

Wstęp do analizy falkowej
Notatki z wykładu

Maciej Paluszyński

13 czerwca 2008

Rozdział 1

Przestrzeń Hilberta

Sygnały

Funkcje (w języku inżynierów - sygnały) które będziemy rozważali na tym wykładzie będą kilku typów

Sygnały ciągłe (analogowe)

-) $L^2(\mathbf{R})$ to funkcje na prostej spełniające warunek

$$\int_{-\infty}^{\infty} |f(x)|^2 dx < +\infty,$$

(czyli funkcje całkowalne z kwadratem). Z punktu widzenia zastosowań to bardzo interesująca przestrzeń sygnałów. Jeżeli $f(x)$ oznacza, na przykład, napięcie jakiegoś przebiegu elektrycznego, to warunek całkowalności z kwadratem oznacza, że reprezentowany przez to napięcie sygnał ma skończoną całkowitą energię - bardzo rozsądne założenie.

-) $L^2([-\pi, \pi])$, przestrzeń funkcji okresowych o okresie 2π (czyli $f(x+2\pi) = f(x)$), całkowalnych z kwadratem w okresie

$$\int_{-\pi}^{\pi} |f(x)|^2 dx < +\infty.$$

Podobnie jak poprzednio ten warunek oznacza, że wartość energii sygnału w okresie jest skończona. Długość okresu nie jest szczególnie ważna, gdyż poprzez proste przeskalowanie można nasze rozważania przenieść na sygnały o innych okresach. Okres 2π wybrany jest dla wygody (to jest okres funkcji trygonometrycznych, których będziemy używali). Przestrzeń tę będziemy też oznaczali przez $L^2(\mathbf{T})$, gdzie $\mathbf{T}$ oznacza okrąg jednostkowy, czyli odcinek

$[-\pi, \pi]$ z utożsamionymi końcami.

-) Ogólniej, $L^2(\mathbf{R}^n)$ to funkcje n zmiennych rzeczywistych, całkowalnych z kwadratem. Szczególnie interesujący jest przypadek $n = 2$, sygnały takie reprezentują obrazy.
-) Ogólniej $L^2(\mathbf{T}^n)$, funkcje n zmiennych, względem każdej zmiennej okresowe o okresie 2π :

$$f(x_1, \dots, x_i + 2\pi, \dots, x_n) = f(x_1, \dots, x_i, \dots, x_n), \quad i = 1, \dots, n,$$

i całkowalne z kwadratem po swoim okresie:

$$\int_{-\pi}^{\pi} \cdots \int_{-\pi}^{\pi} |f(x_1, \dots, x_n)|^2 dx_1 \cdots dx_n < +\infty.$$

Również tutaj najbardziej interesujący (oprócz $n = 1$) jest przypadek $n = 2$.

Sygnały dyskretne (cyfrowe)

-) $L^2(\mathbf{Z})$ to przestrzeń ciągów podwójnie nieskończonych $\{f_k\}$, sumowalnych z kwadratem

$$\sum_{k=-\infty}^{\infty} |f_k|^2 < +\infty.$$

Przestrzeń tą często będziemy też oznaczać ℓ^2 .

-) $L^2(\mathbf{Z}_p)$, $p = 2, 3, \dots$ to przestrzeń ciągów podwójnie nieskończonych $\{f_k\}$ okresowych o okresie p , czyli $f_{k+p} = f_k \forall k \in \mathbf{Z}$. Takie ciągi są, rzecz jasna, automatycznie sumowalne z kwadratem po okresie:

$$\sum_{k=0}^{p-1} |f_k|^2 < +\infty.$$

Tą przestrzeń będziemy też czasem oznaczać ℓ_p^2 .

-) Ogólniej, $L^2(\mathbf{Z}^n)$ i $L^2(\mathbf{Z}_p^n)$ to przestrzenie ciągów n -wymiarowych okresowych ($L^2(\mathbf{Z}_p^n)$) lub nie ($L^2(\mathbf{Z}^n)$), sumowalnych z kwadratem, w przypadku $L^2(\mathbf{Z}_p^n)$ tylko po okresie. Jak poprzednio, najważniejsze przypadki to $n = 2$.

Uwaga 1.1. a) *Sygnały występujące w rzeczywistości (w naturze) są najczęściej ciągłe. Sygnały dyskretne pojawiają się jako wynik próbkowania sygnałów występujących w rzeczywistości, i to one pojawiają się w algorytmach numerycznych. Poznamy fundamentalne (ale bardzo proste) twierdzenie mówiące kiedy sygnał ciągły można całkowicie odtworzyć z ciągu próbek (mówiąc w skrócie, sygnał musi mieć skończone widmo częstotliwościowe, a próbkowanie musi być wystarczająco częste). W przypadku kiedy sygnału nie da*

się zrekonstruować z próbek (na przykład z jakichś powodów próbkowanie jest zbyt rzadkie) dowiemy się jak przygotować sygnał do próbkowania, aby uzyskać najlepszy efekt rekonstrukcji (będzie to tak zwany filtr antyaliasingowy).

b) Zauważmy, że wszystkie rozważane przestrzenie sygnałów mają pewne wspólne cechy, ważne z punktu widzenia teorii. Tworzą je funkcje (o wartościach zespolonych) na jakimś zbiorze, całkowalne lub sumowalne na tym zbiorze z kwadratem. Ta wspólna cecha pozwoli nam wprowadzić w tych przestrzeniach strukturę przestrzeni Hilberta, nasze podstawowe narzędzie. Drugą cechą wspólną rozważanych przestrzeni sygnałów jest to, że zbiór na którym te sygnały są rozważane ($\mathbf{R}^n$, $\mathbf{T}^n$, $\mathbf{Z}^n$, $\mathbf{Z}_p^n$) ma strukturę grupy abelowej, na przykład $\mathbf{Z}_p$ z dodawaniem modulo p . To z kolei pozwala nam korzystać z naszego drugiego podstawowego narzędzia - transformaty Fouriera - w jej wielu wcieleniach, transformaty ciągłej, dyskretnej czy szeregu Fouriera.

Przestrzeń Hilberta

Przypomnijmy krótko pojęcie przestrzeni liniowej, przestrzeni metrycznej i przestrzeni zupełnej.

Przestrzeń liniowa to taka, której elementy można dodawać, odejmować i mnożyć przez skalary (w naszym przypadku skalarami są liczby zespolone). Przestrzeń liniową często nazywa się też przestrzenią wektorową, a jej elementy wektorami.

Przestrzeń metryczna to taka, w której określona jest funkcja odległości $d(x, y)$ (metryka) dzięki której można zdefiniować zbieżność ciągu: $x_n \rightarrow x$ jeżeli $d(x_n, x) \rightarrow 0$. Otrzymujemy przestrzeń topologiczną, można mówić o ciągłości funkcji, czy zbieżności szeregów.

Przestrzeń metryczna zupełna to taka przestrzeń w której każdy ciąg Cauchy'ego jest zbieżny, czyli jeżeli ciąg $\{x_n\}$ elementów przestrzeni spełnia warunek

$$\forall \epsilon > 0 \exists N \in \mathbf{N} \forall n, m \geq N d(x_n, x_m) < \epsilon$$

to w tej przestrzeni istnieje x takie, że $x_n \rightarrow x$. Zupełność przestrzeni jest ważna z punktu widzenia teorii matematycznej. Wszystkie przestrzenie które będziemy rozważać są zupełne.

Przypomnijmy też funkcję wykładniczą zmiennej zespolonej, zdefiniowaną przez szereg potęgowy

$$e^z = \sum_{n=0}^{\infty} \frac{z^n}{n!}, \quad z - \text{liczba zespolona.}$$

Mamy następującą równość

$$e^{i\varphi} = \cos \varphi + i \sin \varphi. \quad (1.1)$$

Otrzymujemy ją wstawiając $i\varphi$ do definicji, korzystając z faktu, że $i^2 = -1$, rozdzielając składniki zawierające i i nie zawierające i (szereg jest zbieżny absolutnie) i korzystając z rozwinięć Taylora funkcji $\sin x$ i $\cos x$.

Definicja 1.2. *Przestrzeń liniowa E nazywa się przestrzenią Hilberta, jeżeli istnieje w niej iloczyn skalarny, to znaczy funkcja $\langle x, y \rangle$ (o wartościach zespolonych) o następujących własnościach*

- (a). $\langle x, y \rangle = \overline{\langle y, x \rangle}$ (jest antysymetryczny),
- (b). $\langle x+y, z \rangle = \langle x, z \rangle + \langle y, z \rangle$, $\langle \alpha x, y \rangle = \alpha \langle x, y \rangle$ (liniowy względem pierwszej zmiennej),
- (c). $\langle x, y+z \rangle = \langle x, y \rangle + \langle x, z \rangle$, $\langle x, \alpha y \rangle = \overline{\alpha} \langle x, y \rangle$ (antyliniowy względem drugiej zmiennej),
- (d). $\langle x, x \rangle \geq 0$ oraz $\langle x, x \rangle = 0 \Leftrightarrow x = 0$.

(x, y, z to dowolne elementy E , α jest dowolną liczbą zespoloną, a $\overline{\alpha}$ jest liczbą sprzężoną do α). Dodatkowo E musi być zupełna, kwestię metryki i zupełności wyjaśnimy za chwilę.

Jeżeli elementy x i y spełniają

$$\langle x, y \rangle = 0,$$

to mówimy, że są prostopadłe lub ortogonalne. Mając w przestrzeni Hilberta iloczyn skalarny wprowadzamy tak zwaną normę (długość) wektorów

$$\|x\| = \sqrt{\langle x, x \rangle}.$$

Zauważmy, że dzięki własności (d) iloczynu skalarnego pierwiastek można zawsze obliczyć.

Twierdzenie 1.3 (Nierówność Schwarz). *Dla dowolnych elementów x, y przestrzeni Hilberta E*

$$|\langle x, y \rangle| \leq \|x\| \|y\|.$$

Dowód. Ustalmy $x, y \in E$ i dowolną liczbę rzeczywistą λ .

$$\begin{aligned}\|x + \lambda y\|^2 &= \langle x + \lambda y, x + \lambda y \rangle \\ &= \langle x, x + \lambda y \rangle + \lambda \langle y, x + \lambda y \rangle \\ &= \langle x, x \rangle + \lambda \langle x, y \rangle + \lambda \langle y, x \rangle + \lambda^2 \langle y, y \rangle \\ &= \lambda^2 \|y\|^2 + 2\lambda \Re \langle x, y \rangle + \|x\|^2,\end{aligned}$$

($\Re$ - część rzeczywista). Rozważane wyrażenie jest więc (dla ustalonych x i y) funkcją kwadratową zmiennej rzeczywistej λ , o współczynnikach $\|y\|^2$, $2\Re \langle x, y \rangle$ i $\|x\|^2$. Wyrażenie nie może być ujemne, więc funkcja kwadratowa może mieć co najwyżej jeden rzeczywisty pierwiastek. A więc wyróżnik funkcji kwadratowej musi być ujemny:

$$(2\Re \langle x, y \rangle)^2 - 4\|y\|^2\|x\|^2 \leq 0,$$

czyli

$$|\Re \langle x, y \rangle| \leq \|x\| \|y\|.$$

Jeżeli $\langle x, y \rangle$ jest liczbą rzeczywistą to dowód jest zakończony. Jeżeli nie, to pozostaje jeszcze jeden trik: niech φ będzie liczbą rzeczywistą taką, że

$$\langle x, y \rangle = e^{i\varphi} |\langle x, y \rangle|.$$

Taką liczbę zawsze można znaleźć, $e^{i\varphi}$ jest „znakiem zespolonym” liczby $\langle x, y \rangle$ (chyba że $\langle x, y \rangle = 0$, ale w takim przypadku nierówność Schwarz’a jest natychmiastowa). Wtedy

$$e^{-i\varphi} \langle x, y \rangle = \langle x, e^{i\varphi} y \rangle$$

jest liczbą rzeczywistą. Wykorzystaliśmy tu równości

$$\frac{1}{e^{i\varphi}} = e^{-i\varphi} = \overline{e^{i\varphi}},$$

które łatwo wynikają z postaci trygonometrycznej (1.1). Z udowodnionej już części twierdzenia wynika, że

$$|\langle x, e^{i\varphi} y \rangle| \leq \|x\| \|e^{i\varphi} y\|.$$

W końcu, skoro, jak łatwo sprawdzić $|e^{i\varphi}| = 1$, mamy

$$|\langle x, y \rangle| = |e^{-i\varphi} \langle x, y \rangle|, \quad \text{ i } \quad \|e^{i\varphi} y\| = \|y\|.$$

□

Uwaga 1.4. Przyglądając się przedstawionemu wyżej dowodowi zauważmy, że równość (w nierówności Schwarza) zachodzi tylko jeżeli x i y są współliniowe (w sensie zespolonym), to znaczy istnieje liczba zespolona α taka, że

$$x = \alpha y.$$

Twierdzenie 1.5 (Własności normy). (a). $\|x\| \geq 0$ oraz $\|x\| = 0 \Leftrightarrow x = 0$,

(b). $\|ax\| = |a|\|x\|$ dla każdej liczby a i elementu $x \in E$,

(c). $\|x + y\| \leq \|x\| + \|y\|$ (nierówność trójkąta).

Dowód. Własności (a) i (b) wynikają wprost z definicji normy i własności iloczynu skalarnego. Sprawdźmy tylko nierówność trójkąta

$$\begin{aligned} \|x + y\|^2 &= \langle x + y, x + y \rangle = \langle x, x \rangle + \langle x, y \rangle + \langle y, x \rangle + \langle y, y \rangle \\ &= \|x\|^2 + \langle x, y \rangle + \langle y, x \rangle + \|y\|^2 \\ &= \left| \|x\|^2 + \langle x, y \rangle + \langle y, x \rangle + \|y\|^2 \right| \\ &\leq \|x\|^2 + |\langle x, y \rangle| + |\langle y, x \rangle| + \|y\|^2 \\ &\leq \|x\|^2 + 2\|x\| \|y\| + \|y\|^2 \\ &= (\|x\| + \|y\|)^2. \end{aligned}$$

Po drodze skorzystaliśmy z nierówności trójkąta dla liczb zespolonych

$$|a + b| \leq |a| + |b|,$$

oraz z nierówności Schwarza. □

Norma umożliwia nam wprowadzenie w E metryki

$$d(x, y) = \|x - y\|.$$

Wymagane własności metryki

(a). $d(x, y) = d(y, x)$,

(b). $d(x, y) \geq 0$ oraz $d(x, y) = 0 \Leftrightarrow x = y$,

(c). $d(x, y) \leq d(x, z) + d(z, y)$,

wynikają wprost z wyszczególnionych powyżej własności normy. Metryka wprowadza w E topologię, można więc mówić o zbieżności w przestrzeni Hilberta ciągów czy szeregów i o ciągłości funkcji. Z nierówności Schwarz'a wynika, że iloczyn skalarny jest ciągłą funkcją dwóch zmiennych. Przestrzeń Hilberta, z definicji, musi być zupełna jako przestrzeń metryczna z tą metryką.

Przykłady: Wszystkie opisane poprzednio przestrzenie sygnałów są przestrzeniami Hilberta. Żeby się o tym przekonać należy w każdej przestrzeni wprowadzić iloczyn skalarny i sprawdzić warunki (a)–(d) definicji. Należy też udowodnić zupełność powstałej przestrzeni metrycznej. W każdym przypadku przy określeniu iloczynu skalarnego będziemy korzystać z następującej nierówności, prawdziwej dla dowolnych liczb zespolonych:

$$2|ab| \leq |a|^2 + |b|^2. \quad (1.2)$$

• $L^2(\mathbf{R}^n)$. Jeżeli $f, g \in L^2(\mathbf{R}^n)$ to funkcja $f \cdot \bar{g}$ jest absolutnie całkowalna na $\mathbf{R}^n$:

$$|f(x)\overline{g(x)}| \leq \frac{|f(x)|^2}{2} + \frac{|\overline{g(x)}|^2}{2},$$

więc

$$\int_{\mathbf{R}^n} |f(x)\overline{g(x)}| dx \leq \frac{1}{2} \int_{\mathbf{R}^n} |f(x)|^2 dx + \frac{1}{2} \int_{\mathbf{R}^n} |g(x)|^2 dx < +\infty.$$

Iloczyn skalarny określamy następująco:

$$\langle f, g \rangle = \int_{\mathbf{R}^n} f(x)\overline{g(x)} dx. \quad (1.3)$$

Jak z tego wynika

$$\|f\| = \sqrt{\int_{\mathbf{R}^n} |f(x)|^2 dx}, \quad (1.4)$$

a odległość dwóch funkcji

$$d(f, g) = \|f - g\| = \sqrt{\int_{\mathbf{R}^n} |f(x) - g(x)|^2 dx}.$$

Zbieżność ciągu funkcji w przestrzeni Hilberta $L^2(\mathbf{R}^n)$ to nie jest to samo, co zbieżność punktowa. Na przykład, niech

$$f_n(x) = \begin{cases} 1 & : x \in [n, n+1], \\ 0 & : \text{poza tym.} \end{cases}$$

Jak łatwo zauważyć,

$$f_n(x) \rightarrow 0, \quad \forall x \in \mathbf{R},$$

czyli $\{f_n\}$ jest zbieżny w każdym punkcie do funkcji stale równej 0. Z drugiej strony, dla dowolnego n

$$\|f_n\|^2 = \int_{\mathbf{R}} |f_n(x)|^2 dx = \int_n^{n+1} dx = 1.$$

Ciąg nie jest więc zbieżny do 0 w $L^2(\mathbf{R})$. Można też podać przykład ciągu zbieżnego w $L^2(\mathbf{R})$, ale nie zbieżnego punktowo. Niech n będzie liczbą naturalną, i niech $n = 2^k + l$, dla pewnego $k = 0, 1, 2, \dots$ i $l = 0, \dots, 2^k - 1$. Mając taki rozkład n określamy

$$f_n(x) = \begin{cases} 1 & : x \in [2^{-k}l, 2^{-k}(l+1)), \\ 0 & : x \notin [2^{-k}l, 2^{-k}(l+1)). \end{cases}$$

Zauważmy, że ciąg $\{f_n\}$ nie jest zbieżny w żadnym punkcie $x \in [0, 1)$, natomiast

$$\|f_n\|^2 = \int_{2^{-k}l}^{2^{-k}(l+1)} dx = 2^{-k}.$$

Łatwo zauważyć, że $k \rightarrow \infty$ gdy $n \rightarrow \infty$, więc

$$\lim_{n \rightarrow \infty} \|f_n\| = 0,$$

ciąg $\{f_n\}$ zbiega więc do 0 w przestrzeni Hilberta.

Dla całkowitej ścisłości trzeba zrobić następującą uwagę, która odnosi się do wszystkich rozważanych przez nas przestrzeni sygnałów ciągłych. Niech

$$f(x) = \begin{cases} 1 & : x = 0, \\ 0 & : x \neq 0. \end{cases} \quad (1.5)$$

f nie jest funkcją zerową, ale $\|f\| = 0$. Iloczyn skalarny wprowadzony wzorem (1.3) nie spełnia warunku (d) definicji dla pewnej grupy specyficznych funkcji takich jak (1.5). Definicję przestrzeni $L^2(\mathbf{R}^n)$ moglibyśmy uściślić, posługując się pojęciem klasy abstrakcji. W przypadku tego kursu aż taka ścisłość nie jest potrzebna. Wystarczy pamiętać, że jeżeli dwie funkcje różnią się na niewielkim zbiorze, na przykład na zbiorze skończonym, to traktujemy je jako tą samą funkcję. Uwagi te nie dotyczą funkcji ciągłych. Jeżeli f i g są ciągłe, oraz

$$\|f - g\| = 0,$$

to $f = g$ wszędzie.

- $L^2(\mathbf{T}^n)$. Iloczyn skalarny wprowadzamy następująco

$$\langle f, g \rangle = \frac{1}{(2\pi)^n} \int_{-\pi}^{\pi} \cdots \int_{-\pi}^{\pi} f(x) \overline{g(x)} dx \quad x = (x_1, \dots, x_n),$$

podobnie jak poprzednio, korzystając z (1.2) pokazujemy, że całka zawsze istnieje.

- $L^2(\mathbf{Z}^n)$. Mając dwa ciągi z tej przestrzeni, $f = \{f_m\}$ i $g = \{g_m\}$, $m = (m_1, \dots, m_n)$ określamy

$$\langle f, g \rangle = \sum_{m_1, \dots, m_n = -\infty}^{\infty} f_m \overline{g_m}.$$

Ze względu na (1.2) szereg jest zbieżny absolutnie.

- $L^2(\mathbf{Z}_p^n)$. Podobnie dla ciągów okresowych

$$\langle f, g \rangle = \frac{1}{p^n} \sum_{m_1, \dots, m_n = 0}^{p-1} f_m \overline{g_m}.$$

W każdym z powyższych przypadków należy sprawdzić warunki (a)–(d) definicji przestrzeni Hilberta. Istotnym punktem do sprawdzenia pozostaje zupełność tak zdefiniowanych przestrzeni. W przypadku przestrzeni dyskretnych zupełność wynika z zupełności zbioru liczb rzeczywistych (każdy ciąg Cauchy’ego jest zbieżny). Do dowodu zupełności $\mathbf{R}$ trzeba dokładnie przyjrzeć się definicji samych liczb rzeczywistych. W przypadku przestrzeni sygnałów ciągłych w dowodzie zupełności korzysta się z konstrukcji całki Lebesgue’a (przestrzenie zdefiniowane przy użyciu całki Riemanna nie są zupełne). Dowód zupełności pomijamy. W dalszej części kursu wystarczy (chciałoby się powiedzieć „w zupełności”) nam sama świadomość tego, że przestrzenie są zupełne.

Bazy i rozpięcia

Zbiór elementów $\{e_n\}$ przestrzeni Hilberta E (skończony lub nieskończony) nazywa się liniowo niezależnym, jeżeli żaden jego element nie jest kombinacją liniową pozostałych. Można to zapisać następująco. Jeżeli dla jakichś liczb zespolonych $\alpha_1, \dots, \alpha_k$ zachodzi

$$\alpha_1 e_1 + \dots + \alpha_k e_k = 0,$$

to $\alpha_1 = \alpha_2 = \dots = \alpha_k = 0$. Na przykład, zbiór funkcji $\{f_n\}_{n=-\infty}^{\infty}$ gdzie

$$\Delta(x) = \begin{cases} x & : x \in [0, 1], \\ 2 - x & : x \in [1, 2], \\ 0 & : x \notin [0, 2], \end{cases} \quad (1.6)$$

oraz $f_n(x) = \Delta(x - n)$ jest liniowo niezależny w $L^2(\mathbf{R})$. Jeżeli $\alpha_1, \dots, \alpha_k$ są jakimiś liczbami zespolonymi, a $n_1, \dots, n_k$ różnymi liczbami całkowitymi, to funkcja

$$\alpha_1 f_{n_1} + \dots + \alpha_k f_{n_k}$$

ma wartość α_l w punkcie całkowitym $n_l + 1$. Jeżeli więc jest równa wszędzie 0, to $\alpha_1 = \alpha_2 = \dots = \alpha_k = 0$.

Jeżeli jednak do tego zbioru dodamy funkcję

$$f(x) = \begin{cases} x & : x \in [0, 2], \\ 6 - 2x & : x \in [2, 3], \\ 0 & : x \notin [0, 3], \end{cases}$$

to powstały zbiór nie jest już liniowo niezależny, bo

$$f(x) = f_0(x) + 2f_1(x),$$

a więc

$$f(x) - f_0(x) - 2f_1(x) = 0$$

a współczynniki 1, -1, -2 nie są zerami.

Zbiór elementów $\{e_n\}$ (znowu, skończony lub nie) nazywa się *ortonormalnym* jeżeli

$$\langle e_n, e_m \rangle = \begin{cases} 0 & : n \neq m, \\ 1 & : n = m. \end{cases}$$

Zbiór ortonormalny składa się więc z elementów o normie 1, wzajemnie ortogonalnych. Zauważmy, że zbiór ortonormalny jest zawsze liniowo niezależny. Niech $\{e_n\}$ będzie ortonormalny, i niech $\alpha_1, \dots, \alpha_k$ będą liczbami takimi, że

$$\alpha_1 e_1 + \dots + \alpha_k e_k = 0.$$

Weźmy dowolne $j = 1, \dots, k$ i obliczmy

$$\begin{aligned} 0 &= \langle \alpha_1 e_1 + \dots + \alpha_k e_k, e_j \rangle \\ &= \alpha_1 \langle e_1, e_j \rangle + \dots + \alpha_k \langle e_k, e_j \rangle \\ &= \alpha_j. \end{aligned}$$

W takim razie $\alpha_1 = \dots = \alpha_k = 0$. Pokazaliśmy więc że istotnie, zbiór ortonormalny jest liniowo niezależny.

Mając zbiór elementów $\{e_n\}$ przestrzeni Hilberta E *rozpięciem liniowym*

$$\text{Lin } \{e_n\}$$

nazywamy zbiór wszystkich kombinacji liniowych elementów $\{e_n\}$. Jest to najmniejsza podprzestrzeń liniowa przestrzeni E zawierająca wszystkie elementy zbioru $\{e_n\}$. Mówimy, że zbiór $\{e_n\}$ *rozpina* tę podprzestrzeń. Domknięcie tego zbioru

$$\overline{\text{Lin } \{e_n\}},$$

(czyli dołączenie do niego granic wszystkich zbieżnych ciągów) nazywamy *domkniętym rozpięciem liniowym* zbioru $\{e_n\}$.

Na przykład, rozpięcie liniowe zbioru funkcji danych przez (1.6) zawiera dokładnie te funkcje $f \in L^2(\mathbf{R})$, które spełniają następujące warunki: f jest ciągła, f jest liniowa na przedziałach postaci $[n, n+1]$, dla $n \in \mathbf{Z}$, i f jest 0 poza pewnym skończonym przedziałem. Z kolei domknięte rozpięcie liniowe tego zbioru to wszystkie funkcje $f \in L^2(\mathbf{R})$, ciągłe i liniowe na przedziałach postaci $[n, n+1]$. Jak łatwo sprawdzić

$$\overline{\text{Lin } \{f_n; n \in \mathbf{Z}\}} = \left\{ g = \sum_{n=-\infty}^{\infty} \alpha_n f_n; \sum_{n=-\infty}^{\infty} |\alpha_n|^2 < \infty \right\}.$$

Innymi słowy, warunkiem koniecznym i dostatecznym na to, żeby szereg

$$\sum_{n=-\infty}^{\infty} \alpha_n f_n$$

był zbieżny w przestrzeni $L^2(\mathbf{R})$ jest sumowalność z kwadratem ciągu współczynników α_n .

Można pokazać, że jeżeli zbiór $\{e_n\}$ jest skończony, to

$$\text{Lin } \{e_n\} = \overline{\text{Lin } \{e_n\}}.$$

Jeżeli mamy przeliczalny zbiór $\{e_n\}$ elementów przestrzeni Hilberta E , to zawsze możemy znaleźć zbiór ortonormalny $\{f_n\}$, o tym samym rozpięciu liniowym

$$\text{Lin } \{e_n\} = \text{Lin } \{f_n\}.$$

Konstrukcja zbioru $\{f_n\}$ nazywa się procedurą Gramma-Schmidta. Procedura jest indukcyjna. Niech elementy zbioru $\{e_n\}$ będą ustawione w ciąg

$e_1, e_2, \dots$. Jeżeli ciąg zawiera elementy zerowe to odrzucimy je – nie wpływa to na rozpięcie liniowe. Niech

$$f_1 = \frac{e_1}{\|e_1\|}.$$

Mamy więc początek procedury indukcyjnej. Teraz opiszemy krok. Załóżmy, że utworzony już został zbiór ortonormalny $\{f_1, \dots, f_k\}$ o następującej własności: istnieje n_k takie, że

$$\text{Lin}\{e_1, \dots, e_{n_k}\} = \text{Lin}\{f_1, \dots, f_k\}. \quad (1.7)$$

Zauważmy, że element f_1 zdefiniowany przed chwilą spełnia powyższy warunek, z $k = 1$, $n_1 = 1$. Żeby wykonać krok indukcyjny znajdziemy element ciągu $\{e_n\}$, następny po e_{n_k} , który nie należy do powyższego rozpięcia (1.7). Jeżeli takiego elementu w ciągu nie znajdziemy, innymi słowy wszystkie pozostałe elementy $e_{n_k+1}, e_{n_k+2}, \dots$ należą do rozpięcia (1.7), to procedura się kończy, i

$$\text{Lin}\{e_n; n = 1, 2, \dots\} = \text{Lin}\{f_1, \dots, f_k\}.$$

W tym wypadku procedura Gramma-Schmidta jest zakończona, a powstały zbiór ortonormalny jest skończony. Jeżeli natomiast znajdziemy element ciągu $\{e_n\}$, następny po e_{n_k} , który nie należy do rozpięcia (1.7) (niech to będzie pierwszy taki element), to nazywamy go $e_{n_{k+1}}$, i definiujemy f_{k+1}

$$f_{k+1} = \frac{e_{n_{k+1}} - \sum_{l=1}^k \langle e_{n_{k+1}}, f_l \rangle f_l}{\left\| e_{n_{k+1}} - \sum_{l=1}^k \langle e_{n_{k+1}}, f_l \rangle f_l \right\|}.$$

Wprost z powyższego wzoru wynika, że $\|f_{k+1}\| = 1$. Niech $j = 1, \dots, k$ i zauważmy, że f_{k+1} i f_j są ortogonalne

$$\begin{aligned} \langle f_{k+1}, f_j \rangle &= \frac{1}{\left\| e_{n_{k+1}} - \sum_{l=1}^k \langle e_{n_{k+1}}, f_l \rangle f_l \right\|} \left\langle \left(e_{n_{k+1}} - \sum_{l=1}^k \langle e_{n_{k+1}}, f_l \rangle f_l \right), f_j \right\rangle \\ &= \frac{1}{\left\| e_{n_{k+1}} - \sum_{l=1}^k \langle e_{n_{k+1}}, f_l \rangle f_l \right\|} \left(\langle e_{n_{k+1}}, f_j \rangle - \sum_{l=1}^k \langle e_{n_{k+1}}, f_l \rangle \langle f_l, f_j \rangle \right) \\ &= \frac{1}{\left\| e_{n_{k+1}} - \sum_{l=1}^k \langle e_{n_{k+1}}, f_l \rangle f_l \right\|} (\langle e_{n_{k+1}}, f_j \rangle - \langle e_{n_{k+1}}, f_j \rangle) \\ &= 0. \end{aligned}$$

Rozszerzony zbiór $\{f_1, \dots, f_{k+1}\}$ jest więc ortonormalny. Pozostaje zauważyć, że

$$\text{Lin}\{e_1, \dots, e_{n_{k+1}}\} = \text{Lin}\{f_1, \dots, f_{k+1}\},$$

co natychmiast wynika z założenia indukcyjnego i konstrukcji elementu f_{k+1} . Krok indukcyjny jest więc wykonany, i procedura Gramma-Schmidta daje w wyniku skończony lub nieskończony zbiór ortonormalny $\{f_n\}$ o tym samym rozpięciu liniowym co $\text{Lin}\{e_n\}$. Zauważmy, że skoro rozpięcia liniowe są identyczne, to także domknięte rozpięcia liniowe

$$\overline{\text{Lin}\{e_n\}} = \overline{\text{Lin}\{f_n\}}.$$

Przykład: Niech

$$e_1(x) = \begin{cases} 1 & : x \in [1, 2), \\ 0 & : x \notin [1, 2), \end{cases}$$

natomiast następne elementy będą dane wzorem

$$e_n(x) = \begin{cases} 1 & : x \in [n-1, n+1), \\ 0 & : x \notin [n-1, n+1), \end{cases} \quad n = 2, 3, \dots$$

Pokażemy, że zbiór $\{e_1, e_2, \dots\}$ jest liniowo niezależny. Niech będą dane współczynniki $\alpha_1, \dots, \alpha_k$ i niech

$$\alpha_1 e_1 + \dots + \alpha_k e_k = 0.$$

Weźmy liczbę całkowitą $j = 1, \dots, k-1$. Wtedy wartość funkcji po lewej stronie w punkcie j jest równa $\alpha_j + \alpha_{j+1}$, a więc

$$\alpha_j + \alpha_{j+1} = 0, \quad j = 1, \dots, k-1,$$

z kolei wartość lewej strony w punkcie k wynosi α_k , czyli $\alpha_k = 0$. Otrzymujemy więc $\alpha_1 = \dots = \alpha_k = 0$. Zbiór wektorów $\{e_n\}$ jest więc liniowo niezależny. Nie jest jednak ortonormalny. Na przykład

$$\begin{aligned} \langle e_1, e_2 \rangle &= \int_{-\infty}^{\infty} e_1(x) e_2(x) dx \\ &= \int_1^2 2 dx \\ &= 1. \end{aligned}$$

Zastosujmy więc procedurę Gramma-Schmidta. $\|e_1\| = 1$, więc $f_1 = e_1$.

$$e_2(x) - \langle e_2, f_1 \rangle f_1(x) = e_2(x) - e_1(x).$$

Otrzymujemy więc

$$f_2(x) = \begin{cases} 1 & : x \in [2, 3), \\ 0 & : x \notin [2, 3). \end{cases}$$

Kontynuując indukcyjnie otrzymujemy

$$f_n(x) = \begin{cases} 1 & : x \in [n, n+1), \\ 0 & : x \notin [n, n+1). \end{cases}$$

Zbiór $\{f_n\}$ jest ortonormalny i rozpiną tą samą podprzestrzeń co zbiór wyjściowy $\{e_n\}$: podprzestrzeń funkcji stałych na przedziałach postaci $[n, n+1)$, równych 0 dla $x < 1$ i $x \geq M$ dla pewnej liczby naturalnej M .

Mówimy, że zbiór $\{e_n\}_{n=1}^{\infty}$ elementów przestrzeni Hilberta E jest *zupełny*, jeżeli

$$\overline{\text{Lin}\{e_n\}} = E,$$

czyli każdy element przestrzeni E jest granicą ciągu kombinacji liniowych elementów zbioru $\{e_n\}$. Innymi słowy,

$$\forall x \in E \forall \epsilon > 0 \exists \alpha_1, \dots, \alpha_k \left\| x - \sum_{n=1}^k \alpha_n e_n \right\| < \epsilon.$$

Ta zupełność jest pojęciem „zpełnienie” związanym z zupełnością – w sensie przestrzeni metrycznej – dyskutowaną wcześniej. Można pokazać następujący fakt, często stosowany w sytuacji, gdy trzeba sprawdzić zupełność jakiegoś zbioru.

Fakt 1.6. *Zbiór $\{e_n\} \subset E$ jest zupełny wtedy i tylko wtedy, gdy jedynym elementem $x \in E$ prostopadłym do wszystkich e_n jest 0.*

Definicja 1.7. *Zbiór $\{e_n\}$ nazywa się bazą przestrzeni Hilberta E jeżeli jest liniowo niezależny i zupełny. $\{e_n\}$ nazywa się bazą ortonormalną jeżeli jest ortonormalny i zupełny.*

Mówiąc luźno, zbiór tworzy bazę jeżeli do każdego elementu przestrzeni E można podejść dowolnie blisko jakąś kombinacją liniową elementów bazy (zupełność), i zbiór nie zawiera żadnych zbędnych elementów (liniowa niezależność).

Uwaga 1.8. (a) Powyższa definicja różni się zasadniczo od pojęcia bazy przestrzeni liniowej wprowadzanego na wykładzie z algebry liniowej. W tamtej definicji każdy element przestrzeni można przedstawić jako kombinację liniową elementów bazy, a w tej definicji wystarczy, żeby do każdego elementu przestrzeni można było dowolnie blisko „podejść” kombinacjami liniowymi elementów bazy. Dla uniknięcia zamieszania tamtą, algebraiczną bazę czasami nazywa się „bazą Hamela”, a tę „bazą topologiczną”. Na tym wykładzie będziemy korzystali tylko z baz topologicznych, i będziemy je po prostu nazywali

bazami. W przypadku przestrzeni skończenie wymiarowych oba pojęcia baz są identyczne.

(b) Można udowodnić, że dwie bazy tej samej przestrzeni są równoliczne. Liczbę elementów bazy (może być nieskończona) nazywamy wymiarem przestrzeni. Rozważane przez nas przestrzenie są zarówno skończenie wymiarowe (na przykład wymiar przestrzeni ℓ_p^2 jest równy p) jak i nieskończenie wymiarowe (ℓ^2 czy przestrzenie sygnałów analogowych). Skończenie wymiarowe przestrzenie Hilberta są czasem nazywane przestrzeniami Euklidesowymi.

(c) Każdy układ liniowo niezależny można rozszerzyć do bazy (uzupełnić). Każdy układ ortonormalny można rozszerzyć do bazy ortonormalnej. Na tym wykładzie będziemy konstruować konkretne bazy w konkretnych przestrzeniach Hilberta.

Bazy ortonormalne są szczególnie wygodne w zastosowaniach. Poniżej przypomnimy na czym polega ta wygoda.

Twierdzenie 1.9. *Jeżeli $\{e_n\}$ jest bazą o.n. przestrzeni E to dla dowolnych liczb zespolonych $\alpha_1, \dots, \alpha_k$ i dowolnego $x \in E$ zachodzi nierówność*

$$\left\| x - \sum_{n=1}^k \langle x, e_n \rangle e_n \right\| \leq \left\| x - \sum_{n=1}^k \alpha_n e_n \right\|. \quad (1.8)$$

Równość zachodzi tylko w przypadku $\alpha_n = \langle x, e_n \rangle$, $n = 1, \dots, k$

Dowód.

$$\begin{aligned} \left\| x - \sum_{n=1}^k \alpha_n e_n \right\|^2 &= \left\| x - \sum_{n=1}^k \langle x, e_n \rangle e_n + \sum_{n=1}^k (\langle x, e_n \rangle - \alpha_n) e_n \right\|^2 \\ &= \left\| x - \sum_{n=1}^k \langle x, e_n \rangle e_n \right\|^2 + \\ &\quad + 2\Re \left\langle x - \sum_{n=1}^k \langle x, e_n \rangle e_n, \sum_{m=1}^k (\langle x, e_m \rangle - \alpha_m) e_m \right\rangle + \\ &\quad + \left\| \sum_{n=1}^k (\langle x, e_n \rangle - \alpha_n) e_n \right\|^2. \end{aligned}$$

Powyższą równość otrzymaliśmy jak zwykle: normę sumy do kwadratu zapisaliśmy jako iloczyn skalarny sumy przez siebie, i skorzystaliśmy z liniowości

iloczynu skalarnego. Rozważmy drugi składnik w uzyskanym wyrażeniu:

$$\begin{aligned} \left\langle x - \sum_{n=1}^k \langle x, e_n \rangle e_n, \sum_{m=1}^k (\langle x, e_m \rangle - \alpha_m) e_m \right\rangle &= \\ &= \sum_{m=1}^k \overline{(\langle x, e_m \rangle - \alpha_m)} \left\langle x - \sum_{n=1}^k \langle x, e_n \rangle e_n, e_m \right\rangle = \\ &= \sum_{m=1}^k \overline{(\langle x, e_m \rangle - \alpha_m)} \left(\langle x, e_m \rangle - \sum_{n=1}^k \langle x, e_n \rangle \langle e_n, e_m \rangle \right). \end{aligned}$$

Zauważmy, że ostatni nawias $= 0$, gdyż baza jest ortonormalna. Pokazaliśmy więc nierówność (1.8). Jeżeli zachodzi równość, to

$$\left\| \sum_{n=1}^k (\langle x, e_n \rangle - \alpha_n) e_n \right\| = 0.$$

Korzystając z własności normy kombinacja liniowa jest więc zerowa

$$\sum_{n=1}^k (\langle x, e_n \rangle - \alpha_n) e_n = 0.$$

Korzystając z liniowej niezależności elementów bazy otrzymujemy

$$\alpha_n = \langle x, e_n \rangle, \quad n = 1, \dots, k.$$

Udowodniliśmy więc ostatnią część twierdzenia. $\square$

Z definicji bazy wynika, że do każdego elementu można dowolnie blisko podejść *jakaś* kombinacją liniową elementów bazy. Powyższe twierdzenie mówi, że jeżeli baza jest ortonormalna, to najlepszymi kombinacjami liniowymi są kombinacje

$$\sum_{n=1}^k \langle x, e_n \rangle e_n.$$

Mamy więc konkretny wzór na współczynniki tych „najefektywniejszych” kombinacji.

