

Multimedia 1: Daten und Formate (188.135)

Ausarbeitung – Prüfungsordner - PO

Datentypen

Erklären Sie kurz die Begriffe Dynamic Range und Dithering. [2] bzw. Erklären Sie kurz Cluster Dithering. Weshalb wird Clustering Dithering eingesetzt? [2.5]

Dynamic Range

Bei achromatischem Licht (nur Grauwert) stellt die Dynamic Range das Verhältnis zwischen der minimalen und maximalen Intensität (Helligkeit) dar. Wie viele Intensitäten in einem Bild zur Darstellung benötigt werden, hängt vom verwendeten Medium ab. Die Intensitäten reichen aus, wenn die Übergänge in der Darstellung kontinuierlich sind. (Hilfe: Dithering)

Dithering (siehe auch Halftone Approximation)

Approximierung von Intensitäten unter der Verwendung von Dot-Patterns (Punktmustern). Dies funktioniert, da das menschliche Auge so feine Punktmuster vermischt (schwarz/weiß → grau) zu einem Grauton und wird vom menschlichen Auge als Schattierungen wahrgenommen.

Dadurch wird bei Bildern mit niedriger Farbtiefe die Illusion einer höheren Farbtiefe erzeugt. Es werden die fehlenden Farben durch eine bestimmte Anordnung aus verfügbaren Farben nachgebildet und dadurch harte Übergänge zwischen den Farben vermieden. Das menschliche Auge nimmt diese als Mischung der einzelnen Farben wahr.

Verwendung: am häufigsten bei Farbverläufen, z.B. auch bei Druckern (4 Farben). So werden z.B. bei Druckern auf die weiße Fläche des Papiers mehr oder weniger große Punkte aufgetragen. Man versucht dabei die Punkte so zu setzen, dass sie keine regelmäßigen Muster erzeugen. Ist dies nicht möglich, so nimmt man diagonale Linien.

Cluster Dithering

Ditherintensität zentriert auf die Matrizenmitte. [Line patterns](#) und andere Artefakte vermeiden, Patterns sollen wachsen (Minimierung des Kontureffekts), Patterns sollen sich vom Zentrum ausbreiten (Effekt wachsender Punktgrößen).

Es gibt viele verschiedene Algorithmen, diese arbeiten oft auch mit einer Fehlerstreuung (Übertragung von Überläufen oder Rundungsfehlern auf benachbarte Bildpunkte für eine feinere Darstellung). Sind Linien nicht zu vermeiden, verwende diagonale Linien.

Was versteht man unter Halftone Approximation, wofür wird sie eingesetzt und worauf ist dabei zu achten? [2]

Das Dithering ist eine Technik in der Computergrafik, um bei Bildern mit geringer Farbtiefe die Illusion einer größeren Farbtiefe zu erzeugen. Approximierung von Intensitäten unter der Verwendung von Dot-Patterns.

Es werden die fehlenden Farben durch eine bestimmte Anordnung aus verfügbaren Farben/ Pixeln nachgebildet und dadurch harte Übergänge zwischen den Farben/ Pixeln vermieden. Das menschliche Auge nimmt diese als Mischung der einzelnen Farben wahr.

Verwendung: am häufigsten bei Farbverläufen

Zu beachten

Linepatterns und andere Artefakte vermeiden, Patterns sollen wachsen (Minimierung des Kontureffekts), Patterns sollen sich vom Zentrum ausbreiten (Effekt wachsender Punktgrößen)

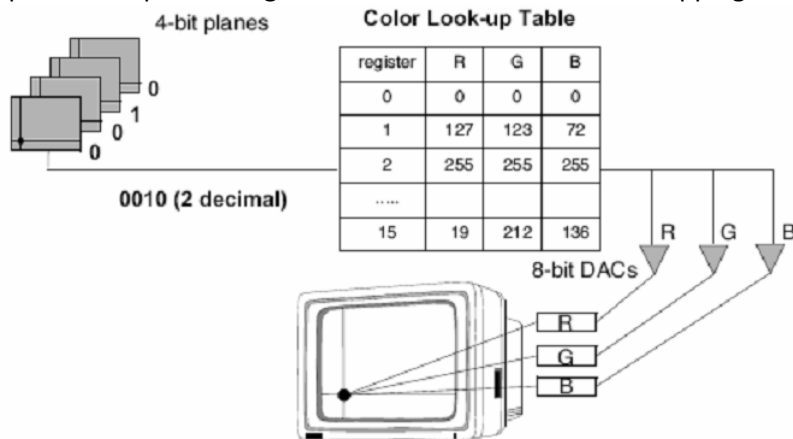
Es gibt viele verschiedene Algorithmen, diese arbeiten oft auch mit einer Fehlerstreuung (Übertragung von Überläufen oder Rundungsfehlern auf benachbarte Bildpunkte für eine feinere Darstellung).

Was sind indizierte Farben (indexed colors) und weshalb werden sie verwendet? [2]

In diesem Farbmodell wird der Farbwert eines Pixels (max. 256 Farben = 2^8 Bit = 1 Byte) durch einen Index in Color Maps oder aus einer CLUT (Color Lookup Table) ausgelesen.

Diese Maps oder Tables sind entweder vordefiniert oder Teil des Bildes. Jedes Pixel verweist auf einen Eintrag in diesem Table der dem jeweiligen Farbwert entspricht.

Es dient der platzsparenden Speicherung von Bildern. Wird beim Color-Mapping verwendet.



Erklären Sie kurz das Prinzip des Zeilensamplings (Line Sampling) für digitales Video. Was bedeuten die n:m:l Angaben im CCIR 601 Videoformat? [3] bzw. Was bedeuten die Komponenten n:m:l in CCIR 601 Video? [1.5]

Diese Angabe bedeutet, dass die Komponente Y (Helligkeit) n-Mal und die Farbkomponenten U und V m- und l-Mal pro Taktperiode gesampled werden. Die Möglichen Taktmultiplikatoren sind: 1,2,3 oder 4.

Sampling

Das Konvertieren von einem analogen zu einem digitalen Video nennt man sampling. Dabei wird das analoge Signal zu bestimmten Zeitpunkten abgetastet und für die jeweiligen Werte eine digitale Repräsentation gesucht. Die einzelnen Frames werden dann durch ein Pixelarray dargestellt. Hier sollte man nach dem Abtasttheorem von Nequist/ Shannon anwenden (Abtastrate = $2 \times$ höchste Frequenz im Signal), um einen Informationsverlust zu vermeiden.

Das YCbCr-Farbmodell

Es gibt Farbmodelle, die eine Farbe nicht durch die additiven Grundfarben Rot, Grün und Blau (kurz RGB), sondern durch andere Eigenschaften ausdrücken. So zum Beispiel das Helligkeit-Farbigkeit-Modell. Hier sind die Kriterien die Grundhelligkeit der Farbe (von Schwarz über Grau bis Weiß), die Farbe mit dem größten Anteil (Rot, Gelb, Grün, Türkis, Blau, Violett bzw. weitere dazwischen liegende reine Farben) und die Sättigung der Farbe, ("knallig" gegenüber blass).

Dieses Farbmodell beruht auf der Fähigkeit des Auges, geringe Luminanz- Unterschiede (Helligkeitsunterschiede) besser zu erkennen als kleine Farbtonunterschiede, und diese wiederum besser als kleine Farbsättigungsunterschiede. So ist ein grau auf schwarz geschriebener Text sehr gut zu lesen, ein blau auf rot geschriebener, bei gleicher Grundhelligkeit der Farben, allerdings sehr schlecht. Solche Farbmodelle nennt man Helligkeit-Farbigkeit-Modelle.

Das YCbCr-Modell ist eine leichte Abwandlung eines solchen Helligkeit-Farbigkeits-Modells. Es wird ein RGB-Farbwert in eine Grundhelligkeit Y und zwei Komponenten Cb und Cr aufgeteilt, wobei Cb ein Maß für die Abweichung von Grau in Richtung Blau ist bzw., wenn es kleiner als 0,5 ist, in Richtung Gelb (Komplementärfarbe von Blau). Cr ist die entsprechende Maßzahl für Abweichung in Richtung Rot bzw. Türkis (Komplementärfarbe von Rot). Diese Darstellung verwendet die Besonderheit des Auges, für grünes Licht besonders empfindlich zu sein. Daher steckt die meiste Information über den Grünanteil (und damit indirekt für dessen Komplementärfarbe Violett) in der Grundhelligkeit Y und man braucht daneben nur noch die Abweichungen beim Rot/Türkis- oder Blau/Gelb-Anteil darzustellen. Die Y-Werte werden dann in den meisten praktischen Anwendungen, etwa auf DVDs, doppelt so fein aufgelöst wie die beiden Werte Cb und Cr (Chroma Subsampling);

CCIR 601 Standard (line sampling)

Stellt eine Norm zur Digitalisierung analoger Videosignale dar. Basiert auf dem YCC - Farbmodell (Y = Helligkeitskanal, CC = 2 Farbkanäle und ist eine Form des YCbCr-Farbmodells). Es gibt mehrere CCIR „Familien“, die sich durch ihre unterschiedlichen Werte für m:n:l unterscheiden.

- m steht hierbei für die Basisrate, mit der Y abgetastet wird.
- n und l bestimmen, mit welcher Rate die beiden Farbkanäle (CC) abgetastet werden.
- Die Basisrate für das Sampling beträgt standardmäßig 3.375MHz.
- m,n und l enthalten nur noch den Multiplikator („Taktmultiplikatoren“) dieser Rate. Dieser kann 1,2,3 oder 4 sein.

Einer der großen Vorteile von YCbCr ist, dass die Abtastrate der Chrominanz-Kanäle niedriger als die des Y Kanals sein kann, ohne dass es zu einer spürbaren Verringerung der zu gewährleistenden Qualität kommt (Chroma Subsampling).

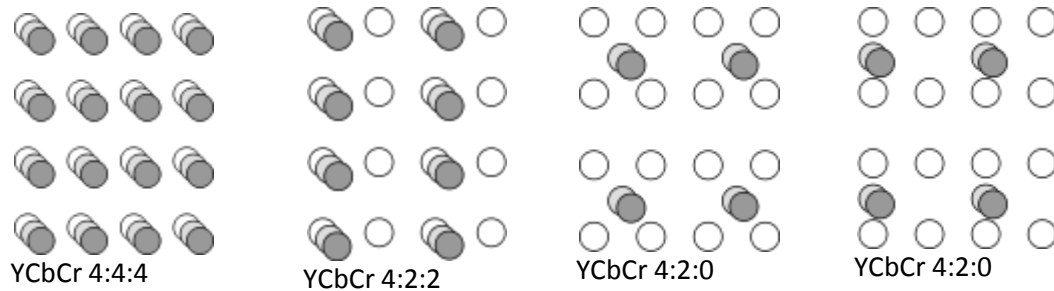
- YCbCr 4:4:4 Chrominanz-Auflösung identisch zur Luma-Auflösung
- YCbCr 4:1:1 Chrominanz-Auflösung horizontal geviertelt und vertikal unverändert

Bei YCbCr 4:2:2 handelt es sich um horizontale Unterabtastung der Farbinformation (Chrominanz) um den Faktor 2, somit wurde die Chrominanz-Auflösung horizontal halbiert. Die 4:2:2-Kodierung entspricht der "Studioqualität" für digitales Video (ITU-R 601-4) und wird von professioneller Aufnahme- und Schnitthardware verwendet.

Bei YCbCr 4:2:0 wird die Chrominanz-Auflösung sowohl horizontal als auch vertikal halbiert. Für Übertragung und für den Heimgebrauch wird typischerweise 4:2:0 verwendet.

Es gibt weiterhin noch Unterschiede über das Zentrum der Chrominanzwerte. Diese können zentriert oder nicht zentriert sein:

Zur Beschreibung wird eine so genannte n:m:C-Notation verwendet. Sie gibt wieder, wie oft Cb und Cr im Vergleich zu Y abgetastet werden.



Weiß Kreise symbolisieren hier die Abtastpositionen für die Luminanz, graue Kreise die Positionen für die Chrominanz.

Beispiel

4:4:4 Pro Pixel werden alle Werte abgetastet. Die ergibt die höchstmögliche Qualität (Studio).

4:2:2 Nur von jedem 2. Pixel pro Zeile werden Farbwerte abgetastet (Broadcast)

4:1:1 Nur von jedem 4. Pixel pro Zeile werden Farbwerte abgetastet (VHS)

Erklären Sie kurz die Begriffe skew, jitter und delay. Nennen Sie 2 Hardwarekomponenten für die Behebung dieser Probleme. [3]

Skew

Mit Skew wird die Zeitdifferenz eines Signals zu seinem korrekten Abspielzeitpunkt (Videosignal) angegeben. Er entsteht beispielsweise aufgrund von Übertragungsdivergenzen (Wenn z.B. das Lesen der Scanline für Video länger dauert, als das für Audio, diese aber synchron abgespielt werden müssen. (Zeitversetztes Lesen, gleichzeitiges Abspielen)). Des Weiteren können solche Divergenzen auch bei der Übertragung über verschiedene Medien entstehen. (Asynchronität zweier Streams).

Jitter

Als Jitter bezeichnet man ein Taktzittern bei der Übertragung von Digitalsignalen, eine leichte Genauigkeitsschwankung im Übertragungstakt (Clock).

Entsteht bei einer unterschiedlichen Verzögerung der Signale eines Datenstroms. Abrupte Störung des Signals in Signalkomponenten z.B. hervorgerufen durch Fehlerwahrscheinlichkeit beweglicher Teile. Lösung: Ein Sync Generator kann helfen, da er einen stabilen Zeittakt vorgibt.

Allgemeiner ist Jitter in der Übertragungstechnik ein abrupter und unerwünschter Wechsel der Signalcharakteristik. Dies kann sowohl Amplitude als auch Frequenz und Phasenlage betreffen.

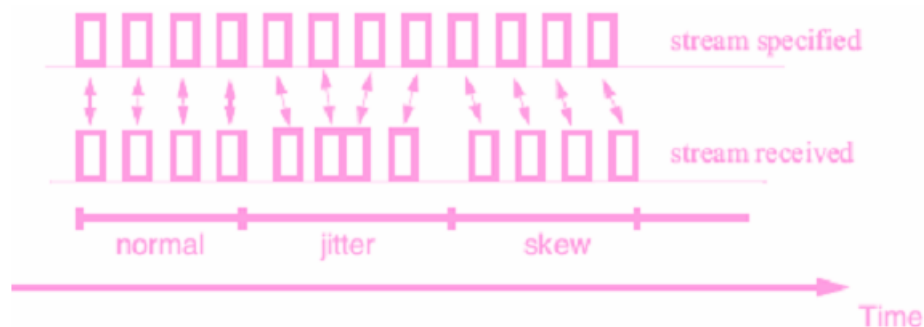
In der Netzwerktechnik wird mit Jitter außerdem die Varianz der Laufzeit von Datenpaketen bezeichnet. Dieser Effekt ist insbesondere bei Multimedia-Anwendungen im Internet (zum Beispiel Video-Streaming und IP-Telefonie) störend, da dadurch Pakete zu spät eintreffen können, um noch zeitgerecht mit ausgegeben werden zu können. Dies wirkt sich wie eine erhöhte Paketverlustrate aus.

Der Effekt wird durch einen sog. Jitterbuffer reduziert, allerdings zum Preis von zusätzlichem Delay, was bei streaming-Anwendungen nichts ausmacht, bei Dialoganwendungen aber ebenfalls stört.

Jitter kann auch bei der Aufzeichnung von Fernsehsignalen auf Videorekordern entstehen. Hier wird der Jitter durch mechanische Toleranzen verursacht. Häufig erscheinen diese Fehler als ein Wackeln des Bildes.

Die Einheit, in der der Jitter durch Messinstrumente gemessen wird, ist die Zeit: die Schwankung der Periodendauer wird typischerweise in μs gemessen.

Jitter entsteht zum Beispiel durch Dämpfung des Hochfrequenzanteils eines Frequenzspektrums auf einer langen Kabelstrecke. Das nennt man dann deterministischen Jitter. Andere Ursachen sind Rauschen (zufälliger Jitter, random Jitter) oder Symbolübersprechen (datenabhängiger Jitter). Um die Qualität einer Übertragungsstrecke beurteilen zu können, versucht man oft, die Ursachen für den vorhandenen Jitter zu finden. Dazu versucht man, den Gesamtjitter in die oben erwähnten Anteile aufzuspalten. Dieses Vorgehen nennt man Jitterseparation.



Delay

Verzögerung des gesamten Streams (im Gegensatz zu Skew, wo die Verzögerung eines Streams gegenüber einem anderen gemessen wird). Audio und Video-Signal sind übertragungsbedingt verzögert. Müssen dadurch wieder synchronisiert werden.

Lösung: Frame Delay-Buffer hilft mehrere externe Sources zu synchronisieren bzw. mittels Sync-Generator Problem zu beheben.

Erklären Sie Abtastfrequenz (sampling frequency), sample size und Quantisierung. Wie lautet das Abtasttheorem? Erklären Sie die Zusammenhänge anhand der Bandbreite von ISDN. [5]

Sample Frequency / Abtastfrequenz

Gibt an, wie oft (mit welcher Frequenz pro Zeitintervall) eine Abtastung eines analogen Signals vorgenommen wird. Je höher der Wert, desto besser ist die Qualität des resultierenden digitalen Signals, da ein geringerer Informationsverlust besteht. Minimale Abtastrate zur verlustlosen Rekonstruktion des Originalsignals: Nyquistfrequenz $\rightarrow$ min. doppelte im Signal höchste vorkommende Frequenz ($2 \times f_{\text{max}}$)

Der Abstand zwischen den Abtastzeitpunkten ist das Abtastintervall. Ist er konstant, heißt die Abtastrate auch Abtastfrequenz oder Samplingfrequenz. Die Wahl einer konstanten Abtastfrequenz vereinfacht die Weiterverarbeitung des Signals.

Sample Size

Die Samplesize bestimmt die Aufteilung der Amplitudenskala bei der Digitalwandlung eines analogen Signals und wird in Bits angegeben. Man kann es also als Speicherplatz, der pro abgetasteten Wert verwendet wird verstehen. Daher gibt an, mit welcher Genauigkeit der Wert gespeichert wird (8bit, 16bit,...).

Ist beispielsweise die Samplesize 8bit, so kann der gesamte im analogen Signal vorkommende dynamische Bereich in 256 diskrete Teile unterteilt werden. Bei der Sampling wird die Amplitude des

analogen Signals durch den nächstgelegenen diskreten Wert (aus den 256 im Falle von 8bit) angenähert. Somit erhält man ein amplitudendiskretes Signal.

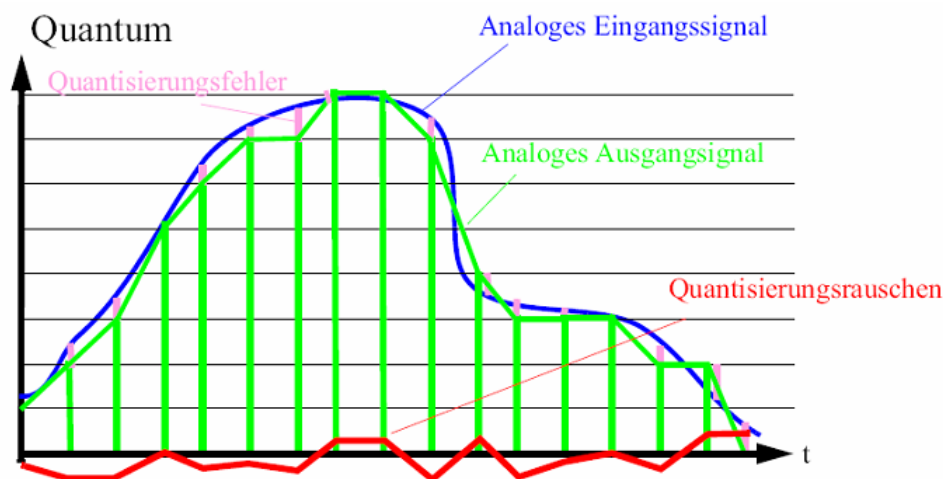
Quantisierung

Näherung des analogen Signals an diskrete Funktionen und Werte. Hierbei werden den abgetasteten Spannungswerten diskrete Zahlenwerte zugeordnet. Der Gesamtspannungsbereich wird in so genannte Quantisierungsintervalle Q unterteilt und der abgetastete Wert dem nächstgelegenen diskreten Zahlenwert zugeordnet. Hierbei entstehen Rundungsfehler, die sich im so genannten Quantisierungsrauschen bemerkbar machen. Sind die Quantisierungsintervalle von konstanter Größe, spricht man von linearer Quantisierung (übliche Methode in der Audiotechnik).

