

ROZDZIAŁ 8. WPROWADZENIE DO KOMPRESJI STRATNEJ

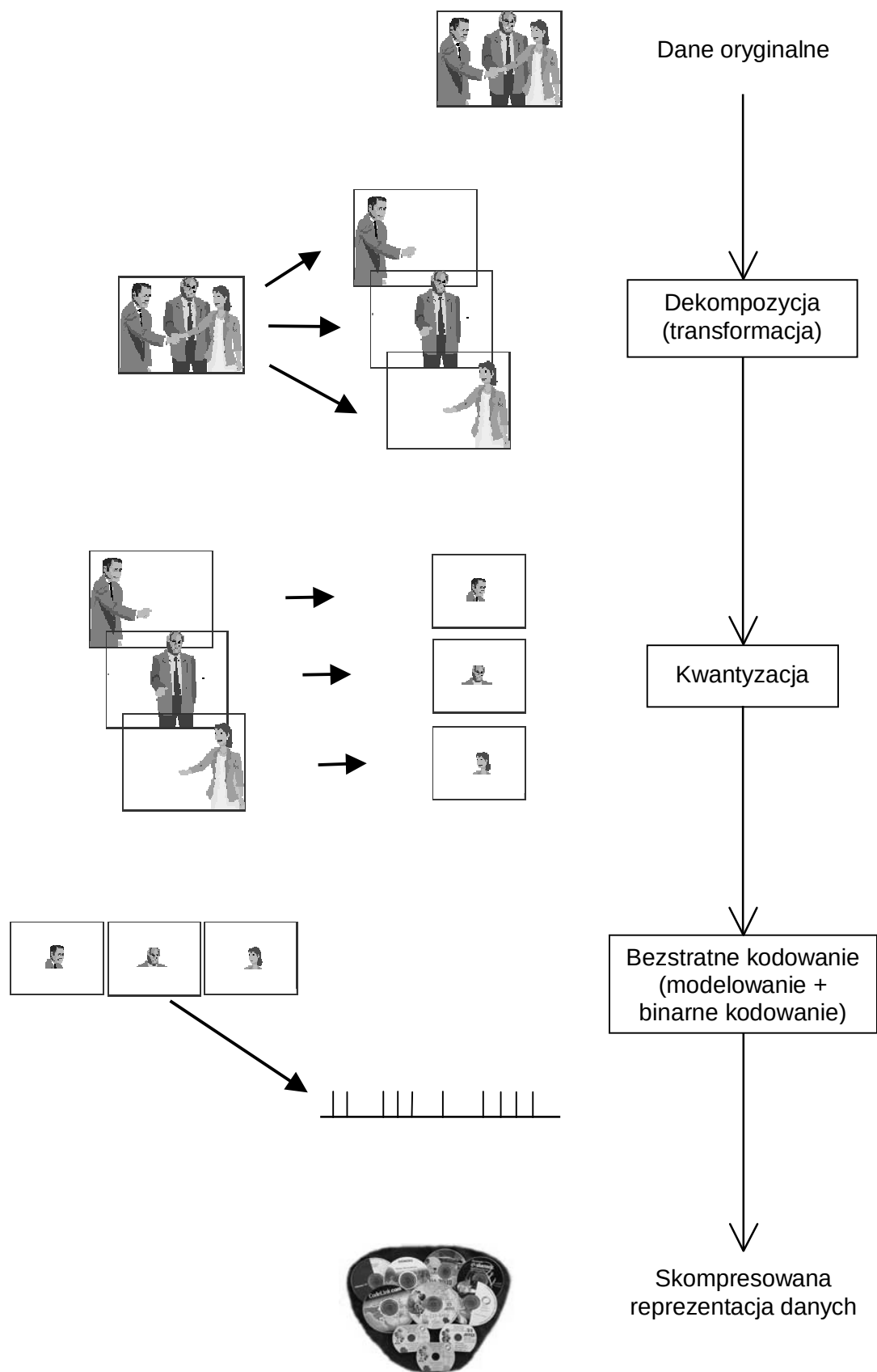
Potrzeba kompresji ogromnych ilości danych zmusza do poszukiwań efektywniejszych od algorytmów bezstratnych rozwiązań. Skuteczność metod odwracalnych w zastosowaniu do kompresji różnego typu zbiorów danych okazuje się ze względu na ograniczenia niewystarczająca w wielu przypadkach. Metody entropijne są często nieefektywne ze względu na brak dobrych modeli statystycznych, szybko adaptujących się do niestacjonarnego charakteru strumienia danych, a wszystkie estymacje globalne przy założeniu stacjonarności źródła informacji zawodzą. Podobnie metody słownikowe są momentami bezradne wobec trudno przewidywalnych zmian własności kompresowanego strumienia. Brak powtarzających się długich, identycznych sekwencji w ciągu danych wejściowych znacznie ogranicza efekty kompresji. Dalsza minimalizacja długości reprezentacji danych możliwa jest poprzez jakościową zmianę koncepcji kompresji i wprowadzenie nieodwracalnego schematu rekonstrukcji sygnału (zbioru, źródła) oryginalnego z uproszczonej reprezentacji wyjściowej koderu stratnego.

Stratne metody kompresji, redukujące informację w stopniu dopuszczalnym dla konkretnych zastosowań (zachowujące np. naturalny charakter danych, istotne szczegóły i wrażenie ciągłości obiektów w obrazach, rozróżnialność poszczególnych słów w zapisie mowy itd.), pozwalają uzyskać znacznie wyższe stopnie kompresji. Stają się więc szansą korzystniejszego rozwiązania problemów dotyczących archiwizacji i transmisji dużych zbiorów danych.

Schemat algorytmów kompresji stratnej w najogólniejszej postaci zawiera dwa kluczowe elementy: kwantyzację oraz bezstratne kodowanie (modelowanie plus binarne kodowanie) kwantowanych wartości. To uzupełnienie procedur kompresji o proces kwantyzacji, wprowadzający pewną redukcję treści czy informacji zawartej w zbiorze danych w celu zwiększenia stopnia kompresji, stanowi centralne zagadnienie przy konstruowaniu algorytmów stratnych. Okazuje się, że w praktycznych realizacjach wygodniej jest poprzedzić kwantyzację wstępną dekompozycją danych, co znacznie ułatwia opracowanie efektywnych kwantyzatorów danych. Stąd bardziej praktyczny schemat stratnej kompresji wygląda jak na rys. 8.1. Znaczenie poszczególnych bloków schematu, sposób ich realizacji, wzajemnych powiązań zmienia się w różnych technikach, określony zbiór istotnych cech jest jednak wspólny niemal każdej metodzie nieodwracalnej kompresji.

Pierwszy element schematu to dekompozycja danych oryginalnych (tj. transformacja lub modelowanie w dziedzinie oryginalnej) do pośredniej reprezentacji, znacznie bardziej podatnej na efektywną kwantyzację i kodowanie. Zasadniczym celem tego etapu jest redukcja nadmiarowości poprzez oszczędny opis danych w nowej bazie reprezentacji pośredniej. Można to uzyskać poprzez upakowanie całej informacji w niewielkim obszarze przestrzeni transformaty, zmniejszenie korelacji pomiędzy wartościami danych, ograniczenie dynamiki danych, zmniejszenie wymiarowości przestrzeni danych, modelowanie wzajemnych zależności treściowych pomiędzy pewnymi obszarami w dziedzinie określoności danych, itp.

Poprzez kwantyzację ogranicza się liczbę symboli alfabetu strumienia danych reprezentacji pośredniej, który jest podawany na koder bezstratny. W konsekwencji następuje nieodwracalne zniekształcenie oryginalnej reprezentacji zbioru danych, co zazwyczaj powoduje redukcję realnej informacji kodowanego źródła. Redukcja ta wykorzystuje często wszelkie dostępne informacje na temat przetwarzanego strumienia danych, charakteru najistotniejszej części informacji, potencjalnych własności informacji nieistotnej w danej aplikacji, itd. Ponadto sposób przeprowadzenia kwantyzacji winien uwzględniać podatność na



Rys.8.1. Ogólny schemat metod stratnej kompresji.

bezstratne kodowanie tworzonego zbioru danych. Bardziej inteligentne rozwiązania uwzględniają więc nie tylko znaczenie jakościowe usuwanej informacji, ale także związany z tym zysk zmniejszenia długości ostatecznej reprezentacji poprzez minimalizację wartości entropii strumienia wyjściowego przy danym poziomie zniekształceń. Rodzaj i poziom kwantyzacji ma więc decydujący wpływ zarówno na uzyskiwany stopień kompresji, jak i na jakość oraz wielkość strat w procesie kompresji. Operator kwantyzacji działa w przestrzeni danych przygotowanej przez wstępną metodę dekompozycji, powinien więc być projektowany z uwzględnieniem specyfiki tej przestrzeni, a czasami wręcz wspólnie z fazą dekompozycji.

Ostatni etap kodowania kwantowanych danych sprowadza się do aplikacji możliwie efektywnej techniki odwracalnej. Zastosowany algorytm zawiera najczęściej fazę modelowania (tym razem w dziedzinie reprezentacji pośredniej, po kwantyzacji), usuwając resztki zależności w zbiorze danych kwantowanych, oraz binarnego kodowania. Rodzaj koderów jest bardzo ściśle związany z własnościami strumienia danych skwantowanych. Może więc być to koder binarny, słowowy, wielostrumieniowy (np. osobno kodowane są znaki i moduły wartości), adaptacyjny lub statyczny, o znanym lub nie modelu statystycznym, itp.

Generalnie, każdy ze składników schematu stratnej kompresji może być zaimplementowany w wersji adaptacyjnej lub statycznej. Algorytm jest adaptacyjny, jeśli struktura jego elementów, bądź ich parametrów zmienia się lokalnie w ramach strumienia danych w oparciu o estymację lokalnej statystyki strumienia. Dynamiczna modyfikacja algorytmu pozwala na zwiększenie efektywności kosztem wzrostu stopnia złożoności schematu kodującego. Może występować w konwencji przyczynowej i nieprzyczynowej (adaptacja wstecz i wprzód - patrz rozdz. 6.2).

W rozdziale tym zdefiniowano zagadnienie kwantyzacji prezentując przy tym podstawowe schematy kwantyzacji skalarnej i omawiając zasadnicze wady i zalety kwantyzacji wektorowej. Następnie krótko scharakteryzowano rolę wstępnej dekompozycji danych oraz odwracalnego kodowania w algorytmach kompresji stratnej. W kolejnych trzech podrozdziałach przedstawiono dla przykładu proste techniki kompresji stratnej, zwracając uwagę na różnorodność rozwiązań problemu kwantyzacji i towarzyszące temu efekty zniekształceń. Rozważania zakończono kilkoma wskazówkami dotyczącymi wyboru optymalnej procedury kompresji nieodwracalnej.

8.1. Kwantyzacja

Kwantyzacja jest sercem algorytmów stratnej kompresji, albo inaczej osią, wokół której następuje koncentracja najlepiej dopasowanych form dekompozycji danych oryginalnych oraz odwracalnego kodowania kwantowanych wartości. Według rozważań teoretycznych narosłych wokół Shannon'owskiej koncepcji kodowania źródeł ze stratą informacji [1], zwanych teorią zniekształceń źródeł informacji, wektorowa kwantyzacja będąca realizacją koderów i dekoderów pozwala zbliżyć się z dowolną dokładnością do granicznej skuteczności stratnej kompresji przy zwiększeniu wymiarowości wektorów do nieskończoności. Gdyby nie trudności realizacyjne, całe zadanie kompresji nieodwracalnej sprowadzałoby się więc do konstrukcji odpowiedniego schematu kwantyzacji wejściowego strumienia danych. Stąd kluczowa rola tego zagadnienia w naszych rozważaniach.

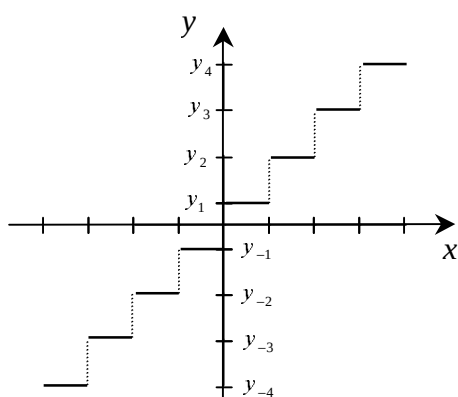
Zadaniem kwantyzacji jest reprezentacja ciągłego zbioru wartości poprzez skończoną, najlepiej niewielką liczbę dyskretnych poziomów kwantyzacji, co pociąga za sobą zmniejszenie ilości informacji wyrażonej poprzez wejściowy ciągły zbiór wartości. Inaczej

kwantyzacja może oznaczać wybór najbardziej istotnej informacji lub też najlepszego jej uogólnienia w znacznie bardziej oszczędnej reprezentacji wyjściowej kwantyzatora. Wyróżnia się też pojęcie kwantyzacji wtórnej, kiedy to dyskretny, skończony zbiór wartości sygnału jest zastępowany poprzez mniejszy zbiór poziomów dyskretnych w sposób możliwie najlepiej zachowujący informację, przy minimalizacji błędu kwantyzacji. Proces odwrotny, wykonywany w dekodерze, polegający na odtwarzaniu (rekonstrukcji) początkowego zbioru wartości sygnału na podstawie poziomów kwantyzacji, nazywany jest dekwantyzacją. Można również mówić w tym przypadku o aproksymacji sygnału wejściowego sygnałem rekonstruowanym. Zasadniczo niemożliwe jest wierne odtworzenie początkowej postaci sygnału, gdyż proces kwantyzacji jest stratny. W znaczeniu ogólnym dla wygody będziemy mówić o procesie kwantyzacji jako łączącym w sobie oba procesy: prosty (kodera) i odwrotny (dekodera).

