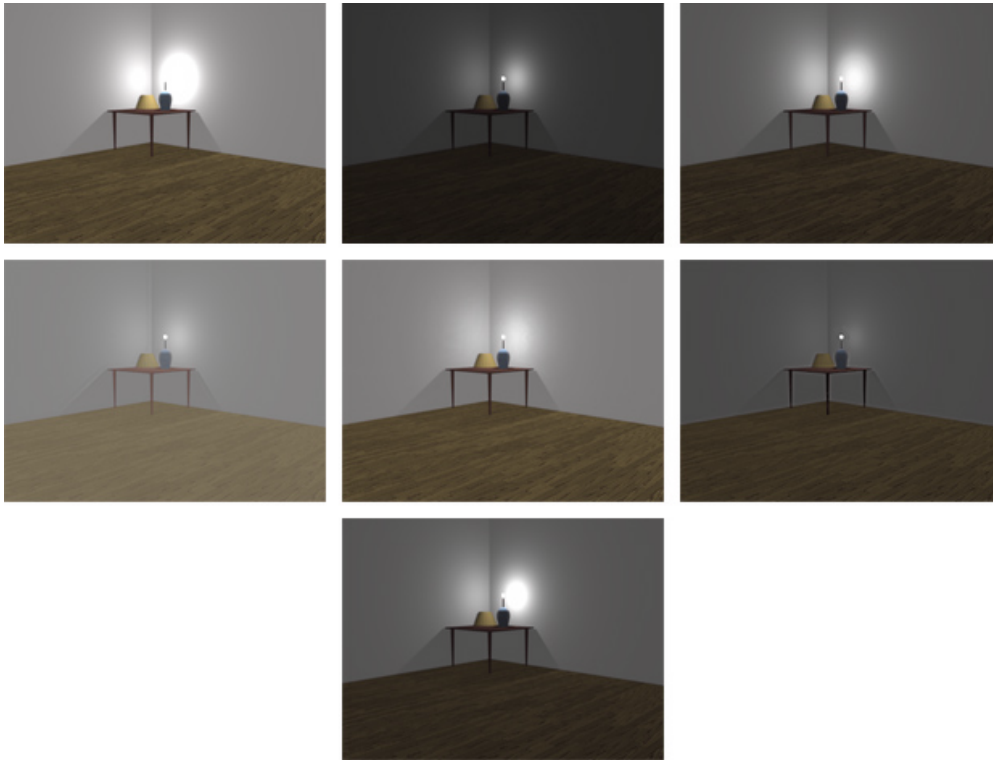
Jak widać obliczanie filtra bilateralnego z definicji (wersja prosta operatora) jest bardzo czasochłonne. Czas obliczeń wersji prostej nie zależy od wartości parametrów Sigma R i Contrast, zależy natomiast istotnie od parametru Sigma S, określającego rozmiar filtra 2.62. Czas w tabeli odpowiada standardowej wartości Sigma S równej 10. Zmniejszenie jego wartości mogłoby przyspieszyć metodę jednak jej autorzy zalecają używanie większych wartości Sigma S a to z kolei mocno wydłużyłoby czas obliczeń. Na szczęście przedstawione przez Duranda i Dorsey metody przyspieszania filtrowania bilateralnego sprawdzają się bardzo dobrze. Wersja przedziałami liniowa jest przynajmniej 20 razy szybsza od wersji prostej, natomiast wersja szybka jest co najmniej 4.5 razy szybsza od wersji przedziałami liniowej.

Czasy działania wersji szybkiej i przedziałami liniowej nie zależą od wartości Sigma S. Na ich działanie wpływ mają parametry Segments i Downsample, ilustrują to rys. 2.27 i rys. 2.29

4.3. Otrzymane obrazy

Wszystkie rezultaty zamieszczone w tym rozdziale powstały przy użyciu programu TM. Nie były one poprawiane po mapowaniu za pomocą opcji Color lub Brightness/Contrast, lub innych programów graficznych.

Najpierw zajmiemy się obrazami otrzymanymi przez zastosowanie metod z parametrami domyślnymi, dzięki czemu będziemy w stanie częściowo stwierdzić jak skomplikowane jest korzystanie z poszczególnych metod. Dowiemy się które części obrazów testowych będą sprawiać najwięcej problemów oraz w jaki sposób objawiać się będą wady poszczególnych metod. Następnie przedstawione zostaną rezultaty mapowań z parametrami dopasowanymi do obrazów testowych, co wykaże pełne możliwości poszczególnych metod. Dodatkowo zmiany w parametrach potrzebne do osiągnięcia dobrych rezultatów uzupełnią naszą wiedzę o tym jak wymagające od użytkownika jest korzystanie z metod.



Rysunek 4.35: Obraz lamp.hdr mapowany wszystkimi zaimplementowanymi operatorami ze standardowymi wartościami parametrów. Rezultaty użycia metod od góry, od lewej: mapowania liniowego z korekcją gamma, Schlicka z 1994 roku, Warda z 1997 roku, Tumblina i Turka z 1999 roku, Reinharda et al. z 2002 roku, Ashikhmina z 2002 roku i Duranda i Dorsey z 2002 roku.