Um eine möglichst gute Qualität zu erreichen soll Q möglichst klein sein, es ergibt sich aus der Länge des Datenwortes. 8bit (256 Intervalle): $Q = V_{pp}/(2^8) = 4 \cdot 10^{-3} V_{pp}$

Man unterscheidet zwischen Linearer Quantisierung (Q konstant, üblich in Audiotechnik) und nichtlinearer Quantisierung (Q von unterschiedlicher Größe $\rightarrow$ kleine Werte: kleiner Quantisierungsfehler, große Werte: größer Quantisierungsfehler)

Bei der nichtlinearen Quantisierung findet eine Unterteilung in unterschiedlich große Intervalle statt. Hierbei gilt, je kleiner die Werte, desto kleiner die Intervalle und umgekehrt. Diese Quantisierung ist somit aufwändiger zu berechnen, erreicht aber bei geringeren Abtastfrequenzen die bessere Qualität und wird daher in Systemen mit geringer Übertragungsbandbreite eingesetzt.



Abtasttheorem nach Nyquist-Shannon

Besagt, dass Signale über der Nyquist-Frequenz (= Hälfte der verwendeten Sampling-Frequenz) nicht mehr fehlerfrei (ohne Informationsverlust) dargestellt werden können. Daher soll man die Sampling-Frequenz auf das doppelte der max. Frequenz des Eingangssignals festlegen, um Artefakte im gesampleten Signal zu vermeiden (Abtastrate = $2 \cdot$ höchste Frequenz im Signal).

ISDN (Integrated Services Digital Network)

Point-to-Point Verbindung für Sprache und Daten. Die Sprache wird umgewandelt in elektrische Signale (AD-Wandler). Die elektrischen Signale bewegen sich in einem Raster mit 256 verschiedenen Zuständen (8bit). Das Signal wird 8000 Mal pro Sekunde abgetastet (Abtastrate 8000Hz $\Rightarrow$ 8kHz). Bitauflösung [bit] x Abtastfrequenz [Hz] = Datenrate [bit/s] $\Rightarrow 8 \times 8000 = 64\,000$ bit/sec (64 kbit/sec).

Bei ISDN bekommt man normalerweise 2-B und 1-D Kanal. Der D-Kanal wird für die Signalisierung/ Steuerkanal (Verbindungsaufbau und -abbau, Rückruf-bei-Besetzt usw.) verwendet mit 16kbit/s.

Beim normalen Telefon werden diese im normalen Kanal übertragen (deshalb kann man das „Wählen“ dort hören). B-Kanal ist Nutzsignal á 64kbit/s.

Erklären Sie kurz die Begriffe “component video” und “composite video”. Was sind Vor- und Nachteile? Welche Rolle spielen “component video” und “composite video” für digitale Videoformate? [3]

Beide dienen der Video Repräsentation

Component Video

Digital Component Video (D1/D5, SMPTE RP125): Trennung der Signale in Helligkeit und Farbe, 27MB/s Datenrate, subsampled Farbsignale 4:2:2.

Composite Video:

Digital Composite Video (D2/D3, SMPTE 244M): 14.3MB/s Datenrate, subsampled Farbsignale 4:2:2.

Component Video

- + mehrere Signale
- benötigt mehr Bandbreite
- Synchronisation aufwändiger
- Senden und Verkabelung komplizierter
- + Jede Primärfarbe und Helligkeit eigenes Signal
- + Bessere Qualität
- + Wird im Profi Bereich z.B. Editing, Encoding,.. verwendet
- + Jedes Signal enthält RGB + luma/ chroma Info

Composite Video

- in einem Signal gemischt
- + einfache Handhabung
- + einfacher zu übertragen
- + einfachere Verkabelung
- Interferenzen können auftreten
- schlechtere Qualität
- ~ wird in Haushalten verwendet (VHS)
- + in einziger Trägerwelle übertragen

Erklären Sie kurz den Begriff Interlacing bzw. Interleaving und führen Sie aus, was er im Kontext der Datentypen Audio, Images und Video bedeutet [3] bzw. Erklären Sie kurz den Begriff Interleaving im Kontext des Datentyps Audio. Welche Vor- bzw. Nachteile haben interleaved Audioformate im Vergleich zu non-interleaved Formaten? [2]

Interlacing ist ein Zeilensprungverfahren und der „Content“ setzt sich in Durchgängen zusammen.

Interlacing Video

Anstatt ein Frame progressiv (daher zeilenweise) aufzubauen werden zwei Halbbilder erzeugt, die jeweils um eine Zeile verschoben sind und nacheinander angezeigt werden. Das geht deshalb, weil der Phosphor im TV-Gerät oder Monitor nachleuchtet und dieser Trick dem menschliche Auge daher nicht auffällt. Zugleich spart man damit gewaltig Bandbreite, hat aber den Nachteil, dass das De-Interlacing sehr aufwändig ist.

Interlacing Image

Als Interlace bzw. Interlacing bezeichnet man ein Speicherverfahren für Grafiken, bei dem das Bild in mehreren Schichten abgespeichert wird. Dies ermöglicht den schnellen Aufbau eines Übersichtsbildes, wenn die Grafik von einem langsamen Medium geladen wird. Die Darstellung wird dann während des Ladevorgangs immer feiner.

Das Bild wird nicht Zeile für Zeile aufgebaut, sondern in mehreren Durchläufen. Beim ersten Durchlauf wird eine bestimmte Anzahl von Zeilen übersprungen, in den nächsten Durchläufen werden die übersprungenen Zeilen dann fortlaufend dargestellt. Sichtbar ist dieser Vorgang insofern, als dass das Bild anfangs unscharf ist und dann mit der Zeit immer „schärfer“ wird.

Interleaving

Der Begriff Verschränkung oder englisch Interleaving (engl. to interleave „verschachteln, überlappen“) bezeichnet den Prozess, mehrere linear durchzählbare Objekte in einer speziellen Reihenfolge anzuordnen.

Interlacing in dem Sinn hat an sich nichts mit der RGB-Codierung zu tun. Das wäre dann wieder Interleaving. Man könnte zwar die drei Farbkanäle in einem File separat speichern, das hätte jedoch Nachteile, da man das File immer komplett geladen haben muss, bevor man es richtig anzeigen kann. Vom Prinzip her würde das dann (vereinfacht) so aussehen:
RRRRRRR...GGGGGG...BBBBBBB...

Beim Interleaving werden linear durchgezählte Objekte (z.B. Video und Audio oder RGB) in eine verschachtelte Form gebracht, um Burstfehler zu vermeiden, indem diese durch das Interleaving zu Einzelbitfehler, die leicht zu korrigieren sind, ausgemerzt werden.

Im Interleaved-Verfahren werden für jeden einzelnen Bildpunkt oder zumindest Bildbereich die einzelnen Kanäle gleich hintereinander gespeichert, also RGBRGBRGB - damit ist es möglich die Bildpunkte sofort auszugeben.

Ohne Interleaving

Fehlerfreie Übertragung: aaaabbbbccccdddeeeeffffgggg

Übertragung mit einem Burstfehler: aaaabbbbcccc____deeeeffffgggg

Nun die gleichen Daten mit Interleaving:

Fehlerfreie Übertragung: abcdefgabdefgabdefgabdefg

Übertragung mit einem Burstfehler: abcdefgabcd____bcdefgabdefg

De- Interleaved Übertragung mit einem Burstfehler

aa_abbbbccccddde_eef_ffg_gg

Jetzt fehlen zwar von a, e, f und g je ein Bit, aber das kann korrigiert werden, weil jeweils nur ein Bit und nicht die ganzen Sequenzen cccc und dddd verloren sind.

Vorteile: Gegen Burstfehler abgesichert. Brauchen für Sicherung keine zusätzlichen Daten.

Nachteile: Latenzzeit erhöht sich.

Interleaving Video

Anstatt die drei Farbkanäle in einem File separat zu speichern (RRRRRRR...GGGGGGGGG...BBBBBBBBBBB...), werden diese interleaved für jeden einzelnen Bildpunkt oder zumindest Bildbereich in die einzelnen Kanäle gleich hintereinander gespeichert (RGBRGBRGBRGB...). Somit muss das File nicht immer komplett neu geladen werden, damit es richtig angezeigt werden kann, sondern die Bildpunkte können sofort ausgegeben werden.

Interleaving Audio

Wird verwendet um die Audiospur mit der Bilderfolge zu einem Stream zu verbinden. Nach einer festgelegten Anzahl Frames wird ein Stück des Audiostreams gespeichert, der die Tonspur der folgenden Frames enthält. Der Abstand zwischen den Audioteilen darf nicht zu groß sein, da sonst das Abspielen ohne Aussetzer nicht gewährleistet ist. Je enger der Abstand ist, desto mehr Systemressourcen werden aber in Anspruch genommen, da der Computer die für das Abspielen des Videos zuständigen Geräte, sehr oft mit eher kleinen Datenmengen beliefern muss. D.h. also, dass Audio und Video in abwechselnden Blöcken gespeichert werden.

Vorteile: Einfache Synchronisation von Audio/Video, weniger Streams.
Nachteile: siehe oben (Abstand)

Interleaving Images

Pixel Interleaving, Raster Interleaving und Frame Interleaving. Die unterschiedlichen Interleaving Schemen mit 4 Reihen und 5 Spalten würde so aussehen:

Pixel interleaved looks like:

RGBRGBRGBRGBRGB	row 1
RGBRGBRGBRGBRGB	row 2
RGBRGBRGBRGBRGB	row 3
RGBRGBRGBRGBRGB	row 4

Raster interleaved looks like:

RRRRRGGGGGBBBBB	row 1
RRRRRGGGGGBBBBB	row 2
RRRRRGGGGGBBBBB	row 3
RRRRRGGGGGBBBBB	row 4

Frame interleaved looks like:

RRRRRRRRRRRRRRRRRRRR	all the red
GGGGGGGGGGGGGGGGGGGG	all the green
BBBBBBBBBBBBBBBBBBBB	all the blue

Nennen Sie 2 Gründe weshalb Video - als einziger Medientyp - in MM-Applikationen auch in analoger Form eingesetzt wird. [2]

- Analoges Video ist kostengünstiger in Bezug auf Speicherplatz und Übertragung (mehrere GB große Videos).
- Die Digitalisierung von Videos benötigt immer noch viel Zeit.
- Analoges Video ist einfacher zu senden.

Erklären Sie kurz die Begriffe chroma keying und navigational video. [2] bzw. Erklären Sie kurz die Begriffe Chroma Keying und Luminance keying. [2]

Luminanz Key

Der Luminanz Key hat seinen Namen aus dem Videosignalbereich in welchem Farbe und Helligkeit getrennt übertragen werden. Man separiert aufgrund von Helligkeitswerten. Bei diesem Verfahren werden die Bildteile eines Videos transparent, die einer bestimmten Helligkeit entsprechen. An diesen Stellen wird ein darunter gelegtes zweites Video/ Bild sichtbar. Das Verfahren ist eher für Effekte interessant. Kann weniger kontrolliert werden, wird aber oft bei Graphikelementen verwendet.

Der Luma Keyer errechnet aus einem farbigen RGB-Bild ein monochromes Graustufenbild. An dieser Stelle kann das Mischverhältnis der drei Farbkanäle justiert werden um Vordergrund und Hintergrund kontraststark zu trennen. Die Stanzmaske wird folglich aufgrund starker Helligkeitsdifferenzen innerhalb des Graustufenbildes erzielt. Um ein gutes Ergebnis zu erreichen sollte der zu keyende Gegenstand sehr hell bis weiß gegenüber einem sehr dunklen oder sogar schwarzen Hintergrund sein; Oder umgekehrt.

Chroma Keying:

Auch der Chroma Key hat seinen Namen vom getrennten Videosignal. Bei Chroma Keying nimmt man als Hintergrund eine bestimmte Farbe (man verwendet gerne blau – weil es eine Komplementärfarbe zu den Hautfarben darstellt - könnte im Prinzip aber jede Farbe sein). Damit kann man dann das oder die Objekte leicht vom Hintergrund separieren und über jeden anderen Hintergrund/ Video darüberlegen. Es wird somit eine Maske (Alpha Kanal) gemacht die besagt, welche Bereiche transparent sind.

Ein Mixer ersetzt eine Farbe in einem Signal, durch Daten eines anderen Signals (z.B.: Blue-Box-Effekt). Im Videobild werden Teile einer ausgewählten Farbe von einer anderen Videosequenz ersetzt. Auch bekannt als Bluebox-Effekt bzw. Bluescreen-Technik für den Austausch von Bild-Hintergründen.

Das gleiche kann man auch genauso im Luminance-Bereich anwenden, da separiert man dann aufgrund von Helligkeitswerten. Durch die Beschränkung auf "eine" Farbe ist diese Methode selektiver und flexibler als der Luma Key. Intern wechselt der Keyer vom RGB Farbraum in HSV (Hue, Saturation, Value), weil er auf Farbnuancen, Sättigungsgrade und die Helligkeit angewiesen ist, was nicht zu den Attributen von RGB zählt. Daher gekeyed wird in HSV, wobei die Justierung ebenso in RGB, YUV, RGBCMYL oder den separaten Kanälen (rot, grün, blau) stattfinden kann.

Navigational Video

Erstellt eine „virtuelle“ Kamerafahrt aus Einzelbildern. Hierbei wird zuerst ein Objekt aus möglichst vielen Richtungen in geringem Abstand zueinander aufgenommen (mit Winkelabständen von beispielsweise 5°). Nun kann man diese Einzelbilder beliebig aneinander reihen und so den Eindruck einer bewegten Kamera erzeugen und einen realistischen Rundumblick bieten (wenn man sie nicht linear abspielt!). Man erhält den Eindruck einer Kamerafahrt und 360°-Ansicht (Panorama).

Erklären Sie kurz den Begriff Multimediasystem im engeren und weiteren Sinne. [2.5]

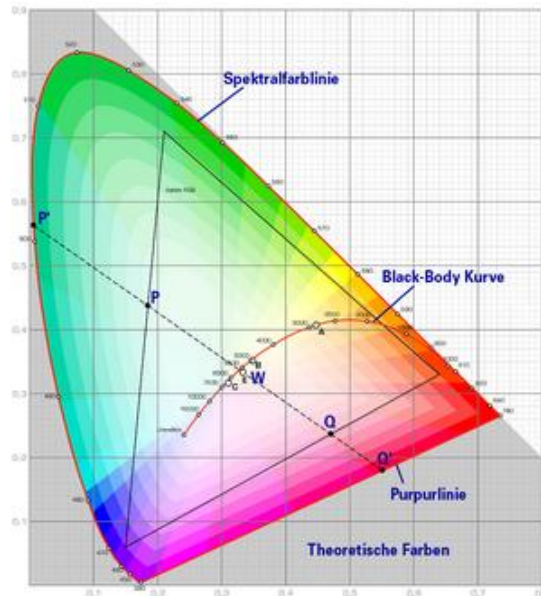
Ein Multimediasystem ist charakterisiert durch computergesteuerte, integrierte Produktion, Manipulation, Präsentation, Speicherung und Kommunikation unabhängiger Information, welche zumindest durch ein kontinuierliches (zeitabhängiges) und ein diskretes (zeitunabhängiges) Medium kodiert ist.

Im weiteren Sinne wird es verwendet um die Verarbeitung individueller Bilder und Texte zu beschreiben, obwohl kein kontinuierliches Medium vorhanden ist.

Erklären Sie kurz den CIE Farbraum. [2]

Der CIE Farbraum wird verwendet um die anderen Farbräume (CMYK, RGB, usw.) zu kalibrieren, da alle in diesem Modell enthalten sind.

Darstellung von Farben beruhend auf dem menschlichen Wahrnehmungsapparat. Additive Farbmischung. Grundlage ist ein sogenannter Normalbeobachter mit einem Sichtfeld von 2° (CIE1931) bzw. 10°(CIE1964). Normfarbwerte X (virtuelles rot), Y (virtuelles grün), Z (virtuelles blau)



Darstellung: Querschnitt durch den Raum; alle Farben befinden sich innerhalb der Fläche.

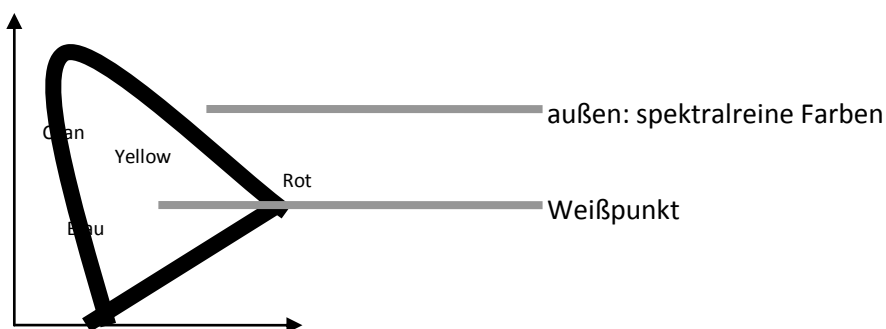
Wellenlänge: gebogene, äußere Linie.

Purpurgerade: gerade Verbindungslinie zwischen Rot und Blau.

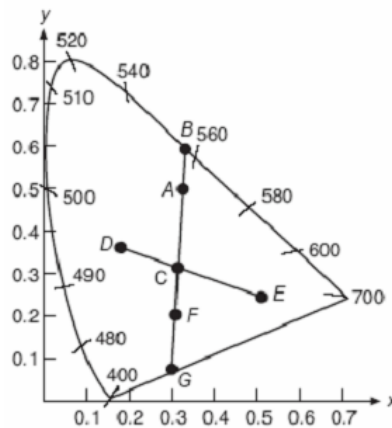
Spektralfarben: höchster Sättigungsgrad; am äußeren Kurvenrand, reines Weiß in Zentrumsnähe (C)

Mittelpunktvalenz: $x = 0,33$, $y = 0,33$ (C)

Wenn A die betrachtete Farbe ist, so wird sie durch die Wellenlänge B dominiert. Wird B mit C gemischt, so erhält man A. Die in F dominierende Wellenlänge G ist definiert als das Komplement der dominierenden Wellenlänge B in A. Werden Komplementärfarben gemischt (E mit D, F mit A), so erhält man weiß. Deshalb kann CIE benutzt werden um Farben zu identifizieren und analysieren, oder um die Leistung eines CRT zu messen (wie viel des CIE Farbraums kann er darstellen)



Auf der äußeren Kurve liegen die reinen Spektralfarben, ca. in der Mitte liegt Weiß(C). Wenn eine Farbe A im Raum zwischen Weiß und der Spektralkurve liegt so bekommt man die dominante Wellenlänge, indem man die Gerade, die durch C und A geht bis zur Kurve verlängert. Der Schnittpunkt zeigt dann die dominante Farbe an. Würde der Schnittpunkt auf der Geraden liegen die den Raum abschließt, so wird die dominante Wellenlänge als das Komplement von A genommen.



Dominante Wellenlänge: zwischen Farbe und Weißpunkt Gerade ziehen und bis Spektralkurve verlängern.

Komplementärfarbe: liegt gespiegelt zum Weißpunkt auf selben Gerade.

Verwendung:

- zur Kalibrierung von anderen Farbmodellen (RGB, CMYK,...).
- zur Identifikation und Analyse von Farben.
- zur Messung für CRT Leisten („Performancemessung“), sprich wie viel vom CIE Farbraum. wird durch CRT Abgedeckt.

Diskutieren Sie kurz 3 Kriterien ihrer Wahl zum Vergleich von Videokomprimierungsverfahren. [3]

Verlustfreie oder verlustreiche Komprimierung

Ist die Bildqualität ausschlaggebend und eine hohe Qualität benötigt, ist es notwendig, auf verlustfreie Kompressionsverfahren zurückzugreifen. Dem stehen allerdings eine niedrigere Kompression und damit eine größere Größe gegenüber. Optimal wäre eine hohe Qualität bei höchstmöglicher Kompression.

Echtzeit Komprimierung

Zeitaufwand für En- und Decodierung. In vielen MM-Anwendungen ist es notwendig, dass die Verarbeitung von Videos in Echtzeit möglich ist (Videoschnitt, aber auch Videokonferenzen,...). Je komplexer ein Verfahren ist, desto schwieriger ist dies zu ermöglichen. Daher muss ein Kompromiss zwischen Qualität und Zeitaufwand gefunden werden.

Editierbarkeit, Navigation: Wie einfach und gut ist bzw. bleibt das komprimierte Video editierbar? Aber auch, ist ein fast forward, fast backward,... möglich?