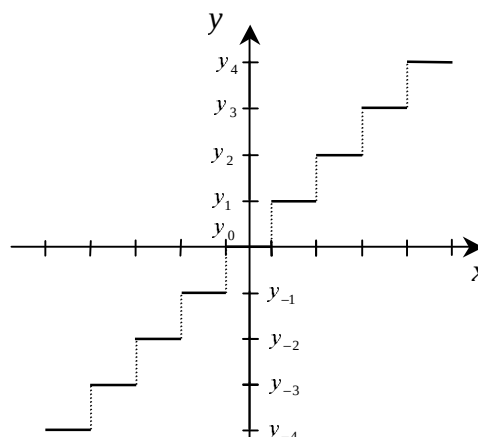
Mechanizm kwantyzacji, przedstawiony na rys. 8.2, rozumiany jest jako złożenie dwu odwzorowań:

- kodera: jako odwzorowanie nieskończonego (skończonego dużego) zbioru wartości rzeczywistych określonego i podzielonego na przedziały wzdłuż osi x , na zbiór całkowitych indeksów kwantyzacji postaci $\{..., -4, -3, -2, -1, 1, 2, 3, 4, ...\}$, które są następnie bezstratnie kodowane;
- dekodera: jako odwzorowanie zbioru indeksów w zbiór rzeczywistych wartości rekonstruowanych: $\{..., y_{-4}, y_{-3}, y_{-2}, y_{-1}, y_1, y_2, y_3, y_4, ...\}$, odpowiadających poszczególnym przedziałom z osi x .

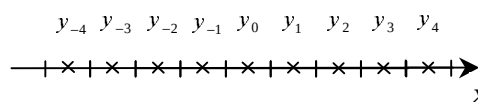
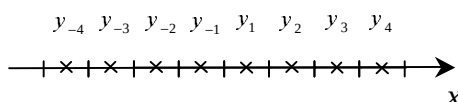
a)



b)



c)



Rys. 8.2. Wykres obrazujący działanie kwantyzatora jako funkcji odwzorowującej nieskończony zbiór wartości (oś x) w skończony zbiór dyskretny (oś y): a) bez zera, b) z zerem. Klasyczny sposób prezentacji odwzorowania argumentu w zbiór wartości funkcji przy pomocy wykresu można w tym przypadku zastąpić jedną osią x z naniesionymi wartościami rekonstrukcji na odpowiednie przedziały kwantyzacji – odpowiednio dla schematów bez zera i z zerem: c).

Kwantyzator jest jednoznacznie zdefiniowany przez zbiór przedziałów kwantyzacji (de facto zbiór punktów granicznych dzielących zakres wartości sygnału wejściowego na te przedziały) oraz zbiór wartości rekonstruowanych. Zbiór wartości rekonstruowanych winien przybliżać zbiór wartości oryginalnych w możliwie 'najlepszy' sposób, tj. minimalizujący zniekształcenie (błąd kwantyzacji) według przyjętej miary. Jeśli operator kwantyzacji oznaczmy przez $Q(\cdot)$, a ciąg wartości wejściowych (oryginalnych) przez $X = \{x_i\}_{i=1}^K$, to w wyniku procesu kwantyzacji tych wartości do M poziomów zbioru wartości rekonstruowanych $\{y_j\}_{j=1}^M$, opisanego zależnością:

$$Q: R \rightarrow R, \text{ tak że } Q(x_i) = \tilde{x}_i, \text{ przy czym } \tilde{x}_i = y_j \text{ jeśli } x_i \in [\beta_{j-1}, \beta_j), \quad (8.1)$$

a $\{\beta_j\}_{j=0}^M$ - granice przedziałów kwantyzacji $B_1, B_2, \dots, B_M$,

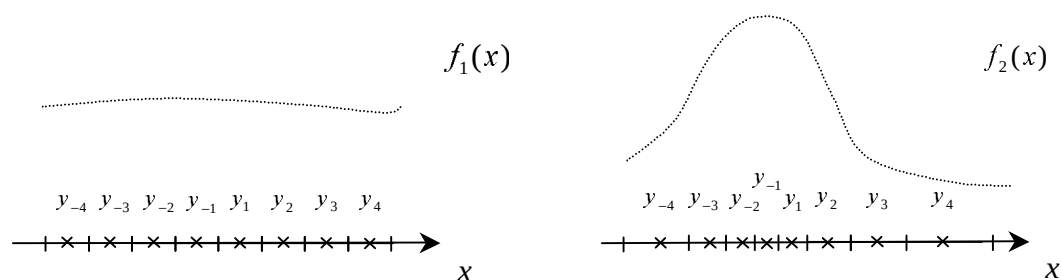
otrzymujemy ciąg wartości rekonstruowanych $\tilde{X} = \{\tilde{x}_i\}_{i=1}^K$ przybliżający sygnał oryginalny. Jakość tego przybliżenia określona jest przez błąd kwantyzacji $D_Q(X, \tilde{X}) = D_Q$. Jedną z najczęściej stosowanych miar błędu kwantyzacji jest błąd średniokwadratowy

$$D_Q = \frac{1}{K} \sum_{i=1}^K (x_i - \tilde{x}_i)^2. \quad (8.2)$$

Przyjmowane najczęściej jako miary jakości procesu kwantyzacji wielkości uśredniające błąd kwantyzacji po wszystkich próbach zmuszają do minimalizacji tego błędu w skali globalnej. Korzystniej jest więc przy założeniu danej liczby przedziałów zagęścić przedziały kwantyzacji w obszarach skali wartości licznie pokrytych przez dane wejściowe (pik na histogramie), zapewniając ich wierniejszą rekonstrukcję, a także umieścić daleko odległe pojedyncze wartości (poziomy) rekonstrukcji w zakresie, gdzie występują bardzo nieliczne wartości danych oryginalnych (łagodne zbocza, niskie tło histogramu). Nawet duże wartości pojedynczych błędów przybliżenia będą mieć wówczas niewielki wpływ na wartość globalnego błędu średniokwadratowego. Uzależnia się więc schemat kwantyzacji od postaci funkcji gęstości prawdopodobieństwa, estymowanej dla zmiennej X na podstawie wartości wejściowych (rys.8.3). Dla w przybliżeniu stałej wartości funkcji gęstości w całym zakresie określoności uzyskujemy stałą szerokość przedziałów kwantyzacji, czyli kwantyzację równomierną. Przy wyraźnym zróżnicowaniu liczby danych wejściowych w poszczególnych obszarach skali wartości pożądanym jest schemat kwantyzacji nierównomiernej, kiedy to szerokość poszczególnych przedziałów kwantyzacji jest różna.

Mozna wykazać, że optymalny model kwantyzacji spełnia następujące warunki:

- mając dane przedziały kwantyzacji (przypisane do funkcji kodera), najlepszym jest odwzorowanie liczb całkowitych indeksów kwantyzacji w zbiór wartości rekonstruowanych (funkcja dekodera), którymi są środki masy (centroidy) tych przedziałów kwantyzacji. Jest to warunek centroidu;
- mając dany zbiór poziomów rekonstrukcji (dekodek), najlepsze określenie przedziałów kwantyzacji polega na wyznaczeniu ich punktów granicznych w środku pomiędzy kolejnymi wartościami rekonstrukcji. Sprowadza się to do przypisania każdej wartości wejściowej x_i najbliższego poziomu rekonstrukcji. Ten warunek nosi nazwę najbliższego sąsiada.



Rys.8.3. Przykłady projektowania efektywnej metody kwantyzacji w zależności od postaci funkcji gęstości prawdopodobieństwa zmiennej opisującej kodowany zbiór danych: kwantyzacja równomierna (po lewej) i nierównomierna (po prawej).

Przy konstrukcji algorytmu kwantyzacji w schemacie kompresji stratnej zasadnicze cele są następujące:

- mała złożoność opisowa kwantyzatora, która nie wymaga przesyłania dużej ilości informacji dodatkowej do dekodera,
- mała złożoność obliczeniowa, aby przyjęta zasada kwantyzacji umożliwiła szybkie implementacje wykonawcze,
- niski poziom błędu kwantyzacji w stosunku do złożoności schematu i uzyskanego stopnia kompresji.

Koncepcja ta może być realizowana zarówno w wersji skalarnej, kiedy to kwantowane są wartości x_i wejściowej wielkości skalnej X , jak również w wersji wektorowej, gdzie na wejściu formowane są kolejno n -wymiarowe realizacje $\mathbf{X}_i = (x_{i1}, x_{i2}, \dots, x_{in})$ wielkości wektorowej $\mathbf{X} = (\mathbf{X}_1, \mathbf{X}_2, \dots, \mathbf{X}_K)$, a składowe wektorów kwantowane są łącznie w przestrzeni n -wymiarowej.

Schemat kwantyzacji skalarnej jest następujący:

$$\tilde{X} = Q(X) = (Q(x_1), Q(x_2), \dots, Q(x_K)) = (\tilde{x}_1, \tilde{x}_2, \dots, \tilde{x}_K), \quad (8.3)$$

podczas gdy wektorową kwantyzację można opisać następującym równaniem:

$$\begin{aligned} \tilde{\mathbf{X}} = Q(\mathbf{X}) = (Q(\mathbf{X}_1), Q(\mathbf{X}_2), \dots, Q(\mathbf{X}_K)) = \\ (Q(x_{11}, x_{12}, \dots, x_{1n}), Q(x_{21}, x_{22}, \dots, x_{2n}), \dots, Q(x_{K1}, x_{K2}, \dots, x_{Kn})) = (\tilde{\mathbf{X}}_1, \tilde{\mathbf{X}}_2, \dots, \tilde{\mathbf{X}}_K) \end{aligned} \quad (8.4)$$

Kwantyzacja skalarna

Ze względu na prostotę schematu aplikacyjnego i dużą efektywność dla źródeł danych w przybliżeniu niezależnych, kwantyzacja skalarna znajduje bardzo duże zastosowanie w metodach stratnej kompresji, praktycznie dla każdego sposobu dekompozycji danych oryginalnych. Ze względu na szerokość przedziałów kwantyzacji w danym operatorze realizującym kwantyzację alfabetu źródła informacji można wyróżnić schemat ze stałą szerokością przedziału - kwantyzację równomierną oraz schemat ze zmienną szerokością przedziału - kwantyzację nierównomierną (zobacz rys. 8.3).

Ogólnie, aby zaprojektować N -poziomowy kwantyzator skalarny trzeba w dziedzinie liczb rzeczywistych określić przedział $[a, b]$ zawierający wartości wszystkich próbek sygnału wejściowego, dokonać podziału tego przedziału na M podprzedziałów

$B_j = [\beta_{j-1}, \beta_j)$, gdzie $j = 1, 2, \dots, M$ oraz $\beta_0 = a$, $\beta_M = b$, a następnie wyznaczyć poziomy rekonstrukcji $y_1, y_2, \dots, y_M$, takie że $\forall_{j=1, \dots, M} y_j \in B_j$. Zbiory wartości $\beta_0, \dots, \beta_M$ oraz $y_1, \dots, y_M$ stanowią jednoznaczną definicję kwantyzatora skalarnego.

Czasami, jako miarę jakości kwantyzacji przyjmuje się stosunek wariancji próbek wejściowych do błędu kwantyzacji określonego równaniem (8.2). Miara ta nazywana jest stosunkiem sygnału do błędu kwantyzacji (ang. signal to quantizing noise ratio - SQNR).

Najprostszym przykładem kwantyzacji jest tzw. kwantyzator wariancji, w którym każda z próbek wejściowych jest kwantowana do jedynej wartości rekonstrukcji równej wartości średniej wejściowego strumienia próbek. Na wyjściu takiego operatora kwantyzacji pojawia się więc wartość średnia $\bar{x} = \frac{1}{K} \sum_{i=1}^K x_i$ zastępująca K wartości próbek x_i . Wartość błędu kwantyzacji jest wtedy dokładnie równa wariancji danych poddanych kwantyzacji.

Kwantyzacja równomierna

Jest to również prosty sposób rozwiązania zagadnienia kwantyzacji. Stosując kwantyzację równomierną, najlepsze efekty można uzyskać dla źródeł informacji o rozkładzie równomiernym. Kwantyzator równomierny może być zdefiniowany w sposób następujący:

- przedział $[a, b]$ określa alfabet wartości wejściowych $X = \{x_i\}_{i=1}^K$, gdzie $\forall_{1 \leq i \leq K} x_i \in [a, b]$;
- liczba przedziałów kwantyzacji M , na które dzielimy przedział $[a, b]$, wynika np. z liczby R bitów kodu przypadających na kwantowaną wartość próbki - wtedy $M = 2^R$;
- funkcja kodera wygląda wtedy jak niżej

$$d_i = \text{round} \left\{ \frac{x_i - a}{b - a} \cdot (M - 1) \right\}, \quad (8.5)$$

przy czym dyskretny alfabet całkowitych indeksów I_Q oznaczonych przez d_i jest następujący: $\{0, 1, \dots, M - 1\}$. d_i jest więc indeksem przedziału kwantyzacji B_j (pomniejszonym o jeden), do którego trafia kolejna dana wejściowa x_i .

- z kolei funkcja dekodera ma postać następującą:

$$\tilde{x}_i = a + \frac{b - a}{M} \cdot \left(d_i + \frac{1}{2}\right), \quad (8.6)$$

umieszczając poziomy rekonstrukcji w środkach M kolejnych przedziałów kwantyzacji, zgodnie z warunkiem centroidu, sprowadzającego się do warunku środka.