4.3.1. Mapowanie z parametrami domyślnymi

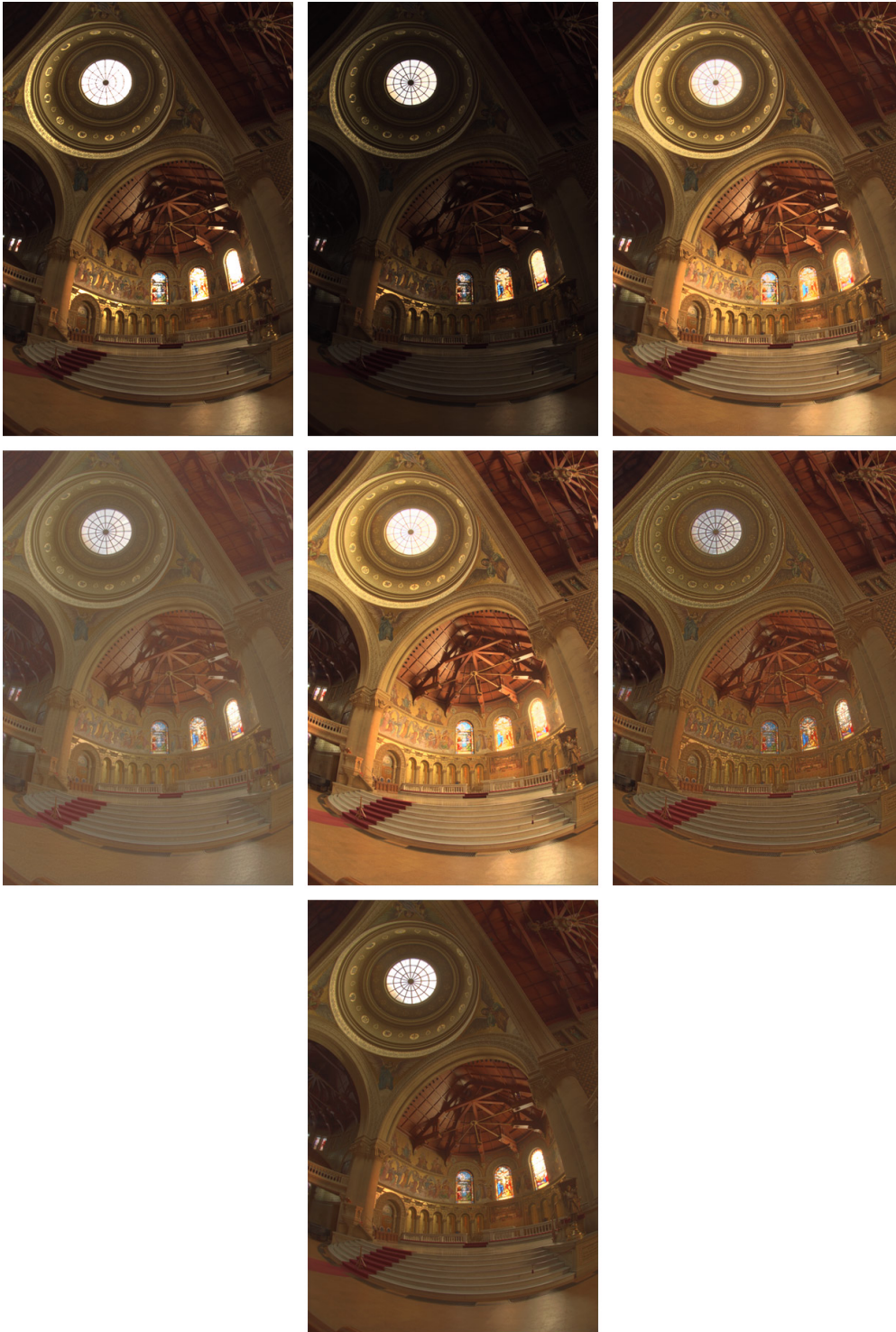
Obraz lamp.hdr

Rezultaty mapowania obrazka lamp.hdr przedstawia rys. 4.35. Znajdują się na nim od góry, od lewej rezultaty użycia metody: mapowania liniowego z korekcją gamma, Schlicka z 1994 roku, Warda z 1997 roku, Tumblina i Turka z 1999 roku, Reinharda et al. z 2002 roku, Ashikhmina z 2002 roku i Duranda i Dorsey z 2002 roku.

Jak widać rezultaty różnią się ogólną jasnością, w czym najbardziej wyróżnia się metoda Schlicka, dająca zdecydowanie najciemniejszy obraz. Mimo tego prawie wszystkie metody są w stanie wystarczająco dobrze skompresować jasności obrazu (również metoda Schlicka), rezultaty ich działania można by właściwie pozostawić bez dodatkowej korekcji parametrów. Najlepsze rezultaty osiągnęły metody Reinharda i Warda. Mimo tego, że pierwsza z nich daje trochę za jasny obraz, a druga tworzy minimalne artefakty na ścianie w pobliżu lampy.

Jedynie metoda mapowania liniowego z korekcją gamma i metoda Duranda i Dorsey nie zachowały detali przy lampie.

4.3. Otrzymane obrazy



Rysunek 4.36: Obraz memorial.hdr mapowany ze standardowymi wartościami parametrów. Kolejność wymików taka jak w rys. 4.35.