Wniosek 1.10. *Jeżeli $\{e_n\}$ jest bazą o.n. przestrzeni E to:*

(a). *dla dowolnego $x \in E$*

$$x = \sum_{n=1}^{\infty} \langle x, e_n \rangle e_n, \quad (\text{rozkład } x \text{ względem bazy})$$

(b). jeżeli dla jakiegoś ciągu współczynników $\alpha_1, \alpha_2, \dots$ zachodzi

$$x = \sum_{n=1}^{\infty} \alpha_n e_n,$$

to współczynniki muszą być iloczynami skalarnymi

$$\alpha_n = \langle x, e_n \rangle \quad (\text{jednoznaczność rozkładu}),$$

(c). dla dowolnego $x \in E$

$$\|x\|^2 = \sum_{n=1}^{\infty} |\langle x, e_n \rangle|^2 \quad (\text{równość Plancherela}).$$

W powyższym wniosku przyjęliśmy, że baza jest nieskończona. Oczywiście, jeżeli jest skończona, to wniosek też jest prawdziwy a wszystkie sumy nieskończone zastępujemy skończonymi

Dowód. (a) Korzystamy z poprzedniego twierdzenia. Niech $k \geq m$ i niech $\alpha_n = \langle x, e_n \rangle$ dla $n = 1, \dots, m$ i $\alpha_n = 0$ dla $n = m+1, \dots, k$. Nierówność z poprzedniego twierdzenia wygląda więc następująco

$$\left\| x - \sum_{n=1}^k \langle x, e_n \rangle e_n \right\| \leq \left\| x - \sum_{n=1}^m \langle x, e_n \rangle e_n \right\|.$$

Dalej, z definicji bazy wiemy, że dla każdego $\epsilon > 0$ istnieją współczynniki $\alpha_1, \dots, \alpha_k$ takie, że

$$\left\| x - \sum_{n=1}^k \alpha_n e_n \right\| < \epsilon.$$

Z poprzedniego twierdzenia i poprzedniej uwagi wynika w takim razie, że dla tego ϵ i tego k mamy

$$\left\| x - \sum_{n=1}^l \langle x, e_n \rangle e_n \right\| \leq \left\| x - \sum_{n=1}^k \langle x, e_n \rangle e_n \right\| \leq \left\| x - \sum_{n=1}^k \alpha_n e_n \right\| < \epsilon \quad \forall l \geq k,$$

a więc szereg

$$\sum_{n=1}^{\infty} \langle x, e_n \rangle e_n,$$

jest zbieżny do x .

(b) Jeżeli

$$x = \sum_{n=1}^{\infty} \alpha_n e_n,$$

to stosując iloczyn skalarny i korzystając z jego ciągłości mamy

$$\langle x, e_n \rangle = \left\langle \sum_{m=1}^{\infty} \alpha_m e_m, e_n \right\rangle = \sum_{m=1}^{\infty} \alpha_m \langle e_m, e_n \rangle = \alpha_n.$$

(c) Korzystając z ciągłości iloczynu skalarnego otrzymujemy

$$\begin{aligned} \|x\|^2 &= \langle x, x \rangle = \left\langle \sum_{n=1}^{\infty} \langle x, e_n \rangle e_n, \sum_{m=1}^{\infty} \langle x, e_m \rangle e_m \right\rangle \\ &= \sum_{n=1}^{\infty} \sum_{m=1}^{\infty} \langle x, e_n \rangle \overline{\langle x, e_m \rangle} \langle e_n, e_m \rangle \\ &= \sum_{n=1}^{\infty} |\langle x, e_n \rangle|^2. \end{aligned}$$

□

Uwaga 1.11. (i) Z powyższego wniosku wynika, że jeżeli baza $\{e_n\}$ jest ortonormalna, to szereg

$$\sum_{n=1}^{\infty} \alpha_n e_n$$

jest zbieżny wtedy i tylko wtedy, gdy szereg współczynników jest sumowalny z kwadratem

$$\sum_{n=1}^{\infty} |\alpha_n|^2 < \infty.$$

(ii) Równości w (a) i (b) są równościami w przestrzeni E i, na przykład, równość w (a) oznacza

$$\lim_{N \rightarrow \infty} \left\| x - \sum_{n=1}^N \langle x, e_n \rangle e_n \right\| = 0.$$

Elementy rozważanych przez nas przestrzeni są funkcjami, a zbieżność szeregu funkcyjnego w normie przestrzeni nie oznacza z reguły zbieżności w każdym punkcie.

(iii) Równość Plancherela

$$\|x\|^2 = \sum_{n=1}^{\infty} |\langle x, e_n \rangle|^2$$

można rozumieć jako uogólnione (z przypadku 2-wymiarowego) twierdzenie Pitagorasa: kwadrat długości wektora jest sumą kwadratów jego składowych w kierunkach elementów bazy ortonormalnej.

Będziemy korzystać z następujących pojęć

Definicja 1.12. Zbiór $\{e_n\}$ elementów przestrzeni Hilberta E (skończony lub nie) nazywamy układem Riesz'a jeżeli istnieją stałe $A, B > 0$ takie, że dla dowolnego ciągu liczb $\{\alpha_n\}$

$$A \sum_{n=1}^{\infty} |\alpha_n|^2 \leq \left\| \sum_{n=1}^{\infty} \alpha_n e_n \right\|^2 \leq B \sum_{n=1}^{\infty} |\alpha_n|^2. \quad (1.9)$$

Jeżeli układ Riesz'a jest dodatkowo zbiorem zupełnym (rozpinają całą przestrzeń E), to nazywamy go bazą Riesz'a.

Warunek (1.9) wystarczy sprawdzić dla sum skończonych. Zauważmy też, że z warunku (1.9) wynika, że szereg

$$\sum_{n=1}^{\infty} \alpha_n e_n$$

jest zbieżny wtedy i tylko wtedy gdy szereg współczynników jest sumowalny z kwadratem

$$\sum_{n=1}^{\infty} |\alpha_n|^2 < \infty.$$

Można pokazać, że jeżeli w powyższej definicji $A = B = 1$, to $\{e_n\}$ jest układem ortonormalnym. Baza Riesz'a jest więc pojęciem ogólniejszym, niż baza ortonormalna. Bazy Riesz'a są często stosowane w praktyce. Posiadają wszystkie główne zalety baz ortonormalnych. Na przykład, można pokazać, że jeżeli baza $\{e_n\}$ jest bazą Riesz'a to istnieje inna baza Riesz'a $\{f_n\}$ taka, że dla każdego $x \in E$

$$x = \sum_{n=1}^{\infty} \langle x, f_n \rangle e_n = \sum_{n=1}^{\infty} \langle x, e_n \rangle f_n.$$

Dla bazy Riesz'a mamy więc, podobnie, jak dla bazy ortonormalnej, jawny wzór na znajdowanie współczynników bazowych, musimy tylko wygenerować wcześniej bazę $\{f_n\}$.

Przykład: Rozważmy ponownie przykład dany przez (1.6), $e_n(x) = \Delta(x - n)$. Wiemy, że elementy e_n , $n = 0, \pm 1, \pm 2, \dots$ są liniowo niezależne, ale nie tworzą zbioru ortonormalnego. Pokażemy, że tworzą zbiór Riesz'a. Weźmy

dowolne współczynniki α_n , $n = N, \dots, M$, gdzie $N, M \in \mathbf{Z}$.

$$\begin{aligned} \int_{-\infty}^{\infty} \left| \sum_{n=N}^M \alpha_n e_n(x) \right|^2 dx &= \int_{-\infty}^{\infty} \sum_{n=N}^M \alpha_n \Delta(x-n) \overline{\sum_{k=N}^M \alpha_k \Delta(x-k)} dx \\ &= \sum_{n,k=N}^M \alpha_n \overline{\alpha_k} \int_{-\infty}^{\infty} \Delta(x-n) \Delta(x-k) dx \\ &= \sum_{n,k=N}^M \alpha_n \overline{\alpha_k} \beta_{k-n}, \end{aligned}$$

gdzie

$$\beta_l = \int_{-\infty}^{\infty} \Delta(x) \Delta(x-l) dx = \begin{cases} \frac{2}{3} & : l = 0, \\ \frac{1}{6} & : l = \pm 1, \\ 0 & : |l| \geq 2. \end{cases}$$

Tak więc

$$\begin{aligned} \left\| \sum_{n=N}^M \alpha_n e_n \right\|^2 &= \beta_0 \sum_{n=N}^M |\alpha_n|^2 + \beta_1 \sum_{n=N}^{M-1} \alpha_n \overline{\alpha_{n+1}} + \beta_{-1} \sum_{n=N+1}^M \alpha_n \overline{\alpha_{n-1}} \\ &= \beta_0 \sum_{n=N}^M |\alpha_n|^2 + 2\beta_1 \Re \left(\sum_{n=N}^{M-1} \alpha_n \overline{\alpha_{n+1}} \right). \end{aligned}$$

Korzystając z nierówności Schwarza szacujemy ostatni składnik

$$\begin{aligned} \left| \Re \left(\sum_{n=N}^{M-1} \alpha_n \overline{\alpha_{n+1}} \right) \right| &\leq \left| \sum_{n=N}^{M-1} \alpha_n \overline{\alpha_{n+1}} \right| \\ &\leq \left(\sum_{n=N}^{M-1} |\alpha_n|^2 \right)^{1/2} \left(\sum_{n=N}^{M-1} |\overline{\alpha_{n+1}}|^2 \right)^{1/2} \\ &\leq \sum_{n=N}^M |\alpha_n|^2. \end{aligned}$$

Podsumowując powyższe otrzymujemy

$$(2/3 - 2/6) \sum_{n=N}^M |\alpha_n|^2 \leq \left\| \sum_{n=N}^M \alpha_n e_n \right\|^2 \leq (2/3 + 2/6) \sum_{n=N}^M |\alpha_n|^2.$$

Zgodnie z definicją, wektory $\{e_n\}_{n=-\infty}^{\infty}$ tworzą układ Riesz ze stałymi $A = 1/3$ i $B = 1$. Używając procedury Gramma-Schmidta możemy wyprodukować z układu Riesz $\{e_n\}_{n=-\infty}^{\infty}$ układ ortonormalny. Jednak nowy układ

ortonormalny nie posiadałby ważnych własności, na przykład tej, że wszystkie jego elementy są przesunięciami jednej funkcji-matki. Można udowodnić, że dla tego układu Riesza „dualny” układ $\{f_n\}_{n=-\infty}^{\infty}$, o którym była mowa powyżej również składa się z całkowitych przesunięć jednej funkcji Θ . Ta jedna funkcja jest ciągła, liniowa pomiędzy sąsiednimi liczbami całkowitymi, a jej wartości w punktach $n \in \mathbf{Z}$ wynoszą

$$\Theta(n) = \sqrt{3}(\sqrt{3} - 2)^{|n|}, \quad n \in \mathbf{Z}.$$

Definicja 1.13. Zbiór elementów $\{e_n\}_{n=1}^{\infty}$ przestrzeni Hilberta E (niekoniecznie liniowo niezależny) nazywamy rozpięciem dokładnym jeżeli dla każdego $x \in E$

$$\|x\|^2 = \sum_{n=1}^{\infty} |\langle x, e_n \rangle|^2.$$

Z definicji wynika, że rozpięcie dokładne musi być zbiorem zupełnym. Jeżeli $\{e_n\}$ jest rozpięciem dokładnym i jest zbiorem liniowo niezależnym, to musi być bazą ortonormalną. Rozpięcie dokładne jest więc pojęciem ogólniejszym od bazy ortonormalnej, często wykorzystywanym w zastosowaniach. Jeżeli $\{e_n\}$ jest rozpięciem dokładnym, to istnieje inne rozpięcie dokładne $\{f_n\}$ takie, że dla każdego $x \in E$

$$x = \sum_{n=1}^{\infty} \langle x, f_n \rangle e_n = \sum_{n=1}^{\infty} \langle x, e_n \rangle f_n.$$

W przypadku rozpięcia dokładnego również mamy więc sytuację, gdzie każdy x można przedstawić jako kombinację elementów rozpięcia, i są jawne wzory na współczynniki tego rozwinięcia.

Przypomnijmy jeszcze kilka pojęć, z których będziemy korzystać. Jeżeli $H \subset E$ jest podprzestrzenią, to *dopełnieniem ortogonalnym* H nazywamy podprzestrzeń

$$H^{\perp} = \{x \in E : x \perp y \ \forall y \in H\},$$

($x \perp y$ oznacza $\langle x, y \rangle = 0$). Dopełnienie ortogonalne podprzestrzeni jest zawsze domknięte, co wynika z ciągłości iloczynu skalarnego. Jeżeli $H \subset E$ jest podprzestrzenią domkniętą to istnieje rzut *prostopadły* na H , czyli przekształcenie liniowe $P : E \rightarrow H$ takie, że

$$\begin{aligned} Px &= x & \forall x \in H, \\ Px &= 0 & \forall x \in H^{\perp}. \end{aligned}$$

P przypisuje dowolnemu $x \in E$ najbliższy mu element podprzestrzeni H . Jeżeli w podprzestrzeni H mamy bazę o.n. $\{e_n\}$, to

$$Px = \sum_{n=1}^{\infty} \langle x, e_n \rangle e_n.$$

Jeżeli G i H są dwoma podprzestrzeniami przestrzeni Hilberta E , to mówimy, że E jest *ortogonalną sumą prostą* podprzestrzeni G i H jeżeli każdy element $x \in E$ można jednoznacznie zapisać jako

$$x = x_1 + x_2, \quad x_1 \in G, \quad x_2 \in H,$$

oraz $G \perp H$ (czyli $x \perp y$ dla dowolnych $x \in G$ i $y \in H$). Piszemy wtedy

$$E = G \oplus H.$$

W takiej sytuacji podprzestrzenie G i H muszą być domknięte.

Uwaga 1.14. (i) Jeżeli $E = G \oplus H$ to żeby skonstruować bazę w E wystarczy osobno skonstruować bazy w G i w H . Ich suma, jako zbiorów, będzie bazą E . Jeżeli te bazy w G i H są ortonormalne, to ich suma też jest ortonormalna. (ii) Jeżeli H jest domkniętą podprzestrzenią E , to

$$E = H \oplus H^\perp.$$

Rozdział 2

Przekształcenie Fouriera

Przypomnimy podstawowe zagadnienia związane z transformatą Fouriera. Transformata Fouriera przekształca wyjściową funkcję w ten sposób, że wartości transformaty nie informują o wartości samej funkcji w jakimkolwiek punkcie, ale informują jaka jest zawartość składowej o danej *częstotliwości* w funkcji wyjściowej. W związku z tym obliczanie transformaty Fouriera nazywa się czasem *analizą częstotliwościową*, *analizą harmoniczną* lub *analizą spektralną* funkcji. Transformata Fouriera stanowi przekształcenie funkcji które zachowuje całą informację, i jest odwracalne. Funkcję wyjściową można odtworzyć z transformaty, przy pomocy transformaty odwrotnej. Transformata Fouriera stanowi więc, intuicyjnie, rozkład funkcji na składowe częstotliwościowe. Nie jest to ściśle stwierdzenie, bo na przykład w przypadku funkcji w $L^2(\mathbf{R})$ transformata jest funkcją zmiennej rzeczywistej, więc musimy rozważać wszystkie częstotliwości będące liczbami rzeczywistymi, a funkcja wyjściowa w ogóle nie jest okresowa.

Transformata Fouriera jest, matematycznie, pojęciem bardzo ogólnym, które występuje w wielu sytuacjach. Dla matematyka nie jest więc niczym dziwnym, że podobne pojęcie występuje w tak wielu formach w praktyce. Jest ciągła transformata Fouriera, jest rozwinięcie w szereg Fouriera, w końcu jest dyskretna transformata Fouriera. Z punktu widzenia zastosowań ogólne podejście do transformaty Fouriera nie jest interesujące, ale kilka faktów warto znać. Jeżeli G jest lokalnie zwartą grupą abelową to rozważamy homomorfizmy grupy G w grupę liczb zespolonych o module 1 (z mnożeniem):

$$f : G \rightarrow \{z \in \mathbf{C} : |z| = 1\}.$$

Każdy taki homomorfizm nazywa się charakterem. Charaktery można mnożyć (tak jak mnożymy funkcje, punktowo), i z tak zdefiniowanym działaniem tworzą grupę abelową. Grupę charakterów nazywamy grupą dualną do G i

oznaczamy $\hat{G}$. Można zdefiniować „transformatę Fouriera”, która jest wzajemnie jednoznaczny izometrycznym przekształceniem

$$\mathfrak{F} : L^2(G) \leftrightarrow L^2(\hat{G}), \quad f \mapsto \hat{f}, \quad \hat{f}(\xi) = \langle f, \xi \rangle,$$

(przypomnijmy, że charakter ξ jest funkcją na G). Nie będziemy zajmować się tą ogólną teorią, w szczególności nie będziemy wyjaśniać szczegółów powyższych wzorów czy całkowania na G i $\hat{G}$. Wspomnijmy jeszcze, że grupą dualną do $\mathbf{R}$ jest też $\mathbf{R}$, a grupą dualną do $\mathbf{T}$ (liczby rzeczywiste z dodawaniem modulo 2π) jest $\mathbf{Z}$, i vice versa.

Transformata Fouriera w $L^2(\mathbf{R})$

Niech funkcja f będzie całkowalna na $\mathbf{R}$. Wtedy funkcja $f(x)e^{-ix\xi}$ też jest całkowalna, dla każdego $\xi \in \mathbf{R}$. Transformatą Fouriera funkcji f nazywamy funkcję

$$\hat{f}(\xi) = \int_{-\infty}^{\infty} f(x)e^{-ix\xi} dx. \quad (2.1)$$

Korzystając z twierdzenia o zbieżności ograniczonej otrzymujemy natychmiast następujący fakt.

Fakt 2.1. *Transformata Fouriera funkcji całkowalnej jest funkcją ciągłą i ograniczoną na $\mathbf{R}$.*

Przykłady: (a) Funkcja charakterystyczna przedziału $[-1/2, 1/2]$;

$$f(x) = \chi_{[-\frac{1}{2}, \frac{1}{2}]}(x) = \begin{cases} 1 & : x \in [-\frac{1}{2}, \frac{1}{2}], \\ 0 & : x \notin [-\frac{1}{2}, \frac{1}{2}]. \end{cases}$$

Niech $\xi \neq 0$. Wtedy

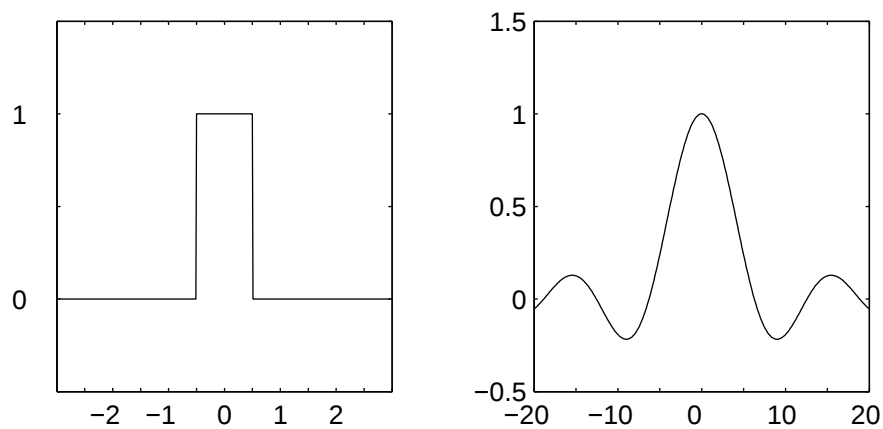
$$\hat{f}(\xi) = \int_{-\infty}^{\infty} \chi_{[-\frac{1}{2}, \frac{1}{2}]}(x)e^{-ix\xi} dx = \int_{-\frac{1}{2}}^{\frac{1}{2}} e^{-ix\xi} dx = \left. \frac{e^{-ix\xi}}{-i\xi} \right|_{-\frac{1}{2}}^{\frac{1}{2}} = \frac{1}{-i\xi} (e^{i\xi/2} - e^{-i\xi/2}) = \frac{\sin(\xi/2)}{\xi/2}.$$

Dla $\xi = 0$ rachunek jest jeszcze prostszy:

$$\hat{f}(0) = \int_{-\infty}^{\infty} \chi_{[-\frac{1}{2}, \frac{1}{2}]}(x) dx = 1.$$

(b) Jądro Gaussa-Weierstrassa

$$w(x) = \frac{1}{\sqrt{2\pi}} e^{-\frac{x^2}{2}}.$$



Rysunek 2.1: Funkcja z przykładu (a) i jej transformata Fouriera

Wiadomo, że

$$\hat{w}(0) = \int_{-\infty}^{\infty} w(x) dx = 1.$$

To jest jedna z powszechnie znanych całek, którą można obliczyć na przykład używając współrzędnych biegunowych w całce podwójnej. $w(x)$ jest funkcją całkowalną na $\mathbf{R}$, i

$$\hat{w}(\xi) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} e^{-\frac{x^2}{2}} e^{-ix\xi} dx.$$

Transformatę tę obliczymy korzystając z następującego triku: niech $F(\xi) = \hat{w}(\xi)$. Wtedy różniczkując pod znakiem całki i całkując przez części otrzy-

mujemy:

$$\begin{aligned}
F'(\xi) &= \frac{d}{d\xi} \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} e^{-\frac{x^2}{2}} e^{-ix\xi} dx \\
&= \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} e^{-\frac{x^2}{2}} \frac{d}{d\xi} (e^{-ix\xi}) dx \\
&= \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} e^{-\frac{x^2}{2}} (-ix) e^{-ix\xi} dx \\
&= \frac{i}{\sqrt{2\pi}} \int_{-\infty}^{\infty} \frac{d}{dx} \left(e^{-\frac{x^2}{2}} \right) e^{-ix\xi} dx \\
&= \frac{i}{\sqrt{2\pi}} e^{-\frac{x^2}{2}} e^{-ix\xi} \Big|_{-\infty}^{\infty} - \frac{i}{\sqrt{2\pi}} \int_{-\infty}^{\infty} e^{-\frac{x^2}{2}} \frac{d}{dx} (e^{-ix\xi}) dx \\
&= 0 - \frac{\xi}{\sqrt{2\pi}} \int_{-\infty}^{\infty} e^{-\frac{x^2}{2}} e^{-ix\xi} dx \\
&= -\xi F(\xi).
\end{aligned}$$

Transformata $F(\xi)$ spełnia więc proste równanie różniczkowe

$$F'(\xi) + \xi F(\xi) = 0.$$

Jest to proste równanie stopnia 1 ze zmiennymi rozdzielonymi, które można łatwo rozwiązać:

$$F(\xi) = c e^{-\frac{\xi^2}{2}}.$$

Podstawiając $\xi = 0$ w końcu otrzymujemy

$$\hat{w}(\xi) = e^{-\frac{\xi^2}{2}}.$$

Jądro Gaussa-Weierstrassa jest więc niezmiennikiem transformaty Fouriera. Jest to jedna z przyczyn, dla których ta funkcja pojawia się w wielu sytuacjach. Dla osób, które nie lubią rozwiązywać równań różniczkowych, a które lubią całkowanie po krzywych na płaszczyźnie w Dodatku obliczymy tą transformatę w inny sposób.

Definicja 2.2. *Splotem funkcji f i g określonych na $\mathbf{R}$ nazywamy funkcję*

$$f * g(x) = \int_{-\infty}^{\infty} f(x-y) g(y) dy, \quad (2.2)$$

o ile całka (2.2) istnieje dla każdego $x \in \mathbf{R}$.

Splot jest przemienny ($f * g = g * f$). Jeżeli f i g są całkowalne to splot istnieje i też jest całkowalny:

$$\begin{aligned} \int_{-\infty}^{\infty} |f * g(x)| dx &= \int_{-\infty}^{\infty} \left| \int_{-\infty}^{\infty} f(x-y) g(y) dy \right| dx \\ &\leq \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} |f(x-y)| |g(y)| dy dx \\ &= \int_{-\infty}^{\infty} |g(y)| \int_{-\infty}^{\infty} |f(x-y)| dx dy \\ &= \int_{-\infty}^{\infty} |g(y)| dy \int_{-\infty}^{\infty} |f(x)| dx. \end{aligned}$$

Sploty i filtry

Pojęcie splotu będzie dla nas narzędziem w dowodach. Warto jednak wspomnieć, że splot jest jednym z podstawowych pojęć w praktyce przetwarzania sygnału. Zauważmy, że zgodnie z powyższą obserwacją, jeżeli funkcja φ jest całkowalna, to splot z φ jest przekształceniem w przestrzeni funkcji całkowalnych:

$$T_{\varphi}(f)(x) = (f * \varphi)(x).$$

Przekształcenie T_{φ} jest liniowe (wynika to z liniowości całki), oraz przemienne z przesunięciami. Oznaczmy przez τ_{x_0} przesunięcie sygnału o x_0

$$(\tau_{x_0}f)(x) = f(x - x_0).$$

Wtedy

$$\begin{aligned} T_{\varphi}(\tau_{x_0}f)(x) &= \int_{-\infty}^{\infty} (\tau_{x_0}f)(x-y) \varphi(y) dy \\ &= \int_{-\infty}^{\infty} f(x-y-x_0) \varphi(y) dy \\ &= \int_{-\infty}^{\infty} f((x-x_0)-y) \varphi(y) dy \\ &= (T_{\varphi}(f))(x-x_0) \\ &= \tau_{x_0}(T_{\varphi}(f))(x). \end{aligned}$$

Innymi słowy

$$T_{\varphi} \cdot \tau_{x_0} = \tau_{x_0} \cdot T_{\varphi}.$$

Splot funkcji całkowalnej g z funkcją $f \in L^2(\mathbf{R})$ jest funkcją z $L^2(\mathbf{R})$, i mamy nierówność

$$\|f * g\| \leq \|f\| \int_{-\infty}^{\infty} |g(x)| dx.$$

To jest prosty fakt wynikający z własności całki (splot nie musi istnieć w każdym punkcie, tak jak jest w przypadku splotu dwóch funkcji całkowalnych, ale istnieje w wystarczająco wielu punktach — w prawie wszystkich punktach — aby określić funkcję w $L^2(\mathbf{R})$). Z tego faktu będziemy korzystać w dowodzie twierdzenia Plancherela. Warto zwrócić uwagę, że powyższa nierówność jest zupełnie naturalna, i można o niej myśleć jako o nierówności trójkąta. Intuicyjnie, możemy przybliżyć splot sumą Riemanna

$$f * g(x) \approx \sum_{k=-M}^M f\left(x - \frac{k}{N}\right) g\left(\frac{k}{N}\right) \frac{1}{N}.$$

Stosując do lewej strony normę $L^2(\mathbf{R})$ i wykorzystując nierówność trójkąta, otrzymujemy

$$\begin{aligned} \|f * g\| &\approx \left\| \sum_{k=-M}^M f\left(\cdot - \frac{k}{N}\right) g\left(\frac{k}{N}\right) \frac{1}{N} \right\| \\ &\leq \sum_{k=-M}^M \left\| f\left(\cdot - \frac{k}{N}\right) \right\| \left| g\left(\frac{k}{N}\right) \right| \frac{1}{N} \\ &= \|f\| \sum_{k=-M}^M \left| g\left(\frac{k}{N}\right) \right| \frac{1}{N} \\ &\approx \|f\| \int_{-\infty}^{\infty} |g(x)| dx. \end{aligned}$$

Kropka zastępuje zmienną całkowania występującą w normie $L^2(\mathbf{R})$ (której nie piszemy). Skorzystaliśmy z faktu, że norma $L^2(\mathbf{R})$ jest niezmiennicza na przesunięcia

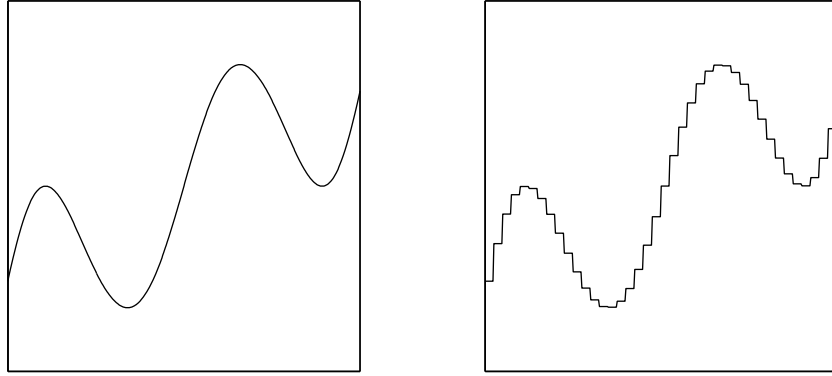
$$\|f(\cdot - t)\| = \|f\|.$$

Przekształcenie sygnału, które jest liniowe, przemienne z przesunięciami, i spełnia jeszcze pewne, niewielkie założenie ciągłości w zastosowaniach nazywa się filtrem. Filtry to właśnie sploty z funkcjami. Widzieliśmy, że splot z funkcją jest przekształceniem liniowym i przemennym z przesunięciami. Założenie ciągłości też jest spełnione, chociaż nie będziemy się zajmować szczegółami. Sploty są więc filtrami. Na odwrót też: okazuje się, że każdy filtr jest splotem. Łatwo się o tym przekonać intuicyjnie. Niech H będzie przekształceniem liniowym, przemennym z przesunięciami. Sygnał f przybliżymy funkcją schodkową:

$$f(x) \approx \sum_{k=-M}^M f\left(\frac{k}{N}\right) \chi_{[k/N, (k+1)/N)}(x),$$

gdzie $\chi_{[a,b]}$ jest funkcją charakterystyczną

$$\chi_{[a,b]}(x) = \begin{cases} 1 & : x \in [a, b) \\ 0 & : x \notin [a, b). \end{cases}$$



Rysunek 2.2: Sygnał i jego przybliżenie funkcją schodkową

Wtedy

$$\begin{aligned} (Hf)(x) &\approx H \left(\sum_{k=-M}^M f \left(\frac{k}{N} \right) \chi_{[k/N, (k+1)/N)} \right) (x) \\ &= \sum_{k=-M}^M f \left(\frac{k}{N} \right) H(\chi_{[k/N, (k+1)/N)})(x) \\ &= \sum_{k=-M}^M f \left(\frac{k}{N} \right) H(\chi_{[0, 1/N)} \left(x - \frac{k}{N} \right) \\ &= \sum_{k=-M}^M f \left(\frac{k}{N} \right) H(N\chi_{[0, 1/N)} \left(x - \frac{k}{N} \right) \frac{1}{N} \\ &\approx \int_{-\infty}^{\infty} f(y) g(x - y) dy \\ &= f * g(x), \end{aligned}$$

gdzie

$$g(x) \approx H(N\chi_{[0, 1/N)})(x).$$

Warunek ciągłości nałożony na filtr H powinien umożliwić przejście graniczne w naszym przybliżeniu, czyli powinien umożliwić zastąpienie $\approx$ przez równość. Widzimy więc, że filtr H splata sygnał z pewną funkcją g , która, w przybliżeniu jest odpowiedzią filtru H na sygnał $N\chi_{[0,1/N)}$ dla dużych N . Z powyższych rozważań można wyciągnąć następujące dwa wnioski:

- (a). Działanie filtru sprowadza się do splotu z pewną funkcją g
- (b). Funkcja g jest odpowiedzią filtru na „impuls jednostkowy”, czyli funkcję dodatnią, niezerową tylko w małym otoczeniu zera, której całka jest 1.

Uwaga 2.3. *Mówiąc o g używamy słowa „funkcja”. Od razu jednak widać, że g może być czymś ogólniejszym od funkcji. Jeżeli H jest filtrem identycznościowym (czyli takim, który nic nie robi, $Hf = f$) to g nie może być funkcją. g , jako odpowiedź filtru na impuls jednostkowy sama jest takim „impulsem jednostkowym”. Inżynierowie g nazywają funkcją uogólnioną, a matematycy dystrybucją. Impuls jednostkowy jest właśnie przykładem funkcji uogólnionej. Osoby, które chciałyby lepiej zrozumieć opisywane tutaj luźno zagadnienia powinny zapisać się na jeden z naszych wykładów z zaawansowanej analizy lub analizy funkcjonalnej, natomiast na tym wykładzie poprzestaniemy na takich, intuicyjnych, uwagach.*

Naszym celem teraz jest udowodnienie twierdzenia Plancherela, do którego będziemy potrzebowali kilka faktów.

Fakt 2.4. *Jeżeli f i g są całkowalne, to*

$$\widehat{(f * g)}(\xi) = \hat{f}(\xi) \hat{g}(\xi). \quad (2.3)$$

Dowód. Podstawiamy do wzorów, i korzystamy z własności funkcji wykładniczej

$$\begin{aligned} \widehat{(f * g)}(\xi) &= \int_{-\infty}^{\infty} f * g(x) e^{-ix\xi} dx \\ &= \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} f(x-y) g(y) e^{-ix\xi} dy dx \\ &= \int_{-\infty}^{\infty} g(y) e^{-iy\xi} \int_{-\infty}^{\infty} f(x-y) e^{-i(x-y)\xi} dx dy \\ &= \int_{-\infty}^{\infty} g(y) e^{-iy\xi} dy \int_{-\infty}^{\infty} f(x) e^{-ix\xi} dx \\ &= \hat{g}(\xi) \hat{f}(\xi). \end{aligned}$$

□

Wróćmy na moment do naszych rozważań o filtrach. Mówiliśmy, że filtry to sploty z funkcjami lub funkcjami uogólnionymi. Powyższy Fakt przenosi się na przypadek splotu z funkcją uogólnioną. Działanie filtru na sygnale, po stronie transformaty Fouriera, sprowadza się więc do mnożenia przez tak zwaną „charakterystykę” filtru. Charakterystyka filtru to transformata Fouriera jego odpowiedzi impulsowej. Bardzo często opisując filtr inżynierowie podają jego charakterystykę. Mamy też następujący prosty wzór

Fakt 2.5. *Jeżeli f jest całkowalna i*

$$f_t(x) = \frac{1}{t} f\left(\frac{x}{t}\right), \quad t > 0, \quad (2.4)$$

to

$$\hat{f}_t(\xi) = \hat{f}(t\xi).$$

Dowód. Korzystamy ze wzoru na całkowanie przez podstawienie

$$\begin{aligned} \hat{f}_t(\xi) &= \int_{-\infty}^{\infty} \frac{1}{t} f\left(\frac{x}{t}\right) e^{-i x \xi} dx \\ &= \int_{-\infty}^{\infty} f(x) e^{-i t x \xi} dx \\ &= \hat{f}(t\xi). \quad \square \end{aligned}$$

□

Korzystając z powyższych dwóch faktów udowodnimy twierdzenie o istnieniu transformaty odwrotnej. Jeżeli funkcja f jest całkowalna, i jej transformata Fouriera $\hat{f}$, dana wzorem (2.1), też jest całkowalna, to funkcję f można odtworzyć z $\hat{f}$ przy pomocy wzoru

$$f(x) = \frac{1}{2\pi} \int_{-\infty}^{\infty} \hat{f}(\xi) e^{i x \xi} d\xi. \quad (2.5)$$

Ten wzór daje nam jeszcze jedną, intuicyjną wskazówkę, co do natury transformaty Fouriera. Funkcja f jest średnią ważoną oscylacji $x \mapsto e^{i x \xi}$, z wagami $|\hat{f}(\xi)|$.

Twierdzenie 2.6. *Jeżeli f i $\hat{f}$ są całkowalne ($\hat{f}$ dana wzorem (2.1)), to zachodzi wzór (2.5). Innymi słowy przekształcenie całkowe*

$$\check{F}(x) = \frac{1}{2\pi} \int_{-\infty}^{\infty} F(\xi) e^{i x \xi} d\xi \quad (2.6)$$

jest przekształceniem odwrotnym do transformaty Fouriera.