Błąd kwantyzacji określony równaniem (8.2) jest ograniczony w tym przypadku od góry poprzez kwadrat połowy szerokości przedziału kwantyzacji

$$D_Q \leq \left(\frac{b - a}{2M} \right)^2 \quad (8.7)$$

Kwantyzator skalarny jest w praktyce zwykle realizowany nieco prościej, z uwzględnieniem przedziałów kwantyzacji rozmieszczonych symetrycznie wokół zera. Jeśli znając maksymalną wartość próbek wejściowego źródła danych X : $x_{\max} = \max_{x_i \in X} (|x_i|)$ chcemy uzyskać zmniejszenie np. czterokrotne dynamiki danych kodowanych, wówczas ustalamy wartość współczynnika kwantyzacji $\Delta = 4$. Współczynnik ten określa szerokość przedziału

kwantyzacji i jest jedynym parametrem takiego kwantyzatora. Operator kwantyzacji w koderze jest wówczas następujący: $d_i = \text{round}\left(\frac{x_i}{\Delta}\right)$, podczas gdy w dekoderyze przybiera postać $\tilde{x}_i = d_i \cdot \Delta$. Uzyskujemy w tym schemacie zerowy przedział kwantyzacji, tj. o zerowym poziomie rekonstrukcji, o szerokości równej wartości Δ wokół zera osi wartości próbek wejściowych, czyli $\tilde{x}_i = 0$ jeśli $-\Delta/2 < x_i < \Delta/2$ (patrz rys. 8.2 b). W skrócie można mówić w tym przypadku o kwantyzatorze z zerem (ang. mid-tread quantizer). Inny podstawowy model kwantyzatora nie zawiera zerowego przedziału kwantyzacji (kwantyzator bez zera - ang. mid-riser quantizer) i jest określony przez następujące funkcje koderki: $d_i = \text{znak}(x_i) \text{round}\left(\frac{|x_i|}{\Delta} + 0.5\right)$ oraz dekodera: $\tilde{x}_i = \text{znak}(d_i)(|d_i| - 0.5)\Delta$. Został on przedstawiony na rys. 8.2 a).

W przypadku źródeł informacji, które nie mogą być przybliżone statystycznym modelem o rozkładzie równomiernym, ze względu na prostotę stosuje się niejednokrotnie również kwantyzację równomierną. Źródła te bardzo często można opisać uogólnionym rozkładem Gaussa. Najprostszy sposób podziału przedziału zmienności wartości danych generowanych przez źródło na równe odcinki - przedziały kwantyzacji - nie daje w tym przypadku najlepszych rezultatów. Lepszym rozwiązaniem w sensie średniokwadratowym będzie podział na równe odcinki jedynie tej części zakresu dynamiki sygnału wejściowego, w której koncentruje się zdecydowana większość informacji sygnału (obszar pików na histogramie danych źródła). Pozostała część zakresu bardzo słabo pokryta punktami danych będzie przybliżana punktem rekonstrukcji leżącym najbliżej, na skraju obszaru uznanego za godny dokładnej rekonstrukcji w procesie kwantyzacji. Takie rozwiązanie spowoduje znaczny wzrost maksymalnego błędu kwantyzacji, ale poprzez wierniejsze odtworzenie zasadniczej informacji zawartej w sygnale wejściowym globalna średnia wartość błędu może okazać się niewielką.

Aby wyznaczyć dla takiego równomiernego kwantyzatora szerokość przedziału kwantyzacji, trzeba wziąć pod uwagę estymatę funkcji gęstości prawdopodobieństwa źródła informacji. Posłużmy się w tym przypadku modelem źródła informacji X o wartościach ciągłych opisanym funkcją gęstości prawdopodobieństwa f_X . W takim przypadku średniokwadratowy błąd kwantyzacji można wyrazić równaniem:

$$D_{Q(X)} = D_Q = \sum_{j=1}^M \int_{\beta_{j-1}}^{\beta_j} (x - y_j)^2 f_X(x) dx. \quad (8.8)$$

Wyrażenie to trzeba przekształcić jako funkcję szerokości przedziału kwantyzacji $D_Q(\Delta)$ w przypadku rozważanego kwantyzatora równomiernego. Przyjmujemy, że projektowany kwantyzator ma M poziomów i jest symetryczny, tak jak funkcja gęstości prawdopodobieństwa opisująca źródła gaussowskie. Jest to wygodne i często użyteczne w praktyce uproszczenie. Można wtedy policzyć błąd jedynie dla wartości dodatnich i podwoić go, zarówno w części informacji znaczącej (pik w estymacie funkcji gęstości) $D_{Q_{zn}}(\Delta)$, jak i informacji nieznaczącej (płaski, niski poziom funkcji gęstości) $D_{Q_{nz}}(\Delta)$. Błąd kwantyzacji informacji znaczącej nazywany jest często szumem ziarnistym, a informacji nieznaczącej - szumem przeładowania (ang. overload). Wokół zera symetrycznie rozmieszczone przedziały kwantyzacji informacji znaczącej są następujące: $[-\frac{M}{2}\Delta, -\frac{M-2}{2}\Delta), \dots, [-2\Delta, -\Delta),$

$[-\Delta, 0), [0, \Delta), \dots, [\frac{M-2}{2}\Delta, \frac{M}{2}\Delta)$, a poziomy rekonstrukcji odpowiednio: $-\frac{M-1}{2}\Delta, \dots, -\frac{3}{2}\Delta, -\frac{1}{2}\Delta, \frac{1}{2}\Delta, \dots, \frac{M-1}{2}\Delta$. Natomiast punkty rekonstrukcji dla obszarów informacji nieznaczącej - dodatniego i ujemnego - wynoszą odpowiednio: $\frac{M-1}{2}\Delta$ oraz $-\frac{M-1}{2}\Delta$.

Wartość błędu kwantyzacji można więc zapisać jako:

$$D_Q(\Delta) = D_{Q_{zn}}(\Delta) + D_{Q_{nz}}(\Delta) = 2 \sum_{j=1}^{\frac{M}{2}-1} \int_{(j-1)\Delta}^{j\Delta} \left(x - \frac{2j-1}{2}\Delta\right)^2 f_X(x) dx + 2 \int_{\frac{(M-1)}{2}\Delta}^{\infty} \left(x - \frac{M-1}{2}\Delta\right)^2 f_X(x) dx \quad (8.9)$$

Wyrażenie na błąd kwantyzacji trzeba teraz zrózniczkować po Δ i przyrównać do zera. Problem ten można numerycznie rozwiązać bez większych problemów, wygodnie jest jednak skorzystać w algorytmie kompresji z gotowych już tablic, wyliczonych dla różnych rozkładów opisujących źródła informacji oraz różnej liczby poziomów kwantyzacji M . Przykład takiej tablicy zamieszczono poniżej.

Tablica 8.1. Wartości optymalnych przedziałów kwantyzacji równomiernego kwantyzatora strumieni danych ze źródeł informacji modelowanych rozkładem Gaussa lub Laplace'a o zerowej wartości średniej i odchyleniu standardowym równym jeden. Dane do tabeli zostały zaczerpnięte z [2].

Liczba poziomów kwantyzacji M	Źródło o rozkładzie Gaussa	Źródło o rozkładzie Laplace'a
	Szerokość przedziału kwantyzacji Δ	Szerokość przedziału kwantyzacji Δ
2	1.5960	1.4140
4	0.9957	1.0873
6	0.7334	0.8707
8	0.5860	0.7309
10	0.4908	0.6334
12	0.4238	0.5613
14	0.3739	0.5055
16	0.3352	0.4609
32	0.1881	0.2799

Kwantyzacja nierównomierna

Celem kwantyzacji nierównomiernej jest dokładniejsza kwantyzacja wartości danych, które występują częściej w strumieniu wejściowym kosztem mniejszej wierności rekonstrukcji danych pojawiających się rzadziej. Zazwyczaj miarą intensywności występowania danych w poszczególnych obszarach przedziału wartości określającego dynamikę źródła wejściowego X jest histogram h_X , estymujący funkcję gęstości prawdopodobieństwa statystycznego modelu źródła danych. Algorytm wyznaczania schematu nierównomiernej kwantyzacji wykorzystujący taką estymatę częstości wystąpień danych wejściowych przedstawiono poniżej. Znany jest on pod nazwą metody Lloyd-Maxa,

opracowanej niezależnie przez S. P. Lloyd'a [3] oraz J. Maxa [4], a której pierwotny zarys można znaleźć w polskim czasopiśmie naukowym [5]. Aby dostosować się do charakteru źródła informacji należy umieścić więcej poziomów rekonstrukcji w obszarach występowania znaczących pików histogramu. Rozmieszczenie daleko-odległych punktów rekonstrukcji w miejscach pokrytych przez znikome ilości próbek daje duże oszczędności w kodowaniu przy niewielkim błędzie kwantyzacji. Dominującym zadaniem w konstrukcji kwantyzatora nierównomiernego jest więc odpowiednie ustawienie punktów rekonstrukcji, na podstawie analizy histogramu, a następnie dopisanie do nich punktów granicznych przedziałów kwantyzacji według zasady najbliższego sąsiada. Algorytm budowania optymalnego modelu nierównomiernej kwantyzacji jest więc następujący:

Algorytm 8.1. Kwantyzacja nierównomierna metodą Lloyd-Maxa

- Inicjujemy M początkowych reprezentantów y_j będących poziomami rekonstrukcji w funkcji dekodera, wokół których rozpinamy przedziały kwantyzacji $B_j = [\beta_{j-1}, \beta_j)$, $j = 0, 1, \dots, M$, pokrywające cały przedział zmienności danych wejściowych tak, że $y_j \in [\beta_{j-1}, \beta_j)$. Zachowując warunek najbliższego sąsiada punkty graniczne wyznaczone są w środku pomiędzy kolejnymi wartościami rekonstrukcji, czyli

$$\beta_j = \frac{y_j + y_{j+1}}{2}. \quad (8.10)$$

- Na podstawie określonych przedziałów B_j obliczamy nowe punkty rekonstrukcji według zależności:

$$y_j = \frac{\sum_{x_i \in \beta_j} h_X(x_i) \cdot x_i}{\sum_{x_i \in \beta_j} h_X(x_i)}. \quad (8.11)$$

- Liczymy błąd kwantyzacji D_Q według wyrażenia $D_Q = \sum_{j=1}^M \sum_{x \in [\beta_{j-1}, \beta_j)} (x - y_j)^2$ i sprawdzamy warunek jakości kwantyzacji (np. względna zmiana wartości błędu w kolejnym kroku iteracji mniejsza od założonej wartości ε , bezwzględna wartość błędu mniejsza od założonej itp.). W przypadku spełnienia przyjętego warunku jakości algorytm jest zakończony. W przeciwnym wypadku określamy nowe przedziały kwantyzacji B_j według zależności (8.10) i powtarzamy wyznaczanie punktów rekonstrukcji zgodnie z wyrażeniem (8.11).

Okazuje się, że w kolejnych iteracjach wartość błędu kwantyzacji spełnia warunek: $D^{(l+1)} \leq D^{(l)}$ i zbiega do lokalnego minimum. Dużej wagi nabiera więc sposób inicjalizacji algorytmu, w którym możliwie najszerzej winna być wykorzystana wiedza dostępna a priori na temat źródła danych, które podlega kwantyzacji. Algorytm ten jest szczególnym, jednowymiarowym przypadkiem metody budowania książki kodowej w wektorowej kwantyzacji, zwanej algorytmem LBG, który został szerzej opisany w rozdziale dziesiątym. Do innych interesujących cech algorytmu Lloyd-Maxa należy zaliczyć:

- wartości średnie sygnału wejściowego oraz sygnału wyjściowego po kwantyzacji są sobie równe;
- wariancja danych na wyjściu kwantyzatora jest nie większa niż wariancja na jego wejściu;
- średniokwadratowy błąd kwantyzacji jest równy:

$$D_{Q(X)} = \sigma_X^2 - \sum_{j=1}^M y_j^2 \Pr(\beta_{j-1} \leq X < \beta_j), \quad (8.12)$$

gdzie σ_X^2 jest wariancją wejściowego sygnału kwantyzatora;

- sygnał wyjściowy kwantyzatora oraz szum kwantyzacji są nieskorelowane.

Kwantyzacja adaptacyjna

Konstruując schemat kwantyzacji zakłada się często stacjonarność źródła informacji lub też optymalizuje się rozwiązanie w skali globalnej, co w przypadku źródeł wyraźnie zmieniających swoje własności w funkcji czasu (chwili generacji symboli) nie pozwala uzyskać optymalnych rozwiązań z punktu widzenia efektywności kompresji tychże źródeł. Sygnał poddawany kwantyzacji bardzo rzadko wykazuje cechy chociażby przybliżonej stacjonarności. Receptą stają się wtedy dopasowane do lokalnej statystyki adaptacyjne rozwiązania kwantyzatorów w dwu podstawowych schematach: adaptacji wprzód i wstecz. W pierwszym przypadku dzielimy strumień danych wejściowych na części o zbliżonych cechach statystycznych (głównym kryterium jest zbliżona wartość wariancji), dla których niezależnie projektujemy lokalnie optymalny kwantyzator. Kosztem dodatkowej informacji przesyłanej do dekodera oraz większych kosztów obliczeniowych uzyskujemy skuteczniejszą kwantyzację. Dużego znaczenia nabiera w tych rozwiązaniach rodzaj klasyfikatora określającego wielkość i różnorodność rozróżnialnych partii strumienia z punktu widzenia statystycznych cech podatności na kwantyzację.

Innym rodzajem kwantyzacji adaptacyjnej jest metoda wstecz wykorzystująca model przyczynowy, zbudowany na wyemitowanych wcześniej danych wyjściowych kwantyzatora. Nie potrzeba w tym przypadku przysyłać dodatkowej informacji dla dekodera. Przykładowym rozwiązaniem jest kwantyzator opracowany przez N. S. Jayanta [6]. Zakłada on kwantyzację równomierną z modyfikowanym przedziałem kwantyzacji na podstawie ostatniej wartości wyjściowej kwantyzatora (model z pamięcią pierwszego rzędu). Zachowując symetryczny model kwantyzatora przyjmuje się następujący sposób modyfikacji szerokości przedziału: im bliższa zeru jest ostatnio uzyskana wartość wyjściowa po kwantyzacji, tym silniej jest zmniejszany przedział kwantyzacji. Dla większych modułów na wyjściu kwantyzatora, ale należących do zakresu informacji znaczącej, przedział również jest zmniejszany, jednak w mniejszym stopniu. Jeśli natomiast skwantowana dana trafia w obszar informacji nieznaczącej, wtedy trzeba zwiększyć nieznacznie przedział kwantyzacji. Opisuje to następujące równanie:

$$\Delta_i = c_{B(i-1)} \Delta_{i-1}, \quad (8.13)$$

gdzie Δ_i oznacza szerokość przedziału kwantyzacji dla i - tej próbki, a $c_{B(i-1)}$ określa współczynnik modyfikujący szerokość przedziału kwantyzacji Δ_{i-1} poprzedniej próbki x_{i-1} w zależności od przedziału $B(i-1)$, w który wpadła poprzednia próbka, tj. $B(i-1) = B_j$, gdzie $x_{i-1} \in B_j$.