Datenrate

Ist die Datenrate hoch, werden hohe Bandbreiten zur Übertragung benötigt, allerdings ist die Qualität auch viel besser. MPEG-1 ist für niedrige Datenraten ausgelegt, MPEG-2 hingegen benötigt hohe Datenrate, hat aber eine hohe Qualität (bis 15Mbit/sec)

Interframe- oder Intraframe-Kompression

Interframe-Kompression erstellt alle Bilder (ausgenommen Keyframes) aufbauend auf das Vorgängerbild. Intraframe komprimiert nur Einzelbilder.

Erklären Sie kurz Schalldruckpegel und Lautstärkepegel. Wozu dienen diese Werte? [3]

Schalldruckpegel

$L = 20 \cdot \log(p/p_0)$ [dB] mit Bezugsschalldruck p_0 (1kHz, gerade noch hörbar $20\mu\text{Pa}$). Die Einheit des Schalldrucks ist Pa und die des Schalldruckpegels dB. Er dient zur wahrnehmungsgemäßen Beschreibung von Lautstärken und wird in Dezibel gemessen.

Der Schalldruckpegel dient zur sinnvollen Darstellung und Beschreibung des Hörbereichs, denn durch die logarithmische Darstellung kann die Spanne von der Hörschwelle bis zur Schmerzgrenze (1: 2 Millionen) in einem kleineren Wertebereich dargestellt werden (0 – 120dB). Außerdem kommt die logarithmische Darstellung dem Menschen insofern zu Gute, weil sich durch Experimente herausstellte, dass das Empfinden von Lautstärkedifferenzen einem logarithmischen Gesetz folgt und somit diese Darstellung dem Empfinden mehr gerecht wird als die Angabe eines absoluten Druckamplitudenwertes.

Somit sind 10 Geigen, die gleichzeitig spielen, doppelt so laut wie nur eine Geige. dB Angaben werden benötigt um einen gigantischen Hörbereich in einen sinnvollen kleineren Wertebereich einzuschränken.

Lautstärkepegel

Der Lautstärkepegel ist ein Vergleichsmaß. Er beschreibt, welchen Schalldruckpegel ein Sinuston mit einer Frequenz von 1000 Hz haben müsste, damit dieser genauso laut empfunden wird wie der betrachtete Schall. Bei einer Schall-Frequenz von 1000 Hz stimmen Schalldruckpegel, gemessen in dB SPL (sound pressure level), und Lautstärkepegel, gemessen in Phon, überein. Für Sinustöne anderer Frequenzen sowie für komplexe Schallereignisse sind dagegen andere Schalldruckpegel erforderlich, um den gleichen Lautstärkeeindruck zu erzielen. Welcher Schalldruckpegel für einen Einzelton bei welcher Frequenz erforderlich ist, um jeweils den gleichen Lautstärkeeindruck zu erzielen, ist in den "Kurven gleicher Lautstärkepegel" (Isophone) beschrieben. Z. B.: Ton mit 100Hz und Schalldruckpegel $L=60\text{dB} \rightarrow 50$ Phon.

Die wahrgenommene Lautstärke ist also nicht nur abhängig vom Schalldruckpegel, sondern auch von der Frequenz. Um so alle Töne, die für uns gleich laut klingen, mit einem gleichen Pegelwert beschreiben zu können, wurde der Lautstärkepegel LN mit der Einheit Phon definiert.

Beschreiben Sie kurz das Prinzip von MIDI und erklären Sie die 2 MIDI-Konzepte Clock und Palette [2.5] bzw. Beschreiben Sie kurz das Prinzip von MIDI und erklären Sie die 2 MIDI-Konzepte Channel und Voice [2.5]

Das Musical Instrument Digital Interface ist ein Protokoll, das die Kommunikation von Computern, Synthesizern, Keyboards und anderen Musikinstrumenten ermöglicht, es stellt sozusagen die Schnittstelle zwischen den Devices dar, über die die Kommunikation untereinander möglich ist.

Clock

- gibt die Taktrate vor, die am MIDI-Gerät eingestellt werden kann.
- Maß: Teile pro Viertelnote
- versieht die Messages mit Timestamps
- Tempo: bpm

Palette

Aufgrund gespeicherter Werte aus einer Palette werden die entsprechenden gespeicherten Sounds wiedergegeben.

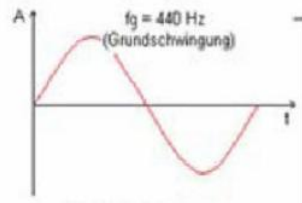
Voice

- ist der Teil des Synthesizers, welcher den Sound produziert.

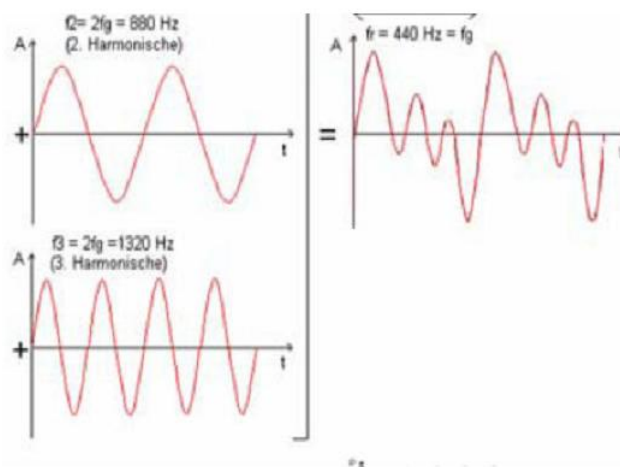
- Synthesizer können mehrere (16, 20, 24, 32, 64,...) voices haben, wobei jede unabhängig und simultan arbeitet um Sound mit verschiedenen Tonlagen und Klangfarben zu produzieren.

Erklären Sie kurz die Begriffe Ton, Klang und Geräusch [1.5]

Als Ton wird eine einzelne Sinusschwingung bezeichnet. $p(t)$ ist eine Sinusschwingung (mit einer bestimmten Frequenz – $1/t$)



Klang entsteht durch Überlagerung eines Grundtons (hörbar als Tonhöhe) und Obertönen ($f = n \cdot f_g$) (hörbar als Klangfarbe) und stellt eine allgemeine periodische Funktion $p(t)$ mit einer Frequenz f_g dar.



Bei einem Geräusch ist kein Grundton mehr erkennbar und nur noch eine Ansammlung von verschiedenen Frequenzen vorliegt. Aperiodische Schwingung (unregelmäßige Schwingung).

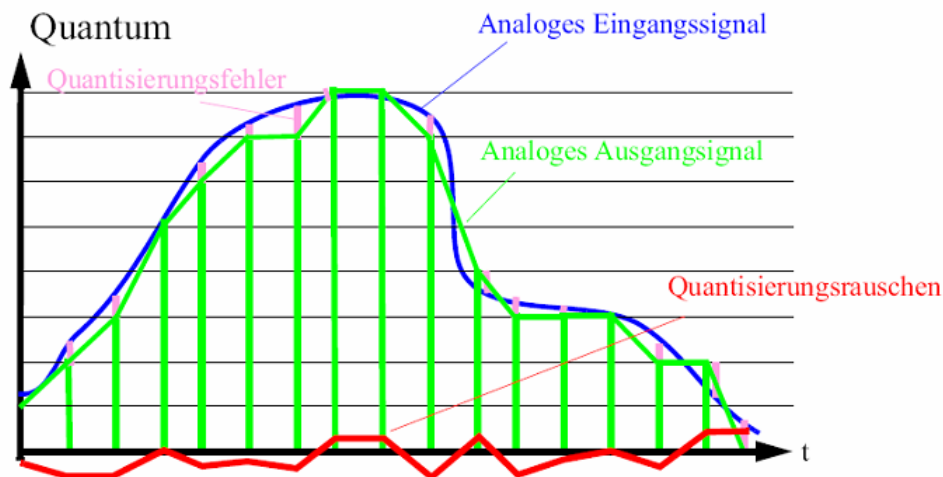


Was ist und wie entsteht Quantisierungsrauschen? Auf welche Weise kann es reduziert werden? [2]

Das Quantisierungsrauschen ist die Folge des Quantisierungsfehlers. Dieser entsteht bei der Quantisierung analoger Signale: Das analoge Signal wird mit einer bestimmten Frequenz abgetastet und der gemessene Amplitudenwert wird auf den nächsten diskreten Wert gerundet (AD-Wandlung). Die Differenz zwischen tatsächlichem analogen Wert und quantisiertem, diskretem Wert bezeichnet man als Quantisierungsfehler, welcher kleiner gleich $0,5 Q$ (Quantisierungsintervall) ist. Wird das diskrete Signal wieder in ein analoges gewandelt, entsteht bei dessen Wiedergabe das Quantisierungsrauschen, welches die akustische Auswirkung des Quantisierungsfehlers ist und den Schalldruckpegel L_{noise} erzeugt.

Reduktion

Das Rauschen kann zum Einen durch eine erhöhte Auflösung beim Quantisieren (höhere Sampleauflösung = Anzahl der Quantisierungspunkte) reduziert werden und andererseits bei der linearen Quantisierung, sofern der Signal-Rausch-Abstand groß genug ist, ausmaskiert werden.



Welchen Zweck haben Distribution Amplifiers und Timebase Correctors? [1.5]

Distribution Amplifier

Ein Distribution Amplifier ist ein NF-Verstärker mit mehreren Ausgängen. Die Ausgänge haben alle das gleiche Ausgangssignal und können an unterschiedliches Equipment angeschlossen werden. Funktional handelt es sich um eine Signalverteilung mit Verstärker (gibt ein Eingangssignal verstärkt an mehreren Ausgängen aus).

Timebase Correctors

Baut Signal wieder auf („Rekonstruiert“), um die Timing-Störungen zu entfernen, die durch VTR/VCR (video cassette recorder) entstehen, zu eliminieren.

Erklären Sie kurz 2 Arten von Video Time Codes Ihrer Wahl. [1.5]

Time Codes notierten auf Band:

- LTC: Longitudinal Time Code: Zeitcode meist im Audiotrack aufgezeichnet; entlang des Bandes
- VITC: Vertical interval TC; in den vertikalen Austastlücken aufgezeichnet; im Nachhinein nicht mehr korrigierbar; kann nur gleichzeitig mit dem Video aufgezeichnet werden

SMPTE time code (HH:MM:SS:FF)

- non-drop frame time code – FF in [0, 29]
- drop frame time code FF in [2, 29], außer bei jeden 10. Minute [10, 29]

Was ist ein Frequenzspektrum? Wie unterscheiden sich die Frequenzspektren von Klängen und Geräuschen? [3]

Frequenzspektrum

Jede periodische Schwingung beliebiger Wellenform (= Klang) kann als eine Überlagerung von Grundschwingung und Oberschwingungen dargestellt werden. Trägt man in einem Diagramm die Amplituden aller beteiligten Schwingungen in Abhängigkeit der Frequenz auf, so erhält man das sogenannte Frequenzspektrum. Diskret/kontinuierlich.

Diskretes Spektrum

Klänge enthalten nur Schwingungen, deren Frequenzen ein ganzzahliges Verhältnis zur Grundfrequenz aufweisen. Kontinuierliches Spektrum: Allgemeine Schallereignisse (Geräusche)

enthalten eine unendliche Anzahl von Einzelschwingungen, deren Frequenzwerte sich kontinuierlich über die x-Achse erstrecken. Es entsteht eine kontinuierliche mathematische Funktion: $p=p(f)$

Was ist eine Color Look-up Table? Weshalb wird sie eingesetzt? [2]

Die LookUp Table (LUT) ist eine hardwaremäßig geschaltete (zwischen Datenspeicher und Darstellungsgerät) Liste von Wertepaaren, die jedem Grauwert r einen neuen Grauwert $s = T(r)$ zuordnet. Sie wird daher bei Punktoperatoren (Grauwertabbildungen) genutzt.

Größe: Anzahl der Farbwerte

Kleiner Bereich genützt: ineffizienter als direkte Indizierung (schlechte Speicherplatzausnutzung) (Speicherung des Farbwertes direkt ohne Tabelle)

Quantisierung – fein –grob

Fein: besser, kleinerer Quantisierungsfehler, genauer, besonders wichtig, wenn sich der Funktionswert nur langsam ändert. In diesen Fällen werden räumlich große Teile auf denselben digitalen (Grau)wert quantisiert und die plötzlichen Änderungen im digitalen Wert (= Sprung von einer Quantisierungsstufe zur nächsten), die an zufälligen Stellen auftreten, neigen dazu, falsche Konturen erscheinen zu lassen und somit Scheinobjekte zu definieren.

Nachteil: höherer Speicherbedarf, höherer Rechenaufwand

Rastern -> fein –grob

Abtasten/Quantisieren

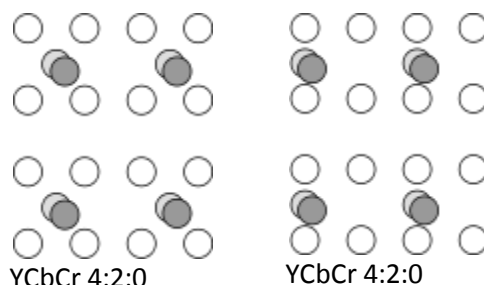
Abtastfrequenz

Frequenz, wie oft ein Signal abgetastet wird. Alle $1/f$ Zeitpunkte wird der Signalwert abgetastet.

Was ist die Besonderheit des Samplings für digitale 4:2:0 Videoformate? [2]

Das Luminanzsignal wird mit 13,5 MHz abgetastet. Bei der Abtastung der Farbdifferenzsignale macht man sich zu nutze, dass das menschliche Auge in vertikaler Richtung eine erheblich geringeres Auflösungsvermögen besitzt als in der Horizontalen. Abwechselnd wird eine Bildzeile im Verhältnis 4:2:2 und die darauffolgende Zeile im Verhältnis 4:0:0 digitalisiert. Zeilen, für die keine Farbinformation aufgezeichnet wurde, übernehmen das Farbsignal der darüber liegenden Zeile.

Oder einfacher ausgedrückt: Hierbei wird das Farbsignal in jeder zweiten Zeile überhaupt nicht abgetastet, und in jeder ersten nur halb so oft wie das Luminanzsignal, was die Datenrate im Vergleich zur 4:2:2-Abtastung noch ein Mal deutlich reduziert.



Die 4:2:0 Abtastung wird im PAL-DV Standard verwendet. Daher benötigt Digital Video trotz einer 5:1 Kompression weniger Speicherplatz (ca. 3,5 MB/sec) als das ähnliche MJPEG-Verfahren. Dieses benötigt bei gleicher Kompression mit 4:2:2 Abtastung ca. 4,4 MB/sec bei PAL-Vollbild (786 x 576 Pixel).

Plattformen

In welchem Kontext spielen Real-Time-Video und Production-Level-Video eine Rolle und was unterscheidet die beiden Formate? [3]

Real-Time-Video und Production-Level-Video sind 2 Varianten der Video-Kompression im Rahmen von DVI (Digital Video Interactive).

Production Level (PLV) bzw. Presentation Level Video

Production Level Video (PLV) repräsentiert eine hohe Qualitätsstufe unter Digital Video Interactive (DVI). Bei PLV werden die Videosequenzen offline von einem Superrechner komprimiert und können dann mit Personal Computer (PC) und DVI-Karte wiedergegeben werden.

- Mit diesem Verfahren werden Kompressionen von 150:1 erreicht. Braucht lange zum Codieren, decodieren dann in Echtzeit.
- Nur I, P-Bilder (keine B-Bilder)
- Luminanz 256x240
- 24bit (true color), mit Zeilen/Spalten-Verdopplung und Interpolation: 512x480;
- Kompressionsalgorithmus nicht freigegeben
- Kompressionszeit/Dekompressionszeit 100:1!

Vorteile

- Hohe Qualität
- Decodieren ist in Echtzeit möglich.
- Kompression viel besser (wird daher eingesetzt, wenn ein Video nur ein Mal komprimiert und oft abgespielt wird (z.B. wenn es als Film verkauft wird)).

Nachteile

- **Lange Übertragungszeiten**
- Hoher Rechenaufwand, daher nicht in Echtzeit komprimiert

Real-Time-Video (RTV)

Kompressionsformat, mit dem Videosignale in Echtzeit digitalisiert, komprimiert, codiert, decodiert und auf der Festplatte gespeichert werden können. Real Time Video (RTV) repräsentiert die einfache Qualitätsstufe unter Digital Video Interactive (DVI). Bei RTV können die Videosequenzen (z.B. bei Videokonferenzen) in Echtzeit komprimiert und mit "Video for Windows" dargestellt werden.

- 128x120, mit Zeilen/Spalten-Verdopplung und Interpolation 256x240
- RTV- Intraframe-Kodierung: Differenz zwischen aktuellem und darüberliegendem Pixel => viele Nullen.
- RTV-Interframe-Kodierung: Differenz von Pixeln gleicher Position in aufeinanderfolgenden Bildern => viele Nullen.
- anschließend RLE+ Huffman
- einfaches, aber schnelles Verfahren
- symmetrisch
- Kompression in Echtzeit

Vorteile

- Multimediaapplikationen zunächst in RTV (= ELV = Edit Level Video), nachfolgend PLV
- codieren und decodieren in Echtzeit. Wichtig bei Videokonferenzen.

Nachteile

- Schlechtere Qualität

- Erzielt nicht so eine gute Kompression (da in Echtzeit codiert und decodiert werden muss)

Was sind die technologischen Unterschiede zwischen Audio-CD und CD-ROM? [2] bzw. Erklären Sie kurz die Unterschiede zwischen CD-Audio und CD-ROM. [1.5]

CD-ROM

- zusätzlicher Layer („Schicht“) für Fehlerdetektion und Fehlerkorrektur
- Unterteilung des Speicherplatzes in aufeinander folgende Datenblöcke (98 frames, 2352 bytes)
- 2 Blockformate
 - Mode1 – 2048 bytes Anwendungsdaten
 - Mode2 – 2333 bytes Anwendungsdaten

CD-Audio

- verwendet nur CLV (constant linear velocity)
- frame mit 588 channel bits, 200 bits nach dem Dekodieren (je AudioCD 20MB nicht nutzbar)
- Enthält nur ausschließlich Audio Daten und keine Multimedialen oder textuellen Zusatzinformationen.
- Audio Daten werden unkomprimiert gespeichert.

Erklären Sie kurz die QuickTime Medienorganisation auf dem konzeptionellen Level. [4] bzw. Erklären Sie kurz die konzeptionelle Medienorganisation von QuickTime. [3]

- **Data entity:** Repräsentiert den tatsächlichen Speicher der für die a/v Daten benutzt wird.
- **Media entity:** ist eine zeitliche Sequenz, in media time gemessen. Jede Entity hat einen Medientyp (entweder a (audio) oder v (video)). Die Elemente referenzieren Speicherregionen.
- **Track entity:** beschreibt die Anordnung von Media Entitäten. Jede Entitäten ist in einer Movie Entity enthalten. Zeit wird in Movie TCS gemessen.
- **Movie entity:** Ist eine Gruppe von Track Entities. Spezifiziert die Zeitskalierung und Dauer. Es beschreibt also die Anordnung, Dauer, Verzögerung („Offset“) von Tracks innerhalb des Movies.

Erklären Sie kurz das Zeitkonzept in QuickTime. [2]

- Zeit ist explizit
- Werte werden als Integer ohne Vorzeichen (Unsigned Integers) repräsentiert
- Die Zeitskalierung wird in units per second angegeben
- Die Dauer wird mit dem Maximalen Zeitwert angegeben.
- Die Timebase („Zeitbasis“) gibt die Playbackrate an.
- Verwendung eines Time Coordinate Systems (TCS) für Skalierung und Dauer

Worin unterscheiden sich die 2 Formate von Videodiscs? Listen Sie kurz Vor- und Nachteile dieser Formate. Welches Format war Vorbild für die CD-Audio? [2.5]

Die Videodisc steht am Scheitelpunkt der digitalen und analogen Medien, ist aber selber ein digitales Medium. Alle Videodiscs sind flache Scheiben und vom Durchmesser mit einer Langspielplatte zu vergleichen. In den siebziger Jahren konkurrierten verschiedene Videodisc-Systeme miteinander. Das Problem aller Systeme bestand darin, gewaltige Informationsmengen kodiert auf relativ begrenzten Flächen unterzubringen.

Das System, das am meisten verbreitet war, ist LV (Laser Vision). Dabei wurde ein Laser verwendet, um kleine Pits (Vertiefungen) in die Oberfläche der einer Master-Disc zu brennen. Die Pits waren in einer langen Spirale, die in der Mitte startete, angeordnet. Von der Master-Disc wurde ein Abdruck gemacht, und aus diesem wurden die Discs für den Verkauf gegossen. Gelesen wurden die Videodiscs von einem Laser, der die Pits abtastete und dadurch in binäre Signale umwandelte. Logischerweise waren diese Discs Read-Only.