Dowód. Wprowadźmy następujące oznaczenia:

$$w_t(x) = \frac{1}{\sqrt{t}} w\left(\frac{x}{\sqrt{t}}\right) = \frac{1}{\sqrt{2\pi t}} e^{-\frac{x^2}{2t}}, \quad \text{a więc, z (2.4),} \quad \widehat{w}_t(\xi) = \widehat{w}(\sqrt{t}\xi) = e^{-\frac{t\xi^2}{2}},$$

oraz

$$f_t(x) = f * w_t(x), \quad \text{a więc, z (2.3),} \quad \widehat{f}_t(\xi) = \widehat{f}(\xi) \widehat{w}_t(\xi).$$

Następnie obliczamy

$$\begin{aligned} \frac{1}{2\pi} \int_{-\infty}^{\infty} \widehat{f}_t(\xi) e^{i\xi y} d\xi &= \frac{1}{2\pi} \int_{-\infty}^{\infty} \widehat{f}(\xi) \widehat{w}_t(\xi) e^{i\xi y} d\xi \\ &= \frac{1}{2\pi} \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} f(x) e^{-ix\xi} dx e^{-\frac{t\xi^2}{2}} e^{i\xi y} d\xi \\ &= \frac{1}{2\pi} \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} f(x) e^{-\frac{t\xi^2}{2}} e^{-i\xi(x-y)} d\xi dx \\ &= \frac{1}{\sqrt{2\pi t}} \sqrt{\frac{t}{2\pi}} \int_{-\infty}^{\infty} f(x) \int_{-\infty}^{\infty} e^{-\frac{t\xi^2}{2}} e^{-i\xi(x-y)} d\xi dx \\ &= \frac{1}{\sqrt{2\pi t}} \int_{-\infty}^{\infty} f(x) \widehat{w}_{\frac{1}{t}}(x-y) dx. \end{aligned}$$

Zauważmy, że

$$\frac{1}{\sqrt{2\pi t}} \widehat{w}_{\frac{1}{t}}(z) = \frac{1}{\sqrt{2\pi t}} e^{-\frac{z^2}{2t}} = w_t(z),$$

a więc, korzystając również z tego, że $w_t(z) = w_t(-z)$, kontynuując, otrzymujemy

$$\begin{aligned} \frac{1}{2\pi} \int_{-\infty}^{\infty} \widehat{f}_t(\xi) e^{i\xi y} d\xi &= \int_{-\infty}^{\infty} f(x) w_t(x-y) dx \\ &= f * w_t(y) \\ &= f_t(y). \end{aligned}$$

Innymi słowy pokazaliśmy, że wzór (2.5) zachodzi dla f_t :

$$f_t(y) = \frac{1}{2\pi} \int_{-\infty}^{\infty} \widehat{f}_t(\xi) e^{iy\xi} d\xi. \quad (2.7)$$

Chcielibyśmy teraz przejść w powyższym wzorze do granicy gdy $t \rightarrow 0$. To jest proste rozumowanie, korzystające z własności całki. Wiemy, że

$$\widehat{f}_t(\xi) = \widehat{f}(\xi) \widehat{w}_t(\xi) \xrightarrow{t \rightarrow 0} \widehat{f}(\xi), \quad \text{gdzie} \quad |\widehat{f}_t(\xi) e^{iy\xi}| \leq |\widehat{f}(\xi)|,$$

przy czym ta ostatnia funkcja jest całkowalna. Korzystając z twierdzenia o zbieżności ograniczonej otrzymujemy

$$f_t(y) = \frac{1}{2\pi} \int_{-\infty}^{\infty} \hat{f}_t(\xi) e^{iy\xi} d\xi \xrightarrow{t \rightarrow 0} \frac{1}{2\pi} \int_{-\infty}^{\infty} \hat{f}(\xi) e^{iy\xi} d\xi, \quad (2.8)$$

dla każdego $y \in \mathbf{R}$. Z drugiej strony, korzystając z faktu, że

$$\int_{-\infty}^{\infty} w_t(x) dx = \int_{-\infty}^{\infty} w(x) dx = 1, \quad \text{ i } \quad w(x) > 0,$$

otrzymujemy

$$\begin{aligned} \int_{-\infty}^{\infty} |f(y) - f_t(y)| dy &= \int_{-\infty}^{\infty} \left| \int_{-\infty}^{\infty} (f(y) - f(y-x)) w_t(x) dx \right| dy \\ &= \int_{-\infty}^{\infty} \left| \int_{-\infty}^{\infty} (f(y) - f(y - \sqrt{t}x)) w(x) dx \right| dy \\ &\leq \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} |f(y) - f(y - \sqrt{t}x)| w(x) dx dy \\ &= \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} |f(y) - f(y - \sqrt{t}x)| dy w(x) dx. \end{aligned}$$

Z własności całki wynika, że przesunięcia są ciągłe względem całki, czyli

$$\lim_{t \rightarrow 0} \int_{-\infty}^{\infty} |f(y) - f(y - \sqrt{t}x)| dy = 0.$$

Korzystając ponownie z twierdzenia o zbieżności ograniczonej otrzymujemy, że $f_t \rightarrow f$ w sensie całki:

$$\lim_{t \rightarrow 0} \int_{-\infty}^{\infty} |f_t(y) - f(y)| dy = 0.$$

Tak więc f_t jest zbieżna do f w sensie całki, i zbieżna w każdym punkcie $y \in \mathbf{R}$ do

$$\frac{1}{2\pi} \int_{-\infty}^{\infty} \hat{f}(\xi) e^{iy\xi} d\xi.$$

Z własności całki wynika, że obie te granice muszą być równe:

$$f(y) = \frac{1}{2\pi} \int_{-\infty}^{\infty} \hat{f}(\xi) e^{iy\xi} d\xi.$$

□

Wnioskiem z powyższego jest najważniejsze twierdzenie, tak zwane twierdzenie Plancherela.

Twierdzenie 2.7 (Plancherel). (a) *Niech $f \in L^2(\mathbf{R})$ będzie także całkowalna. Wtedy $\hat{f}$ też należy do $L^2(\mathbf{R})$, oraz*

$$\|f\|^2 = \frac{1}{2\pi} \|\hat{f}\|^2, \quad (2.9)$$

(b) *Podzbiór $L^2(\mathbf{R})$ składający się z funkcji całkowalnych stanowi gęstą podprzestrzeń $L^2(\mathbf{R})$. Transformata Fouriera rozszerza się z tej podprzestrzeni na całe $L^2(\mathbf{R})$. W przypadku gdy $f \in L^2(\mathbf{R})$ nie jest całkowalna transformatę można obliczyć ze wzoru*

$$\hat{f}(\xi) = \lim_{M \rightarrow \infty} \int_{-M}^M f(x) e^{-i x \xi} dx. \quad (2.10)$$

Tak określone przekształcenie jest wzajemnie jednoznaczne (1-1 i „na”), izometrycznym (z dokładnością do czynnika $\sqrt{2\pi}$, jak we wzorze (2.9)) przekształceniem $L^2(\mathbf{R})$ na siebie. Przekształcenie odwrotne, w przypadku gdy f nie jest całkowalna, dane jest wzorem

$$\check{f}(x) = \lim_{M \rightarrow \infty} \frac{1}{2\pi} \int_{-M}^M f(\xi) e^{i x \xi} d\xi. \quad (2.11)$$

Dowód. (a) Transformata $\hat{f}$ jest ograniczona (Fakt 2.1). Załóżmy najpierw dodatkowo, że $\hat{f}$ też jest całkowalna. Wtedy o f można myśleć jako o transformatie Fouriera funkcji $\hat{f}$, zgodnie z Twierdzeniem 2.6, i w takim razie f też jest ograniczona. Jak łatwo policzyć

$$\begin{aligned} \widehat{\hat{f}}(\xi) &= \int_{-\infty}^{\infty} \overline{\hat{f}(x)} e^{-i x \xi} dx \\ &= \overline{\int_{-\infty}^{\infty} f(x) e^{i x \xi} dx} \\ &= \overline{\hat{f}(-\xi)}, \end{aligned}$$

gdyż $\overline{e^{-ix\xi}} = e^{ix\xi}$. Podstawiamy to do wzoru, i obliczamy

$$\begin{aligned}
\int_{-\infty}^{\infty} |f(x)|^2 dx &= \int_{-\infty}^{\infty} f(x) \overline{f(x)} dx \\
&= \int_{-\infty}^{\infty} f(x) \frac{1}{2\pi} \int_{-\infty}^{\infty} \widehat{f}(\xi) e^{i\xi x} d\xi dx \\
&= \frac{1}{2\pi} \int_{-\infty}^{\infty} \overline{\widehat{f}(-\xi)} \int_{-\infty}^{\infty} f(x) e^{i\xi x} dx d\xi \\
&= \frac{1}{2\pi} \int_{-\infty}^{\infty} \widehat{f}(-\xi) \widehat{f}(-\xi) d\xi \\
&= \frac{1}{2\pi} \int_{-\infty}^{\infty} \widehat{f}(\xi) \overline{\widehat{f}(\xi)} d\xi \\
&= \frac{1}{2\pi} \int_{-\infty}^{\infty} |\widehat{f}(\xi)|^2 d\xi.
\end{aligned}$$

Pozbądźmy się teraz dodatkowego założenia, że $\widehat{f}$ jest całkowalna. Niech więc teraz f będzie w $L^2(\mathbf{R})$ i całkowalna. Podobnie jak w dowodzie poprzedniego twierdzenia, i z tamtymi oznaczeniami, funkcja

$$f_t = f * w_t$$

(splot z przeskalowanym jądrem Gaussa-Weierstrassa) jest całkowalna, w $L^2(\mathbf{R})$ i ma całkowalną transformatę Fouriera ($|\widehat{f}_t(\xi)| \leq c\widehat{w}_t(\xi)$). Fakt, że f_t jest w $L^2(\mathbf{R})$ wynika z własności całki: splot funkcji całkowalnej z całkowalną z kwadratem jest całkowalny z kwadratem (była o tym mowa przy okazji omawiania własności splotów). Dla f_t możemy więc skorzystać z przeprowadzonego już dowodu

$$\|f_t\|^2 = \frac{1}{2\pi} \|\widehat{f}_t\|^2. \quad (2.12)$$

Korzystając z własności całki, podobnie jak w poprzednim dowodzie możemy pokazać, że $f_t \rightarrow f$ w $L^2(\mathbf{R})$ gdy $t \rightarrow 0$, a więc $\|f_t\| \rightarrow \|f\|$ (norma jest funkcją ciągłą). Wynika to z ciągłości przesunięć w $L^2(\mathbf{R})$, (w poprzednim twierdzeniu korzystaliśmy z ciągłości przesunięć w przestrzeni funkcji całkowalnych). Z drugiej strony $\|\widehat{f}_t\| \rightarrow \|\widehat{f}\|$ korzystając z twierdzenia o zbieżności ograniczonej ($|\widehat{f}_t(\xi)|^2 \leq |\widehat{f}(\xi)|^2$). Przechodząc do granicy w (2.12) gdy $t \rightarrow 0$ otrzymujemy (2.9).

(b) Jeżeli $f \in L^2(\mathbf{R})$, to funkcje obcięte

$$f_n(x) = \begin{cases} f(x) & : |x| \leq n \\ 0 & : |x| > n \end{cases} \quad (2.13)$$

tworzą ciąg funkcji całkowalnych, zbieżny do f w $L^2(\mathbf{R})$:

$$\|f - f_n\|^2 = \int_{-\infty}^{\infty} |f(x) - f_n(x)|^2 dx = \int_{|x|>n} |f(x)|^2 dx \xrightarrow{n \rightarrow \infty} 0.$$

Każda funkcja $f \in L^2(\mathbf{R})$ jest więc granicą, w normie $L^2(\mathbf{R})$, ciągu funkcji całkowalnych. Transformatę Fouriera dla funkcji $f \in L^2(\mathbf{R})$, które nie są całkowalne określamy więc następująco. Niech $\{f_n\} \subset L^2(\mathbf{R})$ będzie ciągiem funkcji całkowalnych, zbieżnym do f :

$$f_n \rightarrow f \quad \text{w} \quad L^2(\mathbf{R}).$$

Ciąg $\{f_n\}$ jest ciągiem Cauchy'ego, a więc, zgodnie z (2.9), ciąg $\{\hat{f}_n\}$ też jest Cauchy'ego w $L^2(\mathbf{R})$, a więc jest zbieżny do jakiegoś elementu $F \in L^2(\mathbf{R})$. Ta granica to, z definicji, transformata Fouriera f . Zauważmy, że ta definicja nie zależy od wyboru ciągu $\{f_n\}$ funkcji całkowalnych, zbieżnego do f . Jeżeli dwa różne ciągi są zbieżne do f , to ich różnice tworzą ciąg zbieżny do zera:

$$f_n \rightarrow f, \quad g_n \rightarrow f \quad \text{w} \quad L^2(\mathbf{R}) \quad \Rightarrow \quad f_n - g_n \rightarrow 0 \quad \text{w} \quad L^2(\mathbf{R}).$$

Zgodnie z (2.9) transformaty Fouriera różnic też zbiegają do 0, a więc transformaty Fouriera tych dwóch ciągów mają tę samą granicę. Zauważmy, że oznacza to też, że gdy wyjściowa funkcja f jest całkowalna, to nowa definicja pokrywa się ze starą — jako ciąg funkcji całkowalnych zbieżny do f można wziąć ciąg stały stale równy f . W końcu zauważmy, że skoro ciąg funkcji obciętych (2.13) składa się z funkcji całkowalnych i jest zbieżny do f , to (2.10) wynika z podanej powyżej definicji, zastosowanej do tego konkretnego ciągu. Korzystając z ciągłości normy otrzymujemy (2.9) dla dowolnej funkcji $f \in L^2(\mathbf{R})$. W końcu pozostał do udowodnienia fakt, że transformata Fouriera jest „na” $L^2(\mathbf{R})$, oraz wzór (2.11). Niech f będzie dowolnym elementem $L^2(\mathbf{R})$. Niech f_n zbiega do f w $L^2(\mathbf{R})$, i składa się z funkcji całkowalnych, o całkowalnych transformatach Fouriera. Taki ciąg zawsze istnieje. Można pokazać, że funkcje na przykład różniczkowalne dwukrotnie, równe 0 poza pewnym skończonym przedziałem tworzą podzbiór gęsty $L^2(\mathbf{R})$, oraz mają wymagane własności: są całkowalne, oraz ich transformaty Fouriera są całkowalne. Z definicji $\hat{f}_n \rightarrow \hat{f}$, oraz $\hat{f}_n$ są całkowalne, a więc $\hat{\hat{f}}_n \rightarrow \hat{\hat{f}}$. Wiemy z poprzedniego twierdzenia, że $\hat{\hat{f}}_n(x) = 2\pi f_n(-x) \rightarrow 2\pi f(-x)$. A więc $f = \hat{\hat{F}}$, gdzie

$$F = \frac{1}{4\pi^2} \hat{\hat{f}}.$$

Z tych rachunków wynika (2.11) oraz to, że transformata Fouriera jest „na”. □

Własności transformaty w $L^2(\mathbf{R})$

Przy założeniu, że wszystkie występujące funkcje są w $L^2(\mathbf{R})$ mamy

$$\begin{aligned}g(x) = f(x - c) &\longrightarrow \hat{g}(\xi) = e^{-ic\xi} \hat{f}(\xi), \\g(x) = e^{i\omega x} f(x) &\longrightarrow \hat{g}(\xi) = \hat{f}(\xi - \omega), \\g(x) = f(x/s) &\longrightarrow \hat{g}(\xi) = s \hat{f}(s\xi), \quad s > 0, \\g(x) = f'(x) &\longrightarrow \hat{g}(\xi) = i\xi \hat{f}(\xi), \\g(x) = -ixf(x) &\longrightarrow \hat{g}(\xi) = \hat{f}'(\xi).\end{aligned}$$

Dowody zostawiamy jako ćwiczenie. Pokażemy natomiast zastosowanie transformaty Fouriera do rozwiązywania tak zwanego równania ciepła, lub równania dyfuzji.

Równanie ciepła

Wyobraźmy sobie pręt metalowy leżący na płaszczyźnie wzdłuż osi OX . Niech pręt będzie nieskończenie długi, o pomijalnie małej średnicy, oraz niech będzie zrobiony z materiału przewodzącego ciepło, na przykład jakiegoś metalu. Niech temperatura pręta w punkcie x w czasie t będzie oznaczona przez $u(x, t)$. Można pokazać, że funkcja u spełnia następujące równanie ciepła

$$\frac{\partial^2 u(x, t)}{\partial t^2} = \kappa^2 \frac{\partial^2 u(x, t)}{\partial x^2}. \quad (2.14)$$

Równanie to wynika z praw fizyki, i czasem nazywa się równaniem dyfuzji. Niech rozkład temperatury na pręcie będzie znany w czasie $t = 0$, czyli niech u spełnia *warunek brzegowy*

$$u(x, 0) = f(x),$$

dla zadanej funkcji f . Możemy założyć, że f spełnia jakieś warunki regularności. Nam wystarczy $f \in L^2(\mathbf{R})$. Klasyczny problem sprowadza się do znalezienia funkcji $u(x, t)$, spełniającej równanie ciepła, i zadany warunek brzegowy. Będziemy szukali funkcji u takiej, że dla każdego $t \geq 0$ $u(\cdot, t) \in L^2(\mathbf{R})$, oraz $\frac{\partial u}{\partial t}(\cdot, t) \in L^2(\mathbf{R})$. Okazuje się, że rozwiązanie można znaleźć używając transformaty Fouriera. Teoria transformaty Fouriera powstała właśnie przy okazji badania zagadnienia ciepła. Dla ustalonego $t \geq 0$ zastosujemy transformatę Fouriera, względem zmiennej x , do obu stron równania (2.14). Tak powstałą transformatę oznaczmy przez $F(\xi, t)$. Różniczkując względem t pod znakiem całki z (2.14) otrzymujemy

$$\frac{\partial F(\xi, t)}{\partial t} = -\kappa^2 \xi^2 F(\xi, t). \quad (2.15)$$

Ustalmy teraz ξ i spójrzmy na (2.15) jako na równanie różniczkowe zwyczajne funkcji zmiennej t . Równanie to łatwo rozwiązać, rozwiązanie ma postać

$$F(\xi, t) = c e^{-\kappa^2 \xi^2 t},$$

dla dowolnej stałej c . Podstawiając warunek początkowy (brzegowy), otrzymujemy

$$c = F(\xi, 0) = \hat{f}(\xi),$$

a więc

$$F(\xi, t) = \hat{f}(\xi) e^{-\kappa^2 \xi^2 t}, \quad \text{czyli} \quad u(x, t) = f * w_{\kappa^2 t}(x),$$

gdzie w_t jest znanym nam już jądrem Gaussa-Weierstrassa (zwanym także jądrem ciepła), przeskalowanym zgodnie z (2.4). Mamy więc ilustrację, jak własności transformaty Fouriera pozwalają zastosować ją do rozwiązywania równań różniczkowych. Dla inżynierów to jest główne zastosowanie transformaty Fouriera, i stanowi ona jedno z najważniejszych narzędzi inżynierskich. Poniżej udowodnimy tak zwaną zasadę nieoznaczoności Heisenberga. Zasada ta pokazuje pewien problem występujący w zastosowaniach transformaty Fouriera. Mówiąc bardzo ogólnie to jest właśnie problem, którego chcemy uniknąć zamieniając transformatę Fouriera na transformatę falkową. Zasada nieoznaczoności Heisenberga mówi, że jeżeli rozrzut wartości funkcji wokół jej wartości średniej jest mały (funkcja jest skupiona wokół jakiegoś punktu, i szybko maleje w miarę oddalania się od niego), to rozrzut wartości transformaty musi być duży (transformata Fouriera takiej funkcji nie może być skupiona wokół pewnej częstotliwości). Intuicyjnie to jest łatwe do uzasadnienia. W rzeczywistości jest to ściśle twierdzenie.

Twierdzenie 2.8 (Zasada nieoznaczoności Heisenberga). *Niech $f \in L^2(\mathbf{R})$ będzie unormowana, czyli $\|f\| = 1$, oraz taka, że $xf(x) \in L^2(\mathbf{R})$ oraz $\xi\hat{f}(\xi) \in L^2(\mathbf{R})$. Wtedy*

$$\int_{-\infty}^{\infty} x^2 |f(x)|^2 dx \int_{-\infty}^{\infty} \xi^2 |\hat{f}(\xi)|^2 d\xi \geq \frac{\pi}{2}.$$

Dowód. Wykorzystamy następujący wzór na całkowanie przez części. Jeżeli $h'g$, hg' oraz hg są całkowalne na $\mathbf{R}$, to

$$\int_{-\infty}^{\infty} h'(x) g(x) dx = - \int_{-\infty}^{\infty} h(x) g'(x) dx.$$

Zauważmy, że skoro $x^2|f(x)|^2$ oraz $|f(x)|^2$ są całkowalne, więc $x|f(x)|^2$ też jest całkowalna. Możemy więc skorzystać z powyższego wzoru na całkowanie

przez części:

$$\begin{aligned}
1 &= \int_{-\infty}^{\infty} |f(x)|^2 dx \\
&= - \int_{-\infty}^{\infty} x \frac{d}{dx} |f(x)|^2 dx \\
&= \left| - \int_{-\infty}^{\infty} \left(x f'(x) \overline{f(x)} + x f(x) \overline{f'(x)} \right) dx \right| \\
&\leq 2 \int_{-\infty}^{\infty} |x| |f(x)| |f'(x)| dx \\
&\leq 2 \left(\int_{-\infty}^{\infty} |x|^2 |f(x)|^2 dx \right)^{1/2} \left(\int_{-\infty}^{\infty} |f'(x)|^2 dx \right)^{1/2} \\
&\leq \frac{2}{\sqrt{2\pi}} \left(\int_{-\infty}^{\infty} |x|^2 |f(x)|^2 dx \right)^{1/2} \left(\int_{-\infty}^{\infty} |\xi|^2 |\hat{f}(\xi)|^2 d\xi \right)^{1/2}.
\end{aligned}$$

Przy okazji zauważmy, że równość w powyższym oszacowaniu może zachodzić tylko wtedy, gdy zachodzi równość w wykorzystanej nierówności Schwarza, czyli funkcje $x f(x)$ oraz $f'(x)$ muszą być współliniowe: musi istnieć stała zespolona c taka, że

$$x f(x) = c f'(x).$$

Rozwiązując to równanie różniczkowe otrzymujemy, że f musi być wielokrotnością jądra Gaussa-Weierstrassa. $\square$

Transformata Fouriera w $L^2(\mathbf{R}^n)$

Teoria w $L^2(\mathbf{R}^n)$ jest analogiczna do teorii 1-wymiarowej. Jeżeli f lub F są całkowalne, to

$$\begin{aligned}
\hat{f}(\xi) &= \int_{\mathbf{R}^n} f(x) e^{-i x \cdot \xi} dx, \\
\check{F}(x) &= \frac{1}{(2\pi)^n} \int_{\mathbf{R}^n} F(\xi) e^{i \xi \cdot x} d\xi,
\end{aligned}$$

gdzie $x \cdot \xi$ oznacza zwykły iloczyn skalarny w $\mathbf{R}^n$. Transformata jest wzajemnie jednoznacznym przekształceniem $L^2(\mathbf{R}^n)$ na siebie, a równość Plancherela ma postać

$$\|f\|^2 = \frac{1}{(2\pi)^n} \|\hat{f}\|^2.$$

Wszystkie dowody są takie same, jak w przypadku $L^2(\mathbf{R})$.

Transformata Fouriera w $L^2(\mathbf{T})$ (szeregi Fouriera)

Niech f będzie funkcją na $\mathbf{R}$, okresową o okresie 2π , całkowalną na $[-\pi, \pi]$. Współczynnikami Fouriera f nazywamy ciąg

$$\hat{f}(n) = \frac{1}{2\pi} \int_{-\pi}^{\pi} f(x) e^{-inx} dx, \quad n = 0, \pm 1, \pm 2, \dots \quad (2.16)$$

Przyporządkowanie funkcji f ciągu współczynników $\{\hat{f}(n)\}$ jest też czasem nazywane transformatą Fouriera. W przypadku funkcji okresowych, w odróżnieniu od sytuacji w $L^2(\mathbf{R})$ każda funkcja w $L^2(\mathbf{T})$ jest również całkowalna na przedziale $[-\pi, \pi]$, co wynika z nierówności Schwarza. Dla każdej funkcji w $L^2(\mathbf{T})$ można więc policzyć współczynniki Fouriera ze wzoru (2.16). Zauważmy, że funkcje

$$\{e^{-inx}; n \in \mathbf{Z}\} \quad (2.17)$$

tworzą bazę ortonormalną w $L^2(\mathbf{T})$, więc ciąg współczynników Fouriera stanowi ciąg współczynników bazowych. Twierdzenie Plancherela w przypadku funkcji okresowych jest więc prostsze niż w przypadku $L^2(\mathbf{R})$.

Twierdzenie 2.9 (Plancherel). *Transformata Fouriera jest wzajemnie jednoznaczny przekształceniem $L^2(\mathbf{T})$ na ℓ^2 . Przekształcenie odwrotne dane jest przez tak zwany szereg Fouriera*

$$\mathfrak{F}^{-1} : \ell^2 \rightarrow L^2(\mathbf{T}), \quad \{\alpha_n\}_{n=-\infty}^{\infty} \mapsto f, \quad f(x) = \sum_{n=-\infty}^{\infty} \alpha_n e^{inx}.$$

Szereg Fouriera jest zbieżny w $L^2(\mathbf{T})$. Zachodzą następujące równości

$$\|f\|^2 = \sum_{n=-\infty}^{\infty} |\hat{f}(n)|^2, \quad \langle f, g \rangle = \sum_{n=-\infty}^{\infty} \hat{f}(n) \overline{\hat{g}(n)}.$$

Dowód. Przypomnijmy, że w $L^2(\mathbf{T})$

$$\langle f, g \rangle = \frac{1}{2\pi} \int_{-\pi}^{\pi} f(x) \overline{g(x)} dx.$$

Biorąc pod uwagę, że układ (2.17) stanowi bazę ortonormalną przestrzeni $L^2(\mathbf{T})$, powyższe twierdzenie jest szczególnym przypadkiem Wniosku 1.10 i Uwagi 1.11 z rozdziału o przestrzeni Hilberta. $\square$

Przykład

Niech $f \in L^2(\mathbf{T})$ będzie dana przez

$$f(x) = \begin{cases} -1 & : -\pi < x \leq 0 \\ 1 & : 0 < x \leq \pi. \end{cases} \quad (2.18)$$

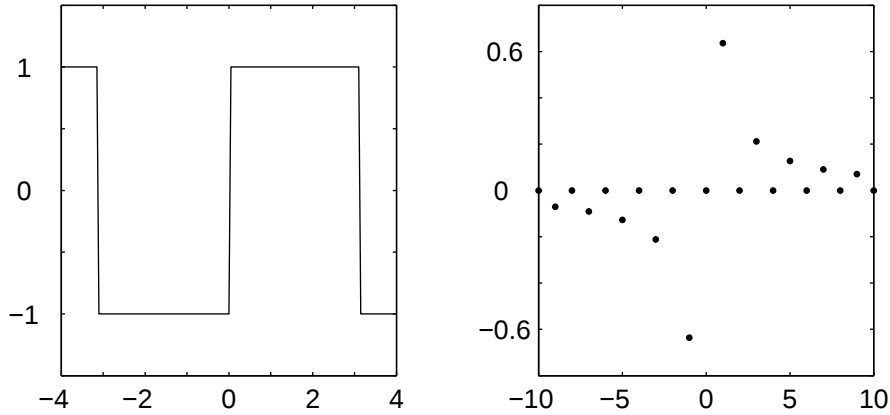
Przypomnijmy, że elementy w $L^2(\mathbf{T})$ można rozważać jako funkcje na całej prostej, okresowe o okresie 2π , albo jako funkcje na przedziale $[-\pi, \pi]$. (2.18) definiuje więc element z $L^2(\mathbf{T})$. W języku inżynierów powyższa funkcja (lub jej 2π okresowe rozszerzenie na $\mathbf{R}$) nazywa się „falą prostokątną”. Fala prostokątna jest typowym przykładem sygnału występującym na przykład w każdym urządzeniu zawierającym mikroprocesor. Policzmy współczynniki Fouriera f

$$\hat{f}(0) = \frac{1}{2\pi} \int_{-\pi}^{\pi} f(x) dx = 0.$$

Niech $n \neq 0$.

$$\begin{aligned} \hat{f}(n) &= \frac{1}{2\pi} \int_{-\pi}^{\pi} f(x) e^{-inx} dx \\ &= -\frac{1}{2\pi} \int_{-\pi}^0 e^{-inx} dx + \frac{1}{2\pi} \int_0^{\pi} e^{-inx} dx \\ &= \frac{1}{2\pi} \int_0^{\pi} (e^{-inx} - e^{inx}) dx \\ &= \frac{-2i}{2\pi} \int_0^{\pi} \sin(nx) dx \\ &= \frac{i}{\pi n} \cos(nx) \Big|_0^{\pi} \\ &= \frac{i}{\pi n} ((-1)^n - 1) \\ &= \begin{cases} \frac{2}{i\pi n} & : n - \text{nieparzyste} \\ 0 & : n - \text{parzyste} \end{cases} \end{aligned}$$

Zgodnie z twierdzeniem Plancherela funkcja f rozwija się więc w następujący szereg Fouriera. Przypomnijmy, że szereg Fouriera jest zbieżny do f w $L^2(\mathbf{T})$, a niekoniecznie w poszczególnych punktach x . W tym konkretnym przypadku szereg Fouriera jest zbieżny w każdym punkcie x , ale w punktach $0, \pm\pi$ jest



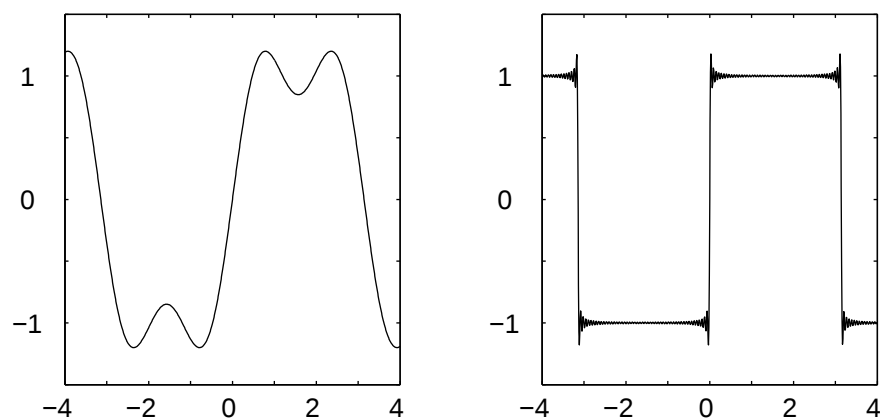
Rysunek 2.3: Funkcja f i jej współczynniki Fouriera

zbieżny do 0, a nie do wartości $f(x)$.

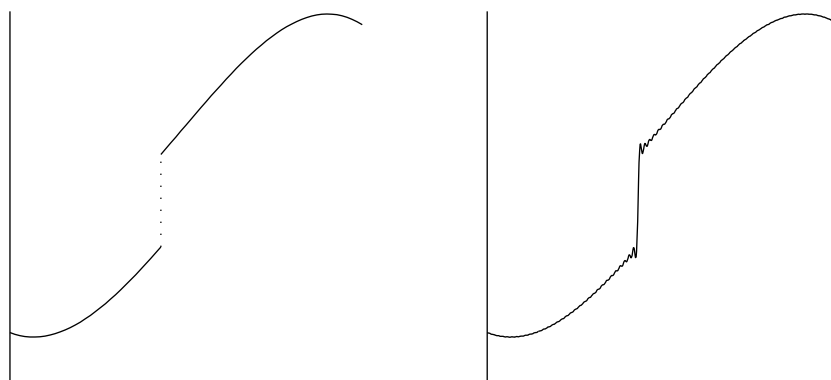
$$\begin{aligned}
 f(x) &= \sum_{\substack{n=-\infty \\ n \text{ nieparzyste}}}^{\infty} \frac{2}{i\pi n} e^{inx} \\
 &= \sum_{\substack{n=1 \\ n \text{ nieparzyste}}}^{\infty} \frac{2}{i\pi n} (e^{inx} - e^{-inx}) \\
 &= \sum_{\substack{n=1 \\ n \text{ nieparzyste}}}^{\infty} \frac{4}{\pi n} \sin nx \\
 &= \frac{4}{\pi} \left(\sin x + \frac{\sin 3x}{3} + \frac{\sin 5x}{5} + \dots \right).
 \end{aligned}$$

Na wykresie 2.4 widać tak zwane zjawisko Gibbsa. W pobliżu nieciągłości skokowej funkcji f (na przykład w pobliżu 0) suma częściowa szeregu Fouriera ma „szpilki”. Szpilki mają zawsze określoną wysokość: suma skończona „przestrzela” wysokość skoku funkcji f o około 9%. Dla coraz dalszych sum częściowych szeregu Fouriera (coraz dokładniejszych przybliżeń f) „szpilki” przysuwają się bliżej nieciągłości, ale zawsze zachowują swoją wysokość. Można to udowodnić. Jako zastosowanie obliczonych powyżej współczynników Fouriera fali prostokątnej obliczymy klasyczną sumę. Dla fali prostokątnej (2.18) mamy

$$\|f\|^2 = \frac{1}{2\pi} \int_{-\pi}^{\pi} |f(x)|^2 dx = 1,$$



Rysunek 2.4: Przybliżenie funkcji f harmonicznymi do 3 i 19 włącznie



Rysunek 2.5: Zjawisko Gibbsa. Z prawej bardzo dokładne przybliżenie funkcji mającej nieciągłość skokową. Szpilki po obu stronach nieciągłości zawsze występują, i zawsze mają określoną wysokość.

a z drugiej strony

$$\sum_{n=-\infty}^{\infty} |\hat{f}(n)|^2 = \sum_{\substack{n=-\infty \\ n \text{ nieparzyste}}}^{\infty} \frac{4}{\pi^2 n^2} = \frac{8}{\pi^2} \sum_{\substack{n=1 \\ n \text{ nieparzyste}}}^{\infty} \frac{1}{n^2}.$$

Podstawiając do równości Plancherela otrzymujemy więc

$$\sum_{\substack{n=1 \\ n \text{ nieparzyste}}}^{\infty} \frac{1}{n^2} = \frac{\pi^2}{8}.$$

Następnie

$$\begin{aligned}
 \frac{\pi^2}{8} &= \sum_{\substack{n=1 \\ n \text{ nieparzyste}}}^{\infty} \frac{1}{n^2} \\
 &= \sum_{n=1}^{\infty} \frac{1}{n^2} - \sum_{\substack{n=1 \\ n \text{ parzyste}}}^{\infty} \frac{1}{n^2} \\
 &= \sum_{n=1}^{\infty} \frac{1}{n^2} - \sum_{k=1}^{\infty} \frac{1}{(2k)^2} \\
 &= \sum_{n=1}^{\infty} \frac{1}{n^2} - \frac{1}{4} \sum_{k=1}^{\infty} \frac{1}{k^2} \\
 &= \frac{3}{4} \sum_{n=1}^{\infty} \frac{1}{n^2},
 \end{aligned}$$

i w końcu

$$\sum_{n=1}^{\infty} \frac{1}{n^2} = \frac{\pi^2}{6}.$$

W podobny sposób, stosując równość Plancherela dla odpowiedniej funkcji możemy obliczyć sumę

$$\sum_{n=1}^{\infty} \frac{1}{n^{\alpha}} \quad (2.19)$$

dla dowolnej liczby całkowitej parzystej α . Jeżeli α nie jest liczbą całkowitą parzystą, to o sumach (2.19) niewiele wiadomo. Można, na przykład, udowodnić, że dla $\alpha = 3$ suma (2.19) jest liczbą niewymierną, ale dowód jest skomplikowany. Twierdzenie Plancherela jest więc bardzo przydatne.

Definicja 2.10. *Splotem dwóch funkcji f i g okresowych o okresie 2π nazywamy funkcję*

$$f * g(x) = \frac{1}{2\pi} \int_{-\pi}^{\pi} f(x-y) g(y) dy, \quad (2.20)$$

o ile całka istnieje dla każdego x . Splotem dwóch ciągów $\alpha = \{\alpha_n\}$, $\beta = \{\beta_n\}$ nazywamy ciąg

$$\alpha * \beta(n) = \sum_{m=-\infty}^{\infty} \alpha_{n-m} \beta_m, \quad (2.21)$$

o ile powyższa suma istnieje dla każdego $n \in \mathbf{Z}$.

Jeżeli f i g są całkowalne na $[-\pi, \pi]$, to spłot $f * g$ istnieje dla (prawie) każdego x , jest okresowy o okresie 2π i jest całkowalny po okresie. Określenie „prawie każdy punkt” ma ścisłe znaczenie, ale nam wystarczy znaczenie intuicyjne: spłot może nie istnieć w jakimś punkcie, albo nawet w wielu punktach, ale tworzących razem zbiór pomijalnie mały z punktu widzenia całkowania. We wzorze (2.20) f i g traktujemy jako funkcje okresowe, o okresie 2π . Możemy również myśleć o nich jako o funkcjach określonych na odcinku $[-\pi, \pi]$, wtedy działanie $x - y$ należy rozumieć modulo 2π . Jeżeli oba ciągi są sumowalne z kwadratem, $\alpha, \beta \in \ell^2$, to spłot istnieje, a jeżeli są absolutnie sumowalne, to spłot istnieje i też jest absolutnie sumowalny. Spłoty (2.20) i (2.21) są przemienne: $f * g = g * f$ i $\alpha * \beta = \beta * \alpha$.

Własności transformaty Fouriera na $L^2(\mathbf{T})$

Transformata w $L^2(\mathbf{T})$ ma własności analogiczne do własności transformaty w $L^2(\mathbf{R})$. Przy odpowiednich założeniach mamy

$$\begin{aligned}\widehat{(f * g)}(n) &= \hat{f}(n) \hat{g}(n), \\ \widehat{(f \cdot g)}(n) &= (\hat{f} * \hat{g})(n), \\ g(x) = f(x - c) &\longrightarrow \hat{g}(n) = e^{-icn} \hat{f}(n), \\ g(x) = e^{imx} f(x) &\longrightarrow \hat{g}(n) = \hat{f}(n - m), \\ g(x) = f'(x) &\longrightarrow \hat{g}(n) = i n \hat{f}(n).\end{aligned}$$

Szereg Fouriera funkcji całkowalnej nie musi być zbieżny do tej funkcji. Znany jest przykład funkcji całkowalnej której szereg Fouriera nie jest zbieżny w żadnym punkcie. Jeżeli funkcja jest w $L^2(\mathbf{T})$, to jej szereg Fouriera jest do niej zbieżny w $L^2(\mathbf{T})$. Jednak nie musi być zbieżny w każdym punkcie. Można podać przykład funkcji ciągłej, której szereg Fouriera jest rozbieżny w jakimś punkcie. Są też dobre wiadomości. Poniższe twierdzenie ma charakter lokalny, to znaczy zbieżność szeregu Fouriera funkcji w punkcie zależy tylko od zachowania tej funkcji w otoczeniu tego punktu. Jest to o tyle ciekawe, że sam szereg Fouriera nie ma charakteru lokalnego. Każdy współczynnik Fouriera zależy od wszystkich wartości funkcji, zawiera całkę po całym okresie $\mathbf{T}$.