Przykładowo, przy symetrycznym kwantyzatorze z M przedziałami kwantyzacji mamy:

$$c_{-1} = c_1, \quad c_{-2} = c_2, \quad \dots, \quad c_{-\frac{M}{2}} = c_{\frac{M}{2}} \quad \text{oraz} \quad c_j \leq c_{j+1} \quad \text{dla} \quad j > 0, \quad c_j \leq 1 \quad \text{dla} \quad |j| \leq \frac{M}{2} - 1 \quad \text{i}$$

$$c_{-\frac{M}{2}} = c_{\frac{M}{2}} > 1 \quad \text{dla} \quad |j| = \frac{M}{2}.$$

Im bardziej współczynniki c_j przypisane poszczególnym przedziałom kwantyzacji różnią się od jedności, tym algorytm szybciej się adaptuje do lokalnych zmian, ale staje się jednocześnie mniej stabilnym. Warunek stabilności tego algorytmu jest następujący:

$$\prod_{j=1}^M c_j^{n_j} = 1, \quad (8.14)$$

gdzie n_j oznacza liczbę próbek wejściowych trafiających do j -tego przedziału kwantyzacji. Aby więc zaprojektować algorytm Jayanta optymalnie, trzeba dokonać wstępnej analizy danych i wyniki tej analizy w postaci zaprojektowanego modelu adaptacji przesłać do dekodera. Jest więc tutaj połączenie cech algorytmu wstecz i wprzód. Przykładowo można zamodelować wstępnie zbiór danych wejściowych rozkładem Gaussa lub Laplace'a, dla założonej liczby poziomów kwantyzacji M przyjąć początkowy przedział kwantyzacji Δ_0 przy pomocy tabeli 8.1, a następnie szacując liczby próbek n_j ustalić wartości współczynników c_j z zachowaniem kryterium stabilności algorytmu adaptacyjnej kwantyzacji.

W przypadku kwantyzacji nierównomiernej bardziej praktycznymi są rozwiązania adaptacji wprzód z podziałem sekwencji wejściowej na wyraźnie różniące się statystycznie fragmenty. Można również wyobrazić sobie schematy hybrydowe, implementując w obszarach strumienia wejściowego o charakterystyce równomiernej kwantyzację ze stałą szerokością przedziału, a dla partii danych o zróżnicowanym, silnie nierównomiernym charakterze stosując optymalne kwantyzatory Lloyd-Maxa.

W adaptacyjnej kwantyzacji można wykorzystywać także wszelką dodatkową informację do modyfikacji schematu podstawowego. I tak w kompresji obrazów, tworząc kwantyzator wstecz, można wykorzystać estymatę znaczenia próbki kwantowanej na podstawie przyczynowego bezpośredniego sąsiedztwa wyższego rzędu w przestrzeni dwuwymiarowej. Przy pomocy kwalifikatora decydującego, czy aktualnie mamy do czynienia z fragmentem krawędzi czy raczej częścią łagodnego obszaru o znikomym poziomie informacji psychowizualnie użytecznej, można regulować wielkość przedziału kwantyzacji w zależności od treści sceny (obrazu). Przykładem bardziej złożonych metod kwantyzacji są rozwiązania stosowane w algorytmach falkowej kompresji, gdzie wykorzystywana jest dodatkowo informacja z różnych skal reprezentacji danych przestrzeń-skala, uzyskanych na skutek wielorozdzielczej dekompozycji obrazów.

Kwantyzacja optymalizowana entropią

Optymalnie zaprojektowany kwantyzator w sensie minimalnego błędu kwantyzacji przy danej liczbie przedziałów kwantyzacji nie musi być najlepszy z punktu widzenia zastosowań w algorytmach kompresji danych. Ostatecznym kryterium jakości jest bowiem uzyskanie możliwie najmniejszego błędu kwantyzacji przy określonej średniej bitowej skompresowanego strumienia danych. Można więc zwiększyć skuteczność kompresji p poprzez kodowanie indeksów przedziałów kwantyzacji, w które wpadają kolejne dane wejściowe, np. przy pomocy kodów o zmiennej długości. Stosując przykładowo efektywny koder entropijny można zniwelować mniejszą skuteczność prostych, łatwych w aplikacji schematów kwantyzacji i w efekcie uzyskać zadawalający poziom kompresji danych. Szczególnie kwantyzacja równomierna sygnałów gaussowskich daje nierównomierny rozkładów indeksów ze względu na większą liczbą indeksów przedziałów wokół wartości zerowej, co stanowi znaczącą nadmiarowość o charakterze statystycznym. Można ją zredukować nawet prostym koderem Huffmana. Po kwantyzacji nierównomiernej rozkład indeksów jest znacznie

bardziej równomierny i w tym przypadku jedynie stosowanie koderów arytmetycznych wyższych rzędów może zauważalnie poprawić skuteczność kompresji całej metody.

Takie schematy kwantyzacji nazywane są skalarnymi kwantyzatorami wymuszonymi entropią (ang. entropy-constrained scalar quantizers -ECSQ). Optymalne schematy ECSQ minimalizują średni poziom zniekształceń dla danego poziomu entropii ze skutecznością, która pozwala się przybliżyć o 1.53 dB do granicznej wartości y wynikającej z teorii stopnia zniekształceń [7].

Innym, potencjalnie lepszym rozwiązaniem jest optymalizacja kwantyzatora z kryterium uwzględniającym entropię strumienia danych wyjściowych. Wymaga ono poruszania się w przestrzeni R-D (ang. rate-distortion) ustalając równowagę pomiędzy uzyskiwaną średnią bitową strumienia (jego entropią) a wartością błędu kwantyzacji. Najlepszym rozwiązaniem jest zmniejszanie równocześnie obu tych wielkości do pewnej optymalnej wielkości, zależnej przede wszystkim od rodzaju kwantyzacji, jak również od sposobu kodowania wartości wyjściowych kwantyzatora. O istniejących ograniczeniach optymalizacji R-D koderów nieco więcej zostało napisane w podrozdziale 8.4.

Zagadnienie entropijnej optymalizacji procesu kwantyzacji jest złożone i sprowadza się do wieloparametrycznej optymalizacji procesu kwantyzacji. Aby je skutecznie rozwiązać, koniecznym jest dokonanie pewnych uproszczeń w celu opracowania niezbyt czasochłonnych algorytmów do zastosowań praktycznych. Jako miarę ilości informacji, będącej wskaźnikiem potencjalnej długości kodu wyjściowego przyjmijmy entropię bezwarunkową źródła indeksów przedziałów kwantyzacji I_Q postaci:

$$H_{I_Q} = -\sum_{j=1}^M P(d_j) \log_2 P(d_j), \quad (8.15)$$

gdzie $P(d_j) = \Pr\{X \in [\beta_{j-1}, \beta_j)\} = \Pr\{\tilde{X} = y_j\}$ oznacza prawdopodobieństwo wystąpienia indeksu d_j , czyli trafienia danej wejściowej x_i do przedziału B_j . Wprowadzenie entropii warunkowej przy jednoczesnym modelowaniu rozmiaru i kształtu kontekstu wprowadziłoby dodatkową komplikację modelu i utrudniło poszukiwanie rozwiązań optymalnych. Przyjmując dalej jako parametr wejściowy liczbę poziomów M trzeba dobrać odpowiednio granice przedziałów kwantyzacji $\{\beta_{j-1}, \beta_j)\}$ oraz wartości poziomów rekonstrukcji $\{y_j\}$. Na uzyskaną entropię strumienia wyjściowego niewątpliwie wpływ mają punkty graniczne przedziałów, decydujące o kolejności wystąpienia i rozkładzie wartości indeksów przedziałów kwantyzacji w tym strumieniu. Nie ma natomiast wpływu dla entropię zbioru wartości punktów rekonstrukcji. Zbiór $\{y_j\}$ jest jednak bardzo istotny ze względu na wartość błędu kwantyzacji, podobnie jak kształt przedziałów kwantyzacji określony przez rozkład punktów granicznych. W procesie optymalizacji kwantyzatora z kryterium minimalnej relacji pomiędzy wartościami entropii i błędu kwantyzacji można kierować się różnymi przesłankami. Przykładowo, można zrealizować zastosować optymalny algorytm Lloyda-Maxa i wyznaczyć punkty graniczne oraz poziomy rekonstrukcji minimalizujące błąd kwantyzacji przy założonej liczbie M . Następnie zakładając pewną wartość entropii, mniejszą lub równą wartości średniej bitowej prostego zapisu indeksów $R_I = \log_2 M$ bitów/symbol, modyfikujemy przedziały kwantyzacji minimalizując błąd kwantyzacji D_Q względem danej wartości entropii. Prowadzi to do rozwiązania układu równań nieliniowych dość złożonej postaci. Ostatecznie można zagadnienie to sprowadzić to iteracyjnego algorytmu będącego uogólnieniem metody Lloyda-Maxa [8].

Zasadniczo, chodzi o rozwiązanie zagadnienia, w którym dla danej wartości bitowego budżetu R_b należy dobrać granice przedziałów kwantyzacji $\{[\beta_{j-1}, \beta_j)\}$ tak, że błąd kwantyzacji D_Q osiąga wartość minimalną dla $H_{I_Q} \leq R_b$. Minimalizując błąd kwantyzacji względem danej wartości entropii równej przyjętemu budżetowi bitów R_b , rozwiązanie znajdujemy przy pomocy następującego funkcjonału, stanowiącego funkcję Lagrange'a kosztu:

$$J = D_Q + \lambda \cdot H_{I_Q}, \quad (8.16)$$

gdzie współczynnik $\lambda \geq 0$ zwany jest mnożnikiem Lagrange'a, minimalizując jego wartość. Mała wartość λ prowadzi do małego poziomu zniekształceń i dużej średniej bitowej wyjściowego strumienia, podczas gdy duża wartość λ daje małą średnią bitową oraz dużą wartość zniekształceń. Rozwiązanie tego problemu jest złożone i sprowadza się do rozwiązania $M-1$ nieliniowych równań (M – liczba przedziałów kwantyzacji), gdzie dobierając wartość λ uzyskuje się założoną R_b [2]. Aby znaleźć schematy praktyczne, stosuje się wobec tego pewne uproszczenia.

Przykładowo, zakładając postać skalarne kwantyzatora równomiernego zadanie optymalizacji można zapisać w sposób następujący:

$$\min_{\{\Delta \in Q\}} [J(\Delta) = D_Q(\Delta) + \lambda H_{I_Q}(\Delta)], \quad (8.17)$$

gdzie $J(\Delta)$ jest dwustronną funkcją kosztu, a nieujemny mnożnik λ ustala równowagę dwóch składników: błędu kwantyzacji i entropii. Jeżeli $\lambda = 0$, wtedy koszt związany jest jedynie z błędem kwantyzacji, podczas gdy dla $\lambda = \infty$ dominująca rola w kształtowaniu wartości funkcji kosztu przypada entropii. Wyznaczenie optymalnego kwantyzatora według zależności (8.17) może więc przebiegać w dwu etapach. W pierwszym z nich następuje minimalizacja funkcji kosztu poprzez poszukiwanie optymalnej wartości przedziału kwantyzacji Δ dla stałej wartości λ oraz stałej wartości entropii. Następnie tak dobieramy wartość mnożnika, że średnia bitowa wyjściowego strumienia danych osiąga dokładnie założoną wartość R_b .

Problem konstrukcji skalarne kwantyzatora optymalizowanego entropią często sprowadza się do zagadnienia minimalizacji pewnego funkcjonału. Optymalny kwantyzator może być poszukiwany poprzez minimalizację entropii dla danego błędu kwantyzacji za pomocą wyrażenia jak niżej:

$$J = H_{\tilde{X}}(Q) + \lambda \cdot D_Q, \quad (8.18)$$

który trzeba zminimalizować dobierając wartość współczynnika λ oraz szerokości kolejnych przedziałów kwantyzacji. Okazuje się, że dla znacznych wartości średniej bitowej, czyli słabszej kompresji z wąskimi przedziałami kwantyzacji (przypadek kwantyzatora wysokorozdzielczego) optymalnym jest kwantyzator równomierny, co upraszcza projektowanie optymalnej kwantyzacji. Ponadto, nawet w przypadku małych wartości średniej bitowej reprezentacji wyjściowej algorytmu kompresji równomierny kwantyzator skalarne pozwala przy pewnych założeniach uzyskać najlepszą skuteczność kwantyzacji z możliwych.

Aby lepiej zrozumieć ideę konstrukcji kwantyzatora skalarne wprowadźmy kilka dodatkowych pojęć. Jeśli przez $f_X(x)$ oznaczymy funkcję gęstości prawdopodobieństwa zmiennej losowej X modelującej źródło informacji o rozkładzie ciągłym (ogólniejszy przypadek), to entropia zmiennej X jest zdefiniowana następująco:

$$H_X = - \int_{-\infty}^{\infty} f_X(x) \log_2 f_X(x) dx \quad (8.19)$$

Kwantyzator jest wysokorozdzielczy, jeśli $f_X(x)$ jest w przybliżeniu stała w każdym przedziale kwantyzacji $[\beta_{j-1}, \beta_j)$ o szerokości $\Delta_j = \beta_j - \beta_{j-1}$. Sytuacja taka występuje w przypadku, kiedy szerokości przedziałów Δ_j są względnie małe w stosunku do stopnia zmienności funkcji gęstości prawdopodobieństwa $f_X(x)$, co występuje przy wysokich średnich bitowych kodu wyjściowego. Następujące twierdzenie (przedstawione w [9], można wykazać także poprzez przekształcenie równania (8.18) jak w [2]) mówi o tym, że kwantyzator równomierny jest optymalny w klasie kwantyzatorów wysokorozdzielczych:

TWIERDZENIE 8.1. Jeśli Q oznacza wysokorozdzielczy kwantyzator względem funkcji $f_X(x)$, to spełniona jest nierówność:

$$H_{I_Q} \geq H_X - \frac{1}{2} \log_2 (12 D_{Q(X)}) \quad (8.20)$$

Nierówność ta staje się równością wtedy i tylko wtedy, kiedy Q jest kwantyzatorem równomiernym, a błąd kwantyzacji wynosi wtedy $D_{Q(X)} = \frac{\Delta^2}{12}$.