LV – Discs werden in 2 Formaten angeboten:

CAV (constant angular velocity)

Angewendet wird dieses Format bei DVD- RAM und ist auch Vorbild für CD-DA. Schallplatten, Laserdisc (LD),...

- konstante Winkelgeschwindigkeit (Platte dreht sich mit konstanter Geschwindigkeit)
- Je weiter der Lesekopf nach außen bewegt wird, desto höher ist die relative Geschwindigkeit zwischen Kopf und Plattenoberfläche
- Die Datenrate pro Umdrehung bleibt gleich
- Dadurch ist die Dichte der Daten nach außen hin immer geringer
- CAV Discs haben ein flexibles Playback: Pause (Still, Bild einfrieren), Play, Forward, Reverse in verschiedenen Geschwindigkeiten.

Vorteil: flexible Abspielmöglichkeiten (freeze frame, vorwärts oder rückwärts abspielen mit verschiedenen Raten) durch Datenanordnung

Nachteil: schlechte Nutzung der Kapazität

CLV (constant linear velocity)

Standard für CD-ROM, Audio Compact Disc (Audio- CD)

- konstante lineare Geschwindigkeit (Absenkung der Drehzahl nach außen hin)
- Je weiter der Lesekopf nach außen bewegt wird, desto mehr wird die Umdrehungsgeschwindigkeit gesenkt.
- Die relative Geschwindigkeit zwischen Lesekopf und Oberfläche bleibt annähernd gleich.
- Datendichte und Datenrate ebenfalls.

Vorteil: doppelte Kapazität (Speicherplatz) im Vergleich zu CAV

Nachteil:

- längere Suche und schlechtere Positionsmöglichkeit (nicht frameadressierbar)
- nicht so gute Playback-Funktionen
- schlechtere Zugriffszeit

Vorbild für die CD- Audio

CD Audio wurde von Philips und Sony entwickelt, in Anlehnung an die LV Videodisc.

Unterschiede:

- Echtes digitales Medium (LV benutzt nur analoge Informationen).
- Die CD verwendet nur eine konstante lineare Geschwindigkeit (CLV).

Beschreiben Sie kurz CD-i. Beschreiben Sie kurz die wichtigsten Unterschiede zwischen CD-ROM und CD-i. [3]

CD-i ist die erste MM-Plattform Anfang der 90er Jahre gewesen. Es ist ein komplettes System im Gegensatz zu CD- Rom. Sowohl logische als auch physikalische Organisation von Multimediainformationen auf der CD. Hardware Interfaces für CD-i Player, Softwareinterfaces für CD-i Programme.

Im Gegensatz zur CD-ROM sind auf der CD-i für jeden Medientyp Formate definiert, der Prozessor ist in den Player integriert → eine CD-i ist auf einem PC nicht mehr abspielbar, jedoch aber CD-ROM XA (extended architecture, die eine offene Form der CD-i ist und beide Formate lesen kann)

Die CD-i gibt dem Abspielgerät zusätzliche Informationen über den auf der CD vorhandenen Medientyp (Video, Bilder, Audio, Daten), die alle in verschiedenen vorgegebenen Standards mit unterschiedlichen Qualitätsstufen gespeichert werden können.

Auf CD-i-Playern lassen sich Audio-, Photo- und VideoCDs als auch interaktive CDs mit Spielen oder Lernsoftware abspielen. Mit ihr waren schließlich erstmals interaktive Videos möglich.

Auf der Cd-i sind folgende Medientypen festgelegt:

- Audio (CD-DA, Level A, Level B, Level C)
- Image (RGB 5:5:5, DYUV, CLUT (256, 128, 16 Einträge), RL (runlength coding, 128 oder 8 Einträge))
- Video (MPEG-1 enkodiert)
- Text und Graphik (applikationsabhängig)

Erklären Sie kurz die Abbildung von Media Time zur Speicheradresse in QuickTime. Welche Datenstrukturen werden dafür benötigt? [3.5]

Zu einer bestimmten Media Time werden entsprechende Samples abgefragt welche Chunks zugeordnet sind, diese beinhalten die Information über die Speicheradresse.

Media Value: Die Media Value weist die Samples den einzelnen Chunks zu.

Sample number: Sample Nummer

Sample Span: Anzahl der einander gereihten Samples der gleichen Größe

Sample Start Time: Anfangszeit des ersten Samples

Sample Duration: Dauer eines einzelnen Samples

Time to Sample: Bei Time to Sample werden die Samples der Zeit zugeordnet.

Sample to Chunk: Effiziente Speicherung der Zuordnung von Samples zu den Chunks

Chunk Number: Chunk Nummer

Chunk Span: Anzahl der einander gereihten Chunks der gleichen Größe

Start Sample in Chunk: Nummer der ersten Samples im jeweiligen Chunk

Samples/ Chunk: Anzahl der Samples im jeweiligen Chunk

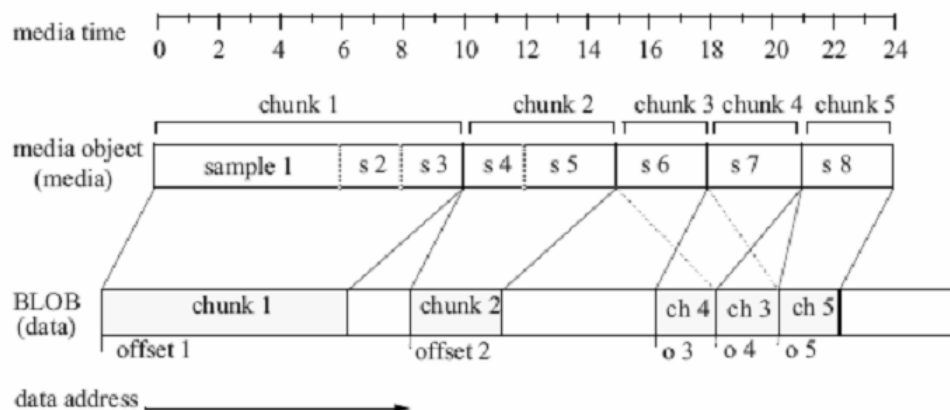
Encoding: Welches Encoding wird hier verwendet

Hierzu benötigt man für die Zeit die Datenstruktur media time, für die Samples media und für die Chunks BLOB (Binary Large Objects)

		Time-to-sample			
		sample number	sample span	sample start time	sample duration
Ch 1	sample 1	1	1	0	6
	sample 2	2	3	6	2
	sample 3	5	4	12	3
Ch 2	sample 4				
Ch 2	sample 5				
Ch 3	sample 6				
Ch 4	sample 7				
Ch 5	sample 8				

Sample-to-chunk				
chunk number	chunk span	start sample in chunk	samples /chunk	encoding
1	1	1	3	E1
2	1	4	2	E1
3	2	6	1	E1
5	1	8	1	E2

Durch diese Tabellen wird Speicherplatz gespart. Die verwendeten Datenstrukturen sind die Media und Data Entity.



Geben Sie eine verbale Beschreibung des Algorithmus, der für einen gegebenen QT-Track und einen gegebenen Zeitpunkt, die Adresse des zu diesem Zeitpunkt abzuspielenden Samples ermittelt. Welche Besonderheit haben die Tabelleneinträge, z. B. bei aufeinanderfolgenden Samples gleicher Länge? [5]

Algorithmus

- 1) Rechne Movie-Time mit Hilfe des TCS des Media-Entities in Media-Time um.
- 2) Bestimme zum diskreten Zeitpunkt das entsprechende Sample mit Hilfe der Time-To-Sample-Tabelle.
- 3) Bestimme zum Sample den Chunk mit Hilfe der Sample-To-Chunk-Tabelle. In diesem Chunk sind die Daten enthalten, die zum gefragten Zeitpunkt abgespielt werden müssen. Die Adresse des Chunks ist bekannt.

Besonderheiten in den Tabellen

- Die Time-To-Sample-Tabelle enthält Samples nicht, wenn sie die gleiche Länge haben des vorherigen.
- Die Sample-To-Chunk-Tabelle enthält Chunks nicht, wenn sie die gleiche Sample-Anzahl wie der vorherige Chunk haben und gleich codiert sind.

Datenkompression

Worin unterscheidet sich die I-Frame-Kodierung des MPEG-Verfahrens von der JPEG-Kodierung [2] bzw. Worin unterscheiden sich die Quantisierungsschritte im JPEG- und MPEG-Verfahren? [2]

JPEG Verfahren

Die JPEG Codierung (Joint Photographic Experts Group) ist ein Komprimierungsverfahren auf Grundlage der Diskreten Cosinus Transformation (DCT). Ziel ist, das hinauswerfen von Bits (Kürzung).

Uniforme Quantisierung: Dividiere die Koeffizientenwerte $S(u, v)$ durch N und runde das Ergebnis. Anhand der Quantisierungstabelle wird entschieden, wie viele Bits gekürzt werden sollen. Sie besteht aus 64 Elementen, wobei jedes 8-bit verwendet: Q_{uv} . Die neuen komprimierten Werte ergeben sich durch die Verwendung der Tabellen: $S_{quv} = S_{uv} / Q_{uv}$ zwei verschiedene Standardtabellen: Helligkeit und Farbe

Sprich: Zuerst wird das Bild in 8x8 (getrennt für Farbkanäle) Pixelblöcke unterteilt => 64 Koeffizienten für die 64 Pixel werden gebildet, Pixelblock wird als Vektor angesehen.

Durch Anwendung der DCT bekommt man den dazugehörigen 8x8 Frequenzblock. Nun erfolgt die Quantisierung. Hierbei wird jeder der Werte aus der Frequenzmatrix durch den zugehörigen Wert in einer vordefinierten Quantisierungsmatrix dividiert und auf Integer gerundet => es entsteht eine Gewichtung.

Zusammenfassung der Schritte beim JPEG Verfahren

Die Kompression erfolgt durch das Anwenden mehrerer Verarbeitungsschritte, von denen nur zwei verlustbehaftet sind.

- Farbraumumrechnung vom (meist) RGB-Farbraum ins YCbCr-Farbmodell (nach IEC 601).
- Tiefpassfilterung und Unterabtastung der Farbabweichungssignale Cb und Cr (verlustbehaftet).
- Blockeinteilung in 8x8 Pixel Blöcke die dann als Vektor der Dimension 64 interpretiert und mit DCT (Diskreter Cosinus Transformation) in den Frequenzbereich transformiert werden.
- Quantisierung des DCT- Ergebnisvektors (verlustbehaftet).
- Umsortierung.
- Entropiekodierung.
- **Codierung des quantisierten Vektors mittel Huffman oder arithmetischer Codierung**

Die Datenreduktion erfolgt durch die verlustbehafteten Verarbeitungsschritte in Zusammenwirken mit der Entropiekodierung.

MPEG Verfahren - I- Frame Kodierung

I Frames werden ohne zusätzliche Informationen bezüglich anderer Einzelbilder kodiert und komprimiert (Intraframekodierung). Sie benötigen am meisten Speicherplatz, lassen sich aber unabhängig von vorangegangenen Bildern dekodieren. Daher sind sie notwendig, um (nahezu) beliebig in einem Video springen zu können. Andere Bilder werden in Abhängigkeit von den anderen Bildern in dem Videostrom kodiert und benötigen dadurch weniger Speicherplatz. I- Bilder bilden somit die Anker für den wahlfreien Zugriff.

Ein I- Bild wird wie ein Standbild behandelt. Hier wurde in MPEG auf die Ergebnisse von JPEG zurückgegriffen. Die Kompression muss jedoch im Gegensatz zu JPEG auch in Echtzeit möglich sein. Die Kompression ist deshalb am geringsten.

MPEG Verfahren - I-Frame Kodierung Ablauf

Hier werden ebenfalls 8x8 Blöcke definiert, die jedoch in einem Makroblock zu finden sind. Auf diesen Blöcken wird die DCT Transformation ausgeführt.

Nach der Anwendung der DCT werden alle Werte in der resultierenden Frequenzmatrix durch einen Quantisierungswert dividiert und gerundet. Die Quantisierung erfolgt also durch einen konstanten Wert für alle Werte. Bei I-Frames ist der Quantisierungswert für alle Makroblöcke der gleiche. Bei P-Frames wird er für jeden Block separat ermittelt.

Die Quantisierung wird an die Datenrate angepasst. Quantisierung hat konstanten Wert für alle DCT Koeffizienten. Wenn nun sich die Datenrate erhöht und einen bestimmten Grenzwert übersteigt, erhöht die Quantisierung die Schrittweite. Wenn aber nun sich die Datenrate verringert, wird die Quantisierung mit feinerer Granularität (Schritten) durchgeführt.

Die MPEG-Quantisierung ist angepasst:

- Wenn die Datenrate über einen Schwellwert steigt, vergrößert die Quantisierung die Schrittweite
- Wenn die Datenrate absinkt wird die Quantisierung mit einer feineren Körnung ausgeführt

Folgerung: JPEG-Quantisierung ist uniform, MPEG-Quantisierung wird der aktuellen Gegebenheit angepasst

Zusammenfassung der Schritte beim MPEG Verfahren - I-Frame Kodierung

- Abtrennung der Farbinformation vom Schwarz-Weiß-Bild und Vergrößerung derselben, da auch das menschliche Auge Farbunterschiede gröber wahrnimmt als Helligkeitsunterschiede.
- Aufteilung des Bildes in 8x8 Pixel große Blöcke, deren Datenbedarf mittels Diskreter Kosinustransformation und anschließender Quantisierung stark verkleinert wird
- Zusammenfassung von je vier Blöcken zu 16x16 Pixel großen sog. Makroblöcken, deren Ähnlichkeit zu Makroblöcken in vorigen und/oder folgenden Bildern ebenfalls datenreduzierend genutzt wird
- Verwendung von Binärcodes mit variabler Länge, um häufigere Bitfolgen kürzer darstellen zu können als seltenere.

Antwort 2

Die I-Frame-Kodierung verwendet zur Run-Length-Kodierung nur Huffman-Kodierung. Um eine konstante Bit-Rate aufrecht zu erhalten, wird die Quantisierungsmatrix von einem Buffer Regulator verwendet.

Antwort 3

Im gelben Skriptum habe ich noch einen weiteren Unterscheid gefunden, der aber ziemlich ins Detail geht. Es wird ja bei den DC-Koeffizienten die Differenz zum vorherigen berechnet. Diese Differenz wird nicht direkt abgespeichert, sondern man nimmt die Anzahl der benötigten Bits für diese Differenz, bestimmt dafür einen Huffman-Code, und speichert schließlich das Tupel (<Huffman code für die Bit-Anzahl>, <diff>). Das ganze gilt bei JPEG. Bei MPEG gibt es einen kleinen Unterschied: hier wird nur für die häufigsten Längen (Bit-Anzahlen) ein Huffman-Code generiert. Die seltenen werden mit einer Escape-Sequenz direkt gespeichert. I-Frames werden mit einer konstanten Quantisierung komprimiert und bei JPEG werden Quantisierungsmatrizen definiert.

Beschreiben Sie kurz die MPEG-4 Videokodierung. [5]

Der MPEG-4-Standard besteht generell aus einem Systemteil, einem Audio-Teil und einem Visual-Teil.

Jeder einzelne Videoframe wird in eine Anzahl von willkürlich geformten Regionen unterteilt, die video object planes (VOP) genannt werden. Aufeinander folgende VOPs die zum selben physikalischen Objekt in einer Szene (z.B. ein Mensch) gehören werden als video objects bezeichnet. Die Form, Bewegung und Textur der VOPs, die zum selben VO gehören, werden jeweils in einem separaten Video object layer (VOL) enkodiert. Relevante Informationen die benötigt werden, um jedes der VOLs zu identifizieren, sowie wie die verschiedenen VOLs zusammengesetzt werden, werden auch kodiert. Dadurch wird dann auch selektives Dekodieren möglich und außerdem erlaubt es uns einzelne object levels getrennt zu skalieren. So kann man z.B. ein Objekt, das als separates VO kodiert wurde vergrößern und in eine neue Umgebung platzieren.

Formulieren Sie in Pseudocode einen Algorithmus, der für die im Array Symbole gespeicherten Zeichen eine Huffman-Kodierung durchführt. Das Feld Symbole enthält in der 1. Komponente das zu kodierende Zeichen, in der 2. Komponente seine Wahrscheinlichkeit und in der 3. Komponente die zu errechnende Kodierung. [6]

Erstelle einen Array e, der die sortieren Symbole[z][p][c] aufsteigend nach p enthält;

Für alle Elemente in e {

Für jedes Zeichen aus e[i]

$c[e[i].\text{charAt}(j)][c] = „1“ + c[e[i].\text{charAt}(j)][c];$

Für jedes Zeichen aus e[i+1]

$c[e[i+1].\text{charAt}(j)][c] = „0“ + c[e[i+1].\text{charAt}(j)][c];$

$e[e[i]+e[i+1]][p] = \text{Symbole}[e[i]][p] + \text{Symbole}[e[i+1]][p];$

lösche e[i] und e[i+1];

sortiere e neu;

}

Symbole.c = c;

Einfachere aber längere Methode

Sortiere aufsteigend, mit dem kleinsten beginnend, das Array nach seiner Wahrscheinlichkeit.

Do solange das Array mindestens 2 Elemente besitzt

{

Nimm ersten 2 Elemente aus dem Array und bilde ein neues Element x

$x.w = \text{Array}[0].w + \text{Array}[1].w$ //Wahrscheinlichkeit des neuen Elements

Bilde nun einen binären Baum, wobei x der Elternknoten und seine Kinderknoten Array [0] und Array [1] sind.

Wähle min (Array [0], Array[1]) als linken und max (Array[0], Array [1]) als rechten NachfolgeKnoten.

Markiere linken Nachfolger mit 1 und rechten mit 0.

Lösche beide ausgewählte Elemente Array [0] und Array [1] aus dem Array.

Sortiere Array aufsteigend nach seiner Wahrscheinlichkeit.

}

Was können Sie über die AC und DC Koeffizienten der DCT aussagen, wenn der zu transformierende 8x8 Block eines Bildes aus einem Schachbrettmuster besteht? [2]

Das erste Pixel im Block (links oben) ist der DC-Koeffizient. Er hat die geringste Frequenz und gibt die Grundfarbe des gesamten Blocks vor. Alle übrigen Pixel (die AC-Koeffizienten) haben höhere Frequenzen und können vom menschlichen Auge nicht mehr so gut unterschieden werden, wie niedrigere.

Der DC Koeffizient liegt an (0,0) und gibt die mittlere Graustufe des Blocks an. Die AC Koeffizienten sind entsprechend hoch, da ein Schachbrettmuster extrem hohe Ortsfrequenzen aufweist. Ungleich Null, konzentriert zu den höheren Frequenzen (entspricht einem idealen Schachbrettmuster).

Was können Sie über die AC und DC Koeffizienten der DCT aussagen, wenn im zu transformierenden 8x8 Block eines Bildes 2 Flächen in einer horizontalen Linie schneiden? [2]

Scharfe Schnittkante → höchste Frequenzen
horizontale Linie → vertikal in der DCT – Darstellung

- DC-Koeffizient: mittel (wenn sich beide Flächen genau in der Mitte schneiden und eine schwarz und eine weiß ist; spiegelt die Grundfarbe des Blocks wider)
- AC-Koeffizienten: erste (und zweite) Spalte von links ungleich 0, die anderen sollten 0 sein (da die Überschneidung der 2 Flächen eine hohe Differenz angeben).

Was können Sie über einen 8x8 Pixelblock aussagen, der nach der DCT-Transformation nur in den ersten beiden Spalten Werte aufweist, die von 0 verschieden sind? [2.5]

Es treten in vertikaler Richtung (im Bild, nicht in der transformierten Darstellung!) nur sehr niedrige Frequenzen $s(1,0)$ auf, in horizontaler Bildrichtung jedoch alle → es dürfte sich um ein Bild handeln, in dem sich zum Beispiel zwei verschiedenfarbige Flächen horizontal mit einer scharfen Kante („horizontale Linie“) schneiden.

Erklären Sie kurz den Begriff Quantisierung im Kontext der Sampling-Theorie und im Kontext von Komprimierungsverfahren. [2]

Sampling

Die abgetasteten analogen Werte werden diskreten Zahlenwerten zugeordnet (quasi eine Rundung). Der gesamte Spannungsbereich wird dazu in sogenannte Quantisierungsintervalle aufgeteilt und die abgetasteten Werte werden dem nächstgelegenen diskreten Wert zugeordnet.

Komprimierung

Bei Komprimierungsverfahren werden benachbarte Werte auf einen neuen Wert abgebildet. Jeder Wert in einer Frequenzmatrix wird durch seinen zugehörigen Wert in der Quantisierungstabelle dividiert und auf Integer gerundet um möglichst viele AC-Koeffizienten abzuschwächen (also auf 0 zu setzen) oder gar zu eliminieren.