Twierdzenie 2.11. (a) *Jeżeli f jest okresowa i całkowalna po okresie $[-\pi, \pi]$ to*

$$\lim_{n \rightarrow \pm\infty} \hat{f}(n) = 0.$$

(b) Jeżeli istnieje pochodna $f'(\theta)$ w jakimś punkcie θ , to szereg Fouriera f jest zbieżny w tym punkcie θ do $f(\theta)$

$$f(\theta) = \sum_{n=-\infty}^{\infty} \hat{f}(n) e^{in\theta}.$$

Dowód. (a) Skorzystamy, jak poprzednio, z ciągłości przesunięć funkcji względem całki. Najpierw zauważmy, że jeżeli funkcja jest okresowa o okresie 2π , to całka z tej funkcji po przedziale $[a, a + 2\pi]$ nie zależy od a . Następnie obliczamy

$$\begin{aligned} \hat{f}(n) &= \frac{1}{2\pi} \int_{-\pi}^{\pi} f(x) e^{-inx} dx \\ &= \frac{1}{2\pi} \int_{-\pi+\pi/n}^{\pi+\pi/n} f(x) e^{-inx} dx \\ &= \frac{1}{2\pi} \int_{-\pi}^{\pi} f\left(x + \frac{\pi}{n}\right) e^{-in(x+\pi/n)} dx \\ &= -\frac{1}{2\pi} \int_{-\pi}^{\pi} f\left(x + \frac{\pi}{n}\right) e^{-inx} dx, \end{aligned}$$

a więc

$$\begin{aligned} |2\hat{f}(n)| &= \frac{1}{2\pi} \left| \int_{-\pi}^{\pi} \left(f(x) - f\left(x + \frac{\pi}{n}\right) \right) e^{-inx} dx \right| \\ &\leq \frac{1}{2\pi} \int_{-\pi}^{\pi} \left| f(x) - f\left(x + \frac{\pi}{n}\right) \right| dx \xrightarrow{n \rightarrow \pm\infty} 0. \end{aligned}$$

(b) Wprowadźmy funkcję pomocniczą

$$g(x) = \frac{f(x) - f(\theta)}{e^{-ix} - e^{-i\theta}}.$$

Zauważmy, że g jest okresowa o okresie 2π , i całkowalna po okresie. Całkowalność wynika z tego, że w pewnym otoczeniu θ jest ograniczona (z istnienia pochodnej $f'(\theta)$), a poza otoczeniem θ mianownik jest ograniczony od dołu. Z (a) wynika, że $\hat{g}(n) \rightarrow 0$ gdy $n \rightarrow \pm\infty$. Zauważmy jednak, że skoro

$$g(x) e^{-ix} - g(x) e^{-i\theta} = f(x) - f(\theta),$$

więc dla wszystkich n

$$\hat{g}(n+1) - \hat{g}(n) e^{-i\theta} = \widehat{f}(n),$$

gdzie $\tilde{f}(x) = f(x) - f(\theta)$. Mnożąc obie strony przez $e^{i(n+1)\theta}$ otrzymujemy

$$\hat{g}(n+1) e^{i(n+1)\theta} - \hat{g}(n) e^{in\theta} = e^{i\theta} \widehat{\tilde{f}}(n) e^{in\theta}.$$

Sumując obie strony po wszystkich $n = -N, \dots, M$, i zwiijając sumę teleskopową po lewej stronie mamy

$$\hat{g}(M+1) e^{i(M+1)\theta} - \hat{g}(-N) e^{i(-N)\theta} = e^{i\theta} \sum_{n=-N}^M \widehat{\tilde{f}}(n) e^{in\theta}.$$

Korzystając z (a) lewa strona $\rightarrow 0$ gdy $M, N \rightarrow \infty$, więc podobnie dzieje się z prawą stroną. Zauważmy, że $\widehat{\tilde{f}}(n) = \hat{f}(n)$ dla $n \neq 0$, oraz $\widehat{\tilde{f}}(0) = \hat{f}(0) - f(\theta)$. Tak więc

$$0 = \sum_{n=-\infty}^{\infty} \widehat{\tilde{f}}(n) e^{in\theta} = \sum_{n=-\infty}^{\infty} \hat{f}(n) e^{in\theta} - f(\theta).$$

Udowodniliśmy więc, że szereg Fouriera f jest zbieżny w θ i jego sumą jest $f(\theta)$. $\square$

Transformata Fouriera w $L^2(\mathbf{T}^n)$ (wielokrotne szeregi Fouriera)

Teoria jest analogiczna do 1-wymiarowej:

$$\hat{f}(k) = \frac{1}{(2\pi)^n} \int_{-\pi}^{\pi} \cdots \int_{-\pi}^{\pi} f(x) e^{-i k \cdot x} dx_1 \cdots dx_n, \quad k = (k_1, \dots, k_n), \quad x = (x_1, \dots, x_n),$$

$$\check{f}(x) = \sum_{k_1, \dots, k_n = -\infty}^{\infty} \hat{f}(k) e^{i k \cdot x},$$

$$\langle f, g \rangle = \frac{1}{(2\pi)^n} \int_{-\pi}^{\pi} \cdots \int_{-\pi}^{\pi} f(x) \overline{g(x)} dx_1 \cdots dx_n = \sum_{k_1, \dots, k_n = -\infty}^{\infty} \hat{f}(k) \overline{\hat{g}(k)},$$

$$\|f\|^2 = \sum_{k_1, \dots, k_n = -\infty}^{\infty} |\hat{f}(k)|^2.$$

Nie ma istotnych różnic w dowodach powyższych faktów dla 1 i wielu wymiarów.

Transformata Fouriera w ℓ_p^2 (dyskretna transformata Fouriera)

W przestrzeni p -elementowych wektorów transformata Fouriera jest wzajemnie jednoznacznym przekształceniem ℓ_p^2 na siebie. Dla $x = (x(0), \dots, x(p -$

1)) $\in \ell_p^2$ mamy

$$\begin{aligned}\hat{x}(l) &= \sum_{j=0}^{p-1} x(j) e^{-i \frac{2\pi l j}{p}}, \quad l = 0, \dots, p-1, \\ \check{x}(j) &= \frac{1}{p} \sum_{l=0}^{p-1} x(l) e^{i \frac{2\pi j l}{p}}, \quad j = 0, \dots, p-1, \\ \langle x, y \rangle &= \sum_{j=0}^{p-1} x(j) \overline{y(j)} = \frac{1}{p} \sum_{l=0}^{p-1} \hat{x}(j) \overline{\hat{y}(j)} = \frac{1}{p} \langle \hat{x}, \hat{y} \rangle, \\ \|x\|^2 &= \frac{1}{p} \|\hat{x}\|^2.\end{aligned}$$

Powyższe wzory wynikają z następującej obserwacji

$$\sum_{l=0}^{p-1} e^{i \frac{2\pi j l}{p}} = \begin{cases} p & : j = 0 \\ 0 & : j \neq 0. \end{cases}$$

Splotem dwóch elementów $x, y \in \ell_p^2$ nazywamy element

$$x * y(k) = \sum_{l=0}^{p-1} x(k-l) y(l). \quad (2.22)$$

Jeżeli o x i y myślimy jako o wektorach p -elementowych to działanie $k-l$ należy rozumieć jako odejmowanie modulo p , i wtedy k przebiega zakres $0, \dots, p-1$ a więc splot też jest wektorem p -elementowym. Jeżeli natomiast o elementach x i y myślimy jako o ciągach p -okresowych, to działanie $k-l$ w (2.22) jest zwykłym odejmowaniem, i splot $x*y$ też jest ciągiem p -okresowym. Jak się łatwo domysleć, zachodzi następujący wzór

$$\widehat{x * y}(k) = \hat{x}(k) \hat{y}(k).$$

Szybka transformata Fouriera

Dyskretna transformata Fouriera występuje często w zastosowaniach. Jeżeli chcemy policzyć numerycznie transformatę Fouriera funkcji, to w rzeczywistości sprowadza się to do policzenia dyskretnej transformaty Fouriera pewnej ilości próbek danej funkcji. Obliczenie dyskretnej transformaty Fouriera p -elementowego wektora $x = (x(0), \dots, x(p-1))$ sprowadza się do pomnożenia

go przez macierz: $\hat{x} = A_p x$, gdzie A_p jest $p \times p$ macierzą

$$A_p = \begin{pmatrix} 1 & 1 & \cdots & 1 \\ 1 & e^{-i 2\pi/p} & \cdots & e^{-i 2\pi(p-1)/p} \\ \vdots & \vdots & \ddots & \vdots \\ 1 & e^{-i 2\pi(p-1)/p} & \cdots & e^{-i 2\pi(p-1)(p-1)/p} \end{pmatrix},$$

o współczynnikach $\{e^{-i 2\pi(k-1)(l-1)/p}\}_{k,l=1}^p$. Postępując naiwnie, do obliczenia transformaty Fouriera długości p będziemy więc musieli wykonać p^2 mnożeń zmiennoprzecinkowych. Istnieje szybszy algorytm obliczania transformaty, wykorzystujący występujące w macierzy A_p symetrie. Jest to tak zwana szybka transformata Fouriera (FFT), która redukuje liczbę mnożeń do $p \log p$. Korzyść jest wielka: jeżeli $p = 10^6$, to szybka transformata Fouriera jest 50 tysięcy razy szybsza od algorytmu naiwnego. Jeżeli nasz komputer, używając algorytmu FFT policzy transformatę w godzinę, to używając algorytmu naiwnego potrzebowałby na to ponad 6 lat. Algorytm FFT odkryto w latach 60 ubiegłego wieku. Sam algorytm jest bardzo prosty, i wkrótce okazało się, że był stosowany już od dawna. Obliczenia z wykorzystaniem algorytmu FFT znaleziono w pracach Gaussa z końca XVIII wieku. Wniosek płynie z tego następujący: matematycy częściej mogliby interesować się zastosowaniami, a inżynierowie częściej mogliby zaglądać do prac teoretyków. I jedni i drudzy mogą znaleźć jakąś niespodziankę.

Algorytm FFT jest bardzo prosty. Zamiast od razu formułować odpowiednie twierdzenie przeprowadźmy proste rachunki, które wszystko wyjaśnią. Niech długość sygnału $x = (x(0), x(1), \dots, x(p-1))$ będzie potęgą 2: $p = 2^q$, gdzie $q = 1, 2, \dots$, wtedy dla $l = 0, 1, 2, \dots, p-1$

$$\begin{aligned} \hat{x}(l) &= \sum_{k=0}^{p-1} x(k) e^{-\frac{2\pi i k l}{p}} \\ &= \sum_{\substack{k=0 \\ n\text{-parzyste}}}^{p-1} x(k) e^{-\frac{2\pi i k l}{p}} + \sum_{\substack{k=0 \\ n\text{-nieparzyste}}}^{p-1} x(k) e^{-\frac{2\pi i k l}{p}} \\ &= \sum_{k=0}^{p/2-1} x(2k) e^{-\frac{2\pi i (2k) l}{p}} + \sum_{k=0}^{p/2-1} x(2k+1) e^{-\frac{2\pi i (2k+1) l}{p}} \\ &= \sum_{k=0}^{p'-1} x'(k) e^{-\frac{2\pi i k l}{p'}} + e^{-\frac{2\pi i l}{p}} \sum_{k=0}^{p'-1} x''(k) e^{-\frac{2\pi i k l}{p'}} \\ &= \hat{x}'(l) + e^{-\frac{2\pi i l}{p}} \hat{x}''(l), \end{aligned}$$

gdzie $p' = p/2$, $x'(k) = x(2k)$, $x''(k) = x(2k+1)$, $k = 0, \dots, p'-1$. Zauważmy, że $\hat{x}'$ i $\hat{x}''$ są okresowe o okresie p' , więc wystarczy je obliczyć dla $l = 0, \dots, p'-1$. Obliczenie transformaty Fouriera sygnału o długości p sprowadza się więc do:

- (a). rozdzielenia parzystych i nieparzystych składowych x ,
- (b). obliczenia transformaty Fouriera osobno dla składowych parzystych i nieparzystych, każda rzędu $p' = p/2$,
- (c). utworzenia kombinacji liniowej:

$$\begin{aligned}\hat{x}(l) &= \hat{x}'(l) + e^{-\frac{2\pi i l}{p}} \hat{x}''(l), \\ \hat{x}(l + p') &= \hat{x}'(l) - e^{-\frac{2\pi i l}{p}} \hat{x}''(l)\end{aligned}$$

dla $l = 0, \dots, p'$.

W ten sposób udowodniliśmy następujące twierdzenie, będące podstawą algorytmu FFT:

Twierdzenie 2.12 (Danielson-Lanczos). *Macierz*

$$A_p = \begin{pmatrix} 1 & 1 & \dots & 1 \\ 1 & e^{-i 2\pi/p} & \dots & e^{-i 2\pi(p-1)/p} \\ \vdots & \vdots & \ddots & \vdots \\ 1 & e^{-i 2\pi(p-1)/p} & \dots & e^{-i 2\pi(p-1)(p-1)/p} \end{pmatrix}$$

można rozłożyć na czynniki

$$A_p = E_p \tilde{A}_{p/2} P_p,$$

gdzie macierz P_p jest postaci

$$P_p = \begin{pmatrix} 1 & 0 & 0 & 0 & \dots & 0 \\ 0 & 0 & 1 & 0 & \dots & 0 \\ & \vdots & & & \ddots & \vdots \\ 0 & 1 & 0 & 0 & \dots & 0 \\ 0 & 0 & 0 & 1 & \dots & 0 \end{pmatrix}.$$

Macierz P_p ma dokładnie jedną jedynkę w każdym wierszu, a jej działanie na wektorze sprowadza się do przestawienia składowych: w pierwszej połowie umieszczone zostają składowe o numerach parzystych, a w drugiej składowe

o numerach nieparzystych: $(x(0), x(2), \dots, x(p-2), x(1), x(3), \dots, x(p-1))$.
Macierz $\tilde{A}_{p/2}$ ma dwie niezerowe, identyczne klatki

$$\tilde{A}_{p/2} = \begin{pmatrix} A_{p/2} & 0 \\ 0 & A_{p/2} \end{pmatrix}.$$

Działanie macierzy $\tilde{A}_{p/2}$ sprowadza się do obliczenia transformaty Fouriera rzędu $p/2$ osobno dla pierwszych i osobno dla ostatnich $p/2$ współczynników sygnału. W końcu macierz E_p ma postać czterech klatek

$$E_p = \begin{pmatrix} I & D \\ I & -D \end{pmatrix},$$

gdzie I są macierzami identycznościowymi rzędu $p/2$, a D jest macierzą diagonalną rzędu $p/2$ ze współczynnikami

$$e^{-\frac{2\pi i 0}{p}}, e^{-\frac{2\pi i 1}{p}}, \dots, e^{-\frac{2\pi i (p/2-1)}{p}}$$

na przekątnej.

Zauważmy, że z powyższego, prostego twierdzenia wynika następujący wniosek:

Wniosek 2.13. *Transformatę Fouriera rzędu p (potęga 2) można obliczyć przy pomocy nie więcej niż $p \log_2 p$ mnożeń.*

Dowód. Dowód jest indukcyjny względem potęgi 2. Transformata rzędu 2^1 to

$$\hat{x}(k) = x(0) \pm x(1), \quad k = 0, 1,$$

a więc występuje tylko 1 mnożenie, przez -1 . Krok indukcyjny używa twierdzenia. Do policzenia transformaty rzędu p potrzeba 2 razy tyle mnożeń co do policzenia transformaty rzędu $p/2$ (macierz $\tilde{A}_{p/2}$) i dodatkowo p mnożeń (macierz E_p). Macierz P_p nie wymaga mnożeń. Korzystając z założenia indukcyjnego otrzymujemy

$$2(p/2 \log_2(p/2)) + p = p(\log_2(p/2) + 1) = p \log_2 p.$$

□

Uwaga: Oszacowaliśmy tylko ilość koniecznych mnożeń, gdyż to mnożenia głównie zajmują czas procesora. Podobnie jak ilość koniecznych mnożeń można oszacować ilość wszystkich koniecznych operacji arytmetycznych.

Następujące twierdzenie daje nam algorytm FFT:

Twierdzenie 2.14. Niech $p = 2^q$. Macierz A_p rozkłada się na iloczyn

$$A_p = E_p \tilde{E}_{p/2} \cdots \tilde{E}_2 \tilde{P}_4 \tilde{P}_8 \cdots \tilde{P}_{p/2} P_p, \quad (2.23)$$

gdzie E_{2^j} oraz P_{2^j} mają to samo znaczenie co w twierdzeniu Danielsona-Lanczosa, $\tilde{E}_{2^j}$ oraz $\tilde{P}_{2^j}$ to macierze $p \times p$ z $p/2^j$ klatkami E_{2^j} i P_{2^j} odpowiednio na przekątnej.

Dowód. Wystarczy zauważyć, że twierdzenie Danielsona-Lanczosa można iterować:

$$A_p = E_p \tilde{A}_{p/2} P_p = E_p \tilde{E}_{p/2} \tilde{A}_{p/4} \tilde{P}_{p/2} P_p = \cdots = E_p \tilde{E}_{p/2} \cdots \tilde{E}_2 \tilde{A}_1 \tilde{P}_2 \cdots \tilde{P}_{p/2} P_p. \quad (2.24)$$

Macierze A_1 i P_2 , a więc także $\tilde{A}_1$ i $\tilde{P}_2$ są macierzami identycznościowymi, więc z (2.24) wynika (2.23). $\square$

Uwagi:(i) Każda z macierzy $\tilde{P}_{2^j}$ jest macierzą permutacji, więc ich iloczyn też jest macierzą permutacji. Wynika z tego, że przekształcenie

$$U = \tilde{P}_4 \tilde{P}_8 \cdots \tilde{P}_{p/2} P_p$$

srowadza się do pewnego przestawienia współczynników wektora x . To przedstawienie jest szczególnie proste do zaimplementowania w praktyce. Jeżeli oznaczymy $x = (x(0), \dots, x(p-1))$ oraz $Ux = x' = (x'(0), \dots, x'(p-1))$, to $x'(k) = x(l)$, gdzie k i l mają wzajemnie symetryczne rozwinięcia w układzie dwójkowym

$$(k)_2 = e_{q-1} \cdots e_0, \quad (l)_2 = e_0 \cdots e_{q-1}, \quad e_j = 0, 1.$$

Jest to tak zwane przekształcenie odwrócenia bitów.

(ii) Przedstawiony powyżej algorytm FFT jest tylko jednym z możliwych. Niektóre programy komputerowe stosują inny algorytm. Zauważmy, że macierz dyskretnej transformaty Fouriera A_p jest symetryczna. Twierdzenie Danielsona-Lanczosa można więc zapisać jako

$$A_p = A_p^t = P_p^t \tilde{A}_{p/2} E_p^t, \quad (2.25)$$

gdzie t oznacza transpozycję, i gdzie wykorzystaliśmy fakt, że A_p i $\tilde{A}_{p/2}$ są symetryczne, a więc niezmiennicze ze względu na transpozycję. Powyższy wzór można udowodnić bezpośrednio, grupując elementy sumy inaczej, niż

robiliśmy to wcześniej. Niech $p' = p/2$.

$$\begin{aligned}
\hat{x}(l) &= \sum_{k=0}^{p-1} x(k) e^{-i \frac{2\pi kl}{p}} \\
&= \sum_{k=0}^{p/2-1} x(k) e^{-i \frac{2\pi kl}{p}} + \sum_{k=p/2}^{p-1} x(k) e^{-i \frac{2\pi kl}{p}} \\
&= \sum_{k=0}^{p'-1} \left(x(k) e^{-i \frac{2\pi kl}{p}} + x(k+p') e^{-i \frac{2\pi (k+p')l}{p}} \right) \\
&= \sum_{k=0}^{p'-1} (x(k) + e^{-i \pi l} x(k+p')) e^{-i \frac{2\pi kl}{p}} \\
&= \sum_{k=0}^{p'-1} (x(k) + (-1)^l x(k+p')) e^{-i \frac{2\pi kl}{p}}
\end{aligned}$$

Jeżeli $l = 2n$, $n = 0, \dots, p' - 1$ jest parzysta, to

$$\hat{x}(l) = \sum_{k=0}^{p'-1} (x(k) + x(k+p')) e^{-i \frac{2\pi kn}{p'}},$$

a jeżeli $l = 2n + 1$ jest nieparzysta, to

$$\hat{x}(l) = \sum_{k=0}^{p'-1} (x(k) - x(k+p')) e^{-i \frac{2\pi kn}{p'}} e^{-i \frac{2\pi kn}{p'}}.$$

Innymi słowy, parzyste współczynniki $\hat{x}$ to transformata rzędu p' wektora x' o współczynnikach

$$x'(k) = x(k) + x(k+p'),$$

a nieparzyste to transformata wektora x''

$$x''(k) = (x(k) - x(k+p')) e^{-i \frac{2\pi k}{p}}.$$

Łatwo sprawdzić, że powyższe obliczenia dają nam dokładnie (2.25). Zauważmy też, że macierz U z poprzedniej uwagi jest symetryczna, więc stosując ten algorytm również dochodzimy do przekształcenia odwrócenia bitów, z tym, że w poprzednim algorytmie był to pierwszy krok szybkiej transformaty, a w omawianym teraz wariancie jest to ostatni krok.

(iii) Omawialiśmy przypadek, gdy sygnał miał długość będącą potęgą 2. W praktyce to jest przypadek najważniejszy. Sygnały o innych długościach są przedłużane do najbliższej potęgi 2, na przykład dodaje się odpowiednią ilość zer. Szybki algorytm można jednak skonstruować niezależnie dla sygnałów o innych długościach.

Transformaty trygonometryczne

Naturalnym językiem transformat Fouriera w ich wszystkich wcieleniach, w tym także dyskretnej transformaty Fouriera jest język liczb zespolonych. W zastosowaniach jest to niepraktyczne. Sygnały rzeczywiste mają transformaty o wartościach zespolonych. Takie sygnały o wartościach zespolonych wymagają innych struktur do przechowywania i manipulacji. Nie jest to wielkim problemem, ale w praktyce wygodniejsze jest posługiwanie się transformacjami, które sygnały rzeczywiste przekształcają na rzeczywiste. Są to tak zwane transformaty trygonometryczne. Transformaty takie powstają przez proste przekształcenie transformaty Fouriera i w związku z tym można stosować do nich szybkie algorytmy. Podstawową obserwacją jest fakt, że jeżeli sygnał rzeczywisty $x = (x(0), \dots, x(p-1))$ jest parzysty, czyli

$$x(p-k) = x(k),$$

to jego transformata Fouriera też jest rzeczywista, a jeżeli jest nieparzysty, czyli

$$x(p-k) = -x(k), \quad x(0) = 0,$$

to transformata jest czysto urojona. Istniejące sygnały można więc odpowiednio przedłużyć i zastosować do nich transformatę Fouriera odpowiedni wysokiego rzędu. W praktyce spotyka się kilka odmian transformat sinusowych i cosinusowych, które otrzymuje się właśnie mniej więcej według powyższego schematu. Na przykład, wyprowadzimy wzór na jedną z transformat cosinusowych. Niech dany będzie sygnał o długości p , $x = (x(0), \dots, x(p-1))$. Niech $\tilde{x}$ będzie przedłużeniem x o długości $2p$, zdefiniowanym następująco:

$$\tilde{x}(k) = \begin{cases} x(k) & : k = 0, \dots, p-1 \\ x(2p-k-1) & : k = p, \dots, 2p-1. \end{cases}$$

Zauważmy, że sygnał $\tilde{x}$ ma pewną symetrię — jest „parzysty” wokół punktu $p-1/2$:

$$\tilde{x}((p-1/2)+l) = \tilde{x}((p-1/2)-l),$$

gdzie l jest liczbą połówkową, $2l = 1, 3, \dots, 2p - 1$. Obliczmy transformatę Fouriera długości $2p$ sygnału rozszerzonego $\tilde{x}$.

$$\begin{aligned}\widehat{\tilde{x}}(l) &= \sum_{k=0}^{2p-1} \tilde{x}(k) e^{-i \frac{2\pi kl}{2p}} \\ &= \sum_{k=0}^{p-1} \tilde{x}(k) e^{-i \frac{2\pi kl}{2p}} + \sum_{k=p}^{2p-1} \tilde{x}(k) e^{-i \frac{2\pi kl}{2p}} \\ &= \sum_{k=0}^{p-1} x(k) e^{-i \frac{\pi kl}{p}} + \sum_{k=p}^{2p-1} x(2p - k - 1) e^{-i \frac{\pi kl}{p}}.\end{aligned}$$

Przenumerujmy składniki drugiej sumy, niech nowy indeks $k' = 2p - k - 1$. Nowy indeks (też go potem nazwiemy k) biegnie od 0 do $p - 1$. Otrzymujemy

$$\begin{aligned}\widehat{\tilde{x}}(l) &= \sum_{k=0}^{p-1} x(k) \left(e^{-i \frac{\pi kl}{p}} + e^{-i \frac{\pi (2p-k-1)l}{p}} \right) \\ &= \sum_{k=0}^{p-1} x(k) \left(e^{-i \frac{\pi kl}{p}} + e^{i \frac{\pi kl}{p}} e^{i \frac{\pi l}{p}} \right) \\ &= \sum_{k=0}^{p-1} x(k) e^{i \frac{\pi l}{2p}} \left(e^{-i \frac{\pi (k+1/2)l}{p}} + e^{i \frac{\pi (k+1/2)l}{p}} \right) \\ &= 2 e^{i \frac{\pi l}{2p}} \sum_{k=0}^{p-1} x(k) \cos \left(\frac{\pi (k + 1/2)l}{p} \right).\end{aligned}$$

Zauważmy, że transformata Fouriera $\tilde{x}$ zależy tylko od współczynników „cosinusowych”

$$\hat{x}_c(l) = \sum_{k=0}^{p-1} x(k) \cos \left(\frac{\pi (k + 1/2)l}{p} \right).$$

Mamy więc

$$\widehat{\tilde{x}}(l) = 2 \cdot e^{i \frac{\pi l}{2p}} \cdot \hat{x}_c(l).$$

Zróbmy jeszcze następujące obserwacje

$$\begin{aligned}\widehat{\tilde{x}}(0) &= 2 \cdot \hat{x}_c(0), \\ \widehat{\tilde{x}}(p) &= e^{i \frac{\pi}{2}} \sum_{k=0}^{p-1} x(k) \cos(\pi(k + 1/2)) = 0,\end{aligned}$$

gdyż $\cos(\pi(k+1/2)) = 0$ dla każdej liczby całkowitej k . W końcu zauważmy, że

$$\widehat{x}(2p-l) = \widehat{x}(l).$$

Rekonstruując $\tilde{x}$, korzystając z odwrotnej transformaty Fouriera, otrzymujemy

$$\begin{aligned}\tilde{x}(k) &= \frac{1}{2p} \sum_{l=0}^{2p-1} \widehat{x}(l) e^{i \frac{2\pi kl}{2p}} \\ &= \frac{1}{2p} \widehat{x}(0) + \frac{1}{2p} \sum_{\substack{l=0 \\ l \neq p}}^{2p-1} \widehat{x}(l) e^{i \frac{2\pi kl}{2p}} \\ &= \frac{1}{p} \widehat{x}_c(0) + \frac{1}{2p} \sum_{l=1}^{p-1} \left(\widehat{x}(l) e^{i \frac{2\pi kl}{2p}} + \widehat{x}(2p-l) e^{i \frac{2\pi k(2p-l)}{2p}} \right) \\ &= \frac{1}{p} \widehat{x}_c(0) + \frac{1}{2p} \sum_{l=1}^{p-1} 2\widehat{x}_c(l) \left(e^{i \frac{\pi l}{2p}} e^{i \frac{2\pi kl}{2p}} + \overline{e^{i \frac{\pi l}{2p}} e^{i \frac{2\pi kl}{2p}}} \right) \\ &= \frac{1}{p} \widehat{x}_c(0) + \frac{2}{p} \sum_{l=1}^{p-1} \widehat{x}_c(l) \cos \left(\frac{\pi(k+1/2)l}{p} \right).\end{aligned}$$

Gdy $k = 0, \dots, p-1$ to $\tilde{x}(k) = x(k)$, więc otrzymaliśmy wzór na odwrotną transformatę cosinusową. Podsumowując: transformata cosinusowa i odwrotna transformata cosinusowa dane są wzorami

$$\begin{aligned}\widehat{x}_c(l) &= \sum_{k=0}^{p-1} x(k) \cos \left(\frac{\pi(k+1/2)l}{p} \right), \quad l = 0, \dots, p-1, \\ \tilde{x}_c(k) &= \frac{1}{p} x(0) + \frac{2}{p} \sum_{l=1}^{p-1} x(l) \cos \left(\frac{\pi(k+1/2)l}{p} \right), \quad k = 0, \dots, p-1.\end{aligned}$$

Transformata cosinusowa związana jest z dyskretną transformatą Fouriera wzorem

$$C_p = R_p A_{2p} Q_p,$$

gdzie C_p jest macierzą odpowiadającą transformacie cosinusowej, natomiast R_p i Q_p są pewnymi macierzami prostokątnymi, zawierającymi w większości zera. Dokładne wzory na R_p i Q_p można otrzymać analizując przeprowadzone powyżej rachunki.

Przypomnijmy, że powyższe wzory stanowią tylko jeden z możliwych wariantów transformat trygonometrycznych.

Inne transformaty

- Lokalne transformaty Fouriera i trygonometryczne
- Transformata Laplace'a
- Transformata Mellina
- Transformata Zaka
- Transformata Radona

Rozdział 3

Dodatek: kilka transformat

Rozdział 4

Analiza Wielorozdzielcza

Bazy falkowe

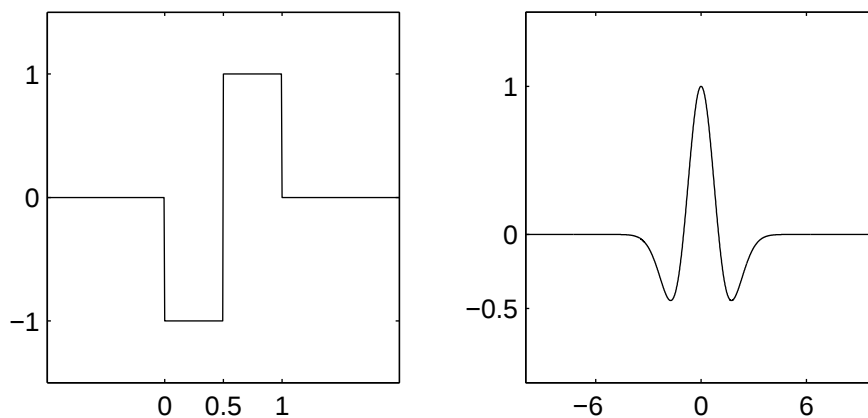
W tym rozdziale przedmiotem naszego zainteresowania będzie konstrukcja szczególnego rodzaju baz w przestrzeni $L^2(\mathbf{R})$, tak zwanych baz falkowych. Baza falkowa to baza ortonormalna w $L^2(\mathbf{R})$ następującej postaci

$$\{2^{\frac{j}{2}}\psi(2^j x - n); n, j \in \mathbf{Z}\}, \quad (4.1)$$

dla pewnej funkcji $\psi \in L^2(\mathbf{R})$. Baza falkowa powstaje z jednej funkcji ψ , w ten sposób, że najpierw tworzymy wszystkie całkowite przesunięcia ψ , a następnie tworzymy wszystkie możliwe przeskalowania powstałej rodziny o współczynniki będące potęgą 2. Jeżeli tak otrzymana rodzina funkcji w $L^2(\mathbf{R})$ jest bazą ortonormalną, to nazywamy ją bazą falkową, a funkcję ψ falką. Nazwa „falka” pochodzi stąd, że ψ jest najczęściej elementarną oscylacją, o wyraźnej lokalizacji w czasie i o określonej częstotliwości. Budując z takiej funkcji bazę według schematu (4.1) widzimy, że funkcje $2^{j/2}\psi(2^j x - n)$ też stanowią taką elementarną oscylację, ale o częstotliwości rzędu 2^j i lokalizacji w okolicy punktu $2^{-j}n$.

Przesunięcie i przeskalowanie takiej funkcji powoduje, że iloczyn skalarny z nią wychwytyje w sygnale szczegóły występujące w konkretnym miejscu i posiadające określoną częstotliwość. Ta własność sprawia, że bazy falkowe odniosły w ostatnich 20 latach ogromny sukces w zastosowaniach. Analiza falkowa jest obecnie obszernym działem matematyki. Z grubsza biorąc można wydzielić dwa rodzaje transformaty falkowej. Po pierwsze mamy ciągłą transformatę falkową

$$W(f)(a, t) = \int_{-\infty}^{\infty} f(x) \sqrt{a} \overline{\psi(ax - t)} dx,$$



Rysunek 4.1: Falka Haara i falka „meksykański kapelusz”.

która funkcji $f \in L^2(\mathbf{R})$ przyporządkowuje funkcję dwóch zmiennych $a, t \in \mathbf{R}$, $a > 0$. Dla każdego $a > 0$ $W(f)(a, t)$ jest, mówiąc intuicyjnie „obrazem” funkcji „na poziomie rozdzielczości a ”. Drugim rodzajem transformaty falkowej jest transformata falkowa dyskretna, czyli rozkład funkcji $f \in L^2(\mathbf{R})$ w bazie falkowej

$$\alpha(f)_{n,j} = \int_{-\infty}^{\infty} f(x) 2^{\frac{j}{2}} \overline{\psi(2^j x - n)} dx, \quad n, j \in \mathbf{Z}.$$

Łatwo zauważyć, że dyskretna transformata falkowa jest rezultatem próbkowania transformaty ciągłej w punktach $(a, t) = (2^j, n)$, $j, n \in \mathbf{Z}$. Jeżeli ψ jest falką, czyli układ funkcji (4.1) stanowi bazę w $L^2(\mathbf{R})$, to funkcję $f \in L^2(\mathbf{R})$ można zrekonstruować z tych próbek $\alpha(f)_{n,j}$. Dla baz falkowych istnieje szybki algorytm numeryczny, tak zwany algorytm Mallata, wyliczający współczynniki bazowe $\alpha(f)_{n,j}$. Algorytm Mallata opiszemy w następującym rozdziale. Problem konstrukcji baz falkowych jest o tyle interesujący, że własności konkretnej falki ψ mają istotne znaczenie dla zastosowań. Dlatego nie wystarczy skonstruować jednej bazy falkowej do wszystkich zastosowań. Istnieje potrzeba konstruowania falek o konkretnych własnościach. W zastosowaniach najczęściej pojawiają się tak zwane falki Daubechies. Jest to rodzina falek o nośniku ograniczonym (czyli $\psi \equiv 0$ poza pewnym przedziałem) i o różnym stopniu gładkości. Falek Daubechies używamy w laboratorium. W tym rozdziale głównie zajmujemy się przypadkiem jednowymiarowym, czyli przestrzenią $L^2(\mathbf{R})$. Pod koniec pokażemy, jak bazy falkowe można budować w przypadku wielowymiarowym, używając skonstruowanych już falek jednowymiarowych.

Analiza wielorozdzielcza MRA

Narzędziem do konstrukcji falek jest tak zwana analiza wielorozdzielcza.

Definicja 4.1. *Analizą wielorozdzielczą (w skrócie: MRA) nazywamy ciąg rosnący podprzestrzeni domkniętych $L^2(\mathbf{R})$*

$$\cdots \subset V_{-2} \subset V_{-1} \subset V_0 \subset V_1 \subset V_2 \subset \cdots \subset L^2(\mathbf{R})$$

spełniających następujące warunki

(a). $\overline{\bigcup_{j=-\infty}^{\infty} V_j} = L^2(\mathbf{R})$,

(b). $\bigcap_{j=-\infty}^{\infty} V_j = \{0\}$,

(c). $f(x) \in V_j \Leftrightarrow f(2x) \in V_{j+1}$,

(d). *istnieje funkcja $\varphi \in V_0$ (tak zwana funkcja skalująca analizy) taka, że zbiór funkcji*

$$\{\varphi(x - n); n \in \mathbf{Z}\}$$

stanowi bazę o.n. (lub bazę Riesz) przestrzeni V_0 .

Dla podkreślenia, czy funkcja skalująca generuje bazę o.n. czy bazę Riesz przestrzeni V_0 czasem mówi się „ortogonalna MRA”, lub „MRA Riesz”. Konstrukcja falek nastąpi w dwóch krokach. Najpierw pokażemy, że mając MRA możemy skonstruować falekę, a następnie pokażemy, jak można skonstruować interesujące nas analizy wielorozdzielcze.

Uwagi: (i) Zauważmy, że funkcja φ całkowicie określa analizę wielorozdzielczą, dla której jest funkcją skalującą. Istotnie, skoro przesunięcia φ stanowią bazę o.n. V_0 , to

$$V_0 = \overline{\text{Lin} \{ \varphi(x - n); n \in \mathbf{Z} \}}. \quad (4.2)$$

Z warunku (c) z kolei wynika, że

$$V_j = \{ f \in L^2(\mathbf{R}); f(2^{-j}x) \in V_0 \}. \quad (4.3)$$

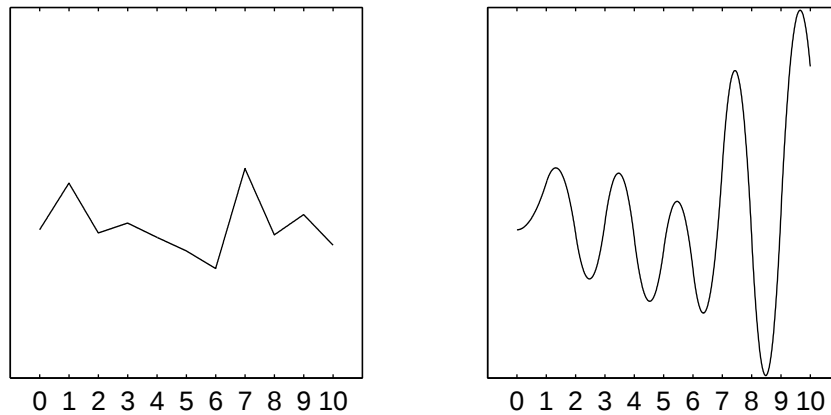
Konstrukcja MRA będzie więc wyglądała następująco. Znajdziemy funkcję φ taką, że jej przesunięcia całkowite stanowią układ ortonormalny. Wtedy przesunięcia będą stanowiły bazę o.n. przestrzeni V_0 zdefiniowanej w (4.2). Następnie przestrzeń V_j zdefiniujemy wzorem (4.3). Znajdziemy dodatkowe warunki na φ , które sprawią, że tak powstały ciąg domkniętych podprzestrzeni jest rosnący (wystarczy, że $V_0 \subset V_1$), oraz spełnione są warunki (a) i (b) definicji.

(ii) Rozróżniamy analizy wielorozdzielcze ortonormalne i Riesz, w zależności od bazy, którą w V_0 generuje funkcja skalująca φ . Pokażemy, że analiza wielorozdzielcza Riesz też jest analizą o.n., tylko trzeba wybrać inną funkcję skalującą. Innymi słowy, dla analizy Riesz, z funkcją skalującą φ można w V_0 znaleźć inną funkcję $\tilde{\varphi}$, której przesunięcia są bazą o.n. V_0 . Pomimo tego, że analiza Riesz jest automatycznie także analizą o.n. warto rozróżniać oba rodzaje analiz. Na przykład, zdarza się, że funkcja $\tilde{\varphi}$, chociaż istnieje, ma bardzo skomplikowaną postać, i wygodniej jest pracować z początkową funkcją φ , chociaż generuje jedynie bazę Riesz, a nie o.n.

Przykłady: (i) Dla $N = 0, 1, 2, \dots$ mówimy, że funkcja f jest splinem rzędu N (nie znam dobrego polskiego tłumaczenia) jeżeli jest różniczkowalna w sposób ciągły $N - 1$ razy oraz jest wielomianem stopnia $\leq N$ na przedziałach postaci $[k, k + 1]$ (pomiędzy sąsiednimi liczbami całkowitymi). Jeżeli $N = 0$ to założenie o różniczkowalności jest puste, funkcja f ma tylko być stała na przedziałach $[k, k + 1]$.

Niech N będzie ustalone. Definiujemy

$$V_0 = \{f \in L^2(\mathbf{R}); f \text{ jest splinem rzędu } N\}.$$



Rysunek 4.2: Spliny rzędu 1 i 2.

Zauważmy, że jest to podprzestrzeń domknięta. Pozostałe przestrzenie definiujemy tak, aby warunek (c) definicji był spełniony

$$V_j = \{f \in L^2(\mathbf{R}); f(2^{-j}x) \in V_0\}.$$

V_j zachowuje wszystkie własności V_0 , które są zachowywane przy przeskalowaniu: jest podprzestrzenią domkniętą, składa się z funkcji różniczkowalnych

$N - 1$ razy w sposób ciągły, które są wielomianami stopnia $\leq N$ na przedziałach postaci $2^{-j}[k, k+1] = [2^{-j}k, 2^{-j}(k+1)]$. W ten sposób skonstruowaliśmy ciąg podprzestrzeni domkniętych w $L^2(\mathbf{R})$, które spełniają warunek (c) definicji MRA. Zauważmy, że $V_j \subset V_{j+1}$. Wynika to z faktu, że każdy przedział postaci $2^{-(j+1)}[k, k+1]$ zawiera się w całości w którymś z dłuższych przedziałów $2^{-j}[k', k'+1]$:

$$2^{-(j+1)}[k, k+1] = 2^{-j}[k/2, (k+1)/2] \subset \begin{cases} 2^{-j}[k/2, k/2+1] & k - \text{parzyste} \\ 2^{-j}[(k-1)/2, (k-1)/2+1] & k - \text{nieparzyste} \end{cases}$$

Jeżeli $f \in V_j$, czyli jest wielomianem na dłuższych przedziałach, to jest także wielomianem na zawartych w nich krótszych przedziałach, czyli $f \in V_{j+1}$. Ciąg podprzestrzeni $\{v_j\}$ jest więc rosnący. Niech φ będzie funkcją Haara, czyli

$$\varphi(x) = \chi_{[0,1]}(x) = \begin{cases} 1 & x \in [0, 1], \\ 0 & x \notin [0, 1], \end{cases}$$

oraz

$$\Delta^N(x) = \varphi * \dots * \varphi(x), \quad (\text{splot } N + 1\text{-krotny}).$$

Fakt 4.2. Δ^N jest splinem rzędu N .