Zakładając powyższą wartość błędu kwantyzacji $D_{Q(X)}$ w schemacie równomiernego kwantyzatora wysokorozdzielczego, określamy przedział kwantyzacji $\Delta = \sqrt{12 D_{Q(X)}}$ uzyskując minimalną średnią bitową na poziomie entropii, równą $R_I = H_X - \log_2 \Delta$. Oczywiście, na podstawie twierdzenia 8.1 można również założyć długość kodu wyjściowego, uzyskiwaną przy pomocy kodera dającego w przybliżeniu graniczną wartość entropii na wyjściu, a następnie wyliczyć szerokość przedziału kwantyzacji jako: $\Delta = 2^{H_X - R_I}$, który daje błąd kwantyzacji równy $D_{Q(X)} = \frac{1}{12} 2^{2[H_X - R_I]}$.

Udowodniono również w [8], że jeżeli nawet warunek wysokiej rozdzielczości kwantyzatora nie jest spełniony, dla szerokiej klasy rozkładów prawdopodobieństwa zmiennych losowych modelujących źródła informacji, włączając w to uogólniony rozkład Gaussa z jego różnymi realizacjami, równomierny kwantyzator skalarny daje błąd kwantyzacji bliski minimalnemu, przy danej wartości entropii kwantowanego strumienia danych, jeśli liczba przedziałów kwantyzacji M jest wystarczająco duża.

Kwantyzacja wektorowa

Wektorowa kwantyzacja jest uogólnieniem kwantyzacji skalarnej na przypadek wielowymiarowych wektorów. Podstawowe zasady konstrukcji kwantyzatora są bardzo podobne, z tym że w miejsce przedziałów kwantyzacji pojawiają się regiony decyzyjne określone w wielowymiarowej przestrzeni wektorów, a zbiór wartości rekonstruowanych zastąpiony jest książką kodów zawierającą wektory rekonstruowane, zwane wektorami kodu. Każdemu wektorowi kodu przyporządkowany jest indeks - kolejne słowo kodowe. Indeksy wskazujące na reprezentantów z książki kodów, przybliżających kolejne wejściowe wektory danych stanowią bitowy strumień nowej, stratnej reprezentacji oryginalnego zbioru danych.

Według teorii informacji łączna kwantyzacja sekwencji zmiennych losowych, pomiędzy którymi występuje zależność, jest zawsze lepsza od niezależnej kwantyzacji każdej

zmiennej. Tak więc wektorowa kwantyzacja (ang. vector quantization - VQ), eliminująca zależności pomiędzy kolejnymi zmiennymi losowymi, jest skuteczniejsza od kwantyzacji skalarnej (ang. scalar quantization - SQ). Jedynie w przypadku, gdy te zmienne losowe są niezależne (jeśli można je np. opisać wielowymiarowym rozkładem Gaussa), zastosowanie SQ daje wyniki porównywalne z VQ. Jednak w praktyce zastosowanie VQ w algorytmach kompresji napotyka na szereg problemów. Mała wymiarowość wektorów nie daje satysfakcjonującej skuteczności kompresji. Przykładowo, bloki o rozmiarach 2×2 czy 4×4 formowane z obrazów, stanowią wektory odpowiednio 4-ro i 16-to elementowe o ograniczonej podatności na kwantyzację. Jeśli natomiast użyjemy bloków o większych rozmiarach, np. 8×8 czy 16×16 , zwiększając wymiarowość wektorów, a co za tym idzie potencjalną skuteczność algorytmu kwantyzacji, wówczas okazuje się, że zbyt duże rozmiary wektorów stają się z kolei mało przydatne w praktycznych aplikacjach VQ ze względu na trudności realizacyjne. Gubione są bowiem szczegóły informacji wejściowej ze względu na zbyt zgrubne przybliżenia wektorami rekonstrukcji. Wynika to z konieczności ograniczenia rozmiarów książki kodowej, a co za tym idzie liczby wektorów kodu. Zbyt duże książki kodowe wymagają bowiem bardzo czasochłonnych algorytmów przeszukiwań oraz olbrzymich pamięci do ich przechowywania. Powodem ograniczenia wymiarowości wektorów w VQ jest także konieczność określenia wielowymiarowej funkcji gęstości prawdopodobieństwa zmiennej losowej, modelującej wartości wektorów w przestrzeni danych, przy konstrukcji optymalnych wektorów kodu. Wiarygodna estymata takiego rozkładu wymaga dużego zbioru treningowego o własnościach zbliżonych do kompresowanych zbiorów danych. Jeśli zwiększamy wymiarowość przestrzeni wektorów, wówczas przy tej samej ilości zbiorów treningowych maleje rzetelność wyznaczanej coraz bardziej złożonej wielowymiarowej funkcji gęstości (zjawisko analogiczne do rozrzedzania kontekstu bezstratnych koderów z modelami statystycznymi wyższych rzędów). Brak jest poza tym praktycznych rozwiązań pozwalających wyznaczyć optymalny zbiór wektorów kodu w skali globalnej całej przestrzeni wektorów danych (wspomniany algorytm LBG znajduje minima lokalne). Stosowane algorytmy wektorowej kwantyzacji ograniczają więc zazwyczaj wymiarowość kwantowanych wektorów oryginalnego zbioru danych do wartości nie przekraczających 16, do 20.

Przykładowo Kuo i Hsu [10] zwiększyli nieznacznie stopień kompresji obrazów przy danym poziomie zniekształceń poprzez zastosowanie kwantyzacji wektorowej zamiast skalarnej (przy podziale obrazu na bloki o rozmiarach 4×4), a Wu i Coll [11] zastosowali VQ w blokach 8×8 , ale jedynie dla wolnozmiennych regionów obrazu. W większości zastosowań VQ do kompresji obrazów wykorzystuje się bloki o rozmiarach 2×2 i 2×4 .

8.2. Dekompozycja danych oryginalnych

Duża złożoność algorytmów wektorowej kwantyzacji, duże koszty obliczeniowe i sprzętowe (konieczność dużej pamięci operacyjnej do realizacji algorytmów) oraz inne ograniczenia realizacyjne powodują problemy z uzyskaniem zadawalającej skuteczności kompresji w dużym obszarze zastosowań. Podejmowano próby wykorzystania różnych schematów wstępnej dekompozycji danych w celu uzyskania możliwie największej redukcji nadmiarowości przed fazą kwantyzacji, przy stosunkowo niewielkich kosztach obliczeniowych i sprzętowych. Do najbardziej popularnych i skutecznych rozwiązań problemu skutecznej dekompozycji danych należy zaliczyć metody wykorzystujące:

- transformaty częstotliwościowe o nieskończonym nośniku,
- przekształcenia fraktalne,

- przekształcenia falkowe,

Bardzo efektywnym rozwiązaniem okazało się wykonanie wstępnej dekompozycji danych przy pomocy różnego typu transformat, a następnie kwantyzacja współczynników tych transformat przy pomocy prostszych rodzajów kwantyzatorów. Dobre metody dekompozycji danych, redukujące znacznie nadmiarowości, w tym przede wszystkim korelacje pomiędzy wartościami współczynników nowej przestrzeni, pozwoliły uzyskać wyższą efektywność całego schematu kompresji przy znacznie mniejszej złożoności i małych kosztach realizacyjnych. Okazuje się, że również w tym przypadku najlepsza postać transformaty w schemacie kodowania transformacyjnego, pozwalająca całkowicie zdekorować wejściowe źródło danych jest bardzo trudna w praktycznej realizacji ze względu na olbrzymie koszty obliczeniowe. Jądro tego przekształcenia jest bowiem zależne od sygnału transformowanego. Można jednak osiągnąć zbliżone poziom dekoracji danych przy pomocy prostszej, niezależnej od sygnału transformaty, która pozwala zrealizować bardzo szybkie algorytmy kompresji.

Nieco inne rozwiązanie, bliższe koncepcji wektorowej kwantyzacji zastosowano w koderach fraktalnych, używanych zasadniczo do stratnej kompresji obrazów. Wykorzystując aparat przekształceń afinicznych do generacji operatorów zwięzających z atraktorami w postaci fraktali przybliżających kolejne porcje obrazu, zbudowano algorytm kompresji oparty na idei samopodobieństwa obrazu. Wyznaczanie dużej księgi kodowej na podstawie licznych wektorów treningowych zastępowane jest procesem definiowania słownika na bazie wzajemnego podobieństwa małych fragmentów kompresowanego obrazu - tak jakby obraz był przybliżany na podstawie samego siebie. Po zdefiniowaniu wielu operatorów zwięzających opisujących obraz można go następnie wygenerować w kilku lub co najwyżej kilkunastu iteracjach, zaczynając od dowolnej postaci pola obrazu.

Jeszcze inna koncepcja wywodzi się z teorii aproksymacji. Wiadomo, że dla oszczędniejszego opisu sygnału przy pomocy jądra liniowej transformaty potrzeba możliwie dużego podobieństwa sygnału i funkcji bazowych tego przekształcenia. Skoro kompresowane dane reprezentują w zdecydowanej większości przypadków źródła silnie niestacjonarne, należy zastosować dobre narzędzie do opisu sygnałów niestacjonarnych. Transformaty Fouriera lub jej podobne (sinusowa, kosinusowa) wykorzystują funkcje sinusoidalne i kosinusoidalne o nieskończonym nośniku przy założeniach stacjonarności (lub pseudostacjonarności) sygnału, a więc nie są dobrym narzędziem w przypadkach sygnałów niestacjonarnych. Wykorzystano więc narzędzie przekształcenia liniowego falek (ang. wavelet) o funkcjach bazowych określonych na skończonym nośniku, o skończonej energii, świetnie nadające się do analizy sygnałów niestacjonarnych tworząc bardzo oszczędną ich reprezentację.

8.3. Kodowanie odwracalne

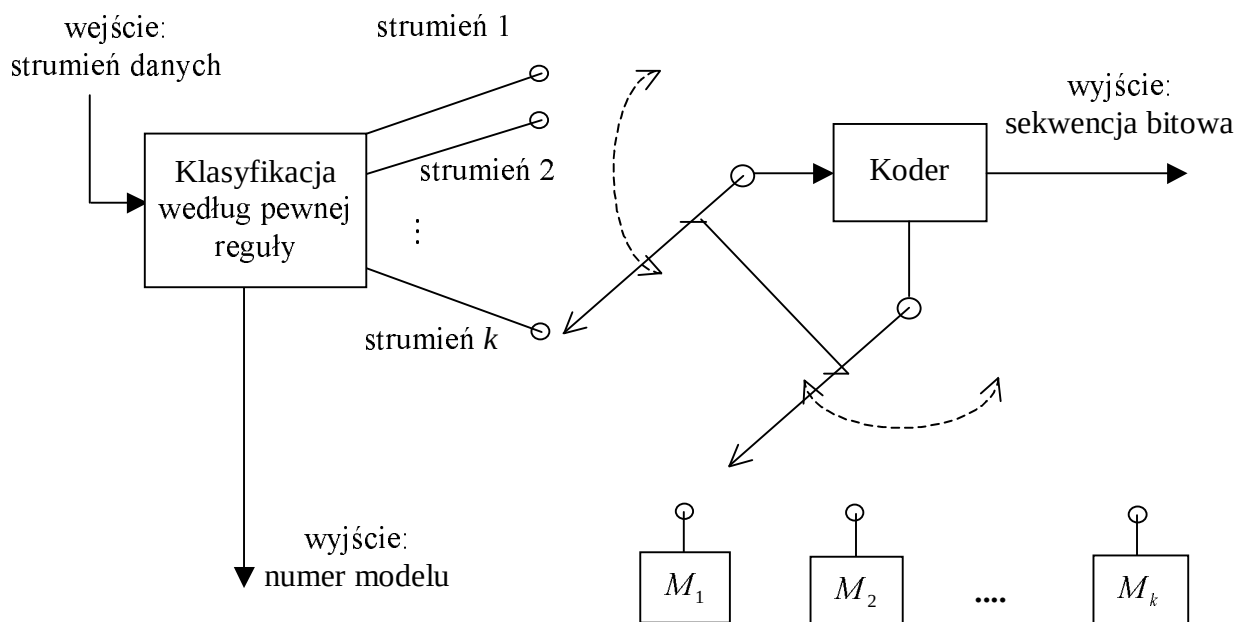
Sposób bezstratnego kodowania powstałych po kwantyzacji zbiorów wartości o zredukowanej dynamice (alfabecie) bardzo silnie zależy od realizowanych wcześniej technik dekompozycji i samej kwantyzacji. Trudno więc tutaj mówić o rozwiązaniach optymalnych w każdym przypadku. Zagadnienie to jest analogiczne do rozważanych w pierwszej części tej pozycji problemów bezstratnej kompresji danych. Wyróżnikiem może być jednak duża wiedza a priori o strumieniach danych wchodzących na wejście kodera ze względu na ukształtowanie tego strumienia przez poprzednie etapy schematu kompresji. Można wyróżnić dwa podstawowe rodzaje tychże schematów. W pierwszym optymalizuje się maksymalnie funkcje dekoracyjne w procesach dekompozycji i kwantyzacji, wyrzucając do kodowania

zbiory danych prawie zupełnie zdekorelowanych czy wręcz niezależnych, o silnie zróżnicowanym alfabecie danych, zmiennej liczbie bitów, itd. W takich przypadkach rola odwracalnego kodowania jest znikoma lub wręcz żadna. Zadanie redukcji nadmiarowości jest w zadawalającym stopniu wykonywane wcześniej, np. w metodzie kwantyzacji, która jest optymalizowana pod kątem minimalnej entropii danych wyjściowych w relacji do poziomu błędu kwantyzacji, a faza kodowania może być pominięta.