Metody Tumblina i Ashikhmina skompresowały kontrast obrazu aż za mocno, co zmniejszyło głębię obrazu. Widać to porównując wielkość i gładkość zmiany jasności owalnego obszaru ściany najbliższej lampy. W przypadku metod Tumblina i Ashikhmina



Rysunek 4.37: Obraz bigfogmap.hdr mapowany wszystkimi zaimplementowanymi operatorami ze standardowymi wartościami parametrów. Kolejność wyników taka jak w rys. 4.35.

obszar ten bardzo słabo wyróżnia się z powierzchni ściany. Dodatkowo metody te tworzą lekkie artefakty "halo", metoda Ashikhmina praktycznie na wszystkich krawędzi-

ach a metoda Tumblina na krawędzi cienia stołu rzucanego na ścianie. Obraz odwzorowany przez metodę Tumblina jest również trochę wyblakły.

Obraz memorial.hdr

Rezultaty mapowania obrazka memorial.hdr zawiera rys. 4.36. Kolejność umieszczonych obrazków jest taka sama jak w rys. 4.35.

Tak jak w przypadku lamp.hdr najbardziej widoczne są różnice w ogólnej jasności rezultatów. Znow najciemniejszy jest wynik metody Schlicka, tym razem jednak jego jasność jest zbyt mała przez co duża część detali z ciemniejszych obszarów sceny jest niewidoczna. Mapowanie liniowe z kompresją gamma nie jest w stanie odwzorować szczegółów zarówno w najjaśniejszych jak i najciemniejszych obszarach.

Pod względem zachowania widoczności jak największej liczby detali, najlepiej wypada metoda Ashikhmina. Rezultat jej działania zachowuje odpowiednią widoczność zdecydowanej większości detali obrazu (poza lewym górnym rogiem). Na szczególną uwagę zasługuje wysoka jakość odwzorowania szklanego zwieńczenia kopuły i witraże. Podobną, choć jednak trochę słabszą, jakość kompresji prezentuje metoda Tumblina, rezultat jej działania jest, tak jak w przypadku obrazu lamp.hdr, trochę wyblakły.

Metody Tumblina i Ashikhmina dobrze zachowują widoczność detali, nie radzą sobie jednak z zachowaniem głębi obrazu. Z tym zadaniem najlepiej uporały się metody Warda i Reinharda, zachowując jednocześnie wysoki poziom detalu.

Obraz bigfogmap.hdr

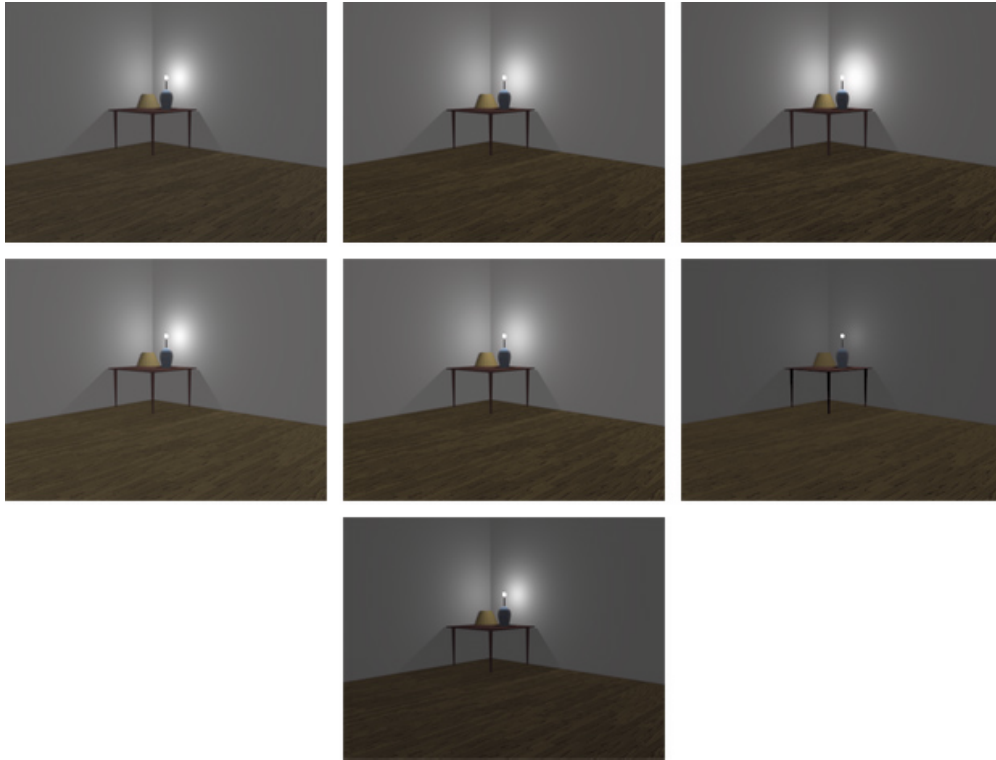
Na rys. 4.37 przedstawione są wyniki mapowania obrazka bigfogmap.hdr. Kolejność umieszczonych rezultatów jest taka sama jak w przypadku rys. 4.35.

Tak jak w dwóch poprzednich obrazach testowych, również teraz obserwujemy wahania ogólnej jasności rezultatów. Wyróżnia się tu szczególnie wynik mapowania liniowego z korekcją gamma. Metoda ta nie radzi sobie zupełnie z obrazem, co sprawia że na jego większej części widać tylko jasną plamę, bez żadnych szczegółów. Najciemniejszy jest wynik metody Schlicka. Widoczność detali w najciemniejszych obszarach sceny jest niewystarczająca. Detale w jasnych częściach obrazu, w tym przypadku część drzewa zasłaniająca lampę, są zbyt mało wyraźne.