Dowód. W przypadku $N = 0$ wynika to od razu z definicji: funkcja Haara jest stała pomiędzy sąsiednimi liczbami całkowitymi. W rozdziale o przestrzeni Hilberta pokazaliśmy, że $\Delta^1 = \varphi * \varphi$ jest ciągła, liniowo rośnie na przedziale $[0, 1]$, liniowo maleje na $[1, 2]$, i jest równa 0 poza tym. Jest więc splinem rzędu 1. Dowód faktu jest indukcyjny. Dla $n \geq 2$

$$\Delta^N(x) = \Delta^{N-1} * \varphi(x) = \int_{-\infty}^{\infty} \Delta^{N-1}(x-y) \varphi(y) dy = \int_0^1 \Delta^{N-1}(x-y) dy = \int_{x-1}^x \Delta^{N-1}(y) dy,$$

gdzie w ostatniej całce zrobiliśmy zamianę zmiennych $y' = x - y$. Korzystając z zasadniczego twierdzenia rachunku różniczkowego mamy

$$(\Delta^N)'(x) = \Delta^{N-1}(x) - \Delta^{N-1}(x-1).$$

Jeżeli więc założymy, że Δ^{N-1} jest splinem rzędu $N - 1$, jest różniczkowalna w sposób ciągły $N - 1$ razy, to Δ^N jest różniczkowalna w sposób ciągły o jeden raz więcej, czyli N razy. Niech $x \in [k, k+1]$. Wtedy $k \in [x-1, x]$, oraz

$$\Delta^N(x) = \int_{x-1}^x \Delta^{N-1}(y) dy = \int_{x-1}^k \Delta^{N-1}(y) dy + \int_k^x \Delta^{N-1}(y) dy.$$

Na obu przedziałach całkowania Δ^{N-1} jest wielomianem stopnia $\leq N - 1$, więc obie całki są wielomianami stopnia $\leq N$ zmiennej x , więc ich suma też. Jeżeli więc Δ^{N-1} jest splinem rzędu $N - 1$ to Δ^N jest splinem rzędu N . $\square$

W rozdziale o przestrzeni Hilberta pokazaliśmy, że dla $N = 0$ przesunięcia całkowite φ stanowią bazę o.n. przestrzeni V_0 , a w przypadku $N = 1$ przesunięcia Δ^1 stanowią bazę Riesz V_0 . Dla dowolnego N można pokazać, że przesunięcia całkowite Δ^N stanowią bazę Riesz V_0 , i warunki (a) i (b) definicji MRA są spełnione. Dowód tego odłożymy do momentu, kiedy udowodnimy wygodną charakteryzację układów Riesz. Dla $N = 0$ tak powstałą analizę nazywamy analizą Haara, a dla $N \geq 1$ nazywamy ją analizą splinową. Są to ważne przykłady analiz, które często spotyka się w praktyce. W oparciu o nie konstruuje się fale Haara i fale splinowe. Funkcje $\Delta^N(x)$ nazywa się splinami podstawowymi rzędu N .

(ii) Analiza Shannona. Niech

$$V_j = \{f \in L^2(\mathbf{R}); \hat{f}(\xi) \equiv 0 \text{ dla } \xi \notin [-2^j\pi, 2^j\pi]\}.$$

V_j składa się z funkcji o „spektrum” ograniczonym do $[-2^j\pi, 2^j\pi]$. Widać więc od razu, że jest to rosnący ciąg podprzestrzeni $V_j \subset V_{j+1}$. Widać też, że są to podprzestrzenie domknięte. Mamy $\widehat{f(2\cdot)}(\xi) = 1/2\hat{f}(\xi/2)$. $\hat{f}(\xi) \equiv 0$ poza $[-2^j\pi, 2^j\pi]$ dokładnie wtedy, gdy $\hat{f}(\xi/2) \equiv 0$ poza $[-2^{j+1}\pi, 2^{j+1}\pi]$, czyli (c) definicji MRA jest spełnione. Pokażemy, że φ zdefiniowane przez swoją transformatę Fouriera

$$\hat{\varphi}(\xi) = \chi_{[-\pi, \pi]}(\xi) = \begin{cases} 1 & : \xi \in [-\pi, \pi] \\ 0 & : \xi \notin [-\pi, \pi], \end{cases}$$

jest o.n. funkcją skalującą. Dowód tego jest bardzo prosty. $f \in V_0$ dokładnie wtedy, gdy $f \in L^2(\mathbf{R})$ i $\hat{f}(\xi) \equiv 0$ dla $\xi \notin [-\pi, \pi]$. Niech $\{\alpha_k\} \in \ell^2$ będą współczynnikami Fouriera funkcji $\hat{f}(\xi)$ na przedziale $[-\pi, \pi]$, czyli

$$\hat{f}(\xi) = \sum_{k=-\infty}^{\infty} \alpha_k e^{ik\xi} \quad \text{w } L^2(\mathbf{T}). \quad (4.4)$$

Ponieważ $\hat{\varphi}(\xi)$ jest równa 1 wszędzie tam, gdzie $\hat{f}(\xi) \neq 0$, więc

$$\hat{f}(\xi) = \hat{\varphi}(\xi)\hat{f}(\xi) = \sum_{k=-\infty}^{\infty} \alpha_k e^{ik\xi} \hat{\varphi}(\xi). \quad (4.5)$$

Powyższa równość, jak łatwo sprawdzić, zachodzi w $L^2(\mathbf{R})$:

$$\begin{aligned} \left\| \hat{f} - \sum_{k=-M}^N \alpha_k e^{ik\cdot} \hat{\varphi} \right\|^2 &= \int_{-\infty}^{\infty} \left| \hat{f}(\xi) - \sum_{k=-M}^N \alpha_k e^{ik\xi} \hat{\varphi}(\xi) \right|^2 d\xi \\ &= \int_{-\pi}^{\pi} \left| \hat{f}(\xi) - \sum_{k=-M}^N \alpha_k e^{ik\xi} \right|^2 d\xi \\ &\xrightarrow{N, M \rightarrow \infty} 0, \end{aligned}$$

korzystając z (4.4). Skoro równość (4.5) zachodzi w $L^2(\mathbf{R})$, to możemy odwrócić transformatę Fouriera, i mamy

$$f(x) = \sum_{k=-\infty}^{\infty} \alpha_{-k} \varphi(x-k) \quad \text{w } L^2(\mathbf{R}). \quad (4.6)$$

Innymi słowy, $f \in V_0$ dokładnie wtedy, gdy istnieją współczynniki $\|\alpha_k\| \in \ell^2$ takie, że zachodzi (4.6). Przesunięcia całkowite φ stanowią więc układ zupełny w V_0 . Łatwo pokazać, że są też układem o.n.:

$$\begin{aligned} \langle \varphi(\cdot - n), \varphi(\cdot - k) \rangle &= \frac{1}{2\pi} \langle \widehat{\varphi(\cdot - n)}, \widehat{\varphi(\cdot - k)} \rangle \\ &= \frac{1}{2\pi} \int_{-\infty}^{\infty} \hat{\varphi}(\xi) e^{-in\xi} \overline{\hat{\varphi}(\xi)} e^{ik\xi} d\xi \\ &= \frac{1}{2\pi} \int_{-\pi}^{\pi} e^{i(k-n)\xi} d\xi \\ &= \begin{cases} 1 & : k = n \\ 0 & : k \neq n. \end{cases} \end{aligned}$$

Warunki (a) i (b) definicji MRA są spełnione automatycznie, co wynika z faktu, udowodnionego przed twierdzeniem ??, dalej w tym rozdziale.

Konstrukcja falki

Niech będzie dana o.n. analiza wielorozdzielcza z funkcją skalującą φ . Niech W_0 będzie dopełnieniem ortogonalnym V_0 w V_1 :

$$W_0 = \{f \in V_1; f \perp g \ \forall g \in V_0\}, \quad \text{czyli} \quad W_0 = V_1 \ominus V_0. \quad (4.7)$$

Ponieważ V_0 jest domknięta, więc

$$V_1 = V_0 \oplus W_0. \quad (4.8)$$

Naszym celem będzie znalezienie takiej funkcji $\psi \in W_0$, że układ

$$\{\psi(x-n); n \in \mathbf{Z}\}$$

stanowi bazę o. n. W_0 . Pokażemy potem, że taka ψ jest falką.

Filtr dolnoprzepustowy

Z definicji MRA wynika, że funkcja $(1/2)\varphi(x/2)$ jest elementem $V_{-1} \subset V_0$. Można ją więc zapisać w bazie jako

$$\frac{1}{2}\varphi\left(\frac{x}{2}\right) = \sum_{n=-\infty}^{\infty} h_n \varphi(x-n), \quad (4.9)$$

gdzie szereg jest zbieżny w $L^2(\mathbf{R})$, współczynniki bazowe dane są przez

$$h_n = \left\langle \frac{1}{2}\varphi\left(\frac{\cdot}{2}\right), \varphi(\cdot-n) \right\rangle,$$

oraz

$$\sum_{j=-\infty}^{\infty} |h_n|^2 = \int_{j=-\infty}^{\infty} \left| \frac{1}{2}\varphi\left(\frac{x}{2}\right) \right|^2 dx = \frac{1}{2} \|\varphi\|^2 = \frac{1}{2}.$$

Zastosujmy transformatę Fouriera do (4.9). Ponieważ szereg jest zbieżny w $L^2(\mathbf{R})$, to możemy z transformatą wejść pod znak sumy.

$$\hat{\varphi}(2\xi) = \sum_{n=-\infty}^{\infty} h_n \hat{\varphi}(\xi) e^{-in\xi}. \quad (4.10)$$

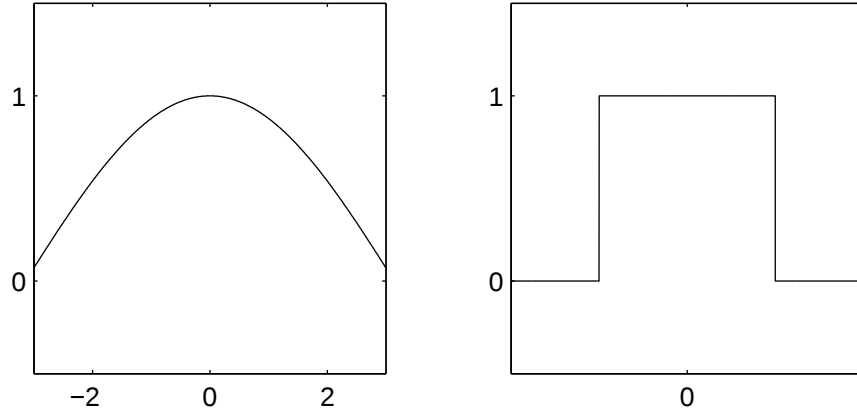
Szereg jest zbieżny w $L^2(\mathbf{R})$. Skoro ciąg $\{h_n\}_{n=-\infty}^{\infty}$ jest sumowalny z kwadratem, to szereg $\sum_{n=-\infty}^{\infty} h_n e^{-in\xi}$ jest zbieżny w $L^2(\mathbf{T})$ do funkcji którą oznaczymy przez m_0 :

$$m_0(\xi) = \sum_{n=-\infty}^{\infty} h_n e^{-in\xi}. \quad (4.11)$$

Łącząc te dwie zbieżności można uzasadnić, że w takim razie

$$\hat{\varphi}(2\xi) = m_0(\xi) \hat{\varphi}(\xi). \quad (4.12)$$

Funkcję $m_0(\xi)$ nazywamy filtrem dolnoprzepustowym. Czasem sam ciąg współczynników $\{h_n\}_{n=-\infty}^{\infty}$ też będziemy nazywali filtrem dolnoprzepustowym. Pokazaliśmy więc, że funkcja $(1/2)\varphi(x/2)$ (lub, co na jedno wychodzi $\varphi(x/2)$) jest wynikiem działania filtru m_0 na φ . Filtr m_0 nazywamy dolnoprzepustowym, gdyż jego wartości (jeżeli jest funkcją ciągłą) w otoczeniu 0



Rysunek 4.3: Filtr dolnoprzepustowy Haara $|m_0|$ i idealny filtr dolnoprzepustowy.

są zbliżone do 1, w czym przypomina charakterystykę typowego filtra dolnoprzepustowego.

Na obrazku 4.3 pokazany jest wykres $|m_0(\xi)|$ dla analizy wielorozdzielczej Haara. Oczywiście nazwanie filtra dolnoprzepustowym jest uproszczeniem. Na przykład funkcja m_0 jest okresowa. Działanie filtra m_0 na funkcjach nie jest splotem z funkcją całkowalną. Łatwo zauważyć, że jest ono splotem z funkcją uogólnioną

$$H(x) = \sum_{k=-\infty}^{\infty} h_k \delta(x - k), \quad \frac{1}{2} \varphi\left(\frac{1}{2}x\right) = (H * \varphi)(x),$$

gdzie symbol δ oznacza impuls jednostkowy (deltę Diraca). Nie będziemy uściślać powyższej uwagi.

W podobny sposób jak (4.12) pokażemy następujące twierdzenie

Twierdzenie 4.3. (i) $f \in V_0 \Leftrightarrow \exists \lambda \in L^2(\mathbf{T})$ taka, że

$$\hat{f}(\xi) = \lambda(\xi) \hat{\varphi}(\xi).$$

Mamy też następującą równość

$$\|f\|^2 = \|\lambda\|^2 = \sum_{k=-\infty}^{\infty} |\hat{\lambda}(k)|^2,$$

(normy w odpowiednich przestrzeniach).

(ii) $f \in V_j \Leftrightarrow \exists \lambda \in L^2(\mathbf{T})$ taka, że

$$\hat{f}(2^j \xi) = \lambda(\xi) \hat{\varphi}(\xi),$$

oraz

$$\|f\|^2 = 2^j \|\lambda\|^2 = 2^j \sum_{k=-\infty}^{\infty} |\hat{\lambda}(k)|^2.$$

Dowód. (i) f należy do V_0 dokładnie wtedy, gdy

$$f(x) = \sum_{n=-\infty}^{\infty} \alpha_n \varphi(x-n) \quad \text{w} \quad L^2(\mathbf{R}),$$

gdzie $\{\alpha_n\}$ jest pewnym ciągiem w ℓ^2 . To z kolei, stosując jak poprzednio transformatę Fouriera, jest równoważne

$$\hat{f}(\xi) = \lambda(\xi) \hat{\varphi}(\xi) \quad \text{gdzie} \quad \lambda(\xi) = \sum_{n=-\infty}^{\infty} \alpha_n e^{-in\xi}.$$

Dodatkowo,

$$\|f\|^2 = \sum_{n=-\infty}^{\infty} |\alpha_n|^2, \quad \text{a} \quad \alpha_n = \hat{\lambda}(-n),$$

a więc otrzymujemy (i). Do (ii) wystarczy zauważyć, że

$$f \in V_j \Leftrightarrow f(2^{-j}x) \in V_0, \quad \text{oraz} \quad \|f\|^2 = 2^{-j} \|f(2^{-j} \cdot)\|^2.$$

□

Kluczowym narzędziem w konstrukcji falki jest następujące twierdzenie

Twierdzenie 4.4. (i) *Dla dowolnej funkcji $f \in L^2(\mathbf{R})$ układ*

$$\{f(x-n); n \in \mathbf{Z}\}$$

jest ortonormalny wtedy i tylko wtedy gdy

$$\sum_{k=-\infty}^{\infty} |\hat{f}(\xi + 2k\pi)|^2 = 1. \quad (4.13)$$

(ii) *Dla dowolnych funkcji $f, g \in L^2(\mathbf{R})$ układy*

$$\{f(x-n); n \in \mathbf{Z}\} \quad \text{i} \quad \{g(x-n); n \in \mathbf{Z}\}$$

są wzajemnie ortogonalne (czyli każda funkcja z jednego zbioru jest ortogonalna do każdej funkcji z drugiego) wtedy i tylko wtedy gdy

$$\sum_{k=-\infty}^{\infty} \hat{f}(\xi + 2k\pi) \overline{\hat{g}(\xi + 2k\pi)} = 0. \quad (4.14)$$

(W każdym przypadku szeregi są zbieżne prawie wszędzie.)

Dowód. Części (i) i (ii) są bardzo podobne. Przeprowadzimy dowód (i).

$$\begin{aligned}\langle f(\cdot - n), f(\cdot - k) \rangle &= \frac{1}{2\pi} \langle \widehat{f(\cdot - n)}, \widehat{f(\cdot - k)} \rangle \\ &= \frac{1}{2\pi} \int_{-\infty}^{\infty} \hat{f}(\xi) e^{-in\xi} \overline{\hat{f}(\xi) e^{-ik\xi}} d\xi \\ &= \frac{1}{2\pi} \int_{-\infty}^{\infty} |\hat{f}(\xi)|^2 e^{-i(n-k)\xi} d\xi.\end{aligned}$$

Całkę po $\mathbf{R}$ zapiszemy jako sumę całek po kolejnych przedziałach $[(2l-1)\pi, (2l+1)\pi]$, $l = 0, \pm 1, \pm 2, \dots$, a następnie zamienimy zmienne.

$$\begin{aligned}&= \frac{1}{2\pi} \sum_{l=-\infty}^{\infty} \int_{(2l-1)\pi}^{(2l+1)\pi} |\hat{f}(\xi)|^2 e^{-i(n-k)\xi} d\xi \\ &= \frac{1}{2\pi} \sum_{l=-\infty}^{\infty} \int_{-\pi}^{\pi} |\hat{f}(\xi + 2l\pi)|^2 e^{-i(n-k)\xi} d\xi,\end{aligned}$$

gdzie w ostatniej całce zamieniliśmy zmienne $\xi \mapsto \xi + 2l\pi$. Funkcja wykładnicza nie zmieniła się, bo jest okresowa. Zamienimy teraz kolejność sumowania i całkowania. W tym wypadku jest to możliwe na mocy twierdzenia o zbieżności ograniczonej, gdyż funkcja

$$F(\xi) = \sum_{l=-\infty}^{\infty} |\hat{f}(\xi + 2l\pi)|^2$$

jest całkowna na $[-\pi, \pi]$, jeżeli $f \in L^2(\mathbf{R})$. Kontynuując rachunki otrzymujemy

$$\begin{aligned}&= \frac{1}{2\pi} \int_{-\pi}^{\pi} \left(\sum_{l=-\infty}^{\infty} |\hat{f}(\xi + 2l\pi)|^2 \right) e^{-i(n-k)\xi} d\xi \\ &= \frac{1}{2\pi} \int_{-\pi}^{\pi} F(\xi) e^{-i(n-k)\xi} d\xi.\end{aligned}$$

Zauważmy, że ostatnie wyrażenie jest współczynnikiem Fouriera rzędu $n-k$ funkcji $F(\xi)$. Funkcja ta jest całkowna na $[-\pi, \pi]$, ale niekoniecznie całkowna z kwadratem. Dla funkcji całkownych też można obliczać współczynniki Fouriera i te współczynniki są jednoznaczne. To znaczy że dwie funkcje całkowne o identycznych współczynnikach Fouriera muszą być równe, prawie wszędzie. Funkcja $F(\xi)$ ma więc współczynniki Fouriera

$$\widehat{F}(n) = \langle f, f(\cdot - n) \rangle.$$

A więc,

$$\langle f(\cdot - n), f(\cdot - k) \rangle = \begin{cases} 1 & n = k, \\ 0 & n \neq k \end{cases}$$

wtedy i tylko wtedy, gdy $F(\xi) \equiv 1$.

Dowód części (ii) wygląda podobnie. Zaczynamy od

$$\langle f(\cdot - n), g(\cdot - k) \rangle$$

i otrzymujemy, że jest to $(n-k)$ -ty współczynnik Fouriera funkcji całkwalnej

$$\sum_{l=-\infty}^{\infty} \hat{f}(\xi + 2l\pi) \overline{\hat{g}(\xi + 2l\pi)}.$$

Wszystkie współczynniki są zerami wtedy i tylko wtedy, gdy sama funkcja jest stale 0. $\square$

Zastosujemy teraz część (i) twierdzenia do funkcji skalującej φ . W następującej sumie, która zgodnie z twierdzeniem jest równa 1 rozdzielamy wyrazy parzyste i nieparzyste, stosujemy (4.12) i korzystamy z okresowości m_0 :

$$\begin{aligned} 1 &= \sum_{k=-\infty}^{\infty} |\hat{\varphi}(2\xi + 2k\pi)|^2 \\ &= \sum_{\substack{k=-\infty \\ \text{parzyste}}}^{\infty} |\hat{\varphi}(2\xi + 2k\pi)|^2 + \sum_{\substack{k=-\infty \\ \text{nieparzyste}}}^{\infty} |\hat{\varphi}(2\xi + 2k\pi)|^2 \\ &= \sum_{k=-\infty}^{\infty} |\hat{\varphi}(2\xi + 2(2k)\pi)|^2 + \sum_{k=-\infty}^{\infty} |\hat{\varphi}(2\xi + 2(2k+1)\pi)|^2 \\ &= \sum_{k=-\infty}^{\infty} |m_0(\xi + 2k\pi)|^2 |\hat{\varphi}(\xi + 2k\pi)|^2 + \\ &\quad + \sum_{k=-\infty}^{\infty} |m_0(\xi + \pi + 2k\pi)|^2 |\hat{\varphi}(\xi + \pi + 2k\pi)|^2 \\ &= |m_0(\xi)|^2 \sum_{k=-\infty}^{\infty} |\hat{\varphi}(\xi + 2k\pi)|^2 + |m_0(\xi + \pi)|^2 \sum_{k=-\infty}^{\infty} |\hat{\varphi}(\xi + \pi + 2k\pi)|^2 \\ &= |m_0(\xi)|^2 + |m_0(\xi + \pi)|^2. \end{aligned}$$

Filtr dolnoprzepustowy spełnia więc tak zwane równanie Barnwella-Smitha

$$|m_0(\xi)|^2 + |m_0(\xi + \pi)|^2 = 1. \quad (4.15)$$

Filtr górnoprzepustowy

Przypomnijmy, że szukamy funkcji $\psi \in V_1$. Korzystając z twierdzenia 4.3 (ii) widzimy, że $\hat{\psi}$ musi mieć postać

$$\hat{\psi}(2\xi) = \lambda(\xi)\hat{\varphi}(\xi) \quad (4.16)$$

dla pewnej funkcji $\lambda \in L^2(\mathbf{T})$. Ponadto przesunięcia całkowite ψ mają stanowić układ ortonormalny, i mają być ortogonalne do przesunięć φ (mają stanowić bazę o.n. W_0). Przeprowadzając rachunek podobny do powyższego, i stosując twierdzenie 4.4 (i) oraz (ii) otrzymujemy

$$|\lambda(\xi)|^2 + |\lambda(\xi + \pi)|^2 = 1, \quad (4.17)$$

$$m_0(\xi)\overline{\lambda(\xi)} + m_0(\xi + \pi)\overline{\lambda(\xi + \pi)} = 0. \quad (4.18)$$

Znalezienie odpowiedniej funkcji λ jest już proste. Zauważmy, że następująca funkcja wstawiona w miejsce λ spełnia (4.17) i (4.18):

$$m_1(\xi) = e^{-i\xi}\overline{m_0(\xi + \pi)}. \quad (4.19)$$

Funkcję m_1 nazywamy filtrem górnoprzepustowym analizy wielorozdzielczej. Ponieważ $\psi \in V_1$, więc $(1/2)\psi(x/2) \in V_0$, a więc istnieje ciąg współczynników $\{g_n\}_{n=-\infty}^{\infty}$, taki, że

$$\frac{1}{2}\psi\left(\frac{x}{2}\right) = \sum_{n=-\infty}^{\infty} g_n\varphi(x - n). \quad (4.20)$$

Współczynniki g_n są współczynnikami filtra górnoprzepustowego

$$m_1(\xi) = \sum_{n=-\infty}^{\infty} g_n e^{-in\xi}, \quad (4.21)$$

a więc możemy je wyliczyć znając współczynniki filtra dolnoprzepustowego.

$$\begin{aligned} m_1(\xi) &= e^{-i\xi}\overline{m_0(\xi + \pi)} \\ &= e^{-i\xi}\overline{\sum_{n=-\infty}^{\infty} h_n e^{-in(\xi + \pi)}} \\ &= \sum_{n=-\infty}^{\infty} \overline{h_n} (-1)^n e^{i(n-1)\xi} \\ &= \sum_{n=-\infty}^{\infty} \overline{h_{1-n}} (-1)^{1-n} e^{-in\xi}. \end{aligned}$$

Z jednoznaczności rozwinięcia w szereg Fouriera mamy

$$g_n = (-1)^{1-n} \overline{h_{1-n}}. \quad (4.22)$$

Sam ciąg współczynników $\{g_n\}$ też czasem nazywa się filtrem górnoprzepustowym analizy.

Filtry QMF

Podstawowe własności pary filtrów m_0 i m_1 to

$$\begin{aligned} |m_0(\xi)|^2 + |m_0(\xi + \pi)|^2 &= 1, \\ |m_1(\xi)|^2 + |m_1(\xi + \pi)|^2 &= 1, \\ m_0(\xi) \overline{m_1(\xi)} + m_0(\xi + \pi) \overline{m_1(\xi + \pi)} &= 0. \end{aligned}$$

Własności te można sformułować bezpośrednio w języku współczynników filtrów:

$$\sum_{k=-\infty}^{\infty} h_{k+2n} \overline{h_k} = \begin{cases} \frac{1}{2} & n = 0 \\ 0 & n \neq 0, \end{cases}$$

i podobnie dla $\{g_n\}$, oraz

$$\sum_{k=-\infty}^{\infty} h_{2n+k} \overline{g_k} = 0, \quad n \in \mathbf{Z}.$$

Parę filtrów (m_0, m_1) spełniających powyższe warunki nazywa się filtrem QMF (quadrature mirror filter). Mają one zastosowanie w teorii przetwarzania sygnału niezależnie od teorii falek.

Sformułujemy teraz ostateczne twierdzenia, które mówią, że ψ istotnie jest falką.

Twierdzenie 4.5. *Niech funkcja ψ będzie dana przez*

$$\hat{\psi}(2\xi) = m_1(\xi) \hat{\varphi}(\xi) = e^{-i\xi} \overline{m_0(\xi + \pi)} \hat{\varphi}(\xi). \quad (4.23)$$

Wtedy zbiór funkcji

$$\{\psi(x - n); n \in \mathbf{Z}\} \quad (4.24)$$

stanowi bazę o.n. przestrzeni $W_0 = V_1 \ominus V_0$.

Dowód. Przypomnijmy, że z definicji $\psi \in V_1$, a przestrzeń V_1 jest niezmienna na przesunięcia całkowite (nawet na przesunięcia połówkowe). Cały zbiór funkcji (4.24) leży więc w V_1 , jest ortonormalny i ortogonalny do V_0

(przypomnijmy warunki (4.18) i (4.19)), a więc jest układem ortonormalnym w W_0 . Pozostaje pokazać, że jest to układ zupełny, to znaczy, że kombinacje liniowe elementów (4.24) leżą gęsto w W_0 . Wystarczy pokazać, że jeżeli $f \in W_0$ i f jest ortogonalna do wszystkich elementów (4.24) to $f \equiv 0$. Niech więc $f \in W_0$. W szczególności $f \in V_1$, więc istnieje $\lambda \in L^2(\mathbf{T})$ taka, że

$$\hat{f}(2\xi) = \lambda(\xi)\hat{\varphi}(\xi).$$

f jako element W_0 jest ortogonalna do wszystkich przesunięć φ . Zakładamy dodatkowo, że jest ortogonalna do wszystkich przesunięć ψ . Oczywiście całkowite przesunięcia f też mają te własności. Korzystając z twierdzenia 4.4 (ii), podobnie jak poprzednio, otrzymujemy

$$m_0(\xi)\overline{\lambda(\xi)} + m_0(\xi + \pi)\overline{\lambda(\xi + \pi)} = 0 \quad (4.25)$$

$$m_1(\xi)\overline{\lambda(\xi)} + m_1(\xi + \pi)\overline{\lambda(\xi + \pi)} = 0. \quad (4.26)$$

Pomnóżmy stronami (4.25) przez $\overline{m_0(\xi)}$. W (4.26) wstawmy wzór na m_1 , pomnóżmy stronami przez $e^{i\xi} m_0(\xi + \pi)$ i uwzględnijmy, że $e^{-i\pi} = -1$. Pozostaje dodać równania stronami, aby otrzymać $\lambda(\xi) = 0$, a ponieważ ξ jest dowolne, to

$$\hat{f}(2\xi) = 0.$$

Widzimy więc, że układ (4.24) jest bazą o.n. przestrzeni W_0 . □

Podsumowując, skonstruowaliśmy funkcję ψ , której całkowite przesunięcia (4.24) stanowią bazę o.n. przestrzeni $W_0 = V_1 \ominus V_0$. Teraz pokażemy, że układ falkowy (4.1) stanowi bazę o.n. całej przestrzeni $L^2(\mathbf{R})$. Niech W_j będzie dopełnieniem ortogonalnym V_j w V_{j+1}

$$W_j = V_{j+1} \ominus V_j, \quad \text{czyli} \quad V_{j+1} = V_j \oplus W_j, \quad j \in \mathbf{Z}. \quad (4.27)$$

Ta definicja rozszerza wcześniejszą definicję W_0 . Mamy następujący fakt

Fakt 4.6. (i) $W_j = \{f \in L^2(\mathbf{R}) : f(2^{-j}x) \in W_0\}$,

(ii) *Układ funkcji*

$$\left\{ 2^{\frac{j}{2}} \psi(2^j x - n) : n \in \mathbf{Z} \right\} \quad (4.28)$$

jest bazą o.n. przestrzeni W_j .

Dowód. (i) Z definicji MRA mamy, że $f \in V_{j+1} \Leftrightarrow f(2^{-j}\cdot) \in V_1$, $g \in V_j \Leftrightarrow g(2^{-j}\cdot) \in V_0$, a przez zamianę zmiennych mamy

$$f \perp g \Leftrightarrow f(2^{-j}\cdot) \perp g(2^{-j}\cdot).$$

Stwierdzenie, że $f \in V_{j+1} \ominus V_j$ jest więc równoważne stwierdzeniu, że $f(2^{-j}\cdot) \in V_1 \ominus V_0$.

(ii) Przez zamianę zmiennych $\xi \mapsto 2^{-j}\xi$ sprowadzamy (ii) do przypadku $j = 0$, który z kolei jest udowodniony w Twierdzeniu 4.5. $\square$

Zauważmy, że podprzestrzenie W_j są do siebie wzajemnie ortogonalne. Jeżeli $j < k$ to $W_j \subset V_{j+1} \subset V_k$, a W_k jest z definicji ortogonalna do V_k . W każdej z tych podprzestrzeni W_j mamy bazę o.n. (4.28). Następujące twierdzenie pokazuje, że nieskończona suma prosta tych podprzestrzeni jest całością $L^2(\mathbf{R})$, a układ (4.1) bazą o.n.

Twierdzenie 4.7. (i) Dla $J \in \mathbf{Z}$ mamy

$$\bigoplus_{j=-\infty}^J W_j = V_{J+1}, \quad (4.29)$$

i dla tej podprzestrzeni układ funkcji

$$\left\{ 2^{\frac{j}{2}} \psi(2^j x - n) : n, j \in \mathbf{Z}, j \leq J \right\} \quad (4.30)$$

stanowi bazę o.n.,

(ii)

$$\bigoplus_{j=-\infty}^{\infty} W_j = L^2(\mathbf{R}), \quad (4.31)$$

i układ funkcji (4.1), czyli

$$\left\{ 2^{\frac{j}{2}} \psi(2^j x - n) : n, j \in \mathbf{Z} \right\}$$

stanowi bazę o.n. (falkową).

Dowód. (i) Jak zauważyliśmy powyżej przestrzenie W_j są do siebie wzajemnie ortogonalne

$$W_j \perp W_k \quad j \neq k,$$

oraz, dla $j \leq J$

$$W_j \subset V_{j+1} \subset V_{J+1}.$$

Przypomnijmy, że nieskończona suma prosta (4.29) jest domknięciem zbioru kombinacji liniowych elementów podprzestrzeni W_j , $j \leq J$. W takim razie

$$\bigoplus_{j=-\infty}^J W_j \subset V_{J+1}.$$

Chcemy pokazać równość. Niech $f \in V_{J+1}$ i niech f będzie ortogonalna do każdej podprzestrzeni W_j , $j \leq J$. Skoro $f \in V_{J+1}$ i $f \perp W_J$ to $f \in V_J$. Skoro $f \perp W_{J-1}$ to $f \in V_{J-1}$. Postępując tak dalej, pokazujemy, że $f \in V_j$ dla każdego $j \leq J+1$, a więc

$$f \in \bigcap_{j=-\infty}^{J+1} V_j = \bigcap_{j=-\infty}^{\infty} V_j,$$

czyli, zgodnie z punktem (b) definicji MRA $f \equiv 0$. Pokazaliśmy więc (4.29). Fakt, że zbiór funkcji (4.30) stanowi bazę o.n. sumy prostej wynika z tego, że jest to zbiór baz o.n. podprzestrzeni składowych sumy.

(ii) Z (i) wynika, że

$$V_J = \bigoplus_{j=-\infty}^{J-1} W_j \subset \bigoplus_{j=-\infty}^{\infty} W_j,$$

a więc

$$\bigcup_{J=-\infty}^{\infty} V_J \subset \bigoplus_{j=-\infty}^{\infty} W_j.$$

Ponieważ suma prosta jest domknięta, więc domknięcie lewej strony też się w niej zawiera. Z punktu (a) definicji MRA mamy, że domknięcie lewej strony jest całością $L^2(\mathbf{R})$. Fakt, że zbiór funkcji (4.1) stanowi bazę o.n. sumy prostej wynika, podobnie jak w części (i) z tego, że jest to zbiór wszystkich baz ortonormalnych przestrzeni składowych W_j , $j \in \mathbf{Z}$. $\square$

Przykłady: (i) MRA Haara ma funkcję skalującą $\varphi = \chi_{[0,1]}$. Widać więc, że

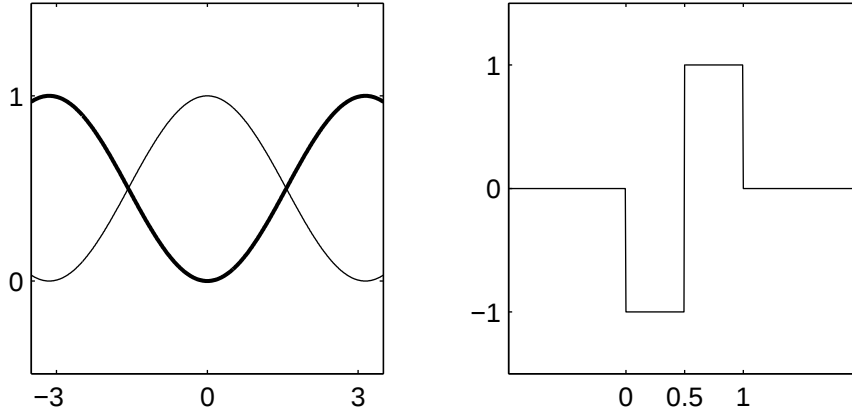
$$\frac{1}{2} \varphi\left(\frac{1}{2}x\right) = \frac{1}{2} \varphi(x) + \frac{1}{2} \varphi(x-1),$$

czyli $h_0 = h_1 = 1/2$, oraz $h_k = 0$ dla $k \neq 0, 1$. W takim razie

$$\begin{aligned} g_0 &= -h_1 = -\frac{1}{2}, & g_1 &= h_0 = \frac{1}{2} & \text{oraz} & & g_k &= 0 \text{ dla } k \neq 0, 1, \\ \frac{1}{2} \psi\left(\frac{1}{2}x\right) &= -\frac{1}{2} \varphi(x) + \frac{1}{2} \varphi(x-1) & \Rightarrow & & \psi(x) &= & -\varphi(2x) + \varphi(2x-1), \\ m_0(\xi) &= \frac{1}{2} + \frac{1}{2} e^{i\xi} = e^{i\xi/2} \cos(\xi/2), & m_1(\xi) &= & -\frac{1}{2} + \frac{1}{2} e^{i\xi} = i e^{i\xi/2} \sin(\xi/2). \end{aligned}$$

(ii) Filtry analizy Shannona wygodniej jest rozważać po stronie transformaty Fouriera. Mamy

$$\hat{\varphi}(\xi) = \chi_{[-\pi, \pi]}(\xi), \quad \hat{\varphi}(2\xi) = \chi_{[-\pi/2, \pi/2]}(\xi),$$



Rysunek 4.4: Wykresy $|m_0(\xi)|^2$ i $|m_1(\xi)|^2$.

więc

$$\begin{aligned} m_0(\xi) &= \hat{\varphi}(2\pi) = \chi_{[-\pi/2, \pi/2]}(\xi) \quad \text{na} \quad [-\pi, \pi], \\ m_1(\xi) &= e^{-i\xi} \chi_{[-\pi, -\pi/2] \cup [\pi/2, \pi]}(\xi) \quad \text{na} \quad [-\pi, \pi], \\ \hat{\psi}(\xi) &= m_1(\xi/2) \hat{\varphi}(\xi/2) = e^{-i\xi/2} \chi_{[-2\pi, -\pi] \cup [\pi, 2\pi]}(\xi). \end{aligned}$$

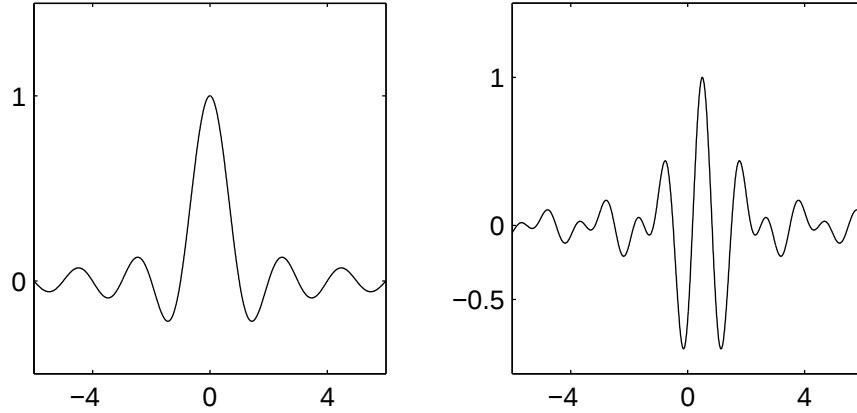
Możemy policzyć φ i ψ analizy Shannona

$$\begin{aligned} \varphi(x) &= \frac{1}{2\pi} \int_{-\pi}^{\pi} \hat{\varphi}(\xi) e^{i\xi x} d\xi = \frac{1}{2\pi} \frac{e^{i\xi x}}{i x} \Big|_{-\pi}^{\pi} = \frac{1}{2\pi i x} (e^{i x \pi} - e^{-i x \pi}) = \frac{\sin(x\pi)}{x\pi}, \\ \psi(x) &= \frac{1}{2\pi} \int_{-2\pi}^{2\pi} \hat{\psi}(\xi) e^{i\xi x} d\xi = \frac{1}{2\pi} \int_{-2\pi}^{-\pi} e^{-i\xi/2} e^{i\xi x} d\xi + \frac{1}{2\pi} \int_{\pi}^{2\pi} e^{-i\xi/2} e^{i\xi x} d\xi \\ &= \frac{1}{2\pi} \frac{e^{i\xi(x-1/2)}}{i(x-1/2)} \Big|_{-2\pi}^{-\pi} + \frac{1}{2\pi} \frac{e^{i\xi(x-1/2)}}{i(x-1/2)} \Big|_{\pi}^{2\pi} \\ &= \frac{1}{2\pi i(x-1/2)} ((e^{-i\pi(x-1/2)} - e^{-i2\pi(x-1/2)}) + (e^{i2\pi(x-1/2)} - e^{i\pi(x-1/2)})) \\ &= \frac{1}{\pi(x-1/2)} (\sin(2\pi(x-1/2)) - \sin(\pi(x-1/2))). \end{aligned}$$

Pakiety falkowe

Zauważmy, że w przestrzeniach V_j mamy w tej chwili wiele różnych baz o.n. Z definicji, mamy bazę

$$\left\{ 2^{\frac{j}{2}} \varphi(2^j x - n); n \in \mathbf{Z} \right\}. \quad (4.32)$$



Rysunek 4.5: Funkcja skalująca i falka Shannona.