Die Quantisierungsmatrix ist deshalb so aufgebaut, dass deren Werte von links oben nach rechts unten zunehmen, da bei der Frequenzmatrix korrespondierend die Frequenzen von links oben nach rechts unten zunehmen und somit mehr komprimiert werden. Hohe Frequenzen werden also stärker komprimiert (eher auf 0 gesetzt), da das menschliche Auge diese schlecht wahrnimmt.

Erklären Sie kurz das Konzept Komposition im MPEG-4 Systems Standard. [3]

Komposition ist ein Modell zur Beschreibung der Zusammenstellung einer komplexen MM-Szene (bestehend aus audiovisuellen Objekten (audio, video, natürlich, synthetisch, 2D, 3D, streamd oder downloaded), welche zusammengehörige MediaObjects kreiert, welche audiovisuelle Szenen

formen) beruhend auf den Konzepten von VRML. Autoren können die Beschreibung im Textformat generieren, möglich auch mit Hilfe eines Authoringtools.

Komposition

- Beschreibt ein Modell zur Komposition von Multimediaobjekten
- Behandelt nur 2D Inhalt
- Gekoppelt mit Streaming Media wie z.B. Video, Audio, Streaming Text.
- Ist auch im Stande Synchronisation durchzuführen
- Die Szenenbeschreibung wird in einer binären Repräsentation kodiert, nämlich in der BIFS (Binary Format for Scene Description) der Effizienz wegen.
- Eine Multimedia Szene beinhaltet hierarchische Strukturen und wird als Graph repräsentiert. Jedes Blatt im Graph repräsentiert ein Mediaobjekt.
- Alle Knoten und Blätter der Graphen der Szene werden an den zugehörigen Stream empfangen oder gesendet. Alle Kanten und Grafenblätter, die Streamingsupport benötigen um Medieninhalte zu erhalten, sind logisch mit den DekodingPipelines verbunden.

Beschreiben Sie kurz die Facial Animation und Facial Definition Parameters im Face Animation-Teil von MPEG-4. [2] bzw. Erklären Sie kurz die Facial Animation und Facial Definition Parameter im Face Animation-Teil von MPEG-4. [2]

Facial Animation beschäftigt sich mit den Parametern für Gesichtsanimation und Gesichtsdefinition.

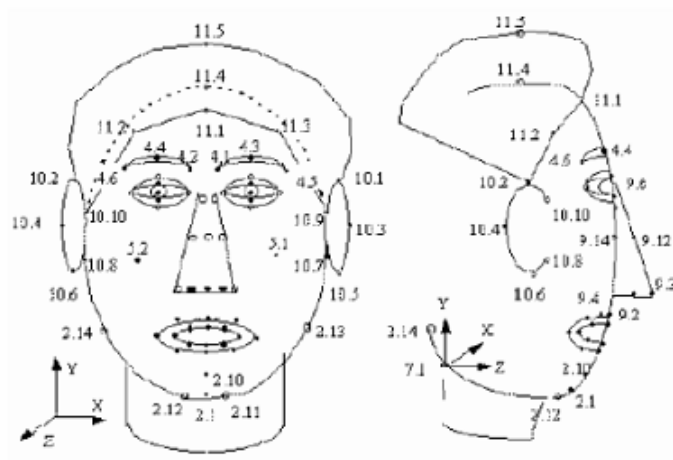
Facial Animation Parameters (FAP)

Erlauben das Anwenden von einem einzelnen Set von Parametern, unabhängig von dem verwendeten Gesichtsmodell das benutzt wird. Die meisten FAPs beschreiben atomare Bewegungen von Gesichtsfeatures (Gesichtszügen), andere (expressions und visemes) definieren viel komplexere Deformationen/ Deformierungen.

Viseme definieren die Position des Mundes zusammenhängend mit den Phonemen. Sie sind auf die Position und Bewegungsabläufe einzelner Gesichtsteile (Mund, Augen...) spezialisiert. FAPs können entweder arithmetisch oder mit DCT **en**kodiert werden

Facial Definition Parameters (FDP)

Werden verwendet um das Receiver Terminal Default Gesicht (Empfängerterminal Standardgesicht) zu kalibrieren (adaptieren oder verändern der Form), oder um eine komplett neue Gesichtsmodelgeometrie und Textur zu erstellen und/oder zu übertragen.



Erklären Sie knapp den Unterschied zwischen Inter- und Intraframe-Synchronisation. Geben Sie 2 Geräte an, die der Synchronisation dienen. [3] bzw. Was versteht man unter Intra-Flow Synchronization und welche Störungen soll sie vermeiden? [2.5]

Inter-Frame-Synchronisation

Inter-Frame-Synchronisation meint die Zusammenführung von 2 Strömen von unterschiedlichen Quellen, so dass die gleichen Frames und die gleichen Zeilen genau zeitgleich ankommen, so dass man sie dann vernünftig verbinden kann (z.B. mit einem video production switcher). Man verwendet dazu einen timebase corrector, der von einem sync generator einen Takt bekommt.

Intra-Frame-Synchronisation

Dagegen meint die Intra-stream-Synchronisation die zeitliche Korrektur eines einzelnen Streams, um die sonst auftretenden Fehler (jitter und skew) zu vermeiden. Dazu verwendet man einen Frame Delay Buffer, der das eingehende Signal zwischenspeichert und genau zu bestimmten Zeitpunkten weitergibt (getriggert von einem sync generator), egal wann es eingetroffen ist.

Zusammengefasst

- Interframe/ Interflow: Zusammenführen und Synchronisieren von 2 Streams (z.B. Audio und Video) um auseinanderlaufen von Ton / Bild zu verhindern.
- Intraframe/ Intraflow: Synchronisation eines Stream (z.B. von Videodaten) untereinander um ruckeln zu vermeiden.

Erklären Sie kurz die verschiedenen Frametypen der MPEG-Komprimierung. [1.5] bzw. Erklären Sie kurz die P- und B-Frame Kodierung der MPEG-Komprimierung. [2]

I-Frames (intracoded Frames)

Ein I-Frame ist ein kodiertes Vollbild. Benötigen keine Referenz auf andere Bilder und werden wie stehende Bilder behandelt. MPEG verwendet JPEG für I-Frames.

Komprimierung muss in Echtzeit passieren, deswegen ist die Komprimierungsrate die geringste innerhalb von MPEG.

I-Frames sind Punkte für den Zufallszugriff in einem MPEG-Stream.

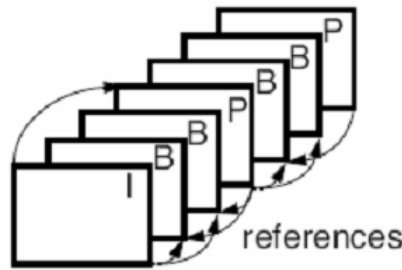
I-Frames verwenden 8x8 Blöcke innerhalb eines Makroblocks, auf diesen wird DCT ausgeführt und für alle DCT Koeffizienten mit konstantem Wert quantisiert.

P-Frame Codierung (predictive-coded Frames)

Sie benötigen Informationen des vorhergehenden I-Frames und/oder vorhergehender P-Frames für Kodierung sowie Dekodierung. Bei der P-Frame Codierung macht man sich aufeinanderfolgende Bilder zu Nutze, in welchen Teile sich nicht verändern bzw. nur „verschoben“ werden. Nur die Differenz zum vorhergehenden Bild wird weiter behandelt. Es ist möglich die Bewegung der Makroblöcke zu berücksichtigen – Motion Compensated Prediction

Zeitliche Redundanz: Bestimme den letzten P- oder I-Frame, der dem aktuellen Block am ähnlichsten ist. Verwende die Motion estimation Methode am Encoder.

Motion Estimation Method: Für jeden Makroblock wird ein Vektor errechnet, der angibt, wo dieser in dem darauffolgenden Referenzbild seine ähnlichste Entsprechung findet. Die beiden Makroblocks werden anschließend voneinander subtrahiert und DCT-kodiert. Die Dekodierung/Kodierung erfolgt also mit Infos bzgl. vorangegangener I- oder P-Bilder und die Kompressionsrate liegt daher auch höher als bei I-Bildern.



B-Frame Codierung (bidirektional predictiv-coded)

Benötigen Information des vorhergehenden und nachfolgenden I und/oder P-Frames. Sind dem P-Frame ähnlich, verwenden zur Kodierung allerdings den vorhergehenden und den folgenden Frame. Wenn Motion Compensation verwendet wird, werden Bewegungsvektoren in die „Vergangenheit“ und in die „Zukunft“ gespeichert.

Motionvektor des vorhergehenden und nachkommenden Referenzframes. Durchschnitt der Blöcke des vorhergehenden und späteren Frames. Für die B-Frame-Kodierung ist es nötig, auch das Vorgängerbild einzubeziehen. A = Block BlockNext, B = Block-BlockLast; Das Ergebnis ergibt sich aus einem Multiplexing von A, B und $(A+B)/2$. Es wird anschließend DCT-kodiert. Die Dekodierung/Kodierung erfolgt also mit Infos bzgl. vorangegangener und später auftretender I- oder P-Bilder und die Kompressionsrate ist daher am höchsten.

D-Frames (DC-kodierte Frames)

Verwendung für fast forward, fast rewind Modi. Nur das Matrixelement $a(00)$ der DCT-Koeffizienten wird verwendet. D-Frames beinhalten somit die niedrigsten Frequenzen des Bildes. DC-Parameter sind DCT-codiert, AC Koeffizienten werden vernachlässigt. Es wird nur ein Typ von Makroblöcken verwendet und nur in MPEG-1.

Beschreiben Sie kurz die Aufgabe von D-Frames im Rahmen von MPEG. Wie werden D-Frames kodiert? [2]

D-Frames sind intraframe-kodiert und werden bei MPEG für schnelles vorwärts bzw. rückwärts spulen verwendet (fast forward, fast rewind Modi). Nur die DC Komponente wird bei D-Frames kodiert, die AC Koeffizienten werden verworfen. Daher bestehen D-Frames nur aus den niedrigsten Frequenzen eines Bildes. Sie verwenden nur eine Art von Makroblocks und werden nur in MPEG-1 benutzt.

Was versteht man unter Group of Pictures? Wie hängen Group of Pictures Muster mit der Datenrate zusammen? [2.5]

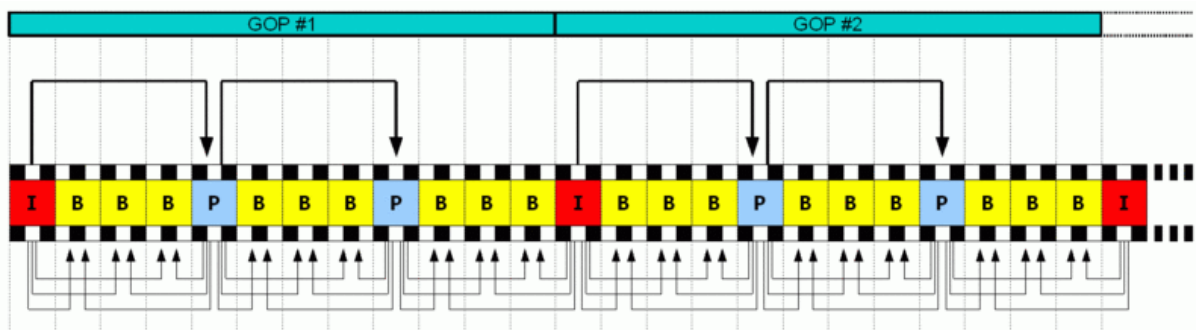
Die Group of Pictures (GOP) ist eine Gruppe von aufeinanderfolgenden MPEG-Bildern, engl. Frames, innerhalb eines MPEG-codierten Films bzw. Videostreams. Jeder MPEG-codierte Film bzw. Videostream besteht aus aufeinanderfolgenden GOPs. Aus den darin enthaltenen MPEG-Bildern werden die sichtbaren Einzelbilder generiert. Eine GOP ist jeweils eine zusammenhängende Folge von I-, P- und B-Frames. Z.B. IPPBPPBPPIPPBPPBPP (den Beginn der neuen GOP hab ich kursiv markiert).

Eine GOP kann folgende Bildtypen enthalten

- I-Bild bzw. I-Frame (engl. intra-coded picture) Referenzbild, entspricht einem Standbild und ist von anderen Bildtypen unabhängig. Jede GOP beginnt mit diesem Bildtyp.
- P-Bild bzw. P-Frame (engl. predictive-coded picture) enthält Differenz-Informationen aus dem vorhergehenden I- oder P-Frame.

- B-Bild bzw. B-Frame (engl. bidirectionally predictive-coded pictures) enthält Differenz-Informationen aus dem vorhergehenden und/oder nachfolgenden I- oder P-Frame.
- D-Bild bzw. D-Frame (engl. DC direct coded picture) dient dem schnellen Vorlauf.

Eine GOP beginnt stets mit einem I-Frame. Danach folgen ein oder mehrere P-Frames, jeweils mit einigen Frames Abstand. In den verbleibenden Lücken befinden sich B-Frames. Mit dem nächsten I-Frame beginnt eine neue GOP. Die folgende Grafik zeigt einen typischen Aufbau von GOPs innerhalb einer MPEG-Datei. Die Pfeile kennzeichnen die Referenzen.



Zusammenhang mit der Datenrate

Bei größeren GOP's werden weniger I-Frames (und die sind ja die aufwendigsten!) benötigt. Da I-Frames aber die geringste Komprimierung erreichen (Intraframe-Kompression) und B-, und P-Frames viel besser komprimiert werden können (und kleiner sind), sinkt mit zunehmender GOP-Größe und damit abnehmender I-Frame-Anzahl die Datenrate.

Erklären Sie kurz die Begriffe Video Object und Video Object Plane. In welchem Kontext werden sie verwendet? [3]

Werden bei der MPEG-4 Codierung verwendet.

Video Object Plane

Jeder Videoframe ist in eine Anzahl willkürlich geformter Regionen, die sogenannten Video Object Planes (VOP), segmentiert.

Video Objects

Aufeinanderfolgende VOPs, die zum selben physikalischen Objekt in einer Szene (z.B. ein Mensch) gehören, werden als Video Objects bezeichnet.

Video Object Layer

in einem Video Object Layer sind die Form-, Bewegungs- und Texturinformationen von VOPs, welche zum selben VO gehören, kodiert. Außerdem sind relevante Informationen, welche zur Identifikation jedes VOPs und wie verschiedenen VOPs zusammengesetzt sind, kodiert. Dies erlaubt eine selektive Dekodierung und bietet Object-Level Skalierbarkeit.

Für welche Einsatzbereiche würden Sie das MotionJPEG-Verfahren dem MPEG-Verfahren vorziehen? [1.5]

- MotionJPEG ist vorzuziehen, wenn man gezielt auf Einzelbilder zugreifen möchte (z.B. nonlinearer Videoschnitt, Navigable Video,...)
- Beim Harddisc-Recording (Verfahren, bei denen viele unterschiedliche Datenraten zum Einsatz kommen).

Formulieren Sie in Pseudocode einen Algorithmus, der in einer quantisierten DCT-Matrix die Runlength-Kodierung durchführt. [3.5]

```
i=0;
Solange (i < 63)
{
    Wenn (Matrix(i) != Matrix(i+1)) (im Zig-Zag-Scan)
    {
        Gib Matrix(i) aus
        i=i+1
    }
    Sonst
    {
        count=i+1
        Solange ((Matrix(i)=Matrix(count) AND (count<63))
        {
            count=count+1
        }
        count=count-i
        z = 0;
        Wenn (count >= 4)
        {

            Gib Matrix(i) aus
            Gib ! aus
            Gib count aus
        }
        Sonst
        {
            solange (z < count) gib Matrix(i) aus
        }
        i=count+1
    }
}
```

Anders formuliert

Für alle Matriceinträge (also 64)
Vergleiche Matriceintrag (i) mit Matriceintrag (i+1)
Wenn beide gleich, dann Zähler um eines erhöhen
Wenn die Anzahl der gleichen Buchstaben größer 3 und nächster Matriceintrag (i+1) nicht gleich sind, dann gib „!“ + Matriceintrag (i) + Anzahl der Buchstaben aus
Sonst wenn Matriceintrag (i+1) nicht gleich gibt Matriceintrag (i) aus
Wenn gleich geh weiter
Erhöhe i um 1

Erklären Sie kurz die Begriffe Profiles und Levels im Kontext von MPEG-2. [2]

Das Angebot unterstützter Codierung ist in Profile und Level unterteilt. Für jedes Profil/Level bietet MPEG-2 die Syntax für den kodierten Bitstream sowie die Dekodierungsanforderungen an.

Profil

Seit MPEG2 dienen Profiles der Definition einer Liste von Codierungs-Möglichkeiten, zusammen mit Levels definieren sie die Mindestanforderungen an eine Decoder-Klasse. Ein Profile legt fest welche

Tools und Funktionen dazu verwendet werden können einen bitstream zu erzeugen und wie dies zu geschehen hat.

Ein Profil ist eine definierte Teilmenge der gesamten Bitstream Syntax. Innerhalb eines Profils ist ein Level als eine Reihe von Bedingungen für die Parameter des Bitstreams (zum Beispiel für die Auflösung, die maximale Bitrate).

Profile und Levels stehen in hierarchischer Beziehung. Die Syntax, die von einem höheren Profil/Level unterstützt wird, inkludiert alle syntaktischen Elemente der niedrigeren Profile/Levels. Dekoder eines bestimmten Profils/Levels sollten demnach auch in der Lage sein Bitstreams niedrigerer Profile/Levels zu dekodieren – Ausnahme: Dekoder des simple Profils am main-Level müssen auch das main Profil für Low-Level Bitstreams dekodieren können (MPEG-1).

Ein Profile definiert aus der Menge der von der Norm zur Verfügung gestellten Verfahren eine bestimmte Anzahl, die in diesem Profile genutzt werden darf. Jeder Decoder, der dieses Profile beherrscht, kann den Datenstrom verarbeiten.

Level

Die Level legen innerhalb eines Profiles wichtige Bildparameter fest. Hierzu gehören u.a. die unterschiedlichen Bildauflösungen (z.B. Mainlevel ML => 720x576), die für die Decodierung benötigten Speichergrößen und die maximalen Größen der Bewegungsvektoren. Die Profiles und Levels sind abwärtskompatibel. Das bedeutet, daß ein höheres Profile oder ein höherer Level alle niederwertigen Profiles und Levels mit beinhaltet.

Beispiele

Profile: Simple, Main, 4:2:2, SNR, spatial, high, multiview

Levels: low (SIF), main (CCIR 601), high-1440, high (HDTV)

Nennen und erläutern Sie kurz 2 Arten von Artefakten bei der Audiokodierung. [2.5] bzw. Beschreiben Sie kurz 2 der häufigsten Artefakte bei der Audiokompression. [2.5]

Pre-Echos

Die Verschlechterung der Zeitauflösung hat zur Folge, dass ein Pre-Echo-Effekt auftreten kann. Als Pre-Echo bezeichnet man solche Geräusche, die vor dem eigentlichen Signal entstehen. Z.B. vor einem lauten plötzlichen Geräusch (zum Beispiel Schlagzeug) sind klirrend-metallische Artefakte hörbar. Um dem Pre-Echo-Effekt vorzubeugen, werden bei der MDCT unterschiedliche Transformationsblocklängen gewählt. Lange Blöcke werden bei stationären, sich also zeitlich wenig ändernden Signalen eingesetzt, um die Kompression zu erhöhen. Kurze Blöcke werden bei sich schnell ändernden Signalen verwendet, um den Bereich, in dem Pre-Echo auftreten kann zu verkleinern. Stärkenänderung beginnen, insbesondere wenn ein Bereich mit sehr niedrigen Amplituden vorausgeht.

roughness, double-speak

Bei niedrigen Bitraten und niedrigeren Sampling- Frequenzen tritt eine Fehlanpassung bzw. ein Ungleichgewicht zwischen der zeitlichen Auflösung des Coders und den Anforderungen, die bestimmte Signale aufgrund ihrer Zeitstruktur stellen, auf.