Słabo wypada również metoda Warda. Nie odwzorowuje dokładnie źródła światła, tworząc do tego duży rozbłysk z jego prawej strony. Detale w ciemnych obszarach obrazu są także słabo widoczne.

Podobne do siebie rezultaty dają metody Tumblina i Duranda. Zamiast źródła światła i przysłaniających je liści widać tylko białą-żółtą plamę, jednak poza tym miejscem obrazu zachowują szczegóły nawet w ciemnych partiach. Metoda Tumblina daje, jeszcze bardziej niż przy lamp.hdr i memorial.hdr, rozjaśniony i wyblakły obraz.

Z obrazem bigfogmap.hdr najlepiej poradziły sobie metody Ashikhmina i Reinharda. Na szczególną uwagę zasługuje bardzo dobre odwzorowanie przez metodę Ashikhmina źródła światła, wyraźnie widoczne są zasłaniające je liście. Lepiej niż w przypadku innych metod widać również odbicia światła w tylnych szybach samochodów pod drzewem. Metoda Reinharda z najtrudniejszymi obszarami sceny radzi sobie trochę gorzej od metody Ashikhmina, jednak lepiej od pozostałych.



Rysunek 4.38: Rezultaty mapowania obrazu lamp.hdr. Kolejność wyników taka sama jak na rys. 4.35.

4.3.2. Mapowanie z dopasowanymi parametrami

Obraz lamp.hdr

Na rys. 4.38 przedstawione są rezultaty mapowania obrazu lamp.hdr z parametrami, których zastosowanie dawało najlepsze wyniki.

Jak widać wszystkie metody dały zadowalające wyniki. Rezultat metody Ashikhmina nie ma już artefaktów przy krawędziach, w dalszym ciągu pozostał jednak zbyt mało kontrastowy. Poza przypadkiem Ashikhmina wszystkie otrzymane obrazy są do siebie bardzo podobne. Występują tylko niewielkie różnice w ogólnej jasności, rozmiarze i gładkości owalnego obszaru ściany najbliższej lampy. Nieznacznie ciemniejsze od innych są efekty działania metod Ashikhmina oraz Duranda.

Do osiągnięcia powyższych rezultatów użyte zostały następujące parametry:

- Mapowanie liniowe z korekcją gamma: Exposure = 0.27, Gamma = 2.7. We wszystkich metodach parametr Gamma odpowiedzialny jest jedynie za korekcję opisaną w rozdziale 1.6. Wyjątkiem od tej reguły jest metoda mapowania liniowego z korekcją gamma, w której podwyższona wartość parametru dodatkowo kompresuje jasności. Możliwości kompresji za pomocą gammy są mocno ograniczone ponieważ zwiększenie parametru szybko powoduje blaknięcie obrazu. W przypadku obrazu lamp.hdr lekkie zwiększenie parametru, razem z odpowiednią wartością Exposure dało bardzo dobre efekty. Przypomnijmy, że najlepsze rezultaty powinna dawać wartość Exposure równa odwrotności luminancji naj-

ciemniejszego miejsca, które w wynikowym obrazie ma być białe.

- Metoda Schlicka z 1994 roku: $DNBG = 4$. Metoda Schlicka powinna działać najlepiej gdy parametr DNBG wybierany jest za pomocą opcji Choose DNBG. Często jednak uzyskuje się lepsze rezultaty przy DNGB trochę mniejszym lub większym od wybranego za pomocą Choose DNBG. W warunkach testowych korzystając z opcji Choose DNBG otrzymywaliśmy wartość 9. Dla sceny lamp.hdr lepiej wyglądały obrazy mapowane z DNBG mniejszym od 9.
- Metoda Warda et al. z 1997 roku: $Sight\ Angle = 170$. Zwykle, gdy opcje odpowiedzialne za symulowanie ludzkiego widzenia są wyłączone, parametr Sight Angle nie ma widocznego wpływu na jakość rezultatu. W przypadku lamp.hdr zwiększona jego wartość eliminuje artefakty ze ściany. Dzieje się tak dlatego, że Sight Angle określa pośrednio liczbę próbek, z których obliczany jest histogram, dzięki czemu jest on dokładniejszy.
- Metoda Tumblina i Turka z 1999 roku: $w\ Color = 0.9$, $LCIS\ Images = 3$ $w0 = 1.0$, $w1 = 1.0$, $w2 = 0.75$, $k0 = 0.0$, $k1 = 0.06$, $k2 = 0.1$. Do osiągnięcia zadowalających rezultatów wystarczą trzy obrazy LCIS. O stopniu kompresji wynikowego obrazu decyduje waga ostatniego z uproszczonych obrazów, w tym przypadku $w2$. Waga zerowego powinna zawsze wynosić 1.0, a wartości wag pośrednich powinny znajdować się między nimi.
- Metoda Reinharda et al. z 2002 roku: $Key\ Value = 0.09$, opcja Is Local wyłączona. Obraz testowy lamp.hdr jest na tyle prosty do odwzorowania, że dobre rezultaty daje już globalna wersja metody Reinharda. Dodatkowo regulacja parametru Key Value umożliwiła dostosowanie jasności obrazu.
- Metoda Ashikhmina z 2002 roku: opcja Is Local wyłączona. Również tu wersja globalna operatora okazała się wystarczająca do odwzorowania obrazu testowego. Niestety metoda Ashikhmina nie umożliwia zmiany ogólnej jasności wyniku.
- Metoda Duranda i Dorsey z 2002 roku: $Contrast = 4.0$. Metoda ta także nie posiada parametru bezpośredniego wpływającego na ogólną jasność. Można co prawda próbować zmienić jasność poprzez zmianę parametrów Contrast, Sigma S i Sigma R, jednak wymaga to wielu prób i nie gwarantuje oczekiwanych zmian.