Dodatkowo, mamy rozkład

$$V_j = V_{j-1} \oplus W_{j-1},$$

i w każdej z przestrzeni składowych bazy o.n.

$$\left\{ 2^{\frac{j-1}{2}} \varphi(2^{j-1}x - n); n \in \mathbf{Z} \right\} \quad \text{oraz} \quad \left\{ 2^{\frac{j-1}{2}} \psi(2^{j-1}x - n); n \in \mathbf{Z} \right\} \quad (4.33)$$

odpowiednio. W przestrzeni V_j mamy więc do wyboru bazę (4.32) lub sumę baz (4.33). Możemy tak postępować dalej, rozkładając V_{j-1} , i widzimy, że w przestrzeniach V_j mamy, dla każdego $J \geq 0$ bazę o.n.

$$\left\{ 2^{\frac{j-J}{2}} \varphi(2^{j-J}x - n), 2^{\frac{j-k}{2}} \psi(2^{j-k}x - n); n \in \mathbf{Z}, 1 \leq k \leq J \right\}.$$

W takiej sytuacji mówimy, że w V_j mamy bibliotekę baz o.n. W rozdziale o pakietach falkowych wrócimy do tego zagadnienia. Będziemy konstruowali jeszcze inne bazy, rozkładając również przestrzeń W_j na sumy proste, używając pary filtrów dolno- i górnoprzepustowego. W ten sposób, mając konkretny sygnał i bibliotekę baz możemy dobrać do niego indywidualną bazę z biblioteki, i w tej bazie go rozłożyć. Kryteria wyboru baz mogą być różne, ale generalnie chodzi o to, żeby w rozkładzie było jak najmniej dużych współczynników.

MRA Riesz

Wspomnieliśmy, że analiza wielorozdzielcza Riesz jest także, po zamianie funkcji skalującej ale dla tych samych podprzestrzeni V_j , analizą ortonormalną. Teraz to uzasadnimy. Będzie nam potrzebna następująca wersja Twierdzenia 4.4 (i).

Twierdzenie 4.8. *Niech $f \in L^2(\mathbf{R})$. Wtedy układ*

$$\{f(x-n); n \in \mathbf{Z}\} \quad (4.34)$$

jest układem Riesz, to znaczy istnieją stałe $A, B > 0$ takie, że

$$A \sum_{n=-\infty}^{\infty} |\alpha_n|^2 \leq \left\| \sum_{n=-\infty}^{\infty} \alpha_n f(\cdot - n) \right\|^2 \leq B \sum_{n=-\infty}^{\infty} |\alpha_n|^2, \quad \{\alpha_n\}_{n=-\infty}^{\infty} \in \ell^2 \quad (4.35)$$

wtedy i tylko wtedy, gdy

$$A \leq \sum_{k=-\infty}^{\infty} |\hat{f}(\xi + 2k\pi)|^2 \leq B. \quad (4.36)$$

Dowód. Dowód przebiega podobnie do dowodu Twierdzenia 4.4. Jeżeli $f \in L^2(\mathbf{R})$ to funkcja

$$F(\xi) = \sum_{k=-\infty}^{\infty} |\hat{f}(\xi + 2k\pi)|^2 \quad (4.37)$$

jest całkowalna na $[-\pi, \pi]$. Niech ciąg $\alpha = \{\alpha_n\}$ będzie skończony, i niech

$$\hat{\alpha}(\xi) = \sum_{k=-\infty}^{\infty} \alpha_k e^{-ik\xi} \quad (4.38)$$

będzie odpowiadającym mu elementem $L^2(\mathbf{T})$ — wielomianem trygonometrycznym. Pokażemy, że

$$\left\| \sum_{k=-\infty}^{\infty} \alpha_k f(\cdot - k) \right\|^2 = \frac{1}{2\pi} \int_{-\pi}^{\pi} |\hat{\alpha}(\xi)|^2 F(\xi) d\xi \quad (4.39)$$

Dowód powyższej równości przeprowadzimy za chwilę, a teraz zauważmy, że twierdzenie wynika z (4.39). Załóżmy, że funkcje (4.34) tworzą układ Riesz, to znaczy zachodzi (4.35). Niech

$$E = \{\xi \in [-\pi, \pi]; F(\xi) > B\},$$

i niech

$$\chi_E(\xi) = \sum_{k=-\infty}^{\infty} \alpha_k e^{-ik\xi},$$

czyli α_k są współczynnikami Fouriera funkcji charakterystycznej χ_E . Ta funkcja nie jest wielomianem trygonometrycznym (ciąg $\{\alpha_k\}$ nie jest skończony),

ale przy założeniu (4.35) równość (4.39) rozszerza się automatycznie na ciągi $\{\alpha_k\} \in \ell^2$. Mamy więc

$$\frac{1}{2\pi} \int_{-\pi}^{\pi} |\chi_E(\xi)|^2 F(\xi) d\xi = \frac{1}{2\pi} \int_{-\pi}^{\pi} \chi_E(\xi) F(\xi) d\xi \geq \frac{1}{2\pi} B \int_{-\pi}^{\pi} \chi_E(\xi) d\xi = \frac{1}{2\pi} B |E|. \quad (4.40)$$

$|E|$ oznacza miarę (długość) zbioru E . Jeżeli $F(\xi)$ jest ciągła, to E jest otwarty, a więc jest sumą rozłącznych odcinków. Wtedy $|E|$ jest ich łączną długością. Miara jest uogólnieniem długości na inne zbiory. Nie wchodząc w szczegóły, pozostajmy przy takim intuicyjnym rozumieniu miary. Zauważmy, że w (4.40) równość może zachodzić tylko jeżeli $|E| = 0$, w przeciwnym wypadku zachodzi nierówność ostra. Łącząc (4.35), (4.39) oraz (4.40) otrzymujemy

$$\begin{aligned} \frac{1}{2\pi} B |E| &\leq \left\| \sum_{k=-\infty}^{\infty} \alpha_k f(\cdot - k) \right\|^2 \\ &\leq B \sum_{k=-\infty}^{\infty} |\alpha_k|^2 \\ &= B \frac{1}{2\pi} \int_{-\pi}^{\pi} |\hat{\alpha}(\xi)|^2 d\xi \\ &= B \frac{1}{2\pi} |E|. \end{aligned}$$

Widzimy więc, że w (4.40) musi zachodzić równość, a więc $|E| = 0$, a więc $F(\xi) \leq B$ prawie wszędzie. Podobnie pokazujemy, że $F(\xi) \geq A$. W drugą stronę jest prościej. Załóżmy (4.36), i mając (4.39) otrzymujemy

$$\begin{aligned} \left\| \sum_{k=-\infty}^{\infty} \alpha_k f(\cdot - k) \right\|^2 &= \frac{1}{2\pi} \int_{-\pi}^{\pi} |\hat{\alpha}(\xi)|^2 F(\xi) d\xi \\ &\leq \frac{B}{2\pi} \int_{-\pi}^{\pi} |\hat{\alpha}(\xi)|^2 d\xi \\ &= B \sum_{k=-\infty}^{\infty} |\alpha_k|^2. \end{aligned}$$

Podobnie z drugą nierównością. Pozostaje pokazać (4.39). Niech $\{\alpha_k\}$ będzie

ciągim skończonym.

$$\begin{aligned}
\left\| \sum_{k=-\infty}^{\infty} \alpha_k f(\cdot - k) \right\|^2 &= \int_{-\infty}^{\infty} \sum_{k,n=-\infty}^{\infty} \alpha_k \overline{\alpha_n} f(x-k) \overline{f(x-n)} dx \\
&= \sum_{k,n=-\infty}^{\infty} \alpha_k \overline{\alpha_n} \int_{-\infty}^{\infty} f(x-k) \overline{f(x-n)} dx \\
&= \sum_{k,n=-\infty}^{\infty} \alpha_k \overline{\alpha_n} \frac{1}{2\pi} \int_{-\infty}^{\infty} \hat{f}(\xi) e^{-ik\xi} \overline{\hat{f}(\xi)} e^{in\xi} d\xi \\
&= \sum_{k,n=-\infty}^{\infty} \alpha_k \overline{\alpha_n} \frac{1}{2\pi} \int_{-\infty}^{\infty} |\hat{f}(\xi)|^2 e^{-i(k-n)\xi} d\xi \\
&= \sum_{k,n=-\infty}^{\infty} \alpha_k \overline{\alpha_n} \frac{1}{2\pi} \sum_{l=-\infty}^{\infty} \int_{\pi(2l-1)}^{\pi(2l+1)} |\hat{f}(\xi)|^2 e^{-i(k-n)\xi} d\xi \\
&= \sum_{k,n=-\infty}^{\infty} \alpha_k \overline{\alpha_n} \frac{1}{2\pi} \int_{-\pi}^{\pi} F(\xi) e^{-i(k-n)\xi} d\xi \\
&= \sum_{k,n=-\infty}^{\infty} \alpha_k \overline{\alpha_n} \hat{F}(k-n) \\
&= \langle \alpha, \alpha * \widehat{\hat{F}} \rangle \\
&= \frac{1}{2\pi} \int_{-\pi}^{\pi} \hat{\alpha}(\xi) \overline{\widehat{(\alpha * \hat{F})}(\xi)} d\xi \\
&= \frac{1}{2\pi} \int_{-\pi}^{\pi} \hat{\alpha}(\xi) \overline{\widehat{\hat{\alpha}}(\xi) \widehat{\hat{\hat{F}}}(\xi)} d\xi \\
&= \frac{1}{2\pi} \int_{-\pi}^{\pi} |\hat{\alpha}(\xi)|^2 \overline{\widehat{\hat{\hat{F}}}(\xi)} d\xi
\end{aligned}$$

Zabawny symbol w ostatniej całce to $F(\xi)$:

$$\begin{aligned}
\overline{\widehat{\hat{\hat{F}}}(\xi)} &= \overline{\sum_{k=-\infty}^{\infty} \overline{\hat{F}(k)} e^{-ik\xi}} \\
&= \sum_{k=-\infty}^{\infty} \hat{F}(k) e^{ik\xi} \\
&= F(\xi),
\end{aligned}$$

a więc udowodniliśmy (4.39). Jest jeszcze drobny szczegół techniczny. Poka-

zując równość

$$\sum_{k,n=-\infty}^{\infty} \alpha_k \overline{\alpha_n} \hat{F}(k-n) = \frac{1}{2\pi} \int_{-\pi}^{\pi} |\hat{\alpha}(\xi)|^2 F(\xi) d\xi \quad (4.41)$$

korzystaliśmy z twierdzenia Plancherela, w ten sposób po cichu zakładając, że $F(\xi) \in L^2(\mathbf{T})$. W rzeczywistości wiemy tylko, że $F(\xi)$ jest całkowalna. Możemy sobie z tym poradzić w standardowy sposób. Niech, dla $N \in \mathbf{N}$ funkcja $F_N(\xi)$ będzie obcięciem $F(\xi)$ do poziomu N :

$$F_N(\xi) = \min\{F(\xi), N\}.$$

$F_N(\xi)$ jako funkcja ograniczona jest w $L^2(\mathbf{T})$. Mamy więc (4.41) z F_N w miejsce F :

$$\sum_{k,n=-\infty}^{\infty} \alpha_k \overline{\alpha_n} \widehat{F_N}(k-n) = \frac{1}{2\pi} \int_{-\pi}^{\pi} |\hat{\alpha}(\xi)|^2 F_N(\xi) d\xi \quad (4.42)$$

Następnie przechodzimy do granicy, gdy $N \rightarrow \infty$. Na mocy twierdzenia o zbieżności ograniczonej prawa strona (4.42) dąży do prawej strony (4.41). Z drugiej strony widać, że

$$\widehat{F_N}(k) \xrightarrow{N \rightarrow \infty} \hat{F}(k),$$

(znowu twierdzenie o zbieżności ograniczonej), a sumy po lewych stronach są skończone, więc z (4.42) wynika (4.41). $\square$

Możemy teraz sformułować wniosek dotyczący analiz wielorozdzielczych Riesz.

Wniosek 4.9. *Jeżeli $\{V_j\}_{j=-\infty}^{\infty}$ i φ tworzą MRA Riesz, to istnieje $\tilde{\varphi} \in V_0$ taka, że $\{V_j\}_{j=-\infty}^{\infty}$ i $\tilde{\varphi}$ tworzą o.n. MRA.*

Dowód. Wiemy, że funkcja

$$F(\xi) = \sum_{k=-\infty}^{\infty} |\hat{\varphi}(\xi + 2k\pi)|^2$$

jest ograniczona od góry i odcięta od 0 od dołu. W takim razie funkcja $F(\xi)^{-1/2}$ jest ograniczona, a więc w $L^2(\mathbf{T})$, a więc istnieją współczynniki $\{\alpha_n\} \in \ell^2$ takie, że

$$\frac{1}{\sqrt{F(\xi)}} = \sum_{n=-\infty}^{\infty} \alpha_n e^{-in\xi}.$$

Niech

$$\begin{aligned}\hat{\tilde{\varphi}}(\xi) &= \frac{\hat{\varphi}(\xi)}{\sqrt{F(\xi)}} \\ &= \sum_{n=-\infty}^{\infty} \alpha_n e^{-in\xi} \hat{\varphi}(\xi),\end{aligned}$$

czyli, jak wiemy,

$$\tilde{\varphi}(x) = \sum_{n=-\infty}^{\infty} \alpha_n \varphi(x - n).$$

Wynika z tego, że $\tilde{\varphi} \in V_0$, a w związku z tym wszystkie całkowite przesunięcia $\tilde{\varphi}$ też należą do V_0 . Tworzą one układ o.n. w V_0 :

$$\begin{aligned}\sum_{k=-\infty}^{\infty} |\hat{\tilde{\varphi}}(\xi + 2k\pi)|^2 &= \sum_{k=-\infty}^{\infty} \frac{|\hat{\varphi}(\xi + 2k\pi)|^2}{F(\xi + 2k\pi)} \\ &= \frac{1}{F(\xi)} \sum_{k=-\infty}^{\infty} |\hat{\varphi}(\xi + 2k\pi)|^2 \\ &= 1.\end{aligned}$$

Korzystając z tego, że $F(\xi)^{1/2} \in L^2(\mathbf{T})$ widzimy, że istnieją również współczynniki $\{\beta_n\} \in \ell^2$, takie, że

$$\varphi(x) = \sum_{n=-\infty}^{\infty} \beta_n \tilde{\varphi}(x - n),$$

a więc domknięte rozpięcia liniowe przesunąć całkowitych $\tilde{\varphi}$ i φ są takie same, równe V_0 . $\square$

Skoro analiza Riesz'a jest także analizą o.n., więc również generuje bazę falkową. W szczególności dla dowolnego $N \in \mathbf{N}$ istnieje falka ψ będąca splinem rzędu N .

Analizy splinowe

Wróćmy na chwilę do analiz splinowych. W rozdziale o przestrzeni Hilberta pokazaliśmy, że przesunięcia całkowite funkcji Haara $\varphi = \Delta^0 = \chi_{[0,1]}$ stanowią układ o.n., a przesunięcia $\Delta^1 = \varphi * \varphi$ stanowią układ Riesz'a ze stałymi $A = 1/3$ i $B = 1$. Korzystając z Twierdzenia 4.8 pokażemy teraz ogólną metodę, przy pomocy której można pokazać, że przesunięcia dowolnego splinu

Δ^N stanowią układ Riesz. Niech

$$F_N(\xi) = \sum_{k=-\infty}^{\infty} |\widehat{\Delta^N}(\xi + 2k\pi)|^2.$$

Wiemy, że

$$\begin{aligned} |\widehat{\Delta^0}(\xi)| &= |\hat{\varphi}(\xi)| = \left| \frac{\sin(\xi/2)}{\xi/2} \right|, \\ F_0(\xi) &= \sum_{k=-\infty}^{\infty} |\hat{\varphi}(\xi + 2k\pi)|^2 = 4 \sin^2(\xi/2) \sum_{k=-\infty}^{\infty} \frac{1}{(\xi + 2k\pi)^2} = 1, \quad (4.43) \\ F_N(\xi) &= \sum_{k=-\infty}^{\infty} |\hat{\varphi}(\xi + 2k\pi)|^{2(N+1)} = 2^{2(N+1)} \sin^{2(N+1)}(\xi/2) \sum_{k=-\infty}^{\infty} \frac{1}{(\xi + 2k\pi)^{2(N+1)}}. \end{aligned}$$

Oszacujemy wartości $F_N(\xi)$ obliczając sumę w ostatnim wierszu. Z (4.43) wynika

$$\sum_{k=-\infty}^{\infty} \frac{1}{(\xi + 2k\pi)^2} = \frac{1}{4 \sin^2(\xi/2)} = \frac{1}{2(1 - \cos(\xi))}. \quad (4.44)$$

Zauważmy, że powyższą sumę można różniczkować wyraz za wyrazem. Wynika to z tego, że szereg, a także szereg pochodnych jest zbieżny jednostajnie na każdym domkniętym podprzedziale otwartego przedziału $(0, 2\pi)$ (w takiej sytuacji mówimy, że szereg jest zbieżny niemal jednostajnie na $(0, 2\pi)$). Różniczkując (4.45) 4 krotnie otrzymujemy

$$\begin{aligned} 6 \sum_{k=-\infty}^{\infty} \frac{1}{(\xi + 2k\pi)^4} &= \frac{2 + \cos(\xi)}{2(1 - \cos(\xi))^2}, \\ 120 \sum_{k=-\infty}^{\infty} \frac{1}{(\xi + 2k\pi)^6} &= \frac{1 + 2 \cos(\xi)}{2(1 - \cos(\xi))^2} + \frac{15 - 12 \cos^2(\xi) - 3 \cos^3(\xi)}{2(1 - \cos(\xi))^4}. \end{aligned}$$

Otrzymujemy więc

$$\begin{aligned} F_1(\xi) &= 16 \sin^4(\xi/2) \sum_{k=-\infty}^{\infty} \frac{1}{(\xi + 2k\pi)^4} = \frac{2 + \cos(\xi)}{3}, \\ F_2(\xi) &= 64 \sin^6(\xi/2) \sum_{k=-\infty}^{\infty} \frac{1}{(\xi + 2k\pi)^6} = \frac{16 + 16 \cos(\xi) + \cos^2(\xi)}{30}. \end{aligned}$$

Pierwszą funkcję łatwo jest oszacować dokładnie od góry i od dołu

$$\frac{1}{3} \leq F_1(\xi) \leq 1,$$

i otrzymujemy te same stałe A i B , które otrzymaliśmy wcześniej, bezpośrednio badając przesunięcia Δ^1 . Dla $F_2(\xi)$ obliczymy pochodną, i znajdziemy maksima i minima. Łatwo można sprawdzić, że $F_2(\xi)$ osiąga swoją wartość maksymalną w punktach $2n\pi$ a minimalną w punktach $(2n+1)\pi$. Wynika stąd, że

$$\frac{1}{30} \leq F(\xi) \leq \frac{11}{10}.$$

W ten sposób udowodniliśmy, że przesunięcia splinu podstawowego $\Delta^2(x)$ stanowią układ Riesz, ze stałymi $1/30$ i $11/10$. Różniczkując (4.45) można uzyskać stałe Riesz ogólnie, dla dowolnego $\Delta^N(x)$.

Konstrukcja falek Daubechies

Falki Daubechies to falki o ograniczonym nośniku i, w zależności od wersji, różnej gładkości. Powstają one z analizy wielorozdzielczej, której funkcja skalująca ma te same własności co falka, czyli ograniczony nośnik i gładkość. Filtr dolnoprzepustowy takiej analizy wielorozdzielczej musi więc być wielomianem trygonometrycznym, czyli jego ciąg współczynników musi być skończony. Konstrukcja takiej analizy wielorozdzielczej rozpoczyna się więc od znalezienia odpowiedniego filtru. Pierwszym krokiem będzie twierdzenie, mówiące, że jeżeli funkcja m_0 spełnia określone warunki, to jest filtrem dolnoprzepustowym pewnej analizy. Dowód twierdzenia jest długi, ale intuicyjnie jasny i dosyć interesujący. Idea jest następująca. Mając funkcję m_0 konstruujemy funkcję skalującą wzorem

$$\hat{\varphi}(\xi) = m_0(2^{-1}\xi)\hat{\varphi}(2^{-1}\xi) = m_0(2^{-1}\xi)m_0(2^{-2}\xi)\hat{\varphi}(2^{-2}\xi) = \cdots = \prod_{j=1}^{\infty} m_0(2^{-j}\xi).$$

Pytania są więc dwa: co trzeba założyć o φ , żeby była to funkcja skalująca pewnej analizy, i co trzeba założyć o m_0 , żeby powyższy iloczyn nieskończony dawał odpowiednie φ . To zostanie rozstrzygnięte w Twierdzeniu 6. Następnie, kiedy już będzie wiadomo jakie warunki ma spełniać m_0 , skonstruujemy ją. Teraz odpowiemy na pierwsze pytanie. Potrzebny nam będzie następujący fakt.

Fakt 4.10. *Niech*

$$\cdots \subset V_{-1} \subset V_0 \subset V_1 \subset \cdots$$

będzie rosnącym ciągiem podprzestrzeni domkniętych $L^2(\mathbf{R})$, i niech będą spełnione warunki (c) i (d) definicji analizy wielorozdzielczej. Wtedy warunek (b) jest spełniony automatycznie, a warunek (a) jest równoważny następującemu

$$\lim_{j \rightarrow \infty} |\hat{\varphi}(2^{-j}\xi)| = 1,$$

dla prawie każdego $\xi \in \mathbf{R}$.

Dowód. Niech $f \in V_j$. Z warunku (c) definicji analizy wynika, że

$$f_j(x) = 2^{j/2} f(2^j x) \in V_0.$$

Jeżeli $f \in \bigcap V_j$ to $f_j \in V_0$ dla każdego $j \in \mathbf{Z}$. Zamieniając zmienne widzimy, że dla każdego $j \in \mathbf{Z}$

$$\|f_j\| = \|f\|.$$

Korzystając z Twierdzenia 4.3 (i) mamy, że istnieją $m_j \in L^2(\mathbf{T})$ takie, że

$$\hat{f}_j(\xi) = m_j(\xi) \hat{\varphi}(\xi), \quad m_j \in L^2(\mathbf{T}), \quad \|m_j\| = \|f_j\| = \|f\|,$$

czyli

$$\hat{f}(\xi) = 2^{j/2} \hat{f}_j(2^j \xi) = 2^{j/2} m_j(2^j \xi) \hat{\varphi}(2^j \xi).$$

Zauważmy, że w takim razie, z nierówności Schwarza

$$\begin{aligned} \int_{2\pi}^{4\pi} |\hat{f}(\xi)| d\xi &= 2^{j/2} \int_{2\pi}^{4\pi} |m_j(2^j \xi)| |\hat{\varphi}(2^j \xi)| d\xi \\ &\leq 2^{j/2} \left(\int_{2\pi}^{4\pi} |m_j(2^j \xi)|^2 d\xi \right)^{1/2} \left(\int_{2\pi}^{4\pi} |\hat{\varphi}(2^j \xi)|^2 d\xi \right)^{1/2}. \end{aligned}$$

Zamieniamy zmienne w obu całkach. m_j jest 2π -okresowa, a przedział $[2^{j+1}\pi, 2^{j+2}\pi]$ składa się z 2^j okresów, więc

$$\begin{aligned} \int_{2\pi}^{4\pi} |m_j(2^j \xi)|^2 d\xi &= 2^{-j} \int_{2^{j+1}\pi}^{2^{j+2}\pi} |m_j(\xi)|^2 d\xi \\ &= \int_{-\pi}^{\pi} |m_j(\xi)|^2 d\xi \\ &= 2\pi \|m_j\|^2 \\ &= 2\pi \|f\|^2. \end{aligned}$$

$$\begin{aligned} \int_{2\pi}^{4\pi} |\hat{\varphi}(2^j \xi)|^2 d\xi &= 2^{-j} \int_{2^{j+1}\pi}^{2^{j+2}\pi} |\hat{\varphi}(\xi)|^2 d\xi \\ &\leq 2^{-j} \int_{2^{j+1}\pi}^{\infty} |\hat{\varphi}(\xi)|^2 d\xi. \end{aligned}$$

Wstawiając powyższe do naszych rachunków

$$\int_{2\pi}^{4\pi} |\hat{f}(\xi)| d\xi \leq \sqrt{2\pi} \|f\| \left(\int_{2^{j+1}\pi}^{\infty} |\hat{\varphi}(\xi)|^2 d\xi \right)^{1/2}.$$

Całka po prawej stronie jest „ogonem” skończonej całki, więc prawa strona $\rightarrow 0$ gdy $j \rightarrow \infty$, a więc

$$\int_{2\pi}^{4\pi} |\hat{f}(\xi)| d\xi = 0 \Rightarrow \hat{f}(\xi) \equiv 0 \quad \text{na} \quad [2\pi, 4\pi].$$

W podobny sposób pokazujemy, że $\hat{f}(\xi) \equiv 0$ dla $\xi \in [-4\pi, -2\pi]$. Następnie możemy przeprowadzić ten sam argument dla funkcji $f(2^k x)$, która też należy do $\bigcap V_j$. Otrzymujemy

$$\hat{f}(\xi) \equiv 0 \quad \text{dla} \quad |\xi| \in 2^{-k}[2\pi, 4\pi] \quad \text{dla każdego} \quad k \in \mathbf{Z},$$

czyli $\hat{f} \equiv 0$. Pokazaliśmy więc, że spełniony jest warunek (b) definicji MRA.

Teraz zajmijmy się warunkiem (a). Przypomnijmy, że niezależnie od (a) (nie korzystając z niego) możemy pokazać istnienie filtru dolnoprzepustowego m_0 takiego, że

$$\hat{\varphi}(2\xi) = m_0(\xi)\hat{\varphi}(\xi), \quad (4.45)$$

$$\sum_{k=-\infty}^{\infty} |\hat{\varphi}(\xi + 2k\pi)|^2 = 1, \quad (4.46)$$

$$|m_0(\xi)|^2 + |m_0(\xi + \pi)|^2 = 1. \quad (4.47)$$

Z (4.47) wynika, że $|m_0(\xi)| \leq 1$ co, biorąc pod uwagę (4.45) daje nam, że ciąg $|\hat{\varphi}(2^{-j}\xi)|$ jest niemalejący

$$|\hat{\varphi}(2^{-j}\xi)| = |m_0(2^{-(j+1)}\xi)| \cdot |\hat{\varphi}(2^{-(j+1)}\xi)| \leq |\hat{\varphi}(2^{-(j+1)}\xi)|.$$

Z (4.46) wynika, że jest to ciąg ograniczony od góry przez 1, więc dla każdego $\xi \in \mathbf{R}$ mamy granicę

$$g(\xi) = \lim_{j \rightarrow \infty} |\hat{\varphi}(2^{-j}\xi)| \quad \text{dla} \quad 0 \leq g(\xi) \leq 1.$$

Wykorzystamy fakt, że w każdej podprzestrzeni V_j mamy bazę o.n. składającą się z funkcji

$$\varphi_{j,k}(x) = 2^{j/2} \varphi(2^j x - k); \quad k \in \mathbf{Z}.$$

Rzut ortogonalny na V_j można więc zapisać wzorem

$$(P_j f)(x) = \sum_{k=-\infty}^{\infty} \langle f, \varphi_{j,k} \rangle \varphi_{j,k}(x) \quad \text{w } L^2(\mathbf{R}),$$

oraz

$$\|P_j f\|^2 = \sum_{k=-\infty}^{\infty} |\langle f, \varphi_{j,k} \rangle|^2.$$

Dowód następującego faktu zostawiamy jako ćwiczenie.

Fakt 4.11. *Mamy następującą równoważność:*

$$\overline{\bigcup_{j=-\infty}^{\infty} V_j} = L^2(\mathbf{R}) \quad \Leftrightarrow \quad \|P_j f\| \xrightarrow{j \rightarrow \infty} \|f\|, \quad \forall f \in L^2(\mathbf{R}).$$

Warunek po prawej można nieco osłabić, zachowując równoważność. Wystarczy żeby zbieżność zachodziła dla f z jakiegoś gęstego podzbioru.

Niech gęstym podzbiorem $L^2(\mathbf{R})$ będzie zbiór funkcji f o ograniczonym spektrum, czyli takich, których transformata Fouriera ma ograniczony nośnik. Niech f będzie taką funkcją.

$$\begin{aligned} \langle f, \varphi_{j,k} \rangle &= \frac{1}{2\pi} \langle \hat{f}, \hat{\varphi}_{j,k} \rangle \\ &= \frac{1}{2\pi} \int_{-\infty}^{\infty} \hat{f}(\xi) 2^{-j/2} \overline{\hat{\varphi}(2^{-j}\xi)} e^{i 2^{-j} \xi k} d\xi \\ &= \frac{2^{j/2}}{2\pi} \int_{-\infty}^{\infty} \hat{f}(2^j \xi) \overline{\hat{\varphi}(\xi)} e^{i \xi k} d\xi. \end{aligned}$$

Niech j będzie dostatecznie duże, tak, aby $\hat{f}(2^j \xi) \equiv 0$ dla $\xi \notin [-\pi, \pi]$. Wtedy

$$= \frac{2^{j/2}}{2\pi} \int_{-\pi}^{\pi} \hat{f}(2^j \xi) \overline{\hat{\varphi}(\xi)} e^{i \xi k} d\xi.$$

Tak więc $\langle f, \varphi_{j,k} \rangle$ jest współczynnikiem Fouriera funkcji z $L^2(\mathbf{T})$:

$$\langle f, \varphi_{j,k} \rangle = \hat{F}(-k), \quad \text{dla} \quad F(\xi) = \hat{f}(2^j \xi) \overline{\hat{\varphi}(\xi)}.$$

Skorzystamy z równości Plancherela w $L^2(\mathbf{T})$

$$\begin{aligned} \|P_j f\|^2 &= \sum_{k=-\infty}^{\infty} |\langle f, \varphi_{j,k} \rangle|^2 \\ &= \frac{1}{2\pi} \int_{-\pi}^{\pi} \left| 2^{j/2} \hat{f}(2^j \xi) \overline{\hat{\varphi}(\xi)} \right|^2 d\xi \\ &= \frac{2^j}{2\pi} \int_{-\pi}^{\pi} |\hat{f}(2^j \xi)|^2 |\hat{\varphi}(\xi)|^2 d\xi \\ &= \frac{1}{2\pi} \int_{-2^j \pi}^{2^j \pi} |\hat{f}(\xi)|^2 |\hat{\varphi}(2^{-j} \xi)|^2 d\xi \\ &\xrightarrow{j \rightarrow \infty} \frac{1}{2\pi} \int_{-\infty}^{\infty} |\hat{f}(\xi)|^2 g(\xi)^2 d\xi, \end{aligned}$$

gdzie w ostatniej linii skorzystaliśmy z twierdzenia o zbieżności ograniczonej. Widzimy więc, że

$$\|P_j\|^2 \xrightarrow{j \rightarrow \infty} \|f\|^2 = \frac{1}{2\pi} \int_{-\infty}^{\infty} |\hat{f}(\xi)|^2 d\xi$$

dokładnie wtedy, gdy

$$\int_{-\infty}^{\infty} |\hat{f}(\xi)|^2 g(\xi)^2 d\xi = \int_{-\infty}^{\infty} |\hat{f}(\xi)|^2 d\xi.$$

Ponieważ $0 \leq g(\xi) \leq 1$ to powyższa równość zachodzi dla każdej funkcji $\hat{f}$ o nośniku ograniczonym wtedy i tylko wtedy gdy $g(\xi) = 1$ prawie wszędzie. $\square$

Skorzystamy teraz z udowodnionego właśnie faktu. Widzimy, że funkcja $\varphi \in L^2(\mathbf{R})$ jest funkcją skalującą pewnej analizy wielorozdzielczej jeżeli spełnia następujące warunki

$$\sum_{k=-\infty}^{\infty} |\hat{\varphi}(\xi + 2k\pi)|^2 = 1, \quad (4.48)$$

$$\text{istnieje } m_0 \in L^2(\mathbf{T}) \text{ taka, że } \hat{\varphi}(2\xi) = m_0(\xi)\hat{\varphi}(\xi). \quad (4.49)$$

$$\lim_{j \rightarrow \infty} |\hat{\varphi}(2^{-j}\xi)| = 1 \text{ dla prawie każdego } \xi \in \mathbf{R}. \quad (4.50)$$

Warunek (4.48) mówi, że przesunięcia całkowite φ stanowią układ o.n., a więc jeżeli V_0 zdefiniujemy jako domknięte rozpięcie liniowe tych przesunięć, to stanowią one bazę o.n. V_0 . Warunek (4.49) z kolei mówi, że jeżeli zdefiniujemy V_j jako odpowiednie przeskalowanie V_0 , to $V_j \subset V_{j+1}$. W końcu, jak wynika z faktu powyżej warunek (4.50) gwarantuje (a) w definicji MRA. Teraz odpowiemy na drugie pytanie, czyli sformułujemy warunki na m_0 , dzięki którym funkcja φ zdefiniowana przy pomocy iloczynu nieskończonego (po stronie transformaty Fouriera), będzie spełniała (37)–(39).

Twierdzenie 4.12. *Niech funkcja $m_0(\xi) \in L^2(\mathbf{T})$ spełnia następujące warunki*

- (i) $|m_0(\xi)|^2 + |m_0(\xi + \pi)|^2 = 1$,
- (ii) m_0 jest różniczkowalna w 0 i $m_0(0) = 1$,
- (iii) m_0 ma wartości rzeczywiste na $[-\pi/2, \pi/2]$, oraz

$$\inf_{\xi \in [-\pi/2, \pi/2]} m_0(\xi) = K > 0.$$

Wtedy m_0 jest filtrem dolnoprzepustowym pewnej analizy wielorozdzielczej. Dodatkowo, jeżeli m_0 jest wielomianem trygonometrycznym (to znaczy ciąg współczynników Fouriera $\{\hat{m}_0(k)\}$ jest skończony), to funkcja skalująca i falka tej analizy mają nośnik ograniczony (poza pewnym skończonym przedziałem są równe 0).

Dowód. Zauważmy, że iloczyn nieskończony

$$\prod_{j=1}^{\infty} m_0(2^{-j}\xi)$$

jest zbieżny dla każdego $\xi \in \mathbf{R}$. Ponieważ $2^{-j}\xi \rightarrow 0$ dla $j \rightarrow \infty$, więc wystarczy pokazać zbieżność iloczynu dla ξ dostatecznie małych. Niech więc $\xi \in [-\pi/2, \pi/2]$. Wtedy, zgodnie z (iii) wszystkie czynniki są dodatnie. Z (i) wynika, że $|m_0(\xi)| \leq 1$, a więc

$$0 < m_0(2^{-j}\xi) \leq 1 \quad j = 1, 2, \dots, \quad \xi \in [-\pi, \pi].$$

Iloczyn częściowe

$$\prod_{j=1}^N m_0(2^{-j}\xi) \tag{4.51}$$

tworzą więc ciąg nierosnący, ograniczony od dołu przez 0, a więc zbieżny. Granicę iloczynu oznaczamy przez $\hat{\varphi}(\xi)$. Pokażemy, że istotnie ta granica jest elementem $L^2(\mathbf{R})$, a więc transformatą Fouriera czegoś, ale obecnie niech to będzie tylko oznaczenie jakiejś funkcji. Mamy więc

$$\hat{\varphi}(\xi) = \prod_{j=1}^{\infty} m_0(2^{-j}\xi). \tag{4.52}$$

Zauważmy, że $\hat{\varphi}$ spełnia (4.49):

$$\begin{aligned} \hat{\varphi}(2\xi) &= \prod_{j=1}^{\infty} m_0(2^{-j+1}\xi) \\ &= \prod_{j=0}^{\infty} m_0(2^{-j}\xi) \\ &= m_0(\xi) \prod_{j=1}^{\infty} m_0(2^{-j}\xi) \\ &= m_0(\xi) \hat{\varphi}(\xi). \end{aligned}$$

Pokażemy teraz, że $\hat{\varphi}$ spełnia (4.50). Z (i) wynika, że $0 \leq |m_0(\xi)| \leq 1$, a więc także $0 \leq |\hat{\varphi}(\xi)| \leq 1$. Zamienimy iloczyn na sumę korzystając z logarytmu. Zauważmy, że $\log m_0(\xi)$ jest określona na $[-\pi/2, \pi/2]$, i jej pochodna w 0 jest 0:

$$\left. \frac{d}{d\xi} \log m_0(\xi) \right|_{\xi=0} = \frac{1}{m_0(0)} m'_0(0) = 0,$$

gdyż $m'_0(0) = 0$. Wynika to z faktu, że m_0 jest różniczkowalna w 0 i ma tam lokalne maksimum. Tak więc dla każdego $\epsilon > 0$ istnieje $\delta > 0$ taka, że

$$\log m_0(\xi) > -\epsilon|\xi| \quad \text{dla} \quad |\xi| < \delta.$$

Niech więc $|\xi| < 2\delta$ i $|\xi| < \pi$, wtedy $\hat{\varphi}(\xi) > 0$ oraz

$$\begin{aligned} \log \hat{\varphi}(\xi) &= \sum_{j=1}^{\infty} \log m_0(2^{-j}\xi) \\ &> \sum_{j=1}^{\infty} -\epsilon 2^{-j} |\xi| \\ &> -\epsilon\pi \sum_{j=1}^{\infty} 2^{-j} \\ &= -\epsilon\pi, \end{aligned}$$

czyli

$$|\hat{\varphi}(\xi)| > e^{-\epsilon\pi} \quad \text{dla} \quad |\xi| < \delta, \quad |\xi| < \pi.$$

Funkcja $|\hat{\varphi}(\xi)|$ jest więc ciągła w 0 i ma tam wartość 1, a więc spełnione jest (4.50). Musimy jeszcze pokazać, że $\hat{\varphi} \in L^2(\mathbf{R})$ oraz (4.48). W tym celu wprowadźmy następujące funkcje

$$\hat{\varphi}_N(\xi) = \chi_{[-2^N\pi, 2^N\pi]}(\xi) \prod_{j=1}^N m_0(2^{-j}\xi). \quad (4.53)$$

Zauważmy, że dla każdego $\xi \in L^2(\mathbf{R})$ mamy

$$\hat{\varphi}(\xi) = \lim_{N \rightarrow \infty} \hat{\varphi}_N(\xi).$$

Przypomnijmy, że iloczyny częściowe (4.51) nie są elementami $L^2(\mathbf{R})$ (choćby dlatego, że są okresowe). Chcemy pokazać, że $\hat{\varphi} \in L^2(\mathbf{R})$, więc utworzyliśmy ciąg (4.53), który należy do $L^2(\mathbf{R})$ i jest zbieżny do $\hat{\varphi}$ w każdym punkcie. Pokażemy teraz, że wszystkie $\hat{\varphi}_N$ mają tę samą normę, równą $\sqrt{2\pi}$.