Bandbreitenverlust

Nennen und erläutern Sie kurz 2 Arten von Scalable Bit Streams in MPEG-2. [3] bzw. Wozu dienen Spatial Scalability und SNR Scalability in MPEG-2? Was haben sie gemein und was

sind die Unterschiede? [3] bzw. Wozu dienen Temporal Scalability and SNR Scalability und was sind die Unterschiede? [2.5]

Es gibt 4 skalierbare Modi – MPEG2 Video wird in verschiedene Layer aufgeteilt, meist mit der Absicht Videodaten zu priorisieren:

- **Spatial scalability:** Kodiert einen Basislayer mit niedrigeren Samplingdimensionen (Auflösungen) als die höheren (darüberliegenden) Layer. Dann hat man die normale Auflösung und die reduzierte
- **Datenpartitionierung:** Teilt den Block von 64 (8x8) quantisierten Koeffizienten in 2 Bitstreams. Der erste, höher priorisierte Bitstream beinhaltet die kritischeren niederfrequenten Koeffizienten und Randinformationen (DC Werte, Bewegungsvektoren)
- **SNR scalability:** Macht auch mehrere Bitstreams nur haben die bei gleicher samplingrate verschiedene Qualität (Bildqualität) weil manche quantisiert werden und andere nicht.
- **Temporal scalability:** Bitstream mit höherer Priorität kodiert Video mit niedrigerer Framerate, mittlere Frames können in einem zweiten Bitstream kodiert werden unter der Verwendung der Rekonstruktion des ersten Bitstreams als Prognose

Unterschiede - Gemeinsamkeiten von Spatial Scalability und SNR Scalability

Beide sind skalierbare Codiermodi von MPEG-2, aber spatial scalability verwendet 2 fixe Layer, SNR ist wirklich frei skalierbar.

Unterschiede - Gemeinsamkeiten von Temporal Scalability und SNR Scalability

???

Erklären Sie kurz die 4 Arten(Modi) der Makroblock Codierung von B-Frames in MPEG-4 [3]

Vorwärts, Rückwärts, Interpolationsmodus und direkter Modus

Ermitteln Sie eine Huffman-Codierung der Symbole A, B, C, D und E; Wahrscheinlichkeiten ihres Auftretens: $p(A)=0.25$, $p(B)=0.41$, $p(C)=0.18$, $p(D)=0.14$, $p(E)=0.02$ [2.5]

Verschiedene Möglichkeiten der Darstellung, je nach Beschriftung der Kanten!

A = 01
B = 1
C = 001
D = 0000
E = 0001

Ermitteln Sie eine Huffman-Codierung der Symbole V, W, X, Y und Z; Wahrscheinlichkeiten ihres Auftretens: $p(V)=0.2$, $p(W)=0.4$, $p(X)=0.18$, $p(Y)=0.15$, $p(Z)=0.07$ [2.5]

V = 01
W = 1
X = 001
Y = 0000
Z = 0001

Was sind die wesentlichsten Unterschiede zwischen MPEG-1 und MPEG-2 Video? [2]

MPEG-2 hat zusätzlich

- zusätzliche Syntax für effiziente Kodierung von interlaced Videos
- 10-bit DCT DC Koeffizienten (MPEG-1 hat 8 bits)

- Skalierbare Erweiterungen, die eine Aufteilung eines kontinuierlichen Videosignals in 2 oder mehrere kodierte Bitstreams erlaubt, die das Video mit verschiedenen Auflösungen, Bildqualitäten oder Bildraten darstellen.
- mehrere Bildformate
- Progressives und interlaced Framekodierung
- zusätzliche Prediction Modi sowie zugehörige Makroblöcke
- 4 Skaliermodi (verschiedene Qualitätsstufen in einem Videostream)
- Verbesserte Bildqualität (Quantisierung, Zig-zag-scan)
- Erhöhung der Genauigkeit der Bewegungsvektoren auf halbe Bildpunkte
- erweiterte Fehlerredundanz durch spezielle Vektoren bei I-Frames
- Präzision der DCT (diskreten Cosinus-Transformation) wählbar

Was versteht man unter Pre-Echo? Wo tritt dieses Phänomen auf? [2]

Pre-Echo

Pre-Echo ist ein Artefakt, das bei der Audio-Kodierung entsteht. Vor einem lauten plötzlichen Geräusch (zum Beispiel Schlagzeug) sind deutliche Artefakte zu hören. Die Verschlechterung der Zeitauflösung hat zur Folge, dass ein Pre-Echo-Effekt auftreten kann. Als Pre-Echo bezeichnet man solche Geräusche, die vor dem eigentlichen Signal entstehen.

Die Verschlechterung der Zeitauflösung hat zur Folge, dass ein Pre-Echo-Effekt auftreten kann. Als Pre-Echo bezeichnet man solche Geräusche, die vor dem eigentlichen Signal entstehen. Z.B. vor einem lauten plötzlichen Geräusch (zum Beispiel Schlagzeug) sind klirrend-metallische Artefakte hörbar. Um dem Pre-Echo-Effekt vorzubeugen, werden bei der MDCT unterschiedliche Transformationsblocklängen gewählt. Lange Blöcke werden bei stationären, sich also zeitlich wenig ändernden Signalen eingesetzt, um die Kompression zu erhöhen. Kurze Blöcke werden bei sich schnell ändernden Signalen verwendet, um den Bereich, in dem Pre-Echo auftreten kann zu verkleinern. Stärkenänderung beginnen, insbesondere wenn ein Bereich mit sehr niedrigen Amplituden vorausgeht.

Post-Echo

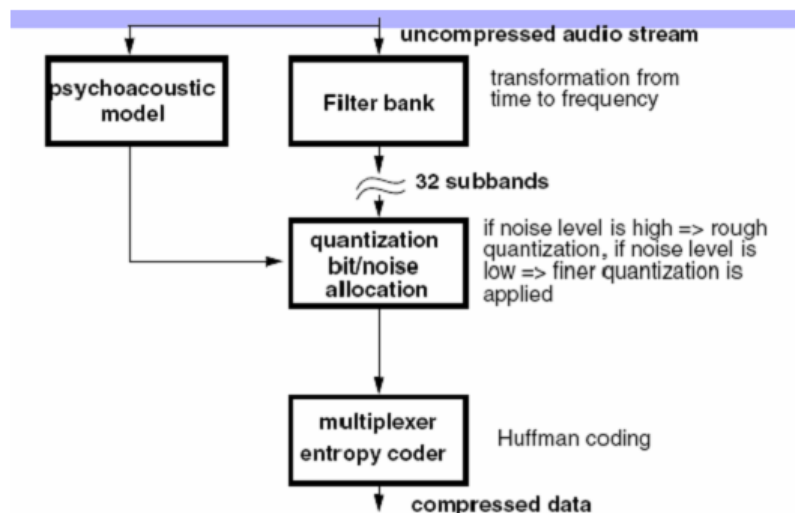
Nach einem plötzlichen Geräusch sind deutliche Artefakte zu hören

Beschreiben Sie kurz die wesentlichen Schritte der MP3-Kodierung. Erklären Sie insbesondere Quantisierung und Kodierung. [5] bzw. Erklären Sie knapp die Schritte der MP3-Kodierung. [5]

MP3 bezeichnet nicht ein unabhängiges Verfahren zur Audiokomprimierung, sondern entspringt der Familie der MPEG-I Codecs. MP3 ist die Abkürzung für MPEG-1-Layer-III, und bietet gegenüber den Layer-I und Layer-II Codecs zusätzliche Funktionalität.

Folgende Algorithmen finden in allen drei Layern Verwendung:

- Im ersten Schritt des Kompressionsprozesses wird das digitale Eingang- Signal in Subbänder geteilt. Dieser Prozess wird auch als Zeit-Frequenz-Zuordnung bezeichnet.
- Der nächste Schritt der Kodierung ist die Implementierung des psychoakustischen Modells
- Im dritten Schritt erfolgt die Quantisierung und Encodierung jedes einzelnen aus Stufe 1 und 2 hervorgegangenen Samples in den Subbändern.
- Der letzte Schritt ist die endgültige Erzeugung der komprimierten Datei, wobei die digitalisierten Audiodaten in einzelne Frames unterteilt werden.



Zusätzlich zu diesen 4 Schritten sieht MPEG1-Layer III noch zwei weitere Schritte vor, die zwar die Zuverlässigkeit des psychoakustischen Modells erhöhen und den Ausgangsdatenstrom so klein wie möglich halten, aber allerdings zu einer längeren Encodierungszeit bei der Erstellung der Ausgangsdatei führen:

- Neben der Implementierung des psychoakustischen Modells aus Schritt 2 kommen noch mehrere Filter- (z.B.: Hoch- Tiefpassfilter) und "Diskrete Cosinus- Transformations"- Algorithmen zum Einsatz.
- In der letzten Phase, in der die Frames vorbereitet werden, ist ein komplexer Algorithmus hinzugefügt, der die Framegröße nach bestimmten Parametern variieren lässt. Es wird auch das sogenannte Reserve Bit gesetzt, dass eventuell auftretende kritische Verzerrungen (Artefakte) bei der Encodierung des Samples löst.

Quantisierung und Kodierung

- Die Spektralkomponenten werden quantisiert und kodiert. Störungen soll unter einem Schwellwert gehalten werden.
- Zwei verschachtelte Iterationsschleifen (rate-loop und Noise-control-loop)
- Quantisiert wird mittels eines potenzförmigen Quantisierers (größere Werte werden automatisch mit weniger Genauigkeit codiert)
- Quantisierte Werte werden nach Huffman codiert

Welchem Zweck dienen BIFS in MPEG-4? Welche Informationen enthalten sie? [2.5]

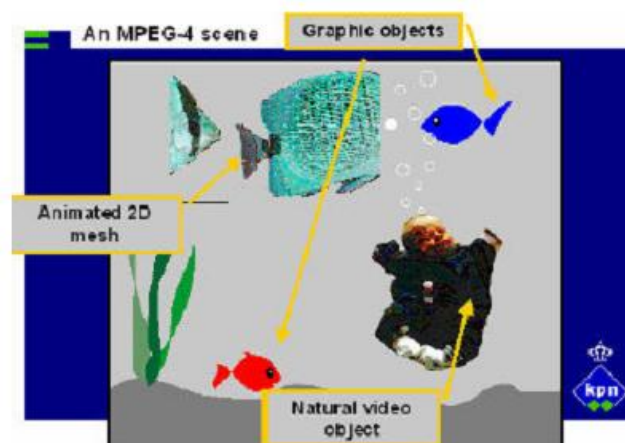
BIFS ist die Abkürzung für Binary Format for Scenes, eine auf VRML97 basierende und in MPEG-4 part 11 (ISO/IEC 14496-11, „Szenenbeschreibung und Anwendungs-Engine“) standardisierte Beschreibungssprache für zwei- und dreidimensionale multimediale audiovisuelle interaktive Inhalte. Sie wird binär codiert und wird bei mpeg4 verwendet um den Szenegraph einer Komposition zu speichern.

Örtlich-zeitliche Zuordnung einzelner Objekte eines Datenstromes können unabhängig voneinander beschrieben werden. Mit BIFS können Objekte beschrieben, hierarchisch geordnet und in einem zwei- oder dreidimensionalen Raum unter Berücksichtigung der Zeitachse gruppiert werden. BIFS setzt auf die virtuelle Modellierungssprache VRML und kann verschiedene Szenen miteinander kombinieren. So könnten beispielsweise die Teilnehmer an einer Videokonferenz aus den verschiedenen Studio-Umgebungen ausgestanzt und in einem virtuellen Studio vereint werden. Bei dieser Technik kommt die Bluescreen-Technik zum Einsatz, die in der Studioteknik als Blue-Keying bekannt ist.

Waren MPEG-1 und MPEG-2 noch Standards, um ausschließlich Audio- und Videodaten zu codieren, war MPEG-4 von Anfang an als "Toolbox" zur Codierung und Übertragung von (nicht notwendigerweise) interaktiven audiovisuellen Inhalten geplant. Wenn im folgenden MPEG-4 erwähnt wird, ist immer der Szenenbeschreibungs-Aspekt gemeint, d.h. die audiovisuellen Inhalte werden als sog. Szene beschrieben, die aus natürlichen (z.B. Videosequenzen oder aufgenommene Audiospuren) und synthetischen Objekten (z.B. 2D- und 3D-Grafiken) bestehen kann. Dazu kommt in MPEG-4 erstmalig das Konzept der objektbasierten Codierung zum Einsatz, bei dem die einzelnen Objekte in einer Szene (z.B. Hintergrund, Darsteller, 3D- und 2D-Objekte, Sprache, Hintergrundmusik etc.) getrennt codiert und übertragen werden. Bei der Wiedergabe wird die Szene dann wieder zusammengesetzt.

Die Vorteile der objektbasierten Codierung und Übertragung

- Sie ermöglicht eine äußerst effiziente Übertragung bzw. Speicherung von multimedialen Inhalten. So reicht es bspw. aus, den Hintergrund der Szene einmalig als Bild zu übertragen (solange die Kamera nicht bewegt wird).
- Für jedes Objekt kann ein passender Codec aus einem der MPEG-Standards verwendet werden (z.B. Sprachcodec, Still-Image-Codec, Videocodec etc.).
- Einzelne Objekte können bei der Produktion von anderen Szenen einfach wiederverwendet werden.
- Die Szene kann im Wiedergabegerät an dessen Fähigkeiten angepasst werden.
- Es sind interaktive Inhalte möglich.



Der Betrachter kann sich in dreidimensionalen Szenen eingeschränkt oder frei bewegen. Weiterhin sind Aktionen (z.B. Start eines Video- und/oder Audio-Clips, Weiterleitung auf einen Webshop, in dem man den betreffenden Artikel erhalten kann etc.) beim Anklicken eines Objektes in der Szene denkbar.

Erklären Sie kurz die Huffman-Kodierung von DC-Koeffizienten. [3.5]

Bei der Huffman-Kodierung werden Zeichenketten innerhalb der zu kodierenden Daten durch kürzere Zeichenketten ersetzt. Häufig vorkommende Zeichenketten werden durch möglichst kurze Symbole ersetzen, selten auftretende Zeichenketten durch längere Symbole.

Die innerhalb von JPEG verwendeten Tabellen beruhen auf der Anzahl der DC-Koeffizienten mit dem Wert Null, die einem Koeffizienten mit einem Wert ungleich Null vorausgehen. Üblicherweise findet die in definierte Huffmantabelle Verwendung.

Diff Values	SSSS (number of bits needed to encode Diff)
0	0 bit
-1, 1	1 bit
-3, ..., 3	2 bits
-7, ..., 7	3 bits
"	"
"	"
"	"
-2047, ..., 2047	11 bits

Hier werden den Diff Values die benötigten Bits zur Encodierung zugeordnet (SSSS). Nun behandelt man die SSSS (Anzahl der für die Diff-Kodierung benötigten bits) wie Huffmancode Symbole, wobei deren Wahrscheinlichkeit p für alle 12 Varianten (0 -11 bits) ausgerechnet und dadurch ein Huffman-Baum erstellt wird.

Für jede Kategorie wird ein zusätzliches Bitfeld angefügt um die Differenz, die in der Kategorie tatsächlich auftrat, eindeutig zu identifizieren.

Beispiel

Wenn SSSS=2 den Huffman Code 001 und die Differenz -3 hat → sende 00100
001: HuffmanCode für 2. 00 entspricht -3 in Zweierkomplementdarstellung

Erklären Sie das Prinzip der Lauflängenkodierung (Run Length Encoding). [1.5]

Die Run-Length Codierung ist ein Kompressionsverfahren zur verlustfreien(!) Kompression von Daten. Es nutzt sich wiederholende Zahlenwerte in Folge, die es mit der Anzahl der Zahlenwerte zusammen als kurzen Code hinterlegt. Es ist ein Entropy-coding Algorithmus, inhaltsabhängige Codierung.

Ersetzung einer Sequenz gleicher aufeinanderfolgender Bytes durch deren Anzahl, gekennzeichnet durch ein spezielles Flag!

Methode

1. Wenn die selben Bytes mindestens 4 mal hintereinander vorkommen, zähle die Anzahl, wie oft sie hintereinander vorkommen
2. Schreibe die komprimierten Daten in der Form: gezähltes Byte + '!' + Vorkommensanzahl

Beispiel

unkomprimiert: ABCCCCCCCCDEFFFFGGG

komprimiert: ABC!9DEF!4GGG

Was sind Quantisierungstabellen, wie sind sie aufgebaut und weshalb setzt man sie ein? [2]

Quantisierungstabelle

Eine Quantisierungstabelle ist eine 8x8 Pixel große Matrix (hat also 64 Elemente), mit einem DC-Koeffizienten (erstes Pixel links oben) und 63 AC-Koeffizienten. Bei der JPEG- Quantisierung wird für jeden Wert dieser Tabelle 8 Bit verwendet.

Die Werte oben links sind kleiner, da sie den niedrigeren Frequenzen zugeordnet werden und damit werden die niedrigeren Frequenzen im Vergleich weniger stark komprimiert. Dies hat den Grund, dass niedrigere Frequenzen vom menschlichen Auge sehr gut wahrgenommen werden, höhere hingegen zunehmend schlechter. Daher werden hohe Frequenzen zunehmend mehr komprimiert.

Weshalb setzt man sie ein?

Zur Komprimierung der Koeffizienten bei JPEG. Die 64 Koeffizienten (=Quantisierung/Frequenzmatrix die durch DCT gewonnenen wurde) werden durch die Werte in der Quantisierungstabelle jeweils dividiert und ergeben die komprimierten Werte der einzelnen Koeffizienten (und werden anschließend auf Integer gerundet).

Was sind die Aufgaben des Synchronization Layers in MPEG-4? [1.5]

Hinzufügen von MPEG-4 spezifischen Informationen für Timing und Synchronisation der kodierten Medien. Elementare Streams werden paketierte, Header mit Timinginformation (Taktreferenzen) und Synchronisationsdaten hinzugefügt.

Die Informationen, die in den Synchronization Layers -Headers enthalten werden, behalten die korrekte Time Base für die grundlegenden Decoder und für den Empfängerterminal, plus die korrekte Synchronisierung in der Darstellung der grundlegenden Medienobjekte in der Szene bei.

Die Clock Referenzen Mechanismus unterstützt Timing des Systems, und das Mechanismus der Time Stamps. Unterstützt Synchronisierung der unterschiedlichen Medien.

Was versteht man unter Perceptual Measurement Techniques und wofür werden sie eingesetzt? [1.5]

Zur Messung der Qualität von Codecs.

- State-of-art ist nicht ausreichend, um große Skala und die gut vorbereitete Hörtests zu überholen.
- Jedoch sind Perceptual Measurement Techniques bis zu dem Punkt gekommen, in dem sie eine sehr nützliche Ergänzung zu den Hörtests sind und sie in einigen Fällen ersetzen können.
- ITU-R Task Group 10/4 (Int. Telecommunications Union, Radiocommunications sector) entwickelte eine Empfehlung für das System PEAQ. (Perceptual Evaluation of AudioQuality)

Welchen Zweck hat die Diskrete Cosinus-Transformation in Kompressionsverfahren (z.B. JPEG)? [2] bzw. Welchen Zweck haben Transformationen in den Frequenzraum im Kontext von Kompressionsverfahren (z.B. DCT bei JPEG)? [2]

Menschen haben nur begrenzte Fähigkeit hohe Frequenzen wahrzunehmen. Bilder die ihre hohe Frequenzen verringern könnten dabei einen hohen Grad an Qualität gewinnen während sie ziemlich beträchtlich zusammengedrückt werden.

Anhand der Darstellung im Frequenzraum können zur Kompression der Daten hohe Frequenzen reduziert werden bei Beibehaltung einer sehr guten Qualität. Niedrige Frequenzen entsprechen den groben Konturen („Outlines“) von Bildobjekten, währenddessen hohe Frequenzkomponenten feine Strukturen darstellen. Dies ist möglich, weil das menschliche Auge nur begrenzt fähig ist hohe Ortsfrequenzen wahrzunehmen.

Mittels Diskreter Cosinus-Transformation werden die 8x8 Pixelblöcke eines unkomprimierten Bildes von der Zeitebene in die Frequenzebene transformiert. Dadurch erhält man Koeffizienten $S(u,v)$, wobei $S(0,0)$ die kleinste Frequenz(0) in beiden Richtungen des Blockes aufweist und damit DC-Koeffizient genannt wird; er stellt die Grundfarbe des Blocks dar. Die anderen ($S(0,1), \dots, S(7,7)$) sind die sogenannten AC-Koeffizienten, wobei die Frequenz von zumindest einer Richtung ungleich 0 ist.

Um die Hochfrequenzen zu verringern müssen die Daten nach Frequenzen sortiert werden->DC

Was sind - abgesehen von Leistungsunterschieden - die wichtigsten Unterschiede zwischen MP3- und AAC-Kodierung? [2]

MP3: MPEG-1 Layer 3

AAC: MPEG-2 Advanced Audio Coding

AAC: bessere Komprimierung, rückwärtskompatible Mehrkanalkodierung fügt die Möglichkeit von vorwärts- und rückwärtskompatibler Codierung von Mehrkanalsignalen (inklusive 5.1 Kanalkonfigurationen)

Zusätzlich erweitert AAC den MPEG-1 Audiostandard um das Kodieren mit niedrigeren Samplingfrequenzen (16kHz, 22.05kHz und 24kHz)

Was versteht man unter Differential Encoding? Nennen Sie Beispiele, wo das Verfahren eingesetzt wird. [2]

- Differential Encoding ist Teil des Sourcecodings.
- Betrachtet man eine Sequenz von Symbolen S1, S2, S3... mit von Null verschiedenen Werten, die sich aber nicht sehr stark unterscheiden, so berechne die Differenz zum vorhergehenden Wert.
- Nullen können mit Lauflängenkodierung kodiert werden.