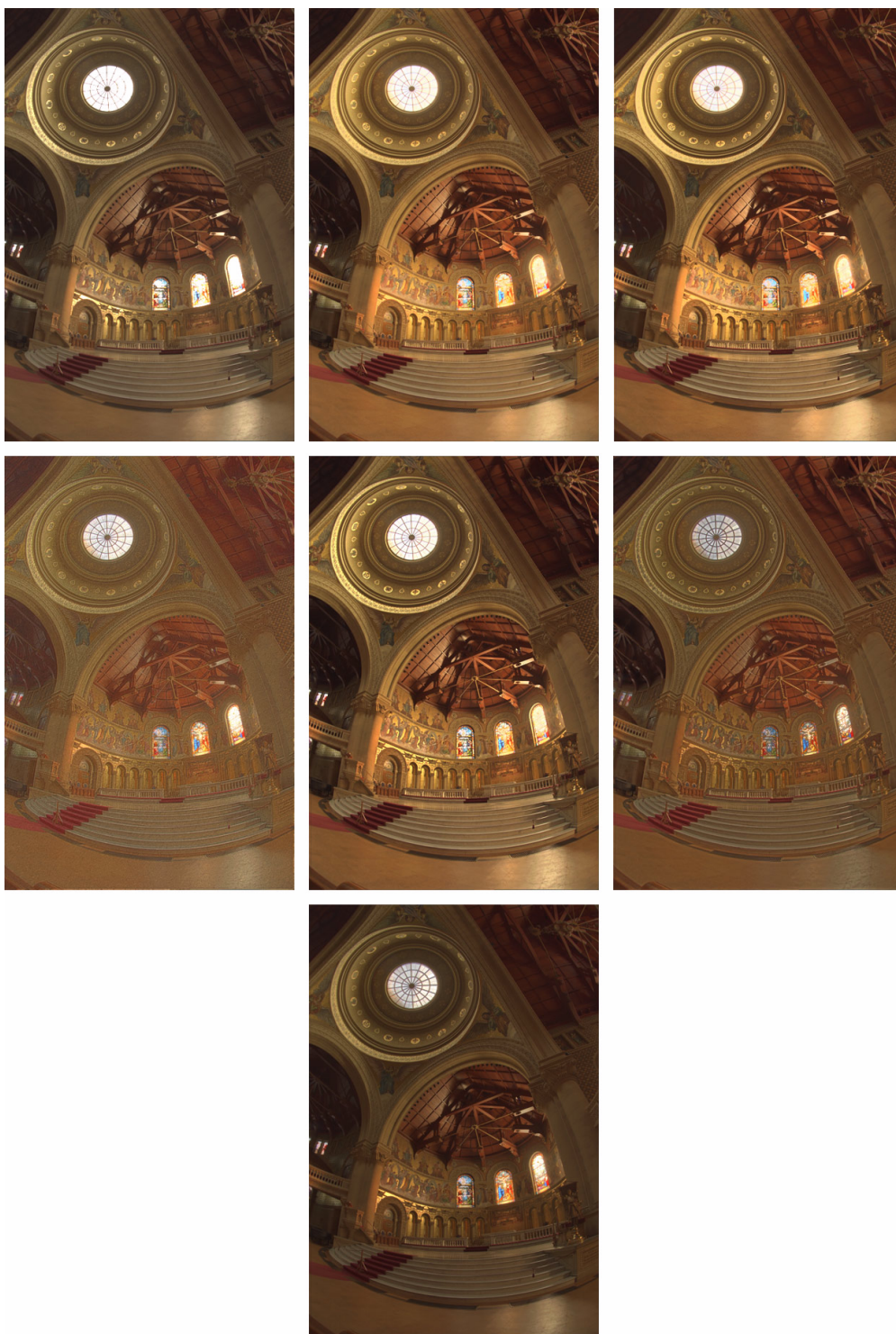
Obraz memorial.hdr

Na rys. 4.39 znajdują się rezultaty mapowania obrazu memorial.hdr z dopasowanymi parametrami.

Różnice w ogólnej jasności dalej są widoczne, jednak znacznie mniej niż w przypadku parametrów domyślnych.

Wśród uzyskanych wyników dość wyraźnie widać podział na metody zachowujące głębię obrazu i metody eksponujące detale. Ogólnie metody z grupy pierwszej słabiej kompresują obraz przez co część detali jest na nich mniej widoczna. Do tej grupy zaliczają się metody: mapowania liniowego z korekcją gamma, Schlicka, Warda, Reinharda i Duranda. Z kolei obrazy generowane przez drugą grupę, czyli metody: Tumblina i Ashikhmina, mają dobrze widoczne detale jednak mogą wydawać się płaskie.