Innym rozwiązaniem jest zmniejszenie czasochłonności i stopnia złożoności algorytmów kwantyzacji poprzez zastąpienie ich rozwiązaniami znacznie prostszymi, bardziej efektywnymi w stosunku do stopnia złożoności. Dają one co prawda mniejszą redukcję nadmiarowości, ale dodatkowa dekorrelacja danych może być wykonana skutecznie w koderze odwracalnym przy znacznie niższych ostatecznych kosztach czasowych. Przykładowo, znając statystykę strumienia danych kwantowanych można zaprojektować efektywny statyczny koder arytmetyczny czy Huffmana o mniej rozbudowanych strukturach i lepszych warunkach początkowych.

Często spotykaną praktyką jest podział zbioru kodowanego na kilka podzbiorów (strumieni) o wyraźnie różnej statystyce i niezależne kodowanie każdego z nich koderem o innych parametrach modelu źródła informacji czy wręcz strukturze, bądź też jednym koderem o modelach przełączanych (rys. 8.4).

Dodatkowo w zależności od zastosowań generowany jest w koderze strumień o cechach kodu sekwencyjnego lub progresywnego, ustalający pewną hierarchię przekazywanej informacji, w niektórych rozwiązaniach również kod zagnieżdżony, kiedy to koder jest zatrzymywany w momencie osiągnięcia założonej długości kodu wyjściowego. Czasami też fazy kwantyzacji i kodowania przeplatają się emitując równolegle ze złożoną procedurą kwantyzacji kolejne partie kwantowanych wartości np. w szybkich zastosowaniach transmisyjnych.



Rys. 8.4. Schemat kodera o kilku modelach $M_1, M_2, \dots, M_k$ przełączanych w zależności od charakterystyki strumienia wejściowego.

8.4. Teoria zniekształceń źródeł informacji

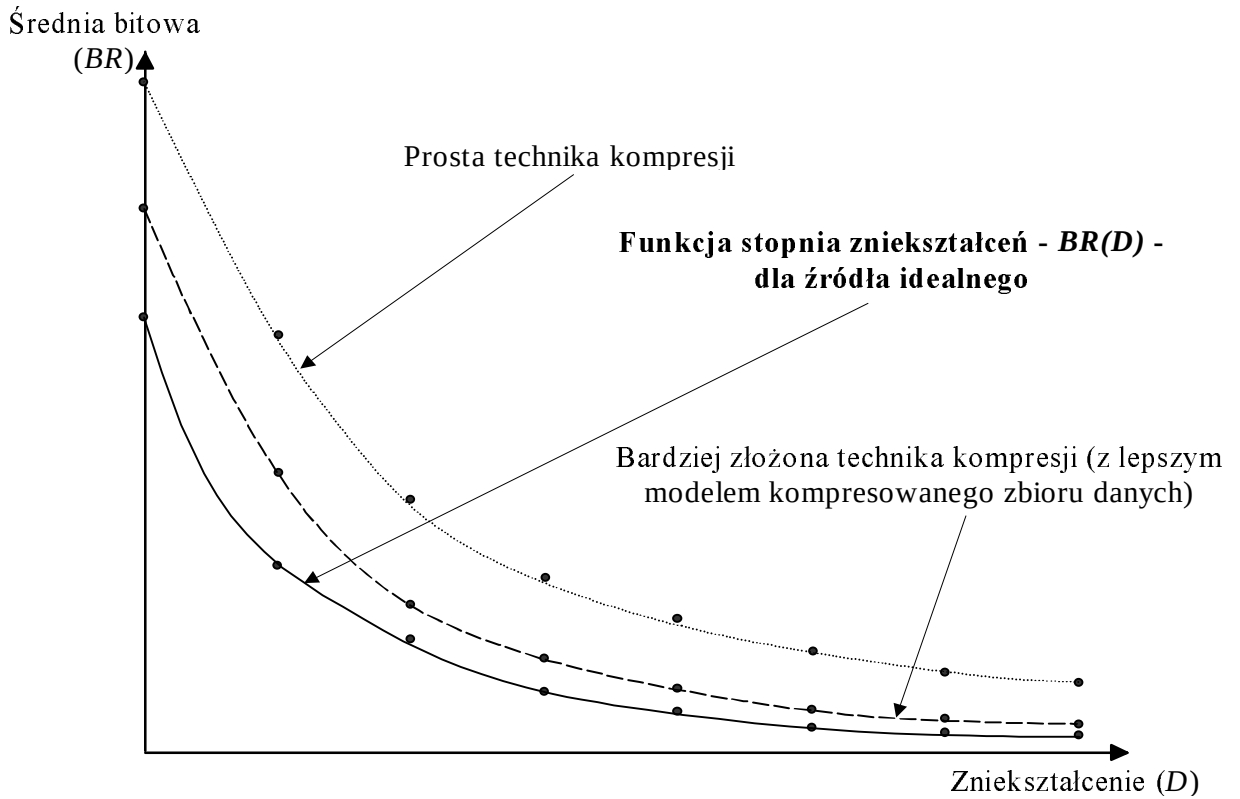
Wyznaczenie granicznej wartości stopnia kompresji informacji zawartej w zbiorze danych poprzez obliczenie entropii źródła modelującego tę informację ma zastosowanie jedynie w przypadku bezstratnych metod kompresji. Przy kompresji stratnej pojawia się również pytanie o wartość graniczną możliwego do uzyskania stopnia kompresji danego obrazu przy poziomie zniekształceń nie przekraczającym pewnej wartości. Odpowiedź na to pytanie adresowana jest w pierwszej kolejności do gałęzi teorii informacji rozważającej stopień zniekształceń źródeł informacji (ang. rate distortion theory). Teoria ta w zastosowaniu do zagadnień kompresji zakreśla teoretyczne granice efektywności technik stratnej rekonstrukcji danych w oparciu o pewien model źródła informacji oraz określone kryterium dokładności. Pozwala na wyznaczenie zależności pomiędzy miarą wielkości kompresji (najczęściej średnią bitową) a przyjętą miarą zniekształceń danych oryginalnych w stratnym procesie kompresji, oznaczoną przez $BR(D)$, która posiada dwie bardzo istotne własności:

- dla dowolnej wartości zniekształceń D , możliwym jest znalezienie algorytmu kodowania w stopniu dowolnie bliskim $BR(D)$ i średnim poziomie zniekształceń dowolnie bliskim wartości D ;
- niemożliwym jest znalezienie reprezentacji kodowej, który pozwala odtworzyć oryginalne źródło informacji ze zniekształceniem D lub mniejszym i stopniem kompresji poniżej $BR(D)$.

$BR(D)$ jest wypukłą, ciągłą i monotonicznie malejącą funkcją D (rys. 8.5). Zasadniczo można stwierdzić, że bardziej wyrafinowane algorytmy kompresji, które lepiej modelują statystykę źródła i dostosowują do niej algorytmy dekompozycji, kwantyzacji i kodowania osiągają efektywność bliższą granicy $BR(D)$. Przedstawiają to dwie funkcje opisujące skuteczność kompresji hipotetycznie różnych technik na rys. 8.5.

Zastosowanie tej teorii do estymacji granicznych krzywych dla stratnych kompresorów działających na konkretnych zbiorach danych nie jest jednak łatwe. Zagadnienie to jest jeszcze matematycznie niezbyt skomplikowane, gdy źródło jest modelowane jako DMS (źródło bez pamięci, gdzie prawdopodobieństwo wystąpienia każdego symbolu z alfabetu opisującego źródło nie zależy od kontekstu) oraz miara zniekształcenia w punktach nie zależy od otoczenia tych punktów w wejściowym strumieniu danych, a kryterium formułowane jest z wykorzystaniem błędu średniokwadratowego lub średniego błędu bezwzględnego.

Takie rozwiązanie w przypadku rzeczywistych zbiorów o różnorodnej strukturze zależności pomiędzy danymi jest jednak nieużyteczne, bo słabo modeluje źródło informacji i wyznaczone wartości graniczne są obarczone dużym błędem. Zarówno model statystyczny źródła jak i miara zniekształcenia winny uwzględniać lokalne korelacje danych w pewnym kontekście, zależnym od charakteru zbioru danych. Stosowane są rozwiązania oparte na Gaussowskim źródle z miarą zniekształceń opartą na ważonym błędzie kwadratowym lub też modele obrazu w postaci dwuwymiarowego źródła Gaussa-Markowa ze współczynnikami korelacji bliskimi 1, co daje znacznie lepszą aproksymację oryginału. Rozwiązania poszukiwane są wówczas dość złożonymi metodami numerycznymi. Skuteczne określenie $BR(D)$ dla modelu źródła, które wiernie przybliży strumień wejściowy oraz dla miary zniekształceń, która dokładnie uwzględnia np. wizualne czy diagnostyczne kryteria oceny jakości danych w kompresji obrazów pozostaje nadal frapującym problemem badawczym.



Rys.8.5. Przykład funkcji stopnia zniekształceń, określającej możliwą skuteczność kompresji w funkcji poziomu zniekształceń - $BR(D)$ oraz przykładowe krzywe określające skuteczność kompresji dwu koderów stratnych.

Ogólną postać wyrażenia określającego wartość zniekształcenia przy statystycznym modelowaniu źródła informacji X (opisującego kompresowany zbiór danych) oraz modelu zrekonstruowanej informacji Y (opisującego zrekonstruowany zbiór danych) można przedstawić następująco:

$$D = \sum_{i=0}^{N_1-1} \sum_{j=0}^{N_2-1} d(x_i, y_j) P(x_i) P(y_j | x_i), \quad (8.21)$$

gdzie N_1 i N_2 oznaczają liczbę symboli w alfabecie źródeł odpowiednio X i Y , $P(x_i)$ to prawdopodobieństwo wystąpienia poszczególnych symboli alfabetu źródła X , a $P(y_j | x_i)$ - prawdopodobieństwo przyporządkowania symbolom źródła X kolejnych symboli alfabetu źródła Y stosując określony mechanizm techniki stratnej kompresji danych. W najprostszych przypadkach (przy jednoznacznym zdeterminowaniu wartości rekonstruowanych) wartość prawdopodobieństwa warunkowego może być określona w następujący sposób:

$$P(y_j | x_i) = \begin{cases} 1 & \text{jeśli } i = j \text{ lub } i = j + 1 \text{ oraz } j = 0, 2, \dots, 10 \\ 0 & \text{w p.p.} \end{cases} \quad (8.22)$$

lub też przy równomiernym rozkładzie prawdopodobieństwa:

$$P(y_j | x_i) = \frac{1}{6} \quad \text{dla } i = 0, 1, \dots, 11 \text{ oraz } j = 0, 2, \dots, 10 \quad (8.23)$$

Wielkość $d(x_i, y_j)$ określa różnicę pomiędzy symbolami obu źródeł i mogą być tutaj stosowane różne miary zniekształcenia, jak chociażby: $d(x_i, y_j) = (x_i - y_j)^2$ lub $d(x_i, y_j) = |x_i - y_j|$.

Tak więc wartości $P(x_i)$ wynikają z charakteru danych wejściowych oraz przyjętego modelu źródła, wartości $P(y_j | x_i)$ charakteryzują przyjęty schemat kompresji, natomiast wartości $d(x_i, y_j)$ wynikają z jednej strony z przyjętej miary zniekształceń, z drugiej zaś z alfabetu modelu danych zrekonstruowanych, a więc zależą także od stosowanej metody kompresji.

Zakładając, że dla rozpatrywanej klasy algorytmów alfabet modelu Y jest jednakowy (może być taki sam jak źródła X) można stwierdzić, że zniekształcenie D jest funkcją jedynie zbioru prawdopodobieństw warunkowych: $D = D(\{P(y_j | x_i)\})$.

Minimalna wartość średniej bitowej dla danego zniekształcenia D^{tar} dana jest, według Shannona [1], poprzez następującą zależność:

$$BR(D) = \min_{\{P(y_j | x_i)\} \in \Lambda} I(X; Y), \quad (8.24)$$

gdzie Λ jest zbiorem prawdopodobieństw warunkowych dla danej metody kompresji (przetwarzania), takim że: $\Lambda = \{\{P(y_j | x_i)\} \text{ takich że } D(\{P(y_j | x_i)\}) \leq D^{tar}\}$, a $I(X; Y)$ jest średnią informacją wzajemną określoną zależnością:

$$I(X; Y) = \sum_{i=0}^{N_1-1} \sum_{j=0}^{N_2-1} P(x_i, y_j) \log \left[\frac{P(x_i | y_j)}{P(x_i)} \right]. \quad (8.25)$$

Jest to ilość informacji statystycznej o zmiennej X zawarta w zmiennej losowej Y .

8.5. Blokowa dwupoziomowa kwantyzacja

Jest to prosta metoda kwantyzacji wykorzystywana w kompresji obrazów. W technice blokowej dwupoziomowej kwantyzacji (ang. Block Truncation Coding - BTC [12]) obraz jest dzielony na rozłączne bloki pikseli o rozmiarach $n \times n$ (najczęściej 4×4). Następnie niezależnie dla każdego bloku przeprowadza się binarną kwantyzację wartości pikseli dobierając próg odpowiednio do lokalnej statystyki bloku. Powstająca wówczas reprezentacja poszczególnych bloków składa się z binarnej mapy $n \times n$, która wskazuje poziom rekonstrukcji dla każdego piksela oraz z dwu wartości rekonstruowanych poziomów. Dekompresja jest więc prostym procesem odtwarzania wartości pikseli według dostarczonej binarnej mapy poziomów.

Najprostszym sposobem określenia progu jest ustalenie go na poziomie średniej wartości pikseli w bloku. Średnie wartości pikseli dwu klas (segmentów) wyznaczonych przez wartość progową mogą posłużyć jako dwie wartości poziomów rekonstrukcji. Z tej podstawowej realizacji wynikają główne zalety i wady techniki BTC.