Beispiele:

- Statische Hintergründe (Videokonferenzen, Nachrichten,...)
- MPEG verwendet die Kodierung (→ motion compensation, motion vectors): 8x8 Blöcke werden verglichen; *areas similar, only shifted, e.g., to the right (motion vector)*

Was versteht man unter Structured Audio in MPEG-4? [2.5]

- Erlaubt die Erstellung von synthetischen Sounds aus kodierten Inputdaten.
- Spezielle Synthesesprache (Structured Audio Orchestra Language – SAOL) erlaubt die Definition eines synthetischen Orchesters, dessen Instrumente Klänge wie reale Instrumente erzeugen oder vorgespeicherte Sounds verarbeiten.
- MPEG-4 standardisiert die Methode der Synthesebeschreibung, nicht die Methoden um Sound zu generieren.
- Download von Noten/Partituren in den Bitstream kontrolliert die Synthese. Noten/Partituren ähneln Skripten in einer speziellen Sprache (Structured Audio Score Language – SASL), sie bestehen aus seinem Set von Befehlen für die verschiedenen Instrumente, wobei Befehle auf verschiedene Instrumente und zu verschiedenen Zeiten angewendet werden können.

Wie erfolgt die zeitliche Komposition (Temporal Composition) in MPEG-4 Datenströmen? [2.5]

Komposition ist die Spezifikation von zeitlichen oder räumlichen Beziehungen zwischen einer Gruppe von Medienobjekten. (Also ist z.B. die räumliche Komposition eine Spezifikation von räumlichen Beziehungen zwischen einer Gruppe von Medienobjekten). Das Resultat einer Komposition ist ein Multimedia-Objekt. Die originalen Elemente werden „Komponenten“ genannt.

- Der composition Stream (BIFS) hat seine eigene Zeitbasis.
- Auch wenn die Zeitbasen für Komposition und elementaren Datenstreams verschieden sind, müssen sie konsistent sein, abgesehen von Translation und Skalierung der Zeitachse.
- Timestamps, welche bei den elementaren Medienstreams hinzugefügt werden, legen fest, zu welcher Zeit die Zugriffseinheit für ein Medienobjekt am Dekoderinput bereit sein muss und wann die Kompositionseinheit am Kompositorinput bereit sein muss. Timestamps des Kompositionstreams bestimmen, wann die Zugriffseinheiten für Komposition am Input des Kompositionsinformationdekoders bereitgestellt sein müssen.

- Felder innerhalb einer Szenenbeschreibung beinhalten auch einen Zeitwert. Er stellt entweder eine Dauer oder einen Moment als Zeit dar.
 - stellen eine relative Zeit zum Timestamp des BIFS Elementarstreams dar.
 - For example, the start of a video clip represents the relative offset between the composition time stamp of the scene and the start of the video display.

Erklären Sie kurz die wesentlichen Charakteristika des MPEG-4 Composition Streams. [3]

- The composition stream is treated differently from others because it provides the information required by the terminal to set up the scene structure and map all other elementary streams to the respective media objects.
- The composition stream (BIFS) has its own time base.
- Time stamps associated to the composition stream specify at what time the access units for composition must be ready at the input of the composition information decoder

Erklären Sie kurz den Skalierfaktor und den Zusammenhang mit Kompressionsverfahren. [2]

Skalierung ist eine Abwandlung der Bandrauschunterdrückungssysteme für Tonbandgeräte (DolbyA, B, C)

- kleine Signalanteile vor Bandaufzeichnung verstärkt
- bei Wiedergabe um selben Faktor gedämpft
- Bandrauschen mit selben Faktor gedämpft

Skalierung bei MPEG Kodierung

- kleine Werte vor Kodierung mit Skalierfaktor multipliziert
- bei Dekodierung mit selben Faktor dividiert
- Quantisierungsrauschen um Skalierfaktor gedämpft

Skalierung verbessert SNR für leise Signale

Erklären Sie kurz die Kritische Abtastung. [2]

Kritische Abtastung digitaler Filterbanken

- Wortlängen bleiben in Subbands unverändert
- Abtastrate in Subbands reziprok zur Anzahl der Subbands
- Reduzierte Abtastrate in Subbands wird als kritische Abtastrate bezeichnet

Kritische Abtastung der Subbands gewährleistet, dass Filterbank keinen Einfluss auf Datenmenge hat

- Es werden zwar mehr Kanäle erzeugt, aber deren Abtastfrequenzen sind dementsprechend geringer.
- Kritische Abtastrate ergibt sich aus Nyquist-Shannon-Theorem

Erklären Sie kurz den Maskierungseffekt und den Zusammenhang mit Kompressionsverfahren. [2]

Zwei Töne erklingen gleichzeitig mit unterschiedlicher Lautstärke und ähnlicher Frequenz.

Maskierungseffekt – Ohr hört nur den lauten Ton

- leiser Ton wird "maskiert"
- der laute Ton mindert Empfindlichkeit des Gehörs im umliegenden Frequenzbereich
- leiser Ton liegt unter neuer Hörschwellenkurve

Maskierungsschwellwert

Maskierungsschwellwert ist minimale Lautstärke, die der leise Ton aufweisen müsste, um nicht maskiert zu werden.

Allgemeines Audiosignal besteht aus vielen Frequenzanteilen

- jeder Frequenzanteil beeinflusst die Hörschwellenkurve
- Audiosignal variiert mit der Zeit
 - Hörschwelle zu jedem Zeitpunkt verschieden

Psychoakustisches Modell

- Berechnet jeweils aktuelle Hörschwellenkurve

Variable Quantisierung

- Frequenzkomponenten so quantisiert, dass Quantisierungsrauschen gerade unter aktueller Hörschwellenkurve
- Wo maskiert wird, grobe Quantisierung, kleine Wortlängen

Erklären Sie kurz Kritische Bandbreite und den Zusammenhang mit Kompressionsverfahren. [2]

- Psychoakustischer Effekt des Gehörs
- Frequenzen gleichzeitig erklingender Töne müssen bestimmten Mindestabstand haben, um als Töne verschiedener Tonhöhe wahrgenommen zu werden
- Kritische Bandbreite – umfasst jeweils einen Frequenzbereich, innerhalb welchem unser Gehör Töne nicht differenzieren kann
 - bezieht sich nur auf gleichzeitig gespielte Töne (bei sequenziell gespielten Tönen ist Frequenzauflösung des Gehörs wesentlich höher)
- elementar für viele psychoakustische Effekte und MPEG Kompression
- Töne unter 500 Hz – kritische Bandbreite konstant 100 Hz
- Töne über 500 Hz – kritische Bandbreite ca. 20% der Frequenz selbst (entspricht ca. kleiner Terz)

Daraus ergibt sich z.B. bei der MP3 Kodierung, dass man die Tatsache nutzt, dass der Mensch nur einen bestimmten Frequenzbereich wahrnehmen kann. Der nicht hörbare Bereich wird weggeschnitten. Dadurch wird die Datengröße reduziert, ohne eine echte Kompression durchgeführt zu haben.

Ähnlich verhält es sich mit der kritischen Bandbreite. Der Frequenzbereich, der vom Menschen nicht getrennt wahrgenommen werden kann, wird bei der Kompression berücksichtigt und weggeschnitten, damit man auch alle Töne getrennt wahrnehmen kann.

Retrieval und Indexing

Nennen Sie ein von geometrischen Transformationen unabhängiges Image-Feature. Diskutieren Sie kurz Vor- und Nachteile. [2]

Farbe- bzw. die Verteilung von Farbe innerhalb von Bildern. Wird in Form von Histogrammen dargestellt. Ein Histogramm gibt die Anzahl der Pixel je vorkommender Farbe an. Die Anzahl der Farben kann dabei frei festgelegt werden.

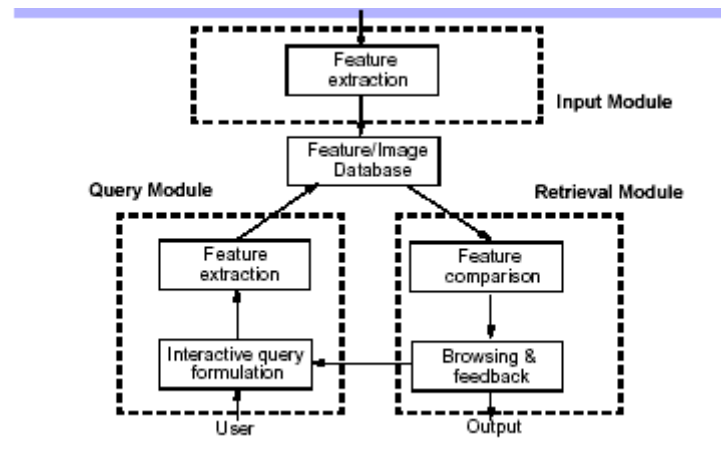
Vorteil der Histogramme ist die Unabhängigkeit gegenüber Translationen und Rotationen entlang der View-Achse. Verändert sich nur wenig, wenn sich der Ansichtswinkel oder der Maßstab verändert.

2 Histogramme können folgendermaßen verglichen werden:

- Intersection: Anzahl der Pixel mit gleicher Farbe, eventuell normalisiert auf eines der beiden Bilder.
- Distance: Alle Farben werden im Histogramm unabhängig voneinander betrachtet. D.h. keine Ähnlichkeiten zwischen Zellen. Lösung über Distanz $D(I, Q) = (I - Q)_t A(I - Q)$.

Beschreibung Sie kurz die Architektur eines Content-Based Image Retrieval-Systems [3]

Ziel ist die Suche in einer Bilddatenbank nach entsprechenden Kriterien.



1. Input Mode

Beim Import der Bilder in die Datenbank sind die entsprechenden Features zu extrahieren und zusammen mit den Bildern in der Datenbank abzulegen.

2. Query Mode

Wird nach Bildern gesucht, hat der Benutzer die entsprechenden Anfragen zu formulieren (RBR, ROA, RSC...). Dies sollte im besten Fall interaktiv geschehen können. Aus diesen Anfragen werden die für die Suche benötigten Features in der Datenbank ermittelt und als Basis für die Suche verwendet.

3. Retrieval Mode

Die von der Benutzeranfrage ermittelten Features werden mit denjenigen der Bilder in der Datenbank verglichen. Dieser Vergleich liefert mitunter eine Menge von Ergebnisbildern, die der Anfrage entsprechen. Daher kann das Ergebnis im nächsten Schritt manuell durchsucht und entsprechend weiter eingeschränkt werden (→ Neue Query). Vorgang solange, bis das gewünschte Ergebnis erreicht wurde.

Was ist die Brodatz Datenbank und weshalb ist sie von Bedeutung? [2] bzw. Beschreiben Sie kurz die Evaluierung von Textur-Features mit Hilfe der Brodatz Datenbank [3.5]

Diese DB ist der de facto Standard für die Evaluierung von Textur- Algorithmen und gibt an wie gut unterschiedliche Algorithmen Texturen in Bildern erkennen können. Sie besteht aus einem Fotoalbum mit 112 Graustufen-Bildern (512x512x8 Bit), welche unterschiedlichen Texturklassen repräsentieren.

Texturauswahl- Algorithmus

9 Subimages werden vom Zentrum jedes Brodatz-Bildes extrahiert. Die Features werden berechnet, indem das Modell aufgefordert wird, jedes der 1008 Subimages abzuschätzen. Die Distanz jedes einzelnen Test- Subimages zu allen anderen wird dann berechnet.

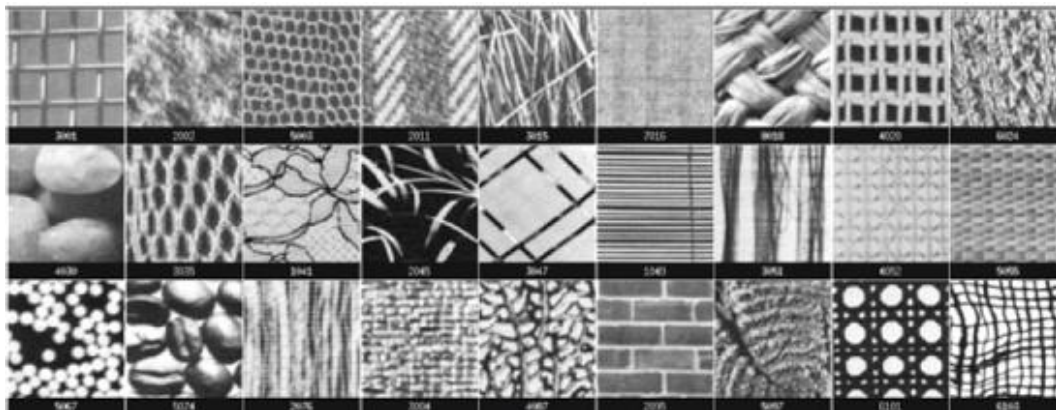
Anders gesagt: Zur Überprüfung des jeweiligen Algorithmus wird nun jedes der 112 Bilder (Klassen) in 9 kleinere Subbilder unterteilt. Man erhält nun in Summe 1008 Einzelbilder wobei jeweils 9 zusammengehören. Dabei wird auf 2 Arten die Qualität des jeweiligen Texturalgorithmus überprüft:

Retrieval Rate

Aus den 1008 Bildern wird nun „eines“ ausgewählt und als Referenz verwendet. Der jeweilige Algorithmus liefert nun eine gewisse Anzahl an Bildern zurück, die dem Referenzbild entsprechen sollen. Unter Retrieval Rate versteht man nun das Verhältnis wie viele der gefundenen Bilder sich einer Klasse (bestehend aus jeweils 9 Subbildern) zuordnen lassen. Beispiel: Werden 10 Bilder gefunden aber nur 7 dieser Bilder entsprechen einer bestimmten Klasse so spricht man von einer Retrieval Rate von 70%.

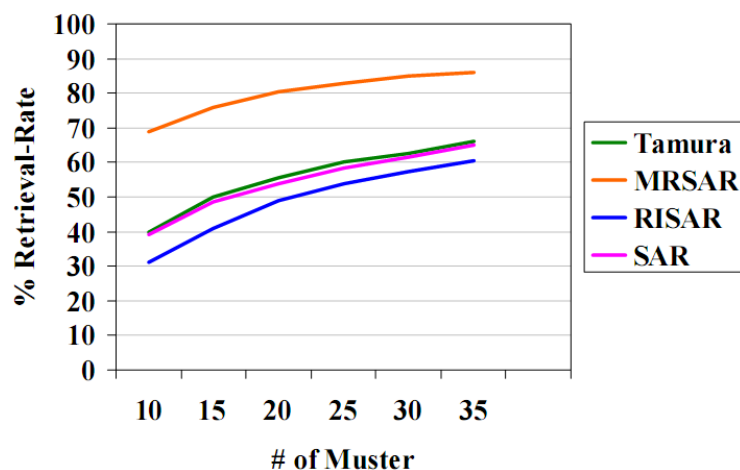
Average Retrieval Rate

Jedes Subimage wird als Test-Subimage berechnet. Also jedes der 1008 Bilder wird als Testbild (Test-Subbild) verwendet und das Ergebnis für verschiedene n-Werte dargestellt. Somit werden „alle“ 1008 Subbilder miteinander verglichen.



Die anderen Modelle im Vergleich

- Tamura:
 - **Vorteil:** Rotations invariant; Die Tamura-Features sind der menschlichen Wahrnehmung ähnlicher.
 - **Nachteil:** Hat die gleiche Genauigkeit, ist aber langsamer als SAR
- MRSAR:
 - **Vorteil:** hat die höchste Genauigkeit
 - **Nachteil:** berechnet niedrigere Auflösungen des Bildes, wenn die Imageregion zu klein ist.



MM Programmierung

Nennen und diskutieren Sie kurz zwei Gründe für das Konzept Ableitung (Derivation) [2]

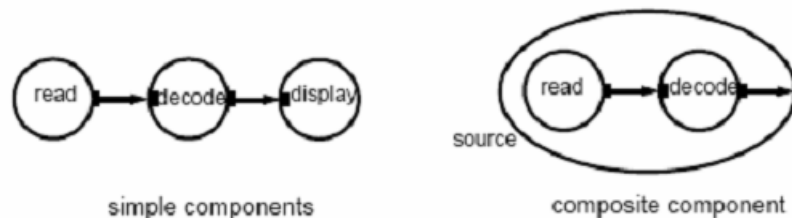
- Lagerungs („Speicher“-) und Netzwerkkosten werden reduziert. (Durch die diversen Schritte der Ableitung, also z.B. Verkettung (Konkatenation) oder Konvertierung (Konversion) wird vor allem Speicherplatz gespart.
- Modifikationen gehen effizienter vonstatten und Ableitung sorgt für Datenunabhängigkeit.

Nennen Sie 3 Gründe, weshalb die Multimedia-Programmierung fast ausschließlich dem objektorientierten Ansatz folgt. [2]

- Kapselung (Encapsulation)
- Keine Software-Altlasten oder –Überbleibsel (völlig neue Ebene)
- Möglichkeit der Erweiterung/Ausdehnung (neue Hardware, neue Media Formate, Standards, Kompressionstechniken etc.)
- Möglichkeit der Cross- Plattform Entwicklung (abstrakte Interfaces reduzieren die Plattformabhängigkeit)

Was versteht man unter dem Begriff Configuration und wie wird das Konzept im vorgestellten Framework realisiert? [2.5]

- Eine Komposition von Komponenten (MPEs), Streams und Medien zu einer Applikation
- Configuration soll das dynamische Ändern von Konfigurationen mit einer beliebigen Anzahl an Komponenten unterstützen
- Weiters soll sie Verteilung (Distribution), das Authoring und das Programmieren unterstützen.



Wozu dienen die Ports von Komponenten? Unter welchen Umständen dürfen Ports von verschiedenen Komponenten verbunden werden? [2]

Components (MPE)

- Streams werden über Ports gesendet und empfangen.
- Verbundene Ports müssen plug-kompatibel sein

Ports dürfen zu verschiedenen Komponenten verbunden sein, wenn

- einer der beiden ein Output-Port und der andere ein Input-Port sind
- beide Ports plug-kompatibel sind
- das Hinzufügen einer Connection nicht das Fan-Limit eines Ports überschreitet.

Worin unterscheiden sich Ihrer Meinung nach die im vorgestellten Framework definierten Components vom allgemeinen Begriff Komponenten wie sonst im Software Engineering verwendet? Was haben die beiden Konzepte gemein? [1.5]

Components unterscheiden sich von Komponenten durch die Ports. Diese können als Schnittstelle gesehen werden, über denen die Streams gesendet bzw. empfangen werden können. Zu jedem Component wird ein oder mehr Ports angebracht.

Was versteht man unter padding und weshalb wird padding eingesetzt? [1]

Padding ist Auffüllen von ungenutztem, leerem Platz mit leeren Daten zwecks Erfüllung von Formatvorgaben.

Anwendungsbereiche – Beispiele

Es wird zum Beispiel beim Schreiben von Audio CDs auf CD-Rs verwendet um sicherzustellen, dass alle Files die korrekte Länge haben.

Ethernet Pakete müssen ≥ 64 byte sein. Will man weniger senden, muss man sinnlose Bytes anhängen.

Was ist die Aufgabe der Klasse Transform im vorgestellten MM-Framework und was ist an den Objekten dieser Klasse bemerkenswert? [3]

Implementiert function objects (Objekte, die sowohl eine Funktion als auch ihre Parameter kapselt).