Jak można się było spodziewać, wśród metod pierwszej grupy najslabiej wypadło mapowanie liniowe z korekcją gamma. Metody Warda i Schlicka dają prawie identyczne rezultaty. Detale w kopule nie są tak dobrze odwzorowane jak w przypadku pozostałych



Rysunek 4.39: Rezultaty mapowania obrazu memorial.hdr. Kolejność wyników taka sama jak na rys. 4.35.

metod, jednak wyniki tych metod świetnie zachowują głębię utrzymując widoczność detali na przyzwoitym poziomie. Metoda Reinharda, tak jak dwie wcześniejsze dobrze oddaje głębię a przy tym większość detali zachowuje lepiej od nich. Metoda Duranda

daje zdecydowanie najciemniejszy rezultat a parametry metody nie umożliwiają jego rozjaśnienia. Mimo tego zarówno w najciemniejszych jak i w najjaśniejszych obszarach sceny detale są całkiem wyraźne. Co więcej, jako jedyna z metod zachowujących głębię bardzo dobrze odwzorowuje zwieńczenie kopuły.

Do uzyskania rezultatów z rys. 4.39 zastosowano następujące parametry:

- Mapowanie liniowe z korekcją gamma: Exposure = 1.0, Gamma = 2.9. Dzięki zwiększeniu Gammy widoczne są detale w ciemniejszych rejonach sceny a dokładniej belki pod dachem z lewej strony. Zwiększając jeszcze bardziej Gamme można by zwiększyć widoczność jasnych detali, jednak spowodowałoby to, że obraz byłby mocno wyblakły.
- Metoda Schlicka z 1994 roku: DNBG = 9, opcja Non Uniform wyłączona. Manipulacja parametrem Non Uniform powoduje niewielkie zmiany w kontraście otrzymanego obrazu. W tym przypadku rezultat wygląda lepiej przy ustawieniu Non Uniform jako wyłączony.
- Metoda Warda et al. z 1997 roku: wszystkie parametry domyślne. Niestety metoda ta nie daje możliwości regulacji widoczności detali. W niektórych obrazach mogą być one widoczne mniej lub bardziej przy różnych wartościach parametru Adjustment, jednak nie w przypadku memorial.hdr.
- Metoda Tumblina i Turka z 1999 roku: Num Steps = 500, Step Length = 0.5, w Color = 0.9, LCIS Images = 5, w0 = 1.0, w1 = 1.0, w2 = 1.0, w3 = 1.0, w4 = 0.34, k0 = 0, k1 = 0.1, k2 = 0.2, k3 = 0.3, k4 = 0.4. Uzyskany obraz jest trochę wyblakły, za jasny i zdecydowanie brak w nim głębi, jednak zachowuje widoczność nawet najdrobniejszych szczegółów. Autorzy metody zalecają używanie Step Length = 1/32, jednak w przypadku tego obrazu użycie większej wartości nie powoduje pojawienia się żadnych artefaktów.
- Metoda Reinharda et al. z 2002 roku: Key Value = 0.09, Sharpness = 13, Convolver = FFT. Convolver FFT daje dokładniejsze rezultaty od wersji Pyramid a większa wartość Sharpness poprawia widoczność detali.
- Metoda Ashikhmina z 2002 roku: Threshold = 0.4, Convolver = FFT.
- Metoda Duranda i Dorsey z 2002 roku: Contrast = 3.5, Sigma R = 1.3. Zwiększenie parametru Sigma R w przypadku tego obrazu zwiększa głębię.

Obraz bigfogmap.hdr

Na rys. 4.40 przedstawione są rezultaty mapowania obrazu bigfogmap.hdr z wartościami parametrów dostosowanymi do obrazu testowego.

Wspomniany przy obrazie memorial.hdr podział na metody zachowujące głębię i zachowujące detale nie jest tu specjalnie widoczny.

Obraz bigfogmap.hdr w sumie sprawił najwięcej problemów a dokładniej najwięcej metod, bo aż trzy, nie dały zadowalających rezultatów. Do metod tych zalicza się mapowanie liniowe z korekcją gamma, metoda Warda oraz metoda Tumblina. Pierwsze dwie poradziły sobie dobrze praktycznie z całą sceną, poza lampą zasłoniętą przez drzewo. Obszar ten wraz z całkiem niemałym otoczeniem został odwzorowany jako biało-żółta plama, z prawie nie rozróżnialnymi detalami. Wynik działania metody Tumblina co prawda zachowuje dobrze szczegóły, również w najtrudniejszym obszarze



Rysunek 4.40: Rezultaty mapowania obrazu bogfogmap.hdr z parametrami dostosowanymi do obrazu. Kolejność wyników taka sama jak na rys. 4.35.

sceny, jednak jest zdecydowanie za jasny i zbyt wyblakły, do tego przy wielu krawędziach występuje "halo".

Pozostałe metody dały dobre rezultaty, zarówno jeśli chodzi o okolice lampy jak

i pozostałe obszary. Niewątpliwie najdokładniej okolice lampy odwzorowała metoda Ashikhmina. Na jej wyniku najwyraźniej widoczne są liście zasłaniające lampę.