Niech $N \geq 2$

$$\begin{aligned}
\|\hat{\varphi}_N\|^2 &= \int_{-\infty}^{\infty} |\hat{\varphi}_N(\xi)|^2 d\xi \\
&= \int_{-2^N\pi}^{2^N\pi} \prod_{j=1}^N |m_0(2^{-j}\xi)|^2 d\xi \\
&= \int_{-2^N\pi}^0 \prod_{j=1}^N |m_0(2^{-j}\xi)|^2 d\xi + \int_0^{2^N\pi} \prod_{j=1}^N |m_0(2^{-j}\xi)|^2 d\xi \\
&= \int_0^{2^N\pi} \left(\prod_{j=1}^N |m_0(2^{-j}(\xi - 2^N\pi))|^2 + \prod_{j=1}^N |m_0(2^{-j}\xi)|^2 \right) d\xi \\
&= \int_0^{2^N\pi} \left(\prod_{j=1}^N |m_0(2^{-j}\xi - 2^{N-j}\pi)|^2 + \prod_{j=1}^N |m_0(2^{-j}\xi)|^2 \right) d\xi.
\end{aligned}$$

m_0 jest 2π -okresowa, więc wszystkie czynniki w obu iloczynach z wyjątkiem N -tego są równe. Możemy je więc wyciągnąć przed nawias, i mamy

$$\begin{aligned}
&= \int_0^{2^N\pi} (|m_0(2^{-N}\xi - \pi)|^2 + |m_0(2^{-N}\xi)|^2) \prod_{j=1}^{N-1} |m_0(2^{-j}\xi)|^2 d\xi \\
&= \int_0^{2^N\pi} \prod_{j=1}^{N-1} |m_0(2^{-j}\xi)|^2 d\xi.
\end{aligned}$$

Funkcja podcałkowa jest okresowa o okresie $2^N\pi$, więc całkę po okresie możemy „przesunąć”:

$$\begin{aligned}
&= \int_{-2^{N-1}\pi}^{2^{N-1}\pi} \prod_{j=1}^{N-1} |m_0(2^{-j}\xi)|^2 d\xi \\
&= \int_{-\infty}^{\infty} |\hat{\varphi}_{N-1}(\xi)|^2 d\xi \\
&= \|\hat{\varphi}_{N-1}\|^2.
\end{aligned}$$

Widzimy więc, że normy wszystkich funkcji $\hat{\varphi}_N$ są równe, $N = 1, 2, \dots$

Policzmy normę $\hat{\varphi}_1$

$$\begin{aligned}
\|\hat{\varphi}_1\|^2 &= \int_{-2\pi}^{2\pi} |m_0(2^{-1}\xi)|^2 d\xi \\
&= \int_{-2\pi}^0 |m_0(2^{-1}\xi)|^2 d\xi + \int_0^{2\pi} |m_0(2^{-1}\xi)|^2 d\xi \\
&= \int_0^{2\pi} (|m_0(2^{-1}(\xi - 2\pi))|^2 + |m_0(2^{-1}\xi)|^2) d\xi \\
&= \int_0^{2\pi} 1 d\xi \\
&= 2\pi.
\end{aligned}$$

Wiemy, że

$$\hat{\varphi}(\xi) = \lim_{N \rightarrow \infty} \hat{\varphi}_N(\xi), \quad \text{oraz} \quad \|\hat{\varphi}_N\|^2 = 1.$$

Skorzystamy teraz z drugiego podstawowego narzędzia teorii całki Lebesgue'a, czyli z lematu Fatou.

$$\int_{-\infty}^{\infty} |\hat{\varphi}(\xi)|^2 d\xi = \int_{-\infty}^{\infty} \lim_{N \rightarrow \infty} |\hat{\varphi}_N(\xi)|^2 d\xi \leq \lim_{N \rightarrow \infty} \int_{-\infty}^{\infty} |\hat{\varphi}_N(\xi)|^2 d\xi = 2\pi.$$

Widzimy więc, że $\hat{\varphi} \in L^2(\mathbf{R})$, a więc istotnie, zgodnie z jej oznaczeniem, jest transformatą Fouriera elementu $\varphi \in L^2(\mathbf{R})$. Pokażemy teraz, że

$$\hat{\varphi}(\xi) = \lim_{N \rightarrow \infty} \hat{\varphi}_N(\xi)$$

nie tylko w każdym punkcie, ale także w $L^2(\mathbf{R})$. W tym celu pokażemy, że istnieje stała c taka, że

$$|\hat{\varphi}_N(\xi)| \leq c|\hat{\varphi}(\xi)|. \quad (4.54)$$

Jeżeli $\xi \notin [-2^N\pi, 2^N\pi]$ to $\hat{\varphi}_N(\xi) = 0$, i (4.54) jest spełnione. Dla $\xi \in [-2^N\pi, 2^N\pi]$ mamy

$$\begin{aligned}
\hat{\varphi}(\xi) &= \prod_{j=1}^{\infty} m_0(2^{-j}\xi) \\
&= \prod_{j=1}^N m_0(2^{-j}\xi) \prod_{j=N+1}^{\infty} m_0(2^{-j}\xi) \\
&= \hat{\varphi}_N(\xi) \prod_{j=1}^{\infty} m_0(2^{-j}(2^{-N}\xi)) \\
&= \hat{\varphi}_N(\xi) \hat{\varphi}(2^{-N}\xi),
\end{aligned}$$

a więc

$$|\hat{\varphi}_N(\xi)| = \frac{|\hat{\varphi}(\xi)|}{|\hat{\varphi}(2^{-N}\xi)|}.$$

Ponieważ dla $\xi \in [-2^N\pi, 2^N\pi]$ mamy $2^{-N}\xi \in [-\pi, \pi]$, więc wystarczy pokazać, że

$$|\hat{\varphi}(\xi)| > \frac{1}{c} \quad \text{dla} \quad \xi \in [-\pi, \pi].$$

Niech więc $\xi \in [-\pi, \pi]$. Wcześniej pokazaliśmy, że dla każdego $\epsilon > 0$ istnieje $\delta > 0$, taka, że

$$\hat{\varphi}(\eta) > e^{-\epsilon\pi} \quad \text{dla} \quad |\eta| < \delta, \quad |\eta| < \pi.$$

Wyberzmy, na przykład, $\epsilon = 1$, i niech $\delta \leq \pi$ będzie zgodna z powyższym. Niech $L \in \mathbf{N}$ będzie wystarczająco duże, tak aby $|2^{-L}\xi| < \delta$. Wtedy dla $j = 1, 2, \dots, L-1$ mamy $2^{-j}\xi \in [-\pi/2, \pi/2]$, oraz

$$\begin{aligned} |\hat{\varphi}(\xi)| &= \prod_{j=1}^{\infty} |m_0(2^{-j}\xi)| \\ &= \prod_{j=1}^{L-1} |m_0(2^{-j}\xi)| \cdot |\hat{\varphi}(2^{-L}\xi)| \\ &\geq \prod_{j=1}^{L-1} K \cdot e^{-\pi} \\ &= K^{L-1} e^{-\pi}. \end{aligned}$$

Niech więc $c = e^{\pi}/K^{L-1}$. Mając (4.54) pokażemy $\varphi_N \rightarrow \varphi$ w $L^2(\mathbf{R})$.

$$\lim_{N \rightarrow \infty} \|\varphi_N - \varphi\|^2 = \lim_{N \rightarrow \infty} \frac{1}{2\pi} \int_{-\infty}^{\infty} |\hat{\varphi}_N(\xi) - \hat{\varphi}(\xi)|^2 d\xi = 0,$$

na mocy twierdzenia o zbieżności ograniczonej, gdyż

$$|\hat{\varphi}_N(\xi) - \hat{\varphi}(\xi)|^2 \leq (|\hat{\varphi}_N(\xi)| + |\hat{\varphi}(\xi)|)^2 \leq |\hat{\varphi}(\xi)|^2 (c+1)^2,$$

a funkcja $|\hat{\varphi}(\xi)|^2$ jest całkowalna. Pokażemy teraz, że funkcje φ_N spełniają (4.48). W tym celu wykonamy rachunek podobny do tego, w którym pokazaliśmy, że wszystkie φ_N mają tą samą normę 1. Niech

$$F_N(\xi) = \sum_{k=-\infty}^{\infty} |\hat{\varphi}_N(\xi + 2k\pi)|^2.$$

Najpierw pokażemy, że $\hat{F}_N(m) = \hat{F}_1(m)$, $m \in \mathbf{Z}$ niezależnie od N . Niech $N \geq 2$

$$\begin{aligned}
\hat{F}_N(m) &= \frac{1}{2\pi} \int_{-\pi}^{\pi} F_N(\xi) e^{-im\xi} d\xi \\
&= \frac{1}{2\pi} \int_{-\pi}^{\pi} \sum_{k=-\infty}^{\infty} |\hat{\varphi}_N(\xi + 2k\pi)|^2 e^{-im\xi} d\xi \\
&= \frac{1}{2\pi} \int_{-\infty}^{\infty} |\hat{\varphi}_N(\xi)|^2 e^{-im\xi} d\xi \\
&= \frac{1}{2\pi} \int_{-2^N\pi}^{2^N\pi} \prod_{j=1}^N |m_0(2^{-j}\xi)|^2 e^{-im\xi} d\xi \\
&= \frac{1}{2\pi} \int_0^{2^N\pi} \left(\prod_{j=1}^N |m_0(2^{-j}(\xi - 2^N\pi))|^2 + \prod_{j=1}^N |m_0(2^{-j}\xi)|^2 \right) e^{-im\xi} d\xi \\
&= \frac{1}{2\pi} \int_0^{2^N\pi} \prod_{j=1}^{N-1} |m_0(2^{-j}\xi)|^2 e^{-im\xi} d\xi \\
&= \frac{1}{2\pi} \int_{-2^{N-1}\pi}^{2^{N-1}\pi} \prod_{j=1}^{N-1} |m_0(2^{-j}\xi)|^2 e^{-im\xi} d\xi \\
&= \frac{1}{2\pi} \int_{-\infty}^{\infty} |\hat{\varphi}_{N-1}(\xi)|^2 e^{-im\xi} d\xi \\
&= \hat{F}_{N-1}(m).
\end{aligned}$$

Funkcje $F_N(\xi)$ mają więc identyczne współczynniki Fouriera, a ponieważ są całkowne na $\mathbf{T}$, więc muszą być sobie równe. Obliczymy współczynniki Fouriera $F_1(\xi)$. Podobnie jak powyżej

$$\begin{aligned}
\hat{F}_1(m) &= \frac{1}{2\pi} \int_{-2\pi}^{2\pi} |m_0(2^{-1}\xi)|^2 e^{-im\xi} d\xi \\
&= \frac{1}{2\pi} \int_0^{2\pi} (|m_0(2^{-1}(\xi - 2\pi))|^2 + |m_0(2^{-1}\xi)|^2) e^{-im\xi} d\xi \\
&= \frac{1}{2\pi} \int_0^{2\pi} e^{-im\xi} d\xi \\
&= \begin{cases} 1 & : \quad m = 0 \\ 0 & : \quad m \neq 0. \end{cases}
\end{aligned}$$

Dla każdego N mamy więc

$$F_N(\xi) \equiv 1,$$

czyli dla każdego N

$$\langle \varphi_N(\cdot - k), \varphi_N(\cdot - l) \rangle = \begin{cases} 1 & : k = l \\ 0 & : k \neq l. \end{cases}$$

Korzystając ze zbieżności $\varphi_N \rightarrow \varphi$ w $L^2(\mathbf{R})$, oraz ciągłości iloczynu skalarnego otrzymujemy (4.48).

Do udowodnienia pozostało jeszcze stwierdzenie o ograniczonym nośniku funkcji skalującej i falki, związanych z analizą wielorozdzielczą, w przypadku, gdy m_0 jest wielomianem trygonometrycznym

$$m_0(\xi) = \sum_{n=-L}^L h_n e^{-in\xi}. \quad (4.55)$$

Do udowodnienia tego ostatniego stwierdzenia najwygodniej jest skorzystać z języka teorii funkcji uogólnionych (dystrybucji). Dlatego tylko naszkicujemy ideę. Iloczyn częściowe zbiegają do $\hat{\varphi}$ w każdym punkcie ξ

$$\prod_{j=1}^N m_0(2^{-j}\xi) \rightarrow \hat{\varphi}(\xi),$$

ale nie w $L^2(\mathbf{R})$, gdyż nie są elementami $L^2(\mathbf{R})$. Są natomiast funkcjami uogólnionymi i jako takie zbiegają do $\hat{\varphi}$, która też jest funkcją uogólnioną. Korzystając z (4.55) widzimy, że iloczyny częściowe mają postać

$$\prod_{j=1}^N \left(\sum_{n=-L}^L h_n e^{-in2^{-j}\xi} \right) = \sum_{n_1, \dots, n_N = -L}^L h_{n_1} \dots h_{n_N} e^{-i(n_1 2^{-1} + \dots + n_N 2^{-N})\xi}.$$

Widzimy więc, że taki iloczyn częściowy jest kombinacją liniową transformat Fouriera impulsów Diraca w punktach $n_1 2^{-1} + \dots + n_N 2^{-N}$. Ponieważ wszystkie współczynniki $|n_i| \leq L$, więc wszystkie impulsy Diraca zlokalizowane są w przedziale $[-L, L]$. Widać więc, przynajmniej intuicyjnie, że funkcja φ , która jako funkcja uogólniona jest granicą funkcji uogólnionych o nośnikach zawartych w $[-L, L]$ też ma tę własność. Z kolei falka jest skończoną kombinacją liniową (gdyż filtr m_1 też jest wielomianem trygonometrycznym) przesunięć φ , a więc też ma nośnik ograniczony. $\square$

Uwaga 4.13. *Oszacowanie rozmiaru nośnika $\varphi(x)$ można zrobić dokładniej. Jeżeli*

$$m_0(\xi) = \sum_{n=M}^N h_n e^{-in\xi}, \quad M < N, \quad M, N \in \mathbf{Z}, \quad (4.56)$$

to nośnik $\varphi(x)$ zawiera się w przedziale $[M, N]$. Podobnie możemy oszacować nośnik falki $\psi(x)$ związanej konstruowaną analizą wielorozdzielczą. Jeżeli filtr dolnoprzepustowy jest postaci (4.56), to filtr górnoprzepustowy jest wielomianem trygonometrycznym postaci

$$m_1(\xi) = \sum_{n=1-N}^{1-M} g_n e^{-in\xi} \quad \text{oraz} \quad \hat{\psi}(\xi) = m_1(2^{-1}\xi) \prod_{j=2}^{\infty} m_0(2^{-j}\xi).$$

Przeprowadzając takie samo rozumowanie jak dla nośnika $\varphi(x)$ otrzymujemy, że nośnik $\psi(x)$ jest zawarty w przedziale $[1/2 - (N-M)/2, 1/2 + (N-M)/2]$.

Konstrukcja falek o nośniku ograniczonym sprowadziłyśmy więc do problemu znalezienia wielomianu trygonometrycznego m_0 spełniającego założenia Twierdzenia 4.12. Konstrukcja takiego wielomianu podzielimy na 2 kroki. W pierwszym znajdziemy wielomian trygonometryczny $g(\xi)$ spełniający

$$g(0) = 1, \quad g(\xi) \geq 0, \quad \text{oraz} \quad g(\xi) > 0 \text{ na } [-\pi/2, \pi/2], \quad (4.57)$$

$$g(\xi) + g(\xi + \pi) = 1. \quad (4.58)$$

W drugim kroku „wyciągniemy pierwiastek” z g , czyli znajdziemy wielomian taki, że

$$|m_0(\xi)|^2 = g(\xi).$$

Filtry Daubechies

Filtry Daubechies mogą mieć dowolne, parzyste długości. Ustalmy pewne $k = 0, 1, 2, \dots$. Niech

$$c_k = \int_0^\pi \sin(t)^{2k+1} dt,$$

oraz

$$g_k(\xi) = 1 - \frac{1}{c_k} \int_0^\xi \sin(t)^{2k+1} dt.$$

Pokażemy, że g_k są wielomianami trygonometrycznymi spełniającymi (4.57) i (4.58). Przede wszystkim $\sin(t)$ jest wielomianem trygonometrycznym

$$\sin(t) = \frac{e^{it} - e^{-it}}{2i},$$

po podniesieniu do potęgi mamy

$$\begin{aligned}\sin(t)^{2k+1} &= \frac{1}{(2i)^{2k+1}} \sum_{l=0}^{2k+1} \binom{2k+1}{l} e^{ilt} e^{-i(2k+1-l)t} \\ &= \frac{1}{(2i)^{2k+1}} \sum_{l=0}^{2k+1} \binom{2k+1}{l} e^{-i(2k+1-2l)t} \\ &== \sum_{\substack{l=-2k-1 \\ l-\text{nieparz.}}}^{2k+1} \alpha_l e^{ilt},\end{aligned}$$

a więc także wielomian trygonometryczny. Wielomian ten nie ma wyrazu wolnego, więc po scałkowaniu w dalszym ciągu jest wielomianem trygonometrycznym:

$$\begin{aligned}\int_0^\xi \sin(t)^{2k+1} dt &= \sum_{\substack{l=-2k-1 \\ l-\text{nieparz.}}}^{2k+1} \alpha_l \int_0^\xi e^{ilt} dt \\ &= \sum_{\substack{l=-2k-1 \\ l-\text{nieparz.}}}^{2k+1} \alpha_l \frac{(e^{i l \xi} - 1)}{i l} \\ &= \sum_{l=-2k-1}^{2k+1} \beta_l e^{i l t}.\end{aligned}$$

Widzimy więc, że $g_k(\xi)$ jest wielomianem trygonometrycznym stopnia $2k+1$. Zakres wartości g_k ustalimy wyznaczając ekstrema.

$$g'_k(\xi) = -\frac{1}{c_k} \frac{d}{d\xi} \int_0^\xi \sin(t)^{2k+1} dt = -\frac{1}{c_k} \sin(\xi)^{2k+1}.$$

Funkcja g_k rośnie więc na przedziałach postaci $[(2n-1)\pi, 2n\pi]$, a maleje na przedziałach postaci $[2n\pi, (2n+1)\pi]$. Ma więc lokalne ściśle maksima w punktach $2n\pi$ i lokalne, ściśle minima w punktach postaci $(2n+1)\pi$. Ponieważ jest 2π -okresowa, więc wystarczy sprawdzić wartości w punktach 0 i π .

$$g_k(0) = 1, \quad g_k(\pi) = 1 - \frac{1}{c_k} c_k = 0.$$

g_k ma więc maksimum 1 w punktach $2n\pi$ i minimum 0 w punktach $(2n+1)\pi$. W innych punktach $0 < g_k(\xi) < 1$, czyli g_k spełnia (??). Udowodnimy teraz (4.58), przy czym wystarczy ograniczyć się do $\xi > 0$, ze względu na

okresowość.

$$\begin{aligned}
g_k(\xi) + g_k(\xi + \pi) &= 1 - \frac{1}{c_k} \int_0^\xi \sin(t)^{2k+1} dt + 1 - \frac{1}{c_k} \int_0^{\xi+\pi} \sin(t)^{2k+1} dt \\
&= 2 - \frac{1}{c_k} \int_0^\xi \sin(t)^{2k+1} dt - \frac{1}{c_k} \int_0^{\xi+\pi} \sin(t)^{2k+1} dt.
\end{aligned} \tag{4.59}$$

Zauważmy, że

$$\int_0^{\xi+\pi} \sin(t)^{2k+1} dt = \int_{-\pi}^\xi \sin(t + \pi)^{2k+1} dt = - \int_{-\pi}^\xi \sin(t)^{2k+1} dt,$$

czyli, kontynuując (4.59) i korzystając z własności całki

$$\begin{aligned}
&= 2 - \frac{1}{c_k} \int_0^\xi \sin(t)^{2k+1} dt + \frac{1}{c_k} \int_{-\pi}^\xi \sin(t)^{2k+1} dt \\
&= 2 + \frac{1}{c_k} \int_{-\pi}^0 \sin(t)^{2k+1} dt \\
&= 2 - \frac{1}{c_k} \int_0^\pi \sin(t)^{2k+1} dt \\
&= 2 - 1 = 1.
\end{aligned}$$

Mamy więc (4.58), czyli dla każdego $k = 0, 1, 2, \dots$ skonstruowaliśmy odpowiednie g_k . Dla każdego k możemy łatwo wyliczyć współczynniki wielomianu g_k . Drugim krokiem w konstrukcji filtrów Daubechies jest następujące twierdzenie, które umożliwia „wyciągnięcie pierwiastka” z g_k .

Twierdzenie 4.14. *Niech $g(\xi)$ będzie wielomianem trygonometrycznym*

$$g(\xi) = \sum_{n=-M}^M \alpha_n e^{in\xi}, \tag{4.60}$$

takim, że $g(\xi) \geq 0$. Wtedy istnieje wielomian trygonometryczny $m_0(\xi)$

$$m_0(\xi) = \sum_{n=0}^M \beta_n e^{-in\xi}$$

taki, że $|m_0(\xi)|^2 = g(\xi)$.

Dowód. Dowód sprowadza się do tego, że pierwiastki g można pogrupować w pary, a następnie te pary można rozdzielić. Niech $P(z)$ będzie wielomianem zespolonym danym wzorem

$$P(z) = \sum_{k=0}^{2M} \alpha_{k-M} z^k,$$

czyli, dla z na okręgu jednostkowym, $z = e^{i\xi}$ mamy

$$P(e^{i\xi}) = (e^{i\xi})^M g(\xi).$$

$P(z)$ spełnia następujące równanie

$$\begin{aligned} z^{2M} \overline{P\left(\frac{1}{\bar{z}}\right)} &= z^{2M} \overline{\sum_{k=0}^{2M} \alpha_{k-M} \bar{z}^{-k}} \\ &= z^{2M} \sum_{k=0}^{2M} \bar{\alpha}_{k-M} z^{-k} \\ &= \sum_{k=0}^{2M} \bar{\alpha}_{k-M} z^{2M-k} \\ &= \sum_{k=0}^{2M} \bar{\alpha}_{M-k} z^k \\ &= \sum_{k=0}^{2M} \alpha_{k-M} z^k \\ &= P(z). \end{aligned}$$

Po drodze skorzystaliśmy z tego, że g ma wartości rzeczywiste, a więc $\alpha_{-k} = \bar{\alpha}_k$. Wielomiany $P(z)$ i $\overline{P(1/\bar{z})}$ mają więc te same pierwiastki (różne od zera), z tymi samymi krotnościami. Niech $r_1, \dots, r_l$ będą różnymi pierwiastkami P , leżącymi wewnątrz okręgu jednostkowego ($|r_i| < 1$, $i = 1, \dots, l$), o krotnościach $n_1, \dots, n_l$. Wtedy $1/\bar{r}_i$, $i = 1, \dots, l$ są różnymi pierwiastkami $\overline{P}$ leżącymi poza okręgiem jednostkowym, o krotnościach $n_1, \dots, n_l$.

$$r_1, \dots, r_l \quad \text{oraz} \quad 1/\bar{r}_1, \dots, 1/\bar{r}_l$$

są wszystkimi pierwiastkami P różnymi od 0 i nie leżącymi dokładnie na okręgu jednostkowym. Niech $s_1, \dots, s_p$ będą różnymi pierwiastkami P leżącymi dokładnie na okręgu jednostkowym, $|s_j| = 1$, $j = 1, \dots, p$. Można uzasadnić, że krotności pierwiastków s_j są liczbami parzystymi. W skrócie, dodajemy $\epsilon > 0$ do naszego wielomianu $g(\xi)$ i staje się on ściśle dodatni. Odpowiadający mu wielomian nie ma więc zer na kole jednostkowym, tylko poza nim, występujące w parach. Gdy $\epsilon \rightarrow 0$ pierwiastki, które przesuną się na okrąg jednostkowy zejdą się na nim parami, jeden przesunie się ze środka, drugi z zewnątrz. Można to uściślić. Krotności pierwiastków na okręgu jednostkowym oznaczmy więc przez $2k_1, \dots, 2k_p$. Wielomian $P(z)$ nie ma pierwiastka w 0. To jest równoważne warunkowi $\alpha_{-M} \neq 0$, czyli temu, że

M w (48) jest najmniejsze. Skoro $P(z)$ nie ma innych pierwiastków oprócz wypisanych powyżej, więc możemy rozłożyć go na czynniki

$$P(z) = A \prod_{i=1}^l (z - r_i)^{n_i} (z - 1/\bar{r}_i)^{n_i} \prod_{j=1}^p (z - s_j)^{2k_j}.$$

Zrobimy następujące przekształcenia

$$(z - s_j) = -zs_j \left(\frac{1}{z} - \bar{s}_j \right), \quad \left(\text{gdzie } \bar{s}_j = \frac{1}{s_j} \right),$$

czyli

$$(z - s_j)^{2k_j} = (-zs_j)^{k_j} (z - s_j)^{k_j} \left(\frac{1}{z} - \bar{s}_j \right)^{k_j}.$$

Podobnie

$$\left(z - \frac{1}{\bar{r}_i} \right) = \frac{-z}{\bar{r}_i} \left(\frac{1}{z} - \bar{r}_i \right),$$

czyli

$$\begin{aligned} P(z) &= A \prod_{i=1}^l \left(\frac{-z}{\bar{r}_i} \right)^{n_i} \prod_{i=1}^l (z - r_i)^{n_i} \left(\frac{1}{z} - \bar{r}_i \right)^{n_i} \prod_{j=1}^p (-zs_j)^{k_j} \prod_{j=1}^p (z - s_j)^{k_j} \left(\frac{1}{z} - \bar{s}_j \right)^{k_j} \\ &= B z^{n_1 + \dots + n_l + k_1 + \dots + k_p} \prod_{i=1}^l (z - r_i)^{n_i} \left(\frac{1}{z} - \bar{r}_i \right)^{n_i} \prod_{j=1}^p (z - s_j)^{k_j} \left(\frac{1}{z} - \bar{s}_j \right)^{k_j} \end{aligned}$$

gdzie $n_1 + \dots + n_l + k_1 + \dots + k_p$ jest połową sumy krotności wszystkich pierwiastków wielomianu $P(z)$, czyli połową stopnia $P(z)$

$$n_1 + \dots + n_l + k_1 + \dots + k_p = \frac{2M}{M} = M.$$

Mamy więc

$$P(z) = z^M B \prod_{i=1}^l (z - r_i)^{n_i} \left(\frac{1}{z} - \bar{r}_i \right)^{n_i} \prod_{j=1}^p (z - s_j)^{k_j} \left(\frac{1}{z} - \bar{s}_j \right)^{k_j}.$$

Dla z na okręgu jednostkowym $z = e^{i\xi}$ mamy więc

$$g(\xi) = B \prod_{i=1}^l (e^{i\xi} - r_i)^{n_i} \left(\frac{1}{e^{i\xi}} - \bar{r}_i \right)^{n_i} \prod_{j=1}^p (e^{i\xi} - s_j)^{k_j} \left(\frac{1}{e^{i\xi}} - \bar{s}_j \right)^{k_j},$$

ale $1/e^{i\xi} = e^{-i\xi} = \overline{e^{i\xi}}$, czyli

$$g(\xi) = B \prod_{i=1}^l |e^{i\xi} - r_i|^{2n_i} \prod_{j=1}^p |e^{i\xi} - s_j|^{2k_j}.$$

Ponieważ $g(\xi) \geq 0$, więc $B \geq 0$. Niech więc

$$m_0(\xi) = \sqrt{B} \prod_{i=1}^l (e^{-i\xi} - \bar{r}_i)^{n_i} \prod_{j=1}^p (e^{-i\xi} - \bar{s}_j)^{k_j}.$$

Zauważmy jeszcze na koniec, że znalezienie m_0 wymaga rozłożenia wielomianu $P(z)$ na czynniki liniowe. Często jest to robione numerycznie. $\square$

Uwaga 4.15. Dla $k = 0, 1, 2, \dots$ skonstruowaliśmy filtry $g_k(\xi)$, spełniające warunki powyższego twierdzenia z $M = 2k + 1$. Filtry $m_0(\xi)$ które otrzymujemy z powyższego twierdzenia są więc postaci

$$m_0(\xi) = \sum_{n=0}^{2k+1} h_n e^{-in\xi}. \quad (4.61)$$

Są to standardowe filtry Daubechies, które służą do konstrukcji standardowych falek Daubechies. Korzystając z (4.61) i uwagi po dowodzie Twierdzenia 4.12 mamy, że dla ustalonego k (odpowiada to filtrowi długości $2k + 2$) nośnik funkcji skalującej $\varphi(x)$ zawiera się w przedziale $[0, 2k + 1]$, a nośnik falki $\psi(x)$ zawiera się w przedziale $[-k, k + 1]$.

Przykłady: Policzmy współczynniki filtrów Daubechies długości 2 i 4. W przypadku filtrów Daubechies rozkład wielomianu $P(z)$ jest trochę łatwiejszy. Wielomiany $g_k(\xi)$ które skonstruowaliśmy mają pierwiastek rzędu $k + 2$ w π , a więc wielomian $P(z)$ ma pierwiastek rzędu co najmniej $k + 2$ w -1 . Niech $k = 0$.

$$\begin{aligned} c_0 &= \int_0^\pi \sin(t) dt = 2, \\ g_0(\xi) &= 1 - \frac{1}{2} \int_0^\xi \sin(t) dt = \frac{1}{2} + \frac{1}{2} \cos(\xi) = \frac{1}{4} e^{-i\xi} + \frac{1}{2} + \frac{1}{4} e^{i\xi} \\ P(z) &= \frac{1}{4} + \frac{1}{2} z + \frac{1}{4} z^2 = \frac{1}{4} (z + 1)^2. \end{aligned}$$

$P(z)$ ma 1 podwójny pierwiastek na okręgu jednostkowym, $z = -1$, czyli $l = 0$, $p = 1$, $k_1 = 2$. Otrzymujemy

$$m_0(\xi) = \frac{1}{2} (e^{-i\xi} + 1) = e^{-i\xi/2} \cos(\xi/2), \quad h_0 = h_1 = \frac{1}{2}, \quad h_k = 0, \quad k \neq 0, 1.$$

To jest filtr dolnoprzepustowy analizy Haara, który jest także najkrótszym z filtrów Daubechies.

Niech $k = 1$. Mamy

$$\begin{aligned} c_1 &= \int_0^\pi \sin^3(t) dt = \frac{4}{3}, \\ g_1(\xi) &= \frac{1}{2} + \frac{3}{4} \cos(\xi) - \frac{1}{4} \cos^3(\xi) = \frac{1}{2} + \frac{9}{16} \cos(\xi) - \frac{1}{16} \cos(3\xi), \\ P(z) &= -\frac{1}{32} + \frac{9}{32} z^2 + \frac{1}{2} z^3 + \frac{9}{32} z^4 - \frac{1}{32} z^6. \end{aligned}$$

Łatwo zauważyć, że $P(z)$ ma pierwiastki rzeczywiste, na okręgu jednostkowym jest pierwiastek -1 krotności 4, oraz poza okręgiem jednostkowym pierwiastki $2 + \sqrt{3}$ i $2 - \sqrt{3}$. W naszych oznaczeniach $l = 1$, $n_1 = 1$, $p = 1$, $k_1 = 4$. Otrzymujemy następujący wzór

$$m_0(\xi) = \frac{1}{\sqrt{32(2 + \sqrt{3})}} \left(-(2 + \sqrt{3}) - (3 + 2\sqrt{3}) e^{-i\xi} - \sqrt{3} e^{-i2\xi} + e^{-i3\xi} \right).$$

W literaturze współczynniki filtrów podaje się najczęściej pomnożone przez $\sqrt{2}$. Jak zobaczymy w następnym rozdziale w ten sposób oszczędza się mnożenie przez ten pierwiastek a algorytmie obliczeniowym.

Funkcji skalujących i falek odpowiadających filtrom Daubechies nie da się zapisać żadnym rozsądnym wzorem. Natomiast istnieje metoda numerycznego generowania ich przybliżonych wykresów. Jest to tak zwany algorytm kaskadowy, który opiszemy w następnym rozdziale.

Falki w wymiarze $n > 1$

W $\mathbf{R}^n$ baza falkowa ma postać

$$\left\{ 2^{\frac{jn}{2}} \psi^l(2^j x - \mathbf{n}); \mathbf{n} \in \mathbf{Z}^n, j \in \mathbf{Z}, l = 1, \dots, 2^n - 1 \right\} \quad (368)$$

Na przykład, w wymiarze $n = 2$ baza falkowa generowana jest przez 3 falki. Dla $\mathbf{R}^n$ można zdefiniować pojęcie analizy wielorozdzielczej MRA i przeprowadzić konstrukcję falek w sposób zupełnie analogiczny do konstrukcji w przypadku $n = 1$. Postępując tak zauważylibyśmy, że w sposób naturalny pojawia się $2^n - 1$ filtrów górnoprzepustowych, przy minimalnych założeniach o regularności funkcji skalującej. Analizę wielorozdzielczą i związane z nią falki w wielu wymiarach można skonstruować używając analizy jednowymiarowej. Mając analizę w $L^2(\mathbf{R})$ z funkcją skalującą $\varphi(x)$ i falką $\psi(x)$ możemy

zdefiniować funkcję skalującą $\Phi(x, y)$ i falki $\Psi^l(x, y)$, $l = 1, 2, 3$ w $L^2(\mathbf{R}^2)$ wzorami

$$\Phi(x, y) = \varphi(x) \varphi(y), \Psi^1(x, y) = \varphi(x) \psi(y), \Psi^2(x, y) = \psi(x) \varphi(y), \Psi^3(x, y) = \psi(x) \psi(y). \quad (4.62)$$

Sformułowanie definicji MRA w $L^2(\mathbf{R}^2)$ i sprawdzenie, że funkcje (4.62) spełniają wszystkie warunki tej definicji zostawiamy jako ćwiczenie.

Uwagi 4.16. (i) w (??) zastosowaliśmy zamianę skali w postaci macierzy

$$A = \begin{pmatrix} 2 & 0 & \cdots & 0 \\ 0 & 2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & 2 \end{pmatrix}.$$

Bazy falkowe można generować również używając innych macierzy A , w których, na przykład, rozciągnięcie połączone jest z jakimś obrotem.

(ii) Bazy falkowe o zmiennych rozdzielonych takie, jak we wzorze (4.62) to tylko jeden szczególny rodzaj baz falkowych w wielu wymiarach. Często stosuje się bazy w których poszczególne falki mają szczególny kształt. Takie falki są użyteczne w wychwytywaniu w sygnale elementów o tego typu kształcie.

Rozdział 5

Algorytmy Numeryczne

Algorytm Mallata

Algorytm Mallata jest podstawą dyskretnej transformaty falkowej (dyskretnej, to znaczy transformaty sygnałów będących ciągami próbek). Algorytm pozwala w sposób efektywny numerycznie obliczać współczynniki bazowe sygnału w jednej bazie na podstawie współczynników w innej bazie. Opiszemy teraz o jakie bazy chodzi.

Przypomnijmy, że dla danej analizy MRA i dowolnego $j \in \mathbf{Z}$ przestrzeń V_j rozkłada się jako ortogonalna suma prosta podprzestrzeni

$$V_j = V_{j-1} \oplus W_{j-1}.$$

Zauważmy, że zgodnie z tym, co wiemy o analizie wielorozdzielczej, w przestrzeni V_j mamy bazę ortonormalną

$$\{\varphi_{j,n}(x); n \in \mathbf{Z}\}, \quad (5.1)$$

gdzie φ jest funkcją skalującą analizy. Dla $j = 0$ jest to część definicji analizy wielorozdzielczej, dla pozostałych $j \in \mathbf{Z}$ wynika z zamiany zmiennych $x \mapsto 2^j x$. Ponieważ j jest dowolne, to także w przestrzeni V_{j-1} mamy bazę o. n.

$$\{\varphi_{j-1,n}(x); n \in \mathbf{Z}\}.$$

W przestrzeni W_{j-1} mamy z kolei bazę o. n.

$$\{\psi_{j-1,n}(x); n \in \mathbf{Z}\},$$

gdzie ψ jest falką związaną z naszą analizą wielorozdzielczą. Przypomnijmy, że używamy następujących oznaczeń

$$\varphi_{j,n}(x) = 2^{j/2} \varphi(2^j x - n), \quad \psi_{j,n}(x) = 2^{j/2} \psi(2^j x - n).$$

W przestrzeni V_j mamy więc do dyspozycji dwie różne bazy o. n., bazę (5.1), oraz bazę

$$\{\varphi_{j-1,n}(x), \psi_{j-1,n}(x); n \in \mathbf{Z}\}. \quad (5.2)$$

Algorytm Mallata pozwala obliczać współczynniki sygnału w jednej z tych baz przy pomocy współczynników w drugiej.

Niech $f \in L^2(\mathbf{R})$, i wprowadźmy następujące oznaczenia na ciągi współczynników f w różnych bazach. Niech $j \in \mathbf{Z}$ będzie dowolne, i niech

$$a_j(n) = \langle f, \varphi_{j,n} \rangle, \quad d_j(n) = \langle f, \psi_{j,n} \rangle, \quad n \in \mathbf{Z}. \quad (5.3)$$

Przypomnijmy, że $\{h_n\}_{n=-\infty}^{\infty}$ i $\{g_n\}_{n=-\infty}^{\infty}$ oznaczają filtry dolno- i górnoprzepustowy analizy wielorozdzielczej. Następujące twierdzenie przedstawia algorytm.

Twierdzenie 5.1 (Mallat). *Transformata:*

$$\begin{aligned} a_{j-1}(k) &= \sqrt{2} \sum_{n=-\infty}^{\infty} \overline{h_{n-2k}} a_j(n), \\ d_{j-1}(k) &= \sqrt{2} \sum_{n=-\infty}^{\infty} \overline{g_{n-2k}} a_j(n). \end{aligned} \quad (5.4)$$

Transformata odwrotna:

$$a_j(n) = \sqrt{2} \sum_{k=-\infty}^{\infty} h_{n-2k} a_{j-1}(k) + \sqrt{2} \sum_{k=-\infty}^{\infty} g_{n-2k} d_{j-1}(k). \quad (5.5)$$

Uwaga 5.2. (a) Zauważmy, że algorytm odwołuje się tylko do filtrów dolno- i górnoprzepustowego. Same wartości funkcji φ i ψ nigdzie nie są potrzebne. Zauważmy też, że wzory powyższe są tym efektywniejsze numerycznie, im krótsze są filtry, a jeżeli są nieskończone to im szybciej maleją.

(b) Zwróćmy uwagę na współczynniki $\sqrt{2}$ występujące w powyższych wzorach. W praktyce po prostu filtry $\{h_n\}$ i $\{g_n\}$ normalizuje się w ten sposób, że $\sqrt{2}$ jest już w nich zawarte. Oszczędza się w ten sposób ciągłego mnożenia przez ten współczynnik. Filtry Daubechies generowane przez polecenie `daubcwf()` pakietu *RWT* są właśnie tak znormalizowane, na przykład filtr Haara to $(\frac{1}{\sqrt{2}}, \frac{1}{\sqrt{2}})$, a nie, tak jak na naszym wykładzie $(\frac{1}{2}, \frac{1}{2})$.