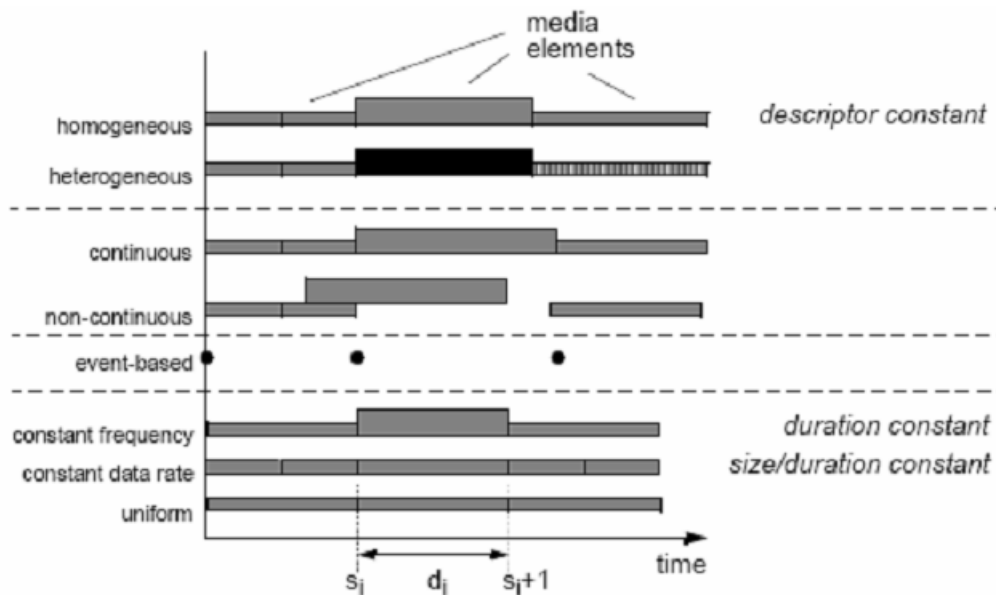
- Instanzvariablen von Transformklassen speichern Parameter.
- Eine Zuweisungsmethode „apply“ löst die Funktion aus und gibt das Ergebnis zurück.
- Alle Medienoperationen (**Medientypen**) können von function objects repräsentiert werden.
- Man kann neue Operationen hinzufügen, ohne existente Klassen und Applikationen zu modifizieren
- Effiziente Lagerung und Bearbeitung der abgeleiteten Objekte.

Nennen und erklären Sie die wichtigsten Eigenschaften, nach denen zeitabhängige Ströme (Timed Streams) klassifiziert werden können. [2.5] bzw. Erklären Sie die wichtigsten Kategorien nach denen zeitabhängige Ströme (timed streams) klassifiziert werden können. [2]

Definition: Ein Timed stream ist eine endliche Folge von Tupeln der Form $\langle e_i, s_i, d_i \rangle, i=1, \dots, n$. Jede Mediasequenz basiert auf einem Medientyp M und einem diskreten Zeitkoordinatensystem D.

Eigenschaften:

- Homogenität (Gleichwertigkeit): Die Elementdeskriptoren sind konstant.
- Kontinuität (Beständigkeit): $s_{i+1} = s_i + d_i$ für $i=1, \dots, n-1$
- Frequenz: $d_i = \text{konstant}$, für $i = 1, \dots, n-1$
- Datenrate: Elementgröße/Dauer = konstant
- Uniformität: Elementgröße und Dauer sind konstant



Was versteht man unter der "dualen Natur" zeitanhängiger Medien? [1]

Zeitbasierte Medien können gesehen werden wie:

- Pools: für medienunabhängige Operationen, z.B. kopieren
- Streams: für medienabhängige Operationen, z.B. bearbeiten („Editing“)

Wichtig für stream „view“:

- Taktinformation für inter und intra- stream Synchronisation
- Informationsgliederung von timed streams

Was versteht man unter Composition und wie wird das Konzept im vorgestellten Framework realisiert? [2] bzw. Was versteht man unter Composition? [1]

Komposition ist die Spezifikation von zeitlichen und/oder räumlichen Beziehungen (Zusammenhängen) zwischen einer Gruppe von Mediaobjekten. Das Ergebnis einer Komposition wird Multimediaobjekt genannt, und die ursprünglichen Elemente werden Komponenten genannt.

Was verstehen Sie unter abgeleiteten Objekten (derived objects), wozu dienen sie und wie werden sie im vorgestellten Framework realisiert? [2] bzw. Was versteht man unter dem Begriff Derivation und wie wird das Konzept im vorgestellten Framework realisiert? [1.5]

Die Ableitung (D) eines Medienobjekts o_1 von einem Medienobjektset O ist die eine Abbildung („Mapping“) der Form $D(O, P_D) \rightarrow o_1$, mit P_D als Set von für D spezifischen Parametern. o_1 wird dann abgeleitetes Objekt genannt. Abgeleitete Objekte verweisen auf die Medienobjekte und verwendeten Parameterwerte.

Erklären Sie kurz den Begriff Quality of Service. [1.5]

Quality of Service (QoS) ist ein gern verwendeter Begriff um die Programm- Erfordernisse bzw. Voraussetzungen für eine vorgegebene Ressource zu repräsentieren. Formal ist QoS eine Menge von Qualitätsanforderungen an das gemeinsame Verhalten bzw. Zusammenspiel von mehreren Objekten.

- Unschärfe oder statistische Beschreibung der Parameter der Service-Übergabe
- Spiegelt Toleranz in menschlicher Wahrnehmung wieder
- Ermöglicht ein Ressourcenmanagement und Trade-offs zwischen sich widersprechenden Zielen

- In allen Levels soll es kompatibel, übertragbar und zugänglich für eine systemweite Optimierung sein

Typische Quality of Service Parameter

- min. und max. Auflösung
- erlaubte Fehlerrate
- akzeptierter Jitter und Delay-Schranken

Manche Applikationen benötigen strenge oder deterministische Garantien für das Service, andere brauchen nur statistische Garantien für eine bestimmte angebotene Performance.

Quality of Service Parameter für Media Server

- **Bandbreitenanforderungen:** Die Zustellung von spezifischen Datenraten (muss nicht konstant sein, durchschnittliche Datenrate)
- **Peak-Bandbreitenanforderungen:** Maximale Dauer der Spitzenbandbreite
- **Burst size oder Workahead:** Maximale Datenmenge, bei der der Stream eine strikt konstante Rate übersteigen kann
- **Verzögerung:** Maximale Verzögerung
- **Verlustwahrscheinlichkeit:** Maximale tolerierbare Verlustwahrscheinlichkeit (Kontrollinformation ist empfindlicher)

Was ist ein Framework und welche Anforderungen werden an ein Multimediaframework gestellt? [2]

Spezielle Art von Klassenbibliothek, die eine Architektur für bestimmte Applikationsklassen bereitstellt. Framework gibt Struktur und Verantwortlichkeit von Klassen/ Objekten, die Interobjekt-Kommunikation und den Kontrollfluss vor.

Hauptprogramm (Anwendungslogik) ist bereits vorgegeben und ruft die Erweiterung des Benutzers auf.

Frameworks bieten eine allgemeine Architektur für die definierte Kooperation und Kommunikation von Objekten, bieten Templates für eine Gruppe von Objekten, welche gemeinsam die Verantwortung für einen bestimmten Problembereich managen, legen vordesignte Lösungen für eine Klasse von Problemen fest. Trennen die generischen und spezifischen Teile einer Lösung und strukturieren generische Parts zu kollaborierten Objekten

Wörtlich übersetzt bedeutet Framework (Programm-)Gerüst, Rahmen oder Skelett. Darin wird ausgedrückt, dass ein Framework in der Regel eine Anwendungsarchitektur vorgibt. Dabei findet eine Umkehrung der Kontrolle statt: Der Programmierer registriert konkrete Implementierungen, die dann durch das Framework gesteuert und benutzt werden, statt – wie bei einer Klassenbibliothek – lediglich Klassen und Funktionen zu benutzen. Wird das Registrieren der konkreten Klassen nicht fest im Programmcode verankert, sondern wird „von außen“ konfiguriert, so spricht man auch von Dependency Injection.

Ein Framework definiert insbesondere den Kontrollfluss der Anwendung und die Schnittstellen für die konkreten Klassen, die vom Programmierer erstellt und registriert werden müssen. Frameworks werden also im Allgemeinen mit dem Ziel einer Wiederverwendung „architektonischer Muster“ entwickelt und genutzt. Da solche Muster nicht ohne die Berücksichtigung einer konkreten Anwendungsdomäne entworfen werden können, sind Frameworks meist domänenspezifisch oder doch auf einen bestimmten Anwendungstyp beschränkt. Beispiele sind Frameworks für grafische Editoren, Buchhaltungssysteme oder elektronische Warenhäuser im World Wide Web.

Bestandteile eines Frameworks

Ein Framework besteht aus Media-, Transform-, Format- und Component-Klassen und stellt vorher designte Lösungen zu Problemen dar.

Anforderungen an ein Multimediaframework

- Wirtschaftlichkeit der Konzepte
- Offen – es sollte möglich sein neue Medientypen, neue Datenrepräsentationen und neue Hardwaremöglichkeiten zu implementieren.
- Queryable – es soll Interfaces für Abfrageumgebungen spezifizieren unter Berücksichtigung derer Möglichkeiten.
- Distribution (Verteilung) – es soll helfen Applikationen zu partitionieren um Verteilung zu ermöglichen.
- Skalierbarkeit – es soll skalierbare Medienrepräsentation unterstützen.
- High-level Interfaces – für Synchronisation, Medienkomposition, Gerätesteuerung, Datenbankintegration und simultane Medienprocessingaktivitäten (gleichzeitig laufende Verarbeitungsaktivitäten).

Was leistet die Programmierungsabstraktion Interpretation? Nennen Sie 2 Sachverhalte, die die Interpretation erschweren. [2.5]

Aufbereitung der Daten sodass man damit arbeiten kann (daher Kapselung der phys. Speicherung). Eine Interpretation eines Streams(BLOB) S, ist ein Abbild von S auf ein Set von Media Objekten (Mediaelementen).

Sachverhalte, die die Interpretation erschweren:

- Uneinheitlichkeit (Heterogenität)
 - o variabel-sortierte Elemente
 - o mehrere Kodierungsmethoden und Parameter
 - o Verschiedengroße Elemente
- die gestörten (out-of-order) Elemente
- interleaving und padding
- das zerstörungsfreie Editing (Bearbeitung) und die Skalierbarkeit

In welchen Klassen des Frameworks werden function objects implementiert? Welchem Zweck dienen diese Klassen? [2] bzw. Was sind function objects? In welchen Klassen des Frameworks werden sie implementiert?

Die Transformklasse implementiert function objects.

Function objects sind Objekte, die sowohl Funktionen als auch Parameter umschließen. Jede Art von Media-Operationen (Konvertierungen, generische oder zeitliche Transformationen etc.) können durch function objects repräsentiert werden.

Function Objects sollten mit Vorsicht verwendet werden, denn Leistungsverlust und möglicher Verlust des Type-Checkings können die Folge sein.

Beschreiben Sie kurz die allgemeinen Kategorien, mit denen zeitabhängige Ströme unterschieden werden können. Ordnen Sie diese Kategorien MPEG-1 Video und CD-Audio zu. [3.5]

???

MPEG-7, MPEG-21 und Content Description

Erklären Sie kurz die Konzepte Rights Data Dictionary (RDD) und Rights Expression Language (REL). In welchem Kontext treten sie auf und wie spielen sie zusammen? [3.5]

Rights Data Dictionary (RDD)

Ein Verzeichnis mit Schlüsselbestimmungen, die benötigt werden um Rechte von denen, die Distribution steuern zu beschreiben, inklusive intellektueller Property Rechte. Rechte sind eindeutig ausgedrückt durch die Verwendung einer standardisierten syntaktischen Konvention und kann auf alle Bereiche angewendet werden. RDD spezifiziert wie weit Bestimmungen definiert werden unter Steuerung/Herrschaft einer Registrierungsautorität (registration authority).

RDD System unterstützt außerdem die Abbildung und Transformation von Metadaten einer Terminologie eines Namespaces in das eines anderen.

Rights Expression Language (REL)

Maschinenlesbare Sprache, die Rechte und Permissions („Erlaubnis“) unter Verwendung von Bedingungen („Terme“), wie sie im RDD definiert sind, deklariert

Bietet flexible und interoperable Mechanismen zur Unterstützung von transparenten und erweiterten Verwendung von digitalen Ressourcen. Schützt digitalen Content und anerkennt die Rechte, Konditionen und Kosten, die für digitalen Content spezifiziert sind. Beabsichtigt weiters das Unterstützen von Zugangsspezifikation und Verwendungskontrollen für digitalen Content in Fällen, wo finanzieller Austausch/Wechsel/Markt nicht Teil der Nutzungsbestimmungen ist und das Unterstützen des Austauschs sensiblen und privaten digitalen Contents.

Erklären Sie kurz die Ziele von MPEG-7 [2]

- Weniger als Kompressionsformat gedacht
- Beschreibung von Multimedia-Inhalten durch Metadaten, um das Suchverfahren, Anzeigestrategien, aber auch Konsistenzprüfungen zu verbessern (z.B. Vision wäre es Videos nach Schlagwörtern zu suchen. Wenn man z.B. als Suchwort „gewaltfreie Szenen“ eingibt, liefert eine Suchmaschine alle Videos, die dem Kriterium entsprechen, deswegen werden auch Metadaten benötigt)
- Flexibilität im Datenmanagement
- Globalisierung und Interoperabilität von Datenressourcen

MPEG-7 bestrebt folgende Punkte zu standardisieren

- Ein Set von Description Schemes und Descriptors um Daten zu beschreiben
- Eine Sprache um Description Schemes zu spezifizieren, wie die Description Definition Language (DDL)
- Ein Schema zur Kodierung der Description.

Erklären Sie kurz die Ziele von MPEG-21. [2]

- Ermöglichung von transparenter und erweiterter Verwendung von MM-Ressourcen über eine große Auswahl von Netzwerken und Devices
- Kontrolle/Betrachtung von Elementen, die entweder existieren oder in Entwicklung sind um eine Infrastruktur für die Übermittlung und den Konsum von MM-Content aufzubauen

Ziele

- Verstehen, wie Elemente zusammenpassen

- Bestimmen neuer Standards, die benötigt werden, wenn Lücken in der Infrastruktur bestehen.
- Ausführen der Integration von verschiedenen Standards

Erklären Sie kurz die Konzepte Deskriptor, Description Scheme und DDL im Kontext von MPEG-7. [3]

Eine Description besteht aus einem Description Scheme (Struktur) und einer Reihe von Description Werten (Umschreibungen), die die Daten beschreiben. Eine Description beinhaltet oder zeigt auf Eine vollständige oder teilweise umschriebene Description Scheme.

Deskriptor(D)

Ein Deskriptor ist eine Repräsentation eines Features (Feature ~ bestimmte Eigenschaft der Daten, die etwas für irgendjemanden bedeutet) und wird dazu verwendet die unterschiedlichen Eigenschaften (Features) von multimedialen Inhalten zu beschreiben. Er definiert die Syntax und Semantik der Featurerepräsentation. Ein Deskriptor muss die Semantik des Features, den assoziierten Datentyp, erlaubte Werte und eine Interpretation der Deskriptorwerte präzise definieren. Mehrere Deskriptoren können ein einzelnes Feature repräsentieren/ beschreiben, ebenso wie multiple Deskriptoren möglich sind (z.B. Histogramm). Ein Deskriptor kann wie folgt aussehen: <name, typekind, spec> value </name>

Beispiele:

- Color: string.
- RGB-Color: [integer, integer, integer]

Description Scheme (DS)

Ein Description Scheme spezifiziert die Struktur und Semantik (Bedeutung) von Beziehungen zwischen seinen Komponenten, welche sowohl Deskriptoren als auch Description Schemes sein können. Der Unterschied zwischen Deskriptor und Description Scheme ist, dass ein Deskriptor sich mit der Repräsentation (Beschreibung) eines Features beschäftigt, das Description Schema jedoch mit der Struktur einer Featurebeschreibung.

Description Definition Language (DDL)

Beschreibt die Syntax der Description Scheme. Description Definition Language ermöglicht Erstellung neuer Description Schemas und Deskriptoren. Außerdem können damit bestehende Description Schemas erweitert und modifiziert werden. DDL basiert auf XML, ist aber keine Modellierungssprache! Sie stellt den Kern von MPEG-7 dar.

Folgende Anforderungen werden abgedeckt

- **Compositional capabilities:** Möglichkeiten der Entwicklung: Es müssen neue Description-Schemata und Deskriptoren entwickelt werden können, aus vielen alten müssen neue generiert werden, die auch dann alle MPEG-7 kompatibel sein sollten.
- **Transformational capabilities:** Wiederverwendung, Erweiterung und Vererbung von existierenden DS und Deskriptoren.
- **Unique identification:** DDL bietet Mechanismen zur eindeutigen Identifizierung von D-Schemata und Deskriptoren für eindeutige Verweise.
- **Data types:** Datentypen müssen definiert sein: Primitives bereitstellen - Mechanismus um Deskriptoren mit verschiedenen Medientypen zu kombinieren: z.B. Text, integer, time/time index. Muss einen Mechanismus anbieten um Deskriptoren Daten mehrerer Medientypen mit eigener Struktur (audio, video, av-presentations,...) zuzuordnen oder auch histogram, graph, RGB, audio, audio-visual presentation, triggers,....
- **Relationships** innerhalb von DS und zwischen DS

- **Relationship** zwischen Beschreibung und Daten.

Erklären Sie sehr kurz die Terminal Architektur von MPEG-7. [4]

Das Terminal ist die Einheit, welche kodierte Repräsentationen von MM-Content nutzt. Dabei kann es sich um eine stand-alone-Applikation oder um einen Teil eines Applikationssystems handeln:

- **Application**
- Transmission/Storage medium: verweist auf die niedrigeren Layer der Übermittlungsinfrastruktur (Netzwerk Layer und darunter genauso wie Speicherung); diese Layer übermitteln gemultiplexte Streams zum Delivery Layer.
- Delivery Layer: umfasst Mechanismen für Synchronisation, Framing und Multiplexen von MPEG-7 Content.
- Compression Layer: der Fluss von Access Units (entweder textuell oder binär kodiert) wird geparts und die Content description wird rekonstruiert.

Infos

Über die Ausarbeitung

Die Fragen stammen aus dem [Informatik Forum](#), [MTB](#), [LVA Homepage](#), von Kollegen die die Prüfung machten bzw. aus einer Sammlung die ich mir über die Jahre anlegte.

Ich habe die Ausarbeitung so gut es geht gemacht, aber trotzdem können sich Fehler einschleichen! Falls man welche findet, bitte per [E-Mail](#) oder [PM](#) an mich weiter leiten damit ich sie ausbessere!

Die Antworten stammen teilweise aus Ausarbeitungen von anderen Studenten (die richtigen Namen standen nicht immer dabei. Also bitte einfach bei mir melden, damit ich die Namen dann unten dazu geben kann zu den anderen!) und Ausarbeitung von mir. Ich habe aber oft die Antworten aus diversen Ausarbeitungen, meine Antworten,... zusammengefasst, sodass man das ganze Wissen konzentriert in der Ausarbeitung hat. Den Rest der Fragen die unbeantwortet sind, versuche ich von Version zu Version weiter auszuarbeiten und zu beantworten.

Da ich es gut finde alle POs, Lösungen, Ausarbeitungen,... in einer Datei gesammelt zu haben, damit man nur die lernen braucht und auch keine Redundanzen hat, machte ich diese Ausarbeitung.

Bei einigen Sachen war ich mir nicht ganz sicher ob sie so stimmen. Deswegen habe ich diese Sätze rot geschrieben (zum leichten erkennen das man da „aufpassen“ sollte, da es womöglich nicht stimmt). Grün ist wenn es mehrere Antworten gibt und die noch zusammengefasst werden müssen, oder der Text auf Englisch ist.

Inkludierte Prüfungsangaben

5. Februar 2003, 20. März 2003, 6. Juni 2003, 2. Oktober 2003, 29. Januar 2004, 29. April 2004, 29. Oktober 2004, 10. Dezember 2004, 25. Januar 2005, 7. April 2005, 23. Juni 2005, 20. Oktober 2005, 27. Januar 2006, 24. März 2006, 30. Mai 2006, 9. Oktober 2006, 26. Jänner 2007, 2. April 2008, 7. Mai 2008.

Falls jemand Angaben / Fragen hat die hier nicht zu finden sind, wäre es auch sehr nett sie mir zukommen zu lassen!

Zusätzliche Informationen

Version: 0.9
LVA Webseite: http://www.ims.tuwien.ac.at/teaching_detail.php?ims_id=188135
Neuste Version: <http://stud4.tuwien.ac.at/~e0402913/uni.html>
Ausarbeitung: Martin Tintel (mtintel)
Ausarbeitungen auf die Marcel Nürnberg (multimedia%20i%20fragenkatalog.pdf)
aufgebaut wird: Elisabeth Wetzinger (mm1_po ausarbeitung.pdf)
Martin Tintel (Pruefungordner 2004- 2005 leicht überarbeitet.doc)
Antworten von [Christoph R.](#) aus dem Forum
Quellen: <http://de.wikipedia.org/wiki/YCbCr-Farbmodell>
<http://www.slashcam.de/multi/Glossar/-Buchstabe--4.html#2>
<http://de.wikipedia.org/wiki/BIFS>
<http://de.wikipedia.org/wiki/MPEG-2>
<http://de.wikipedia.org/wiki/JPEG>