Parametry użyte przy mapowaniu obrazu `bigfogmap.hdr` są następujące:

- Mapowanie liniowe z korekcją gamma: Exposure = 0.035, Gamma = 2.7.
- Metoda Schlicka z 1994 roku: DNGB = 1, opcja Non Uniform wyłączona.
- Metoda Warda et al. z 1997 roku: Adjustment = Linear. W tym przypadku liniowe ograniczenie histogramu dawało najlepszą kompresję detalu w najjaśniejszym obszarze sceny. Niestety w otrzymanym obrazie przejawiało się to tylko zmniejszeniem żółtej plamy w pobliżu lampy.
- Metoda Tumblina i Turka z 1999 roku: Num Steps = 250, Step Length = 0.025, w Color = 0.7, LCIS Images = 4, w0 = 1.0, w1 = 1.0, w2 = 1.0, w3 = 0.175, k0 = 0.0, k1 = 0.05, k2 = 0.1, k3 = 0.15. W tym obrazie testowym użycie większych wartości parametrów Step Length i k powodowało pojawianie się dosyć wyraźnych artefaktów przy krawędziach.
- Metoda Reinharda et al. z 2002 roku: Key Value = 0.14, Sharpness = 12, Convolver = FFT.
- Metoda Ashikhmina z 2002 roku: Threshold = 0.5, Convolver = FFT.
- Metoda Duranda i Dorsey z 2002 roku: Contrast = 4.

4.4. Podsumowanie

Możnaby odnieść wrażenie, że głównym problemem i różnicą między porównywanymi metodami mapowania tonów jest ogólna jasność w generowanych przez nie rezultatach. Problem ten jest jednak tylko pozorny. Można go łatwo wyeliminować stosując do wyników standardowe przekształcenia obrazów nisko-kontrastowych. To samo dotyczy również saturacji i kontrastu. Trzeba jednak pamiętać, że zarówno korekcja jasności jak i nasycenia czy kontrastu metodami nisko-kontrastowymi nie da tak dobrych rezultatów jakie dałyby podobne przekształcenia wykonane przy przetwarzaniu obrazu wysoko-kontrastowego. Dlatego ważne jest aby metoda mapowania tonów umożliwiała łatwą kontrolę nad jasnością sceny, kontrastem i poziomem widoczności detali. W niektórych przypadkach przydatna byłaby także możliwość regulacji nasycenia.

Wśród siedmiu metod przedstawionych w tej pracy żadna nie umożliwiała jednoczesnej kontroli nad wszystkimi własnościami obrazu wyjściowego. Chyba dlatego nie można wybrać spośród nich najlepszej. Każda z metod ma swoje dobre strony i dla pewnych obrazów daje bardzo dobre rezultaty. Jednak dla każdej można również znaleźć obraz, który inne metody odwzorują lepiej.

Podsumowując cechy charakterystyczne wszystkich metod:

Metoda mapowania liniowa z korekcją gamma jest "najsłabszą" z metod. Jest jednak w stanie dobrze odwzorować obraz o średnim zakresie dynamicznym. Do tego jest bardzo łatwa w implementacji, ma najkrótsze czasy obliczeń oraz jest prosta w obsłudze. Mając do czynienia z prostymi scenami, metoda ta może okazać się wystarczająca.

Metoda Schlicka przy minimalnie dłuższych czasach obliczeń i niewiele większym skomplikowaniu zapewnia dobry poziom kompresji tonów, zachowując głębię obrazu.

Pozwala też łatwo modyfikować jakność rezultatu. Niestety nie umożliwia sterowania widocznością detali a w niektórych obrazach jej kompresja okazuje się niewystarczająca do wyraźnego odwzorowania najtrudniejszych detali.

Metoda Warda et al. jest również szybka, w większości przypadków daje dobre mapowanie, nie wymagając praktycznie ingerencji użytkownika. Jednak przy trudniejszych scenach może nie odwzorowywać dobrze detali. Nie ma również żadnej możliwości zmiany jasności czy widoczności szczegółów. Metoda ta jednak jako jedyna umożliwia symulację efektów charakterystycznych dla ludzkiego widzenia.

Metoda Tumblina i Turka potrafi bardzo dobrze wyeksponować nawet subtelne szczegóły z trudnych obrazów. Niestety osiągnięcie dobrych rezultatów wymaga ustawienia od ośmiu do czternastu parametrów co przy bardzo długich czasach działania metody, sprawia że korzystanie z niej jest zdecydowanie trudniejsze od innych metod. Do tego metoda nie umożliwia sterowania jasnością rezultatu a przy tym często daje za jasne wyniki.

Metoda Reinharda et al. wprawdzie nie jest w stanie najlepiej odwzorować najjaśniejszych detali w najbardziej kontrastowych obrazach. W każdej sytuacji zapewnia jednak dobry poziom oddania szczegółu oraz głębi rezultatu. W połączeniu z niezłą szybkością i prostymi w obsłudze parametrami umożliwiającymi intuicyjną zmianę ogólnej jasności i ekspozycji detali, stanowi zdecydowanie jedną z najlepszych metod przedstawionych w tej pracy

Metoda Ashikhmina jest trochę wolniejsza od metody Reinharda, słabo oddaje głębię. Nie ma jednak żadnych problemów z odwzorowywaniem szczegółów nawet najbardziej kontrastowych scen. Do tego wymaga manipulowania tylko jednym parametrem Threshold, określającym widoczność detali. Nie umożliwia zmiany jasności wyniku, jednak przy większości obrazów nie jest to potrzebne ponieważ metoda sama ustala jasność na odpowiednim poziomie.