Główną zaletą jest zachowanie (dokładne odtworzenie) ostrych krawędzi w zrekonstruowanym obrazie, podczas gdy większość stratnych metod kompresji powoduje rozmycie ostrych krawędzi. Wady związane są przede wszystkim z wprowadzaniem artefaktów polegających na zniekształcaniu łagodnych konturów wewnątrz bloków,

spowodowanych gwałtownym skokiem pomiędzy wartościami poziomów rekonstrukcji, oraz wprowadzaniu nieciągłości krawędzi na granicach bloków wynikających z niezależnego przetwarzania sąsiednich bloków. Jakość rekonstrukcji klasycznego algorytmu BTC pokazuje przykład 8.1.

PRZYKŁAD 8.1. Stosując technikę BTC z podziałem obrazu na bloki 4×4 , progiem równym wartości średniej w bloku i średnimi segmentów pikseli z wartościami powyżej oraz poniżej progu jako poziomami rekonstrukcji dokonano kwantyzacji następujących trzech bloków:

$$X_1 = \begin{bmatrix} 45 & 48 & 46 & 47 \\ 44 & 45 & 44 & 46 \\ 45 & 47 & 43 & 45 \\ 47 & 46 & 44 & 47 \end{bmatrix}, \quad X_2 = \begin{bmatrix} 145 & 148 & 146 & 146 \\ 44 & 145 & 144 & 146 \\ 45 & 47 & 143 & 45 \\ 47 & 46 & 45 & 47 \end{bmatrix}, \quad X_3 = \begin{bmatrix} 49 & 53 & 57 & 59 \\ 48 & 51 & 54 & 53 \\ 48 & 47 & 52 & 50 \\ 47 & 46 & 49 & 50 \end{bmatrix}.$$

Progi dla poszczególnych bloków wynoszą odpowiednio: $t_1 = 46$, $t_2 = 96$, $t_3 = 51$. Na ich podstawie przeprowadzono segmentację w każdym z bloków uzyskując następujące mapy bitowe:

$$M_1 = \begin{bmatrix} 0 & 1 & 1 & 1 \\ 0 & 0 & 0 & 1 \\ 0 & 1 & 0 & 0 \\ 1 & 1 & 0 & 1 \end{bmatrix}, \quad M_2 = \begin{bmatrix} 1 & 1 & 1 & 1 \\ 0 & 1 & 1 & 1 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix}, \quad M_3 = \begin{bmatrix} 0 & 1 & 1 & 1 \\ 0 & 1 & 1 & 1 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix}.$$

Wartości poziomów rekonstrukcji w segmentach jak wyżej są następujące: $y_1^L = 44$, $y_1^H = 47$, $y_2^L = 46$, $y_2^H = 145$, $y_3^L = 48$, $y_3^H = 54$. Daje to postać bloków zrekonstruowanych jak niżej:

$$\tilde{X}_1 = \begin{bmatrix} 44 & 47 & 47 & 47 \\ 44 & 44 & 44 & 47 \\ 44 & 47 & 44 & 44 \\ 47 & 47 & 44 & 47 \end{bmatrix}, \quad \tilde{X}_2 = \begin{bmatrix} 145 & 145 & 145 & 145 \\ 46 & 145 & 145 & 145 \\ 46 & 46 & 145 & 46 \\ 46 & 46 & 46 & 46 \end{bmatrix}, \quad \tilde{X}_3 = \begin{bmatrix} 48 & 54 & 54 & 54 \\ 48 & 54 & 54 & 54 \\ 48 & 48 & 54 & 48 \\ 48 & 48 & 48 & 48 \end{bmatrix}.$$

Efekt stratności procesu kompresji techniką BTC widać wyraźnie przy porównaniu bloków oryginalnych X_1, X_2 i X_3 z ich zrekonstruowanymi obrazami $\tilde{X}_1, \tilde{X}_2$ i $\tilde{X}_3$. Ostra krawędź w bloku X_2 została zachowana, natomiast w przypadku bloku X_3 bardzo łagodnie rosnące zbocze zostało sztucznie podzielone na dwa obszary i powstała niczym nieuzasadniona krawędź obiektu o przypadkowym kształcie. W jednorodnym, nieco zaszumionym obszarze bloku X_1 pojawił się z kolei bardziej widoczny szum ziarnisty.

Optymalizacja techniki BTC może być prowadzona w trzech zasadniczych kierunkach. Po pierwsze, konstruowane są bardziej złożone modele kwantyzacji przy wykorzystaniu trzech stopni swobody (próg i dwie wartości poziomów rekonstrukcji). Niektóre z nich zachowują trzy centralne momenty statystyczne (wartość średnia, wariancja, asymetria), inne wprowadzają element minimalizacji błędu średniokwadratowego.

Ponadto stosuje się wiele technik kodowania map bitowych oraz poziomów rekonstrukcji. Może to być proces wprowadzający dodatkowe zniekształcenia, zwiększający skuteczność kompresji kosztem większej złożoności algorytmu. W najprostszym rozwiązaniu

można poddać równomiernej kwantyzacji wartości poziomów rekonstrukcji zapisując je na mniejszej liczbie bitów. Czasami zamiast wartości poziomów koduje się wartość średnią i odchylenie standardowe bloku kwantyzując je łącznie (ang. joint quantization). Kodowanie mapy binarnej, podobnie jak wariancji, można w ogóle pominąć w przypadku, gdy wariancja danego bloku jest mała. Wówczas do strumienia wyjściowego zostaje przesłana jedynie wartość średnia bloku. Dla większych wartości wariancji kodowanie mapy bitowej może odbywać się przy pomocy różnego rodzaju technik. Stosowane są przykładowo techniki wyznaczające w mapie i kodujące jedynie tzw. bity niezależne [13]. W mapie bitowej pewne bity ustalane są jako niezależne, a pozostałe (zależne) ustalane są na podstawie predefiniowanych logicznych zależności. Zależności te określają typ struktur, które będą zachowane w procesie stranej kompresji pozwalając na deformacje pozostałych. Zwiększając liczbę niezależnych bitów uzyskuje się lepszą jakość rekonstrukcji obrazu. Ponadto, do kodowania map bitowych wykorzystywane są typowe techniki wektorowej kwantyzacji [14].

Opracowano także metody adaptacyjne w odniesieniu do wielkości bloków. Przystosowując się do lokalnych statystycznych własności obrazów można zmniejszać lub zwiększać rozmiary bloków zwiększając skuteczność kompresji. Użycie bloków o większych rozmiarach w obszarach łagodnych (o małej wariancji wartości danych w blokach) oraz bloków mniejszych w obszarach aktywnych (o dużej wariancji) pozwala oszczędniej kompresować obrazy o zróżnicowanej treści. Wprowadza się także próby likwidacji niekorzystnych efektów powstających na granicy bloków.

8.6. Metody ekstrakcyjne

Techniki ekstrakcyjne (ang. message extraction techniques) w kompresji stratnej odgrywają niekiedy bardzo duże znaczenie, szczególnie w przypadku wybranych, specjalistycznych zastosowań o znanej charakterystyce kompresowanych zbiorów danych, przy czym najczęściej są to obrazy. Zasadniczą ideą jest dokonanie wstępnej analizy sceny i pewna klasyfikacja prezentowanej tam informacji na bardziej i mniej istotną. Metody te wykorzystywane już w latach sześćdziesiątych są awangardą coraz powszechniej stosowanych, szczególnie w kontekście kompresji danych multimedialnych, współczesnych metod kompresji drugiej generacji z obiektową analizą informacji użytecznej (np. w standardzie MPEG4).

Można wyróżnić trzy podstawowe kategorie metod ekstrakcyjnych [15]. Pierwsza to algorytmy wykorzystujące lokalne cechy regionów. Przykładowo w kodowaniu piramidalnym operatory lokalne wydzielają składowe niskoczęstotliwościowe (informacje o kontraście, często estymowane przez średnią wartość pikseli z danego obrazu) i wysokoczęstotliwościowe (informacje o krawędziach, drobnych fragmentach struktur), które są kodowane oddzielnie. Dekompresor łączy z powrotem te dwa zbiory danych, tworząc zrekonstruowany obraz. Dopuszczalne stopnie kompresji tą metodą wynoszą około 10:1, podczas gdy przy zastosowaniu nieco bardziej złożonej metody anizotropowego niestacjonarnego predyktywnego kodowania, bazującej na analizie spektrum mocy obrazu przy określaniu energii niskich i wysokich częstotliwości oraz kierunkowo-zorientowanej energii obrazu, można uzyskać stopnie równe nawet 35:1.

W drugiej grupie algorytmów obraz jest kompresowany poprzez zakodowanie jego konturów i tekstur. Wpierw wykrywane są krawędzie, na podstawie których tworzone są kontury, a następnie dokonuje się estymacji poziomów szarości w poszczególnych regionach objętych konturami. Przykładem takiej techniki jest zaproponowana przez Chena [16] metoda

kompresji z aproksymacją wartości pikseli w każdym regionie oddzielnie za pomocą funkcji wielomianowych.

Druga grupa algorytmów i duża część pierwszej mogą być zastosowane jedynie do kompresji prostych obrazów (o niezbyt złożonych scenach), natomiast kodowanie bardziej skomplikowanych obrazów o dużej dynamice i różnorodności prezentowanych struktur wymaga bardzo dużego stopnia złożoności tych algorytmów, co czyni je mało efektywnymi i nieprzydatnymi w praktyce.

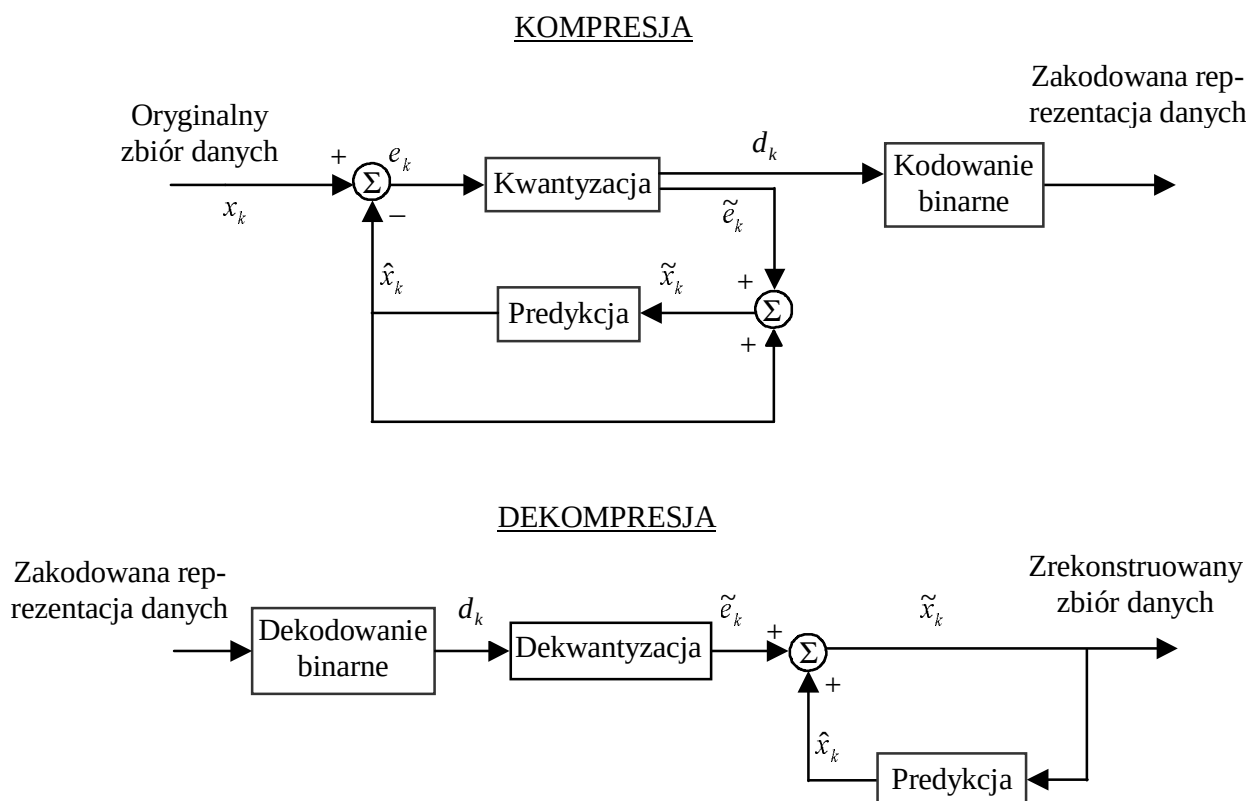
Algorytmy trzeciej grupy polegają na wyznaczeniu w obrazie nieistotnych obszarów, które są następnie z niego usuwane, a pozostałe ważne obszary są kompresowane przy zastosowaniu technik odwracalnych. W badaniach przeprowadzonych w University of Kansas Medical Center (UKMC) zastosowano dwie takie techniki: eliminację tła i kodowanie obwodu (ang. perimeter coding). Eliminacja tła została użyta przy kompresji obrazów CT głowy, gdzie usunięto część obrazu poza powierzchnią skóry, a pozostałą część zakodowano metodą Huffmana uzyskując stopnie kompresji około 5:1. Technikę kodowania obwodu zastosowano do kompresji rentgenowskich obrazów kończyn i kręgosłupa (stopień kompresji 3:1), a polega ona na zaznaczeniu na filmie prostokątnych obszarów o znaczeniu diagnostycznym. Dalej tylko te obszary są przetwarzane na postać cyfrową i kompresowane.

W UKMC podjęto także próbę kompresji obrazów CR poprzez ograniczenie kontrastu za pomocą adaptacyjnej kwantyzacji histogramu (CLAHE), a następnie uśrednienie wartości pikseli w blokach 2×2 . Osiągnięto akceptowalne (według testu ROC scharakteryzowanego w następnym rozdziale) stopnie kompresji równe 6:1.

8.7. Stratne kodowanie predykcyjne

W krótkim przeglądzie prostych metod kompresji stratnej tego rozdziału ważne miejsce zajmują metody predykcyjne, stosowane najczęściej do mało-stratnej kompresji obrazów oraz w kompresji mowy. Ograniczymy się tutaj do najpopularniejszych metod predykcji liniowej. W algorytmach odwracalnej DPCM wyznaczana przy pomocy modelu predykcji wartość przewidywana odejmowana jest od rzeczywistej wartości dla kolejnych danych w strumieniu tworząc strumień wartości różnicowych, tzw. błędu predykcji o zmniejszonym poziomie korelacji danych w stosunku do zbioru oryginalnego. W stratnej DPCM różnicowy zbiór danych podlega procesowi kwantyzacji, który ma decydujący wpływ zarówno na stopień kompresji jak i na jakość rekonstruowanego obrazu.