Metoda Duranda i Dorsey jest najszybszą z metod lokalnych obok metody Reinharda. Dla większości obrazów w celu uzyskania bardzo dobrej widoczności zarówno w ciemnych jak i jasnych partiach sceny wystarczy manipulowanie tylko jednym parametrem Contrast, przy użyciu dwóch innych parametrów Sigma S i Sigma R możliwa jest też zmiana głębi obrazu kosztem widoczności szczegółów. Metodę tę, obok metod Reinharda et al. i Ashikhmina, należy zaliczyć do najlepszych wśród wszystkich przedstawionych w pracy.

Bibliografia

- [1] Adams, A., The negative. The Ansel Adams Photography series. Little, Brown and Company, 1981.
- [2] Ashikhmin, M., A tone mapping algorithm for high contrast images. In 13th Eurographics Workshop on Rendering. Eurographics, June 2002.
- [3] Ashdown, I., Radiosity: A Programmer's Perspective. New York, NY: John Wiley & Sons, Inc., 1994.
- [4] Barash, D., Bilateral Filtering and Anisotropic Diffusion: Towards a Unified Viewpoint, Hewlett-Packard Laboratories Technical Report, HPL-2000-18(R.1), 2000.
- [5] Burt, P., E. Adelson, The Laplacian Pyramid as a Compact Image Code, IEEE Transactions on Communications, Vol. Com-31, No. 4, pp. 532-540, April 1983.
- [6] Curcio, C.A., Kimberly, A.A., Sloan, K.R., Lerea, C.L., Hurley, J.B., Klkock, I.B., Milan, A.H., Distribution and morphology of human cone photoreceptors stained with anti-blue opsin. Journal of Comparative Neurology 312, 610-624, 1991.
- [7] Debevec, P., Reinhard, E., Ward, G., Pattanaik, S., High Dynamic Range Imaging, SIGGRAPH 2004, Course 13, August 9, 2004.
- [8] Debevec, P., Malik, J., Recovering high dynamic range radiance maps from photographs. Proceedings of ACM SIGGRAPH 1997, pp. 369-378, 1997.
- [9] Devlin, K., A review of tone reproduction techniques, Technical Report CSTR-02-005, Department of Computer Science, University of Bristol, November 2002.
- [10] Durand, F., Dorsey, J., Fast bilateral filtering for the display of high-dynamic-range images, ACM Transactions on Graphics, Proceedings of ACM SIGGRAPH 2002, San Antonio, Texas, vol. 21(3), pp. 249-256, 2002.
- [11] Ferwerda, J. A., Pattanaik, S., Shirley, P., Greenberg, D. P. A model of visual adaptation for realistic image synthesis. In Siggraph 96 Conference Proceedings, AddisonWesley, H. Rushmeier, Ed., Annual Conference Series, ACM SIGGRAPH 1996, 249-258, 1996.
- [12] Rozszerzenie do formatu plików tiff umożliwiające zapisywanie obrazów HDR, <http://www.anywhere.com/gward/pixformat/tiffluv.html>.

- [13] Matković, K., Tone Mapping Techniques and Color Image Difference in Global Illumination, Institut für Computergraphik eingereicht an der Technischen Universität Wien Technisch-Naturwissenschaftliche Fakultät, PhD Dissertation, December, 1997.
- [14] Badanie ostrości wzroku, http://www.mediweb.pl/data/print_bd.php?id=25.
- [15] Nitzberg, M., Shiotu, T., Nonlinear image filtering with edge and corner enhancement. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 14(8):826–833, August 1992.
- [16] Format pliku HDR firmy Industrial Light & Magic, <http://www.openexr.com/>.
- [17] Peli, E., Contrast in complex images. *Journal of the Optical Society of America*, A 7, 10, 2032–2040, October 1990.
- [18] Perona, P., Malik, J., Scale-space and edge detection using anisotropic diffusion. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 12(7):629–639, July 1990.
- [19] Renderer Radiance, <http://radsite.lbl.gov/radiance/>.
- [20] Reinhard, E., Stark, M., Shirley, P. and Ferwerda, J., Photographic tone reproduction for digital images, *ACM Transactions on Graphics, Proceedings of ACM SIGGRAPH 2002*, San Antonio, Texas, vol. 21(3), pp. 267-276, 2002.
- [21] Schlick, C., Quantization techniques for visualization of high dynamic range pictures. *5th Eurographics Workshop on Rendering*, 7–20, 1994.
- [22] Spencer, G., Shirley, P., Zimmerman, K., Greenberg, D., Physically-based glare effects for computer generated images, *Proceedings of ACM SIGGRAPH 1995*, pp. 325-334, 1995.
- [23] Tomasi, C., and Manduchi, R., Bilateral filtering for gray and color images. *Proceedings of IEEE Int. Conf. on Computer Vision*, 836–846, 1998.
- [24] Tumblin, J , and Turk, G., LCIS: A boundary hierarchy for detail-preserving contrast reduction. In *Siggraph 1999, Computer Graphics Proceedings*, Addison Wesley Longman, Los Angeles, A. Rockwood, Ed., Annual Conference Series, 83–90, 1999.
- [25] Ward, G., Rushmeier, H., Piatko, C., A visibility matching tone reproduction operator for high dynamic range scenes. *IEEE Transactions on Visualization and Computer Graphics* 3, 4, 291-306, December 1997.
- [26] Wikipedia, Wolna Encyklopedia, <http://pl.wikipedia.org> .