W procesie dekompresji predykcja może być oparta jedynie na zrekonstruowanych wartościach pikseli, które mogą różnić się od wartości rzeczywistych wskutek redukcji informacji przeprowadzonej podczas kwantyzacji. By zapewnić identyczną predykcję podczas kompresji i dekompresji, w czasie kodowania także korzysta się ze zrekonstruowanych wartości pikseli do określania przybliżeń, co prowadzi do umieszczenia etapu kwantyzacji w pętli predykcji. Przedstawia to rys.8.6.



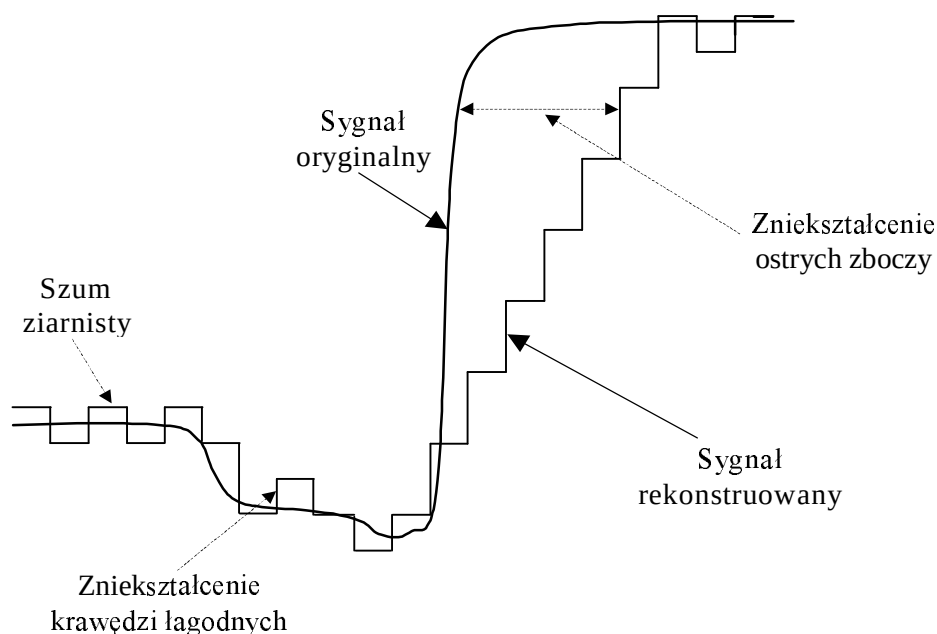
Rys.8.6. Blokowy schemat stratnej techniki DPCM. Oznaczenia jak w podrozdziale 8.1.

Projektowanie efektywnych algorytmów DPCM sprowadza się do optymalizacji procesu predykcji oraz kwantyzacji. Ze względu na umieszczenie kwantyzacji w pętli predykcji istnieje silna zależność pomiędzy błędem predykcji a błędem kwantyzacji, ale najlepsza w tym przypadku łączna optymalizacja prowadzi do złożonych modeli. Można tego uniknąć poprzez oddzielną optymalizację obu procesów. Rozwiązanie takie w wielu przypadkach może być dobrą aproksymacją optymalizacji łącznej według kryterium błędu średniokwadratowego.

Stosuje się predykcję różnych rzędów, jedno- i dwuwymiarową, globalną, lokalną i adaptacyjną. Predykcja 2-D (dwuwymiarowa) daje lepsze rezultaty przy kompresji obrazów w stosunku do metod 1-D (jednowymiarowych) zarówno w ocenie ilościowej (stosunek sygnału do szumu rośnie o około 3 dB), jak też w wyraźnej subiektywnej poprawie jakości obrazów, wyrażającej się głównie w mniejszych zniekształceniach ukośnych krawędzi.

Proces kwantyzacji może być projektowany przy użyciu wizualnych (obrazy), słuchowych (mowa) bądź statystycznych kryteriów. Do typowych zniekształceń wprowadzanych w procesie kwantyzacji wartości różnicowych dla danych obrazowych należą szum ziarnisty oraz zniekształcenia ostrych (ang. slope overload) i łagodnych krawędzi (ang. edge busyness). Szum ziarnisty pojawia się w jednolitych obszarach, gdzie niewielkie zmiany wartości sąsiednich pikseli powodują fluktuacje na wyjściu kwantyzatora. Zniekształcenia konturów w miejscach występowania ostrych zboczy sygnału 2-D, takie jak ich rozmycie i nieraz dodatkowa degradacja obrazu w ich sąsiedztwie w wyniku kwantyzacji, związane jest z faktem, iż gradient przyrostu wartości wyjściowych kwantyzacji nie nadąża za dużym gradientem danych obrazowych. Natomiast zniekształcenia łagodnych krawędzi (przesuwanie, odkształcanie) spowodowane są fluktuacją procesu kwantyzacji oraz

sztucznym tworzeniem zwiększonych gradientów. Rysunek 8.7 prezentuje schematycznie te zniekształcenia.



Rys.8.7. Typowe zniekształcenia występujące w technice stratnej DPCM.

Adaptacyjność technik DPCM może być realizowana w dziedzinie predykcji, kwantyzacji lub w obydwu. Adaptacyjna predykcja zwykle redukuje błąd predykcji przed fazą kwantyzacji, co powoduje że przy tym samym stopniu kompresji zmniejszona dynamika sygnału na wejściu kwantyzatora powoduje mniejszy błąd kwantyzacji i lepszą jakość rekonstruowanego obrazu. Przykłady adaptacyjnych modeli predykcji można znaleźć w rozdziale 6.

Z drugiej strony adaptacyjna kwantyzacja ma na celu zmniejszenie błędu kwantyzacji bezpośrednio przez dobór progów kwantyzacji i poziomów rekonstrukcji zgodnie z lokalną statystyką obrazu. W schematach adaptacyjnej kwantyzacji wprzód zbiór próbek sygnału wejściowego jest dzielony we wstępnej analizie na fragmenty o zbliżonych własnościach statystycznych, dla których procedura kwantyzacji konstruowana jest niezależnie. Dodatkowa informacja przesyłana do dekodera dotyczy podziału strumienia wejściowego na bloki oraz wartości parametrów poszczególnych schematów kwantyzacji. Jest to więc duże obciążenie skuteczności metody kompresji, często niwelujące efekty lokalnej poprawy jakości kwantyzacji. Stosuje się więc w takich przypadkach schemat adaptacji wstecz, podobne do rozwiązań opisanych w punkcie 8.1. Wykorzystuje się równomierną kwantyzację skalarną zmieniając liczbę poziomów kwantyzacji (tj. szerokość przedziału) w zależności od przedziału kwantyzacji, w który wpadła poprzednio kodowana próbka (podobnie jak w metodzie Jayanta). Ponadto wykorzystywane są informacje o perceptualnych skutkach kwantyzacji w przypadku kompresji obrazów. Szerokość przedziału kwantyzacji jest modyfikowana jedynie w zakresie, w którym maksymalny błąd kwantyzacji nie przekracza progu widoczności (powyżej tego progu błąd kwantyzacji powoduje zauważalne zmiany w obrazie). W algorytmie adaptacji określa się rodzaj obszaru, w którym znajduje się piksel, przyporządkowując go do klasy fragmentów płaskich, łagodnych krawędzi lub krawędzi wyraźnych i dobierając odpowiedni kwantyzator, np. wierniej zachowując w kwantyzacji

krawędzie. Rodzaj obszaru w obrazie wyznacza się przy pomocy estymatorów gradientów w różnych kierunkach.

Choć techniki predykcyjne typu DPCM były historycznie pierwszymi stosowanymi technikami kompresji ze stratą informacji, a wprowadzone później metody transformat ortogonalnych okazały się efektywniejsze, to jednak predykcja jako metoda modelowania w koderach stratnych nie straciła swego znaczenia. Zwiększa ona efektywność hybrydowych technik transformacyjnych, kodowania pasmowego i jest szczególnie skuteczna w tzw. prawie bezstratnym kodowaniu, kiedy to kosztem niewielkich odkształceń sygnału oryginalnego, zupełnie nierozpoznawalnej w ocenie subiektywnej, można znacząco zwiększyć uzyskiwany stopień kompresji danych [17].

8.8. Wybór optymalnej techniki kompresji stratnej

Lista czynników wpływających na wybór optymalnego algorytmu kompresji jest długa i może być rozszerzana w zależności od zastosowań. Ponadto waga poszczególnych czynników zmienia się w zależności od aplikacji. I tak przykładowo dla nieodwracalnej kompresji obrazów szczególnie istotne mogą być: czułość na typ obrazu wejściowego (dynamika, szумы, korelacja pomiędzy wartościami pikseli, itd.), zamierzony stopień kompresji (relacja jakość obrazu, a stopień kompresji), stały stopień kompresji przy zadanej jakości obrazów, asymetria koder/dekoder, możliwość progresywnej kompresji, warunki konkretnej implementacji (prostota, mała czasochłonność), tolerancja błędu kanału, itd.

Choć pytanie o najlepszy algorytm stratnej kompresji jest często stawiane, to niestety nie ma na nie jednoznacznej odpowiedzi. Mogą istnieć najodpowiedniejsze algorytmy dla konkretnych aplikacji, przy czym zarówno termin odpowiedniości jak i dokładne określenie zakresu aplikacji zależy od wielu czynników. Przykładowo, kiedy kompresja jest stosowana w systemie transmisji obrazów, najlepiej by było, gdyby operacje kompresji i dekompresji były wykonywane w czasie rzeczywistym, a złożoność struktury kodu odporna na zakłócenia występujące w kanale transmisji. Istotna jest także możliwość buforowania strumienia danych wyjściowych z koderu w związku ze zmieniającą się często prędkością transmisji w kanale oraz dopuszczalny poziom błędu rekonstrukcji po stronie odbiornika (np. w telemedycynie). Są to wymagania zupełnie inne niż w przypadku kompresji obrazów w celu zminimalizowania zajmowanej przez nie pamięci w obszernej bazie danych obrazowych.

Jeszcze raz warto przypomnieć tutaj o rozwiązaniu problemu kwantyzacji jako decydującym o przydatności danego koderu w konkretnej aplikacji. Określenie dopuszczalnego poziomu zniekształceń, sposobu redukcji informacji 'niepotrzebnej' lub 'mniej potrzebnej', zachowanie szczegółów informacji decydujących o jakości kompresowanego zbioru danych, stworzenie wyjściowego strumienia indeksów podatnego na efektywne binarne kodowanie to zasadnicze wymagania stawiane procesowi kwantyzacji w niemal każdym algorytmie kompresji. Pewne zadania skutecznej kwantyzacji może przejąć dobry algorytm wstępnej dekompozycji danych oraz binarnego kodowania odpowiednio formowanych strumieni danych, zmniejszając nieraz znacznie stopień złożoności metody kompresji oraz jej adaptacyjność. Przykłady takich rozwiązań zostały przedstawione w kolejnych rozdziałach.

Bibliografia:

1. C. E. Shannon, *Coding Theorems for a Discrete Source with a Fidelity Criterion*, IRE /Nat. Conv. Rec., 4:142-163, 1959.
2. K. Sayood, *Introduction to Data Compression*, Morgan Kaufmann Publishers, 1996, rozdz. 8.
3. J. Max, *Quantizing for minimum distortion*, IRE Trans. Inform. Theory, IT-6:7-12, 1960.
4. S. P. Lloyd, *Least squares quantization in PCM*, IEEE Trans. Inform. Theory, IT-28:129-137, 1982. Jest to reprodukcja manuskryptu raportu technicznego Bell Laboratories z 1957 roku.
5. J. Łukaszewicz, and H. Steinhaus, *O mierzeniu przez kalibrowanie*, Zastosowania matematyki, s. 225-231, 1955.
6. N. S. Jayant, *Adaptive Quantization with One Word Memory*, Bell System Technical Journal, 52:1119-1144, 1973.
7. H. Gish, J. N. Pierce, *Asymptotically efficient quantizing*, IEEE Trans. Inform. Theory, IT-14:676-683, 1968.
8. N. Farvardin, and J. Modestino, *Optimum Quantizer Performance for a Class of Non-Gaussian Memoryless Sources*, IEEE Trans. Information Theory, IT-30:485-497, 1984.
9. A. Gersho, and R. M. Gray, *Vector Quantization and Signal Compression*, Kluwer Academic Publishers, Boston, 1992, rozdz. 2.
10. C. J. Kuo, and Y. F. Hsu, *Performance Comparison of Quantization Efficiency*, Proceedings of the National Science Council of ROC, 18(3):257-262, 1994.
11. Y. Wu, and D. C. Coll, *BTC-VQ-DCT Hybrid Coding of Digital Images*, IEEE Trans. Commun., 39(9):1283-1287, 1991.
12. E. J. Delp, and O. R. Mitchell, *Image Compression using block truncation coding*, IEEE Trans. Commun., COM-27(9): 1135-1142, 1979.
13. O. R. Mitchell, and E. J. Delp, *Multilevel graphics representation using block truncation coding*, Proc. IEEE, 68(7): 868-873, 1980.
14. V. R. Udpikar, and J. P. Raina, *BTC image coding using vector quantization*, IEEE Trans. Commun., COM-35(3):352-356, 1987.
15. J. M. Bramble, H.K. Huang, M. D. Murphy, *Image Data Compression*, Investigative Radiology, 23:707-712, 1988.
16. L-H. Chen, Y. K. Chen, *A High-Compression Image Coding Method Based on a New Segmentation Technique*, Proceedings of the National Science Council of ROC, 16(5):403-421, 1992.
17. K. Chen, T. V. Ramabadran, *Near-Lossless Compression of Medical Images Through Entropy-Coded DPCM*, IEEE Trans. Medical Imaging, 13(3):538-548, 1994.