Dowód (twierdzenia Mallata). Przypomnijmy wzory

$$\begin{aligned} h_n &= \left\langle \frac{1}{2} \varphi \left(\frac{\cdot}{2} \right), \varphi(\cdot - n) \right\rangle, \\ g_n &= \left\langle \frac{1}{2} \psi \left(\frac{\cdot}{2} \right), \varphi(\cdot - n) \right\rangle. \end{aligned} \quad (5.6)$$

Liczba całkowita n jest dowolna, więc w pierwszym ze wzorów (5.6) wpiszmy w jej miejsce $n - 2l$, gdzie $n, l \in \mathbf{Z}$.

$$\begin{aligned}
h_{n-2l} &= \left\langle \frac{1}{2} \varphi\left(\frac{\cdot}{2}\right), \varphi(\cdot - (n - 2l)) \right\rangle \\
&= \int_{-\infty}^{\infty} \frac{1}{2} \varphi\left(\frac{x}{2}\right) \overline{\varphi(x - n + 2l)} dx \\
&= \int_{-\infty}^{\infty} \frac{1}{2} \varphi\left(\frac{x - 2l}{2}\right) \overline{\varphi(x - n)} dx \\
&= \int_{-\infty}^{\infty} \frac{1}{2} \varphi\left(\frac{x}{2} - l\right) \overline{\varphi(x - n)} dx.
\end{aligned}$$

Dokonujemy zamiany zmiennych $x \mapsto 2^j x$. Granice całkowania nie zmieniają się

$$\begin{aligned}
h_{n-2l} &= \int_{-\infty}^{\infty} \frac{1}{2} \varphi(2^{j-1}x - l) \overline{\varphi(2^j x - n)} 2^j dx \\
&= \int_{-\infty}^{\infty} \frac{1}{\sqrt{2}} 2^{\frac{j-1}{2}} \varphi(2^{j-1}x - l) 2^{\frac{j}{2}} \overline{\varphi(2^j x - n)} dx \\
&= \int_{-\infty}^{\infty} \frac{1}{\sqrt{2}} \varphi_{j-1,l}(x) \overline{\varphi_{j,n}(x)} dx \\
&= \frac{1}{\sqrt{2}} \langle \varphi_{j-1,l}, \varphi_{j,n} \rangle.
\end{aligned}$$

Podobnie wyprowadzamy wzór dla g_{n-2l} i otrzymujemy

$$\begin{aligned}
\sqrt{2} h_{n-2l} &= \langle \varphi_{j-1,l}, \varphi_{j,n} \rangle, \\
\sqrt{2} g_{n-2l} &= \langle \psi_{j-1,l}, \varphi_{j,n} \rangle.
\end{aligned} \tag{5.7}$$

Obliczyliśmy współczynniki bazowe funkcji $\varphi_{j-1,l}$ i $\psi_{j-1,l}$ w bazie (5.1), więc funkcje te mają rozwinięcia

$$\begin{aligned}
\varphi_{j-1,l} &= \sqrt{2} \sum_{n=-\infty}^{\infty} h_{n-2l} \varphi_{j,n}, \\
\psi_{j-1,l} &= \sqrt{2} \sum_{n=-\infty}^{\infty} g_{n-2l} \varphi_{j,n}.
\end{aligned} \tag{5.8}$$

Wzory te wstawiamy do definicji naszych ciągów

$$\begin{aligned}
a_{j-1}(k) &= \langle f, \varphi_{j-1,k} \rangle \\
&= \left\langle f, \sqrt{2} \sum_{n=-\infty}^{\infty} h_{n-2k} \varphi_{j,n} \right\rangle \\
&= \sqrt{2} \sum_{n=-\infty}^{\infty} \overline{h_{n-2k}} \langle f, \varphi_{j,n} \rangle \\
&= \sqrt{2} \sum_{n=-\infty}^{\infty} \overline{h_{n-2k}} a_j(n).
\end{aligned}$$

Podobnie obliczamy $d_{j-1}(k)$. W ten sposób pokazaliśmy pierwszą część twierdzenia.

Wzory (5.7) możemy też potraktować jako wzory na współczynniki bazowe funkcji $\varphi_{j,n} \in V_j$ w bazie (5.2). W takim razie

$$\begin{aligned}
\varphi_{j,n} &= \sum_{k=-\infty}^{\infty} \langle \varphi_{j,n}, \varphi_{j-1,k} \rangle \varphi_{j-1,k} + \sum_{k=-\infty}^{\infty} \langle \varphi_{j,n}, \psi_{j-1,k} \rangle \psi_{j-1,k} \\
&= \sqrt{2} \sum_{k=-\infty}^{\infty} \overline{h_{n-2k}} \varphi_{j-1,k} + \sqrt{2} \sum_{k=-\infty}^{\infty} \overline{g_{n-2k}} \psi_{j-1,k}.
\end{aligned} \tag{5.9}$$

Możemy więc obliczyć współczynniki bazowe f

$$\begin{aligned}
a_j(n) &= \langle f, \varphi_{j,n} \rangle \\
&= \left\langle f, \sqrt{2} \sum_{k=-\infty}^{\infty} \overline{h_{n-2k}} \varphi_{j-1,k} + \sqrt{2} \sum_{k=-\infty}^{\infty} \overline{g_{n-2k}} \psi_{j-1,k} \right\rangle \\
&= \sqrt{2} \sum_{k=-\infty}^{\infty} h_{n-2k} \langle f, \varphi_{j-1,k} \rangle + \sqrt{2} \sum_{k=-\infty}^{\infty} g_{n-2k} \langle f, \psi_{j-1,k} \rangle \\
&= \sqrt{2} \sum_{k=-\infty}^{\infty} h_{n-2k} a_{j-1}(k) + \sqrt{2} \sum_{k=-\infty}^{\infty} g_{n-2k} d_{j-1}(k),
\end{aligned}$$

co kończy dowód. □

Przyjrzyjmy się wzorom. Niech ciągi $\{\tilde{h}_n\}$ i $\{\tilde{g}_n\}$ mają współczynniki

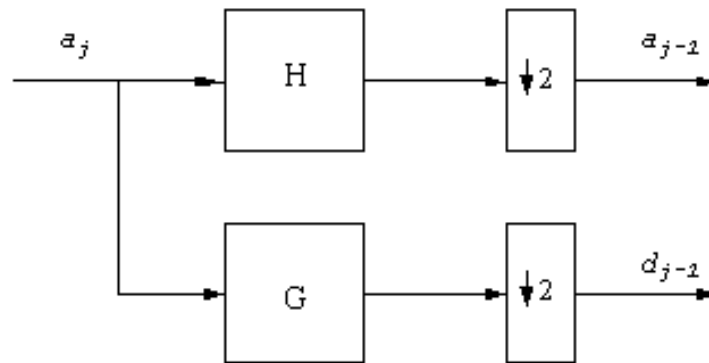
$$\tilde{h}_n = \overline{h_{-n}}, \quad \tilde{g}_n = \overline{g_{-n}}.$$

Wtedy, zgodnie z twierdzeniem Mallata

$$a_{j-1}(k) = \sqrt{2} \sum_{n=-\infty}^{\infty} \tilde{h}_{2k-n} a_j(n) = \sqrt{2} (\tilde{h} * a_j)(2k),$$

$$d_{j-1}(k) = \sqrt{2} \sum_{n=-\infty}^{\infty} \tilde{g}_{2k-n} a_j(n) = \sqrt{2} (\tilde{g} * a_j)(2k).$$

Obliczenie ciągów a_{j-1} i d_{j-1} sprowadza się więc do splotu wyjściowego ciągu a_j z ciągami $\tilde{h}$ i $\tilde{g}$, a następnie wybrania co drugiego elementu powstałych ciągów (i pomnożenia przez $\sqrt{2}$). Odrzucenie co drugiego wyrazu po angielsku nazywa się *downsampling*. W teorii przetwarzania sygnału operacje na sygnale często opisuje się przy pomocy schematów blokowych. Naszą operację możemy przedstawić schematycznie następująco (operatory H i G to sploty z filtrami $\{\tilde{h}_n\}$ i $\{\tilde{g}_n\}$, oraz pomnożenie przez $\frac{1}{\sqrt{2}}$)



Rysunek 5.1: Jeden krok dyskretnej transformaty falkowej.

Podobnie przyjrzyjmy się drugiej części twierdzenia Mallata. Wprowadźmy następujące oznaczenie. Jeżeli dany jest ciąg $\{\alpha_n\}$, to przez $\{\alpha'_n\}$ oznaczamy jego „rozrzedzenie”:

$$\alpha'_n = \begin{cases} \alpha_{n/2} & : n - \text{parzyste} \\ 0 & : n - \text{nieparzyste.} \end{cases}$$

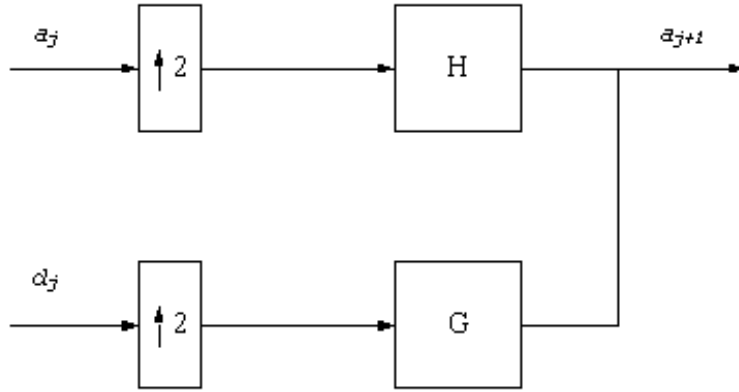
Ciąg $\{\alpha_n\}$ został więc „przepleciony” zerami:

$$\dots, \alpha_{-1}, \alpha_0, \alpha_1, \dots \quad \Rightarrow \quad \dots, \alpha_{-1}, 0, \alpha_0, 0, \alpha_1, 0, \alpha_2, \dots$$

Wzór (5.5) można wtedy zapisać

$$\begin{aligned}
a_j(n) &= \sqrt{2} \sum_{k=-\infty}^{\infty} h_{n-2k} a_{j-1}(k) + \sqrt{2} \sum_{k=-\infty}^{\infty} g_{n-2k} d_{j-1}(k) \\
&= \sqrt{2} \sum_{\substack{k=-\infty \\ k\text{-parzyste}}}^{\infty} h_{n-k} a_{j-1}(k/2) + \sqrt{2} \sum_{\substack{k=-\infty \\ k\text{-parzyste}}}^{\infty} g_{n-k} d_{j-1}(k/2) \\
&= \sqrt{2} \sum_{k=-\infty}^{\infty} h_{n-k} a'_{j-1}(k) + \sqrt{2} \sum_{k=-\infty}^{\infty} g_{n-k} d'_{j-1}(k) \\
&= \sqrt{2} (h * a'_{j-1})(k) + \sqrt{2} (g * d'_{j-1})(k).
\end{aligned}$$

Zauważmy więc, że 1 krok odwrotnej transformaty falkowej sprowadza się najpierw do rozrzedzenia sygnałów wejściowych a_{j-1} i d_{j-1} (po angielsku *upsampling*), a następnie splotu każdego sygnału z odpowiednim ciągiem, i w końcu złożenia (sumy). Sploty są w zasadzie z tymi samymi ciągami co w transformacie. Odwrotną transformatę możemy przedstawić w postaci schematu blokowego (tutaj operatory H i G to sploty z filtrami $\{h_n\}$ i $\{g_n\}$, oraz pomnożenie przez $\frac{1}{\sqrt{2}}$)



Rysunek 5.2: Jeden krok odwrotnej dyskretnej transformaty falkowej.

Sygnały dyskretne

Sygnały dyskretne to ciągi wartości $\{f_n\}$, skończone lub nie (tymczasem rozważamy sygnały 1-wymiarowe). Najczęściej taki ciąg wartości to ciąg próbek jakiegoś sygnału ciągłego, na przykład dźwięku. Analiza falkowa takiego sygnału polega wyborze analizy wielorozdzielczej (czyli wyborze falki),

następnie umieszczeniu tego sygnału w przestrzeni V_0 , a następnie obliczeniu współczynników bazowych w bazie falkowej (ciągów $\{d_j\}$). Numerycznie cały ten proces sprowadza się do prostego i szybkiego algorytmu rekurencyjnego, w którym wartości funkcji skalującej ani falki w ogóle się nie pojawiają. Jedynym elementem analizy wielorozdzielczej występującym w obliczeniach są filtry dolno- i górnoprzepustowy. Algorytm opiera się na twierdzeniu Mallata i nazywa się algorytmem Mallata. Niech naszym sygnałem będzie ciąg $\{f_n\}$. Wybrawszy analizę wielorozdzielczą przyporządkujemy sygnałowi funkcję $f \in V_0$

$$f(x) = \sum_{n=-\infty}^{\infty} f_n \varphi(x - n). \quad (5.10)$$

Takie przyporządkowanie jest bardzo naturalne. Najczęściej wartości f_n są próbkami jakiejś ciągłej wartości fizycznej, pobieranymi w regularnych odstępach czasu. Czujnik pobierający próbkę z reguły pobierając ją dokonuje uśrednienia wartości w jakimś przedziale czasu. Cały proces pobrania próbek i przyporządkowania im elementu $f \in V_0$ jest więc rzutem wyjściowego sygnału ciągłego na V_0 . W sumie bardzo naturalna i łagodna operacja, jeżeli analizę wielorozdzielczą dobierzemy właściwie do charakteru badanego sygnału.

Dyskretna transformata falkowa dyskretnego sygnału $\{f_n\}$ to zbiór współczynników bazowych

$$\{d_j(n) = \langle f, \psi_{j,n} \rangle; j, n \in \mathbf{Z}, j \leq -1\}. \quad (5.11)$$

Widzimy więc, że obliczenie transformaty sprowadza się do rekurencyjnego stosowania twierdzenia Mallata, z $a_0 = \{f_n\}$. Sygnał można odtworzyć z jego transformaty falkowej (5.11) stosując rekurencyjnie drugą część twierdzenia Mallata. W praktyce sygnał analizowany jest zawsze skończony i algorytm Mallata ma skończoną liczbę kroków.

Sygnały skończone

W przypadku sygnału skończonego składającego się z N próbek

$$\{f_0, f_1, \dots, f_{N-1}\}$$

algorytm wygląda następująco. Algorytm ma zastosowanie do sygnałów, których długość jest potęgą 2, więc najpierw ewentualnie wydłużamy sygnał, poprzez dodanie zer, tak, aby jego długość była potęgą 2, niech $N = 2^J$. Następnie określamy wyjściowy ciąg a_0 , jako N -periodyzację sygnału:

$$a_0(n) = f_n, \quad n = 0, \dots, N - 1,$$

oraz

$$a_0(n + kN) = a_0(n), \quad k \in \mathbf{Z}, \quad n = 0, \dots, N - 1.$$

Łatwo zauważyć, że zastosowanie 1 kroku transformaty falkowej do nam 2 ciągi

$$\{a_{-1}(n)\}_{n=-\infty}^{\infty} \quad \text{i} \quad \{d_{-1}(n)\}_{n=-\infty}^{\infty},$$

które są okresowe o okresach dwukrotnie krótszych

$$a_{-1}(n + N/2) = a_{-1}(n), \quad d_{-1}(n + N/2) = d_{-1}(n).$$

Wynika to wprost ze wzorów, na przykład

$$\begin{aligned} a_{-1}(n + N/2) &= \sqrt{2} \sum_{k=-\infty}^{\infty} \overline{h_{k-2(n+N/2)}} a_0(k) \\ &= \sqrt{2} \sum_{k=-\infty}^{\infty} \overline{h_{k-2n-N}} a_0(k) \\ &= \sqrt{2} \sum_{k=-\infty}^{\infty} \overline{h_{k-2n}} a_0(k + N) \\ &= \sqrt{2} \sum_{k=-\infty}^{\infty} \overline{h_{k-2n}} a_0(k) \\ &= a_{-1}(n). \end{aligned}$$

Ze względu na tę okresowość sygnały a_{-1} i d_{-1} są o połowę krótsze, bo wystarczy zapamiętywać tylko jeden okres każdego ciągu. Iterując procedurę, czyli stosując twierdzenie Mallata kolejno do ciągów a_j okresowych o coraz krótszych okresach, po J krokach uzyskujemy kompletną transformatę falkową wyjściowego sygnału $\{f_0, \dots, f_{N-1}\}$:

$$a_{-J}(0), d_{-J}(0), d_{-J+1}(0), d_{-J+1}(1), \dots, d_{-1}(0), \dots, d_{-1}(N/2 - 1).$$

Odwrotną transformatę obliczamy podobnie. Każdy ze skończonych ciągów d_j przedłużamy okresowo do ciągu nieskończonego, następnie stosujemy J razy twierdzenie Mallata. W każdym kroku rekonstruujemy ciąg a_j , o coraz dłuższym okresie.

Algorytm kaskadowy

Przypomnijmy, że analizy wielorozdzielcze Daubechies, a, co za tym idzie, falki Daubechies skonstruowaliśmy konstruując odpowiednie filtry dolnoprzepustowe m_0 . Dzięki własnościom tych konkretnych filtrów mogliśmy udowodnić ważne własności funkcji skalujących i falek, takie jak na przykład

ograniczony nośnik (to znaczy, że funkcje znikają poza skończonym przedziałem), oraz różniczkowalność. Wspominaliśmy, że funkcje skalujących i falek Daubechies nie da się zapisać żadnym jawnym wzorem (z wyjątkiem przypadku filtru długości 2, dla którego funkcje te są funkcjami Haara). Nie ma to znaczenia dla algorytmu Mallata, w którym występują tylko współczynniki filtrów. Bardzo łatwo można jednak obliczyć przybliżone wartości funkcji skalującej i falek, dla dowolnej analizy wielorozdzielczej, dla której te funkcje są wystarczająco regularne (czyli różniczkowalne odpowiednią ilość razy). Dzięki temu, na przykład, można narysować wykresy tych funkcji. Do przybliżonego obliczenia wartości służy tak zwany algorytm kaskadowy, oparty na algorytmie Mallata.

Zilustrujemy teraz algorytm kaskadowy, rysując wykresy funkcji skalujących i falek Daubechies. Funkcja skalująca Daubechies jest ciągła, ma ograniczony nośnik, oraz

$$1 = \hat{\varphi}(0) = \int_{-\infty}^{\infty} \varphi(x) dx, \quad \int_{-\infty}^{\infty} |\varphi(x)| dx < \infty$$

(ograniczmy się do przypadku filtrów dłuższych niż 2). Mamy następujący fakt

Fakt 5.3. *Jeżeli funkcja $f(x)$ jest ciągła, to dla każdego $x \in \mathbf{R}$*

$$\int_{-\infty}^{\infty} f(x+y) 2^j \overline{\varphi(2^j y)} dy \xrightarrow{j \rightarrow \infty} f(x).$$

Jeżeli $f(x)$ jest ciągła jednostajnie, to zbieżność jest jednostajna, a jeżeli $f(x)$ jest ciągła w sensie Höldera, czyli

$$|f(x) - f(y)| \leq C|x - y|^\alpha, \quad \text{dla jakiegoś } 0 < \alpha \leq 1,$$

to zbieżność jest wykładnicza

$$\left| f(x) - \int_{-\infty}^{\infty} f(x+y) 2^j \overline{\varphi(2^j y)} dy \right| \leq C' 2^{-j\alpha}.$$

Dowód. Niech $[-R, R]$ będzie przedziałem, poza którym $\varphi(x) \equiv 0$. Na przykład, z dowodu twierdzenia ?? z rozdziału o analizie MRA wynika, że dla filtru o długości $2k$ mamy $R = 2k - 1$. Ustalmy x oraz $\epsilon > 0$ i niech $N \in \mathbf{N}$ będzie takie, że

$$\forall j \geq N \quad |y| < 2^{-j} R \Rightarrow |f(x) - f(x+y)| < \epsilon.$$

Dla $j \geq N$ mamy wtedy oszacowanie

$$\begin{aligned}
\left| f(x) - \int_{-\infty}^{\infty} f(x+y) 2^j \overline{\varphi(2^j y)} dy \right| &= \left| \int_{-\infty}^{\infty} (f(x) - f(x+y)) 2^j \overline{\varphi(2^j y)} dy \right| \\
&\leq \int_{-\infty}^{\infty} |f(x) - f(x+y)| 2^j |\varphi(2^j y)| dy \\
&= \int_{-\infty}^{\infty} |f(x) - f(x+2^{-j} y)| |\varphi(y)| dy \\
&= \int_{-R}^R |f(x) - f(x+2^{-j} y)| |\varphi(y)| dy \\
&\leq \epsilon \int_{-R}^R |\varphi(y)| dy \\
&\leq C\epsilon,
\end{aligned}$$

gdzie

$$C = \int_{-\infty}^{\infty} |\varphi(y)| dy.$$

Zauważmy, że wszystkie tezy faktu wynikają z powyższego oszacowania. $\square$

Funkcja skalująca i falka Daubechies są różniczkowalne, a więc spełniają warunek Höldera z $\alpha = 1$. Zbieżność w powyższym fakcie jest więc dla nich wykładnicza, czyli szybka. Niech x będzie liczbą diadyczną, czyli liczbą postaci $x = 2^{-J} n$, $J \in \mathbf{Z}$. Wtedy

$$\begin{aligned}
\varphi(x) &= \lim_{j \rightarrow \infty} 2^j \int_{-\infty}^{\infty} \varphi(2^{-J} n + y) \overline{\varphi(2^j y)} dy \\
&= \lim_{j \rightarrow \infty} 2^{j/2} \int_{-\infty}^{\infty} \varphi(y) 2^{j/2} \overline{\varphi(2^j y - 2^{-J+j} n)} dy \\
&= \lim_{j \rightarrow \infty} 2^{j/2} \langle \varphi, \varphi_{j, 2^{-J+j} n} \rangle,
\end{aligned}$$

czyli, w punktach diadycznych $2^{-J} n$, dla odpowiednio dużego $j \geq J$

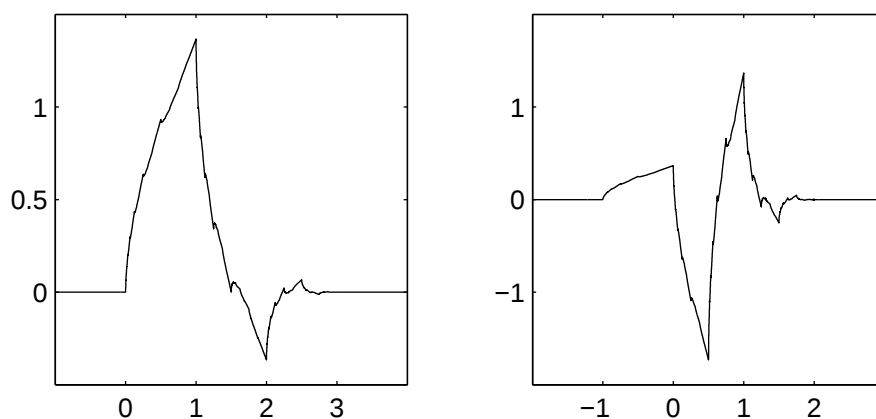
$$\varphi(2^{-J} n) \simeq 2^{j/2} \langle \varphi, \varphi_{j, 2^{-J+j} n} \rangle = 2^{j/2} a_j(2^{j-J} n), \quad (5.12)$$

gdzie ciąg a_j z twierdzenia Mallata jest obliczony dla $f = \varphi$. Ciąg ten możemy obliczyć dla dowolnego $j > 0$ korzystając z drugiej części twierdzenia Mallata. Wiemy, że dla $f = \varphi$ ciąg a_0 składa się z 1 w zerze i poza tym samych zer, a ciągi d_j składają się z samych zer, dla $j \geq 0$. Aby narysować wykres $\varphi(x)$ postępujemy więc następująco. Wybieramy $J \in \mathbf{Z}$ odpowiednio duże, w zależności od potrzebnej dokładności. Wartość J decyduje o tym, jak gęsta

jest siatka liczb postaci $2^{-J}n$, ale także o dokładności przybliżenia (5.12), w którym weźmiemy $j = J$. Następnie stosujemy twierdzenie Mallata J razy. Podobnie postępujemy w celu narysowania wykresu falek. Wychodzimy od ciągów

$$a_0(k) \equiv 0, \quad d_0(k) = \delta_0(k), \quad (\text{jedynka w } 0 \text{ i same } 0 \text{ poza tym}),$$

oraz $d_j(k) \equiv 0$ dla $j \geq 1$. Następnie stosujemy twierdzenie Mallata J razy. Na obrazkach pokazujemy kilka przykładów. W każdym przypadku stosowaliśmy $J = 15$.



Rysunek 5.3: Funkcja skalująca i falka Daubechies, filtry długości 4.

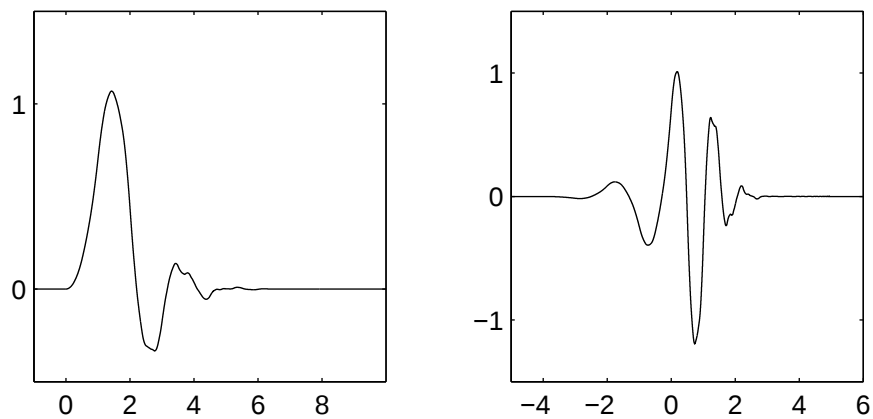
Wykresy funkcji skalujących zostały wygenerowane przy pomocy następującego skryptu w Matlabie. Należy wybrać N w 1 linijce (długość filtra), oraz J w linijce 5 (dokładność).

```
N=20;%długość filtra, musi być parzysta
A=-1;%lewy zakres wykresu funkcji skalującej, musi być ujemny
B=N;%prawy zakres wykresu funkcji skalującej, musi być > N-1
h=daubcwf(N);
J=15;
a=2^J*(N-1);
dx=2^(-J);
X=[floor(A):dx:floor(B)];
phi1=zeros(1,size(X,2));
offset=-floor(A)/dx;
phi=zeros(1,a+1);
temp=zeros(1,a+1);
```

```

phi(1)=1;
for j=1:J
    temp=phi;
    for n=0:a
        c=0;
        par=n-2*floor(n/2);
        for l=par:2:N-1
            ind=(n-l)/2;
            if (ind>=0)\&(ind<=a)
                c=c+h(l+1)*temp(ind+1);
            end;
        end;
        phi(n+1)=c;
    end;
end;
for i=1:a+1
    phi1(offset+i)=phi(i);
end;
phi1=sqrt(2^J)*phi1;
plot(X,phi1,'black');

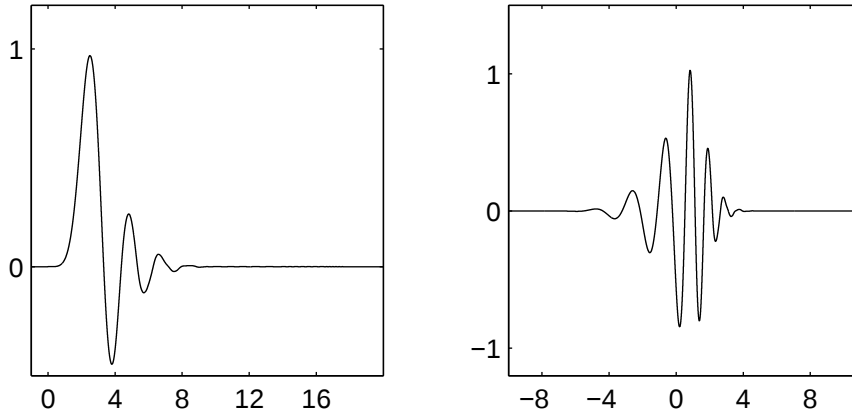
```



Rysunek 5.4: Funkcja skalująca i falka Daubechies, filtry długości 10.

Próbkowanie i kwantyzacja

Sygnały dyskretne (czyli skończone lub nieskończone ciągi wartości) w praktyce powstają jako wynik próbkowania sygnałów ciągłych. Naturalne pyta-



Rysunek 5.5: Funkcja skalująca i falka Daubechies, filtry długości 20.

nie, jakie się pojawia, to czy i w jakim stopniu taki ciąg próbek reprezentuje wyjściowy sygnał. Na przykład, czy i w jaki sposób sygnał ciągły można z takich próbek odtworzyć. Podstawowym twierdzeniem związanym z próbkowaniem jest następujące twierdzenie.

Twierdzenie 5.4 (Shannon, Whittaker). *Jeżeli funkcja $f \in L^2(\mathbf{R})$ ma ograniczone spektrum częstotliwościowe, to znaczy dla pewnego T_0*

$$\hat{f}(\xi) \equiv 0 \quad \text{dla} \quad \xi \notin [-T_0, T_0],$$

to funkcja ta jest całkowicie reprezentowana przez ciąg próbek (w jednakowych odstępach) $\{f(np)\}_{n \in \mathbf{Z}}$, jeżeli próbkowanie jest wystarczająco gęste, a dokładnie jeżeli $p \leq \frac{\pi}{T_0}$. Wartości funkcji można odtworzyć z próbek przy pomocy wzoru

$$f(x) = \sum_{n=-\infty}^{\infty} f(np) \operatorname{sinc}\left(\frac{\pi}{p}x - \pi n\right), \quad (5.13)$$

gdzie funkcja interpolująca $\operatorname{sinc}(\cdot)$ dana jest wzorem

$$\operatorname{sinc}(x) = \frac{\sin(x)}{x}.$$

Uwaga 5.5. (i) *Zauważmy, że funkcja interpolująca*

$$\operatorname{sinc}\left(\frac{\pi}{p}x - \pi n\right)$$

jest równa 0 w punktach postaci $x = kp$, dla $k \in \mathbf{Z}$, $k \neq n$, oraz jest równa 1 dla $x = np$. Od razu więc widać, że prawa strona (5.13) zgadza się z lewą w

punktach postaci kp , czyli prawa strona jest funkcją interpolującą wartości pomiędzy próbkami.

(ii) Powyższe twierdzenie mówi nam, że interpolacja wartości pomiędzy próbkami odtwarza funkcję wyjściową, jeżeli ta ma odpowiednio ograniczone spektrum częstotliwościowe. Większość funkcji występujących w praktyce jest tego typu, i dlatego twierdzenie Shannona-Whittakera jest ważne. Można jednak badać inne przestrzenie funkcji (na przykład spliny, dla których funkcja interpolująca, żeby odtwarzać dokładnie wyjściowy sygnał analogowy, musi mieć inną postać.

(iii) Ciekawe jest pytanie o odtwarzanie funkcji z próbek pobieranych nieregularnie. W tym przypadku również można udowodnić ważne twierdzenia, wyjaśniające sprawę, i dające inżynierom praktyczne narzędzia.

(iv) Powyższe twierdzenie jest sformułowane w ten sposób, że podaje jak często trzeba próbkować sygnał wyjściowy, o spektrum częstotliwościowym ograniczonym do T_0 , żeby móc sygnał odtworzyć. Można to przeformułować: jeżeli próbki pobierane są w odstępach p , to jaka jest maksymalna częstotliwość sygnału, która zostanie odtworzona bez zniekształceń. Jest to tak zwana częstotliwość Nyquista i, jak łatwo odczytać z twierdzenia Shannona-Whittakera, wynosi $\frac{\pi}{p}$. Uwaga: częstotliwość, o której mówimy tutaj, to w języku inżynierów tak zwana częstotliwość kołowa. Częstotliwość, którą najczęściej posługują się inżynierowie (czyli 1/okres) to nasza częstotliwość podzielona przez 2π . W języku inżynierów częstotliwość Nyquista wynosi więc $\frac{1}{2p}$.

(v) Można zastanowić się, co się dzieje, jeżeli sygnał próbkowany jest za rzadko, w stosunku do zakresu swoich składowych częstotliwościowych. Prawa strona (5.13) jest wtedy zaledwie przybliżeniem lewej strony, a zniekształcenia (czyli błąd tego przybliżenia) mają charakterystyczną postać, i noszą nazwę aliasingu. Wrócimy jeszcze do zjawiska aliasingu.

Dowód twierdzenia Shannona-Whittakera. Załóżmy, że

$$\hat{f}(\xi) \equiv 0 \quad \text{dla} \quad \xi \notin [-T_0, T_0],$$

Niech $T \geq T_0$, i niech funkcja pomocnicza $g(\xi)$ będzie dana wzorem

$$g(\xi) = \hat{f}\left(\frac{T}{\pi}\xi\right), \quad \xi \in [-\pi, \pi],$$

a następnie tak zdefiniowaną w przedziale $[-\pi, \pi]$ funkcję g przedłużamy jako okresową o okresie 2π na całą prostą $\mathbf{R}$. Zauważmy, że $g \in L^2(\mathbf{R})$. To proste, $\hat{f}$ jest całkowalna z kwadratem na prostej, a więc też na przedziale

$[-T_0, T_0]$, a całka kwadratu g na przedziale $[-\pi, \pi]$ jest jej równa, po zamianie zmiennych. Funkcję g rozwijamy w szereg Fouriera

$$g(\xi) = \sum_{n=-\infty}^{\infty} \hat{g}(n) e^{in\xi}, \quad \text{w } L^2(\mathbf{T}),$$

a współczynniki Fouriera mają następującą postać:

$$\begin{aligned} \hat{g}(n) &= \frac{1}{2\pi} \int_{-\pi}^{\pi} g(\xi) e^{-in\xi} d\xi \\ &= \frac{1}{2\pi} \int_{-\pi}^{\pi} \hat{f}\left(\frac{T}{\pi} \xi\right) e^{-in\xi} d\xi \\ &= \frac{1}{2T} \int_{-T}^T \hat{f}(\xi) e^{-in \frac{\pi}{T} \xi} d\xi \\ &= \frac{1}{2T} \int_{-\infty}^{\infty} \hat{f}(\xi) e^{-in \frac{\pi}{T} \xi} d\xi \\ &= \frac{\pi}{T} f\left(-n \frac{\pi}{T}\right). \end{aligned}$$

Zauważmy, że wszystkie całki „po drodze” są absolutnie zbieżne, więc możemy skorzystać ze wzoru na odwrotną transformatę Fouriera. Wykorzystaliśmy też fakt, że $\hat{f}$ jest zerem poza przedziałem $[-T, T] \supseteq [-T_0, T_0]$, więc całka z $\hat{f}$ po przedziale $[-T, T]$ i po całej prostej są sobie równe. Z definicji funkcji g mamy

$$\begin{aligned} \hat{f}\left(\frac{T}{\pi} \xi\right) &= g(\xi) \cdot \chi_{[-\pi, \pi]}(\xi) \\ &= \sum_{n=-\infty}^{\infty} \hat{g}(n) \chi_{[-\pi, \pi]}(\xi) e^{in\xi} \\ &= \frac{\pi}{T} \sum_{n=-\infty}^{\infty} f\left(-n \frac{\pi}{T}\right) \chi_{[-\pi, \pi]}(\xi) e^{in\xi}. \end{aligned}$$

W końcu obliczymy wartość $f(x)$ wykorzystując odwrotną transformatę Fo-

uriera.

$$\begin{aligned}
f(x) &= \frac{1}{2\pi} \int_{-\infty}^{\infty} \hat{f}(\xi) e^{ix\xi} d\xi \\
&= \frac{1}{2\pi} \int_{-\infty}^{\infty} \hat{f}\left(\frac{T}{\pi}\xi\right) e^{ix\frac{T}{\pi}\xi} \frac{T}{\pi} d\xi \\
&= \frac{T}{2\pi^2} \int_{-\infty}^{\infty} g(\xi) \chi_{[-\pi, \pi]}(\xi) e^{ix\frac{T}{\pi}\xi} d\xi \\
&= \frac{T}{2\pi^2} \int_{-\pi}^{\pi} \sum_{k=-\infty}^{\infty} \hat{g}(k) e^{ik\xi + ix\frac{T}{\pi}\xi} d\xi \\
&= \frac{T}{2\pi^2} \sum_{k=-\infty}^{\infty} \hat{g}(k) \int_{-\pi}^{\pi} e^{i\xi(k+x\frac{T}{\pi})} d\xi \\
&= \frac{1}{2\pi} \sum_{k=-\infty}^{\infty} f\left(-\frac{\pi}{T}k\right) \int_{-\pi}^{\pi} e^{i\xi(k+x\frac{T}{\pi})} d\xi \\
&= \sum_{k=-\infty}^{\infty} f\left(\frac{\pi}{T}k\right) \frac{1}{2\pi} \int_{-\pi}^{\pi} e^{i\xi(x\frac{T}{\pi}-k)} d\xi \\
&= \sum_{k=-\infty}^{\infty} f\left(\frac{\pi}{T}k\right) \operatorname{sinc}\left(\pi\left(x\frac{T}{\pi}-k\right)\right) \\
&= \sum_{k=-\infty}^{\infty} f\left(\frac{\pi}{T}k\right) \operatorname{sinc}(Tx - \pi k)
\end{aligned}$$

Niech teraz $p = \frac{\pi}{T}$, i koniec dowodu. □

Rozdział 6

Materiały na laboratorium

Odczytywanie, zapisywanie i wyświetlanie obrazków w Matlabie

Obrazki w postaci cyfrowej występują w kilku rodzajach. Są tak zwane formaty indeksowane, gdzie każdemu pikselowi (elementowi obrazu) przypisywana jest liczba (najczęściej 1 bajt) będąca numerem koloru na liście. Lista kolorów, tak zwana paleta, najczęściej składa się z 256 kolorów, które reprezentowane są przez 3 bajty, odpowiadające zawartości składowej czerwonej, zielonej i niebieskiej. Przykładami formatów indeksowanych są formaty **bmp** i **gif**. Format **gif** zawiera obraz oraz paletę w postaci skompresowanej (bezzstratnie). W tym formacie w jednym pliku obrazów może być kilka, i każdy może mieć własną paletę. Jeżeli obrazów jest kilka, to całość może być animowana. Jeden z kolorów może być zadeklarowany jako przezroczysty. Plik w formacie **bmp** zawiera obraz i paletę, najczęściej w postaci nieskompresowanej. Sam format dopuszcza prostą kompresję, ale w praktyce nie spotyka się skompresowanych plików w formacie **bmp**. Drugim rodzajem obrazów cyfrowych są obrazy w odcieniach szarości. Obrazy tego typu najczęściej zapisywane są w jednym z formatów indeksowanych, z użyciem standardowej palety. Standardowa paleta zawiera równomiernie rozłożone odcienie szarości (3 bajty kolorów składowych równe sobie), od koloru (0,0,0), czyli czerni (indeks 0) do (255,255,255), czyli bieli (indeks 255). Trzecim rodzajem obrazów cyfrowych są tak zwane obrazy *true color*, w których każdemu pikselowi bezpośredni przypisane są 3 bajty określające zawartość 3 kolorów składowych. Obrazy tego typu zapisywane są w formatach **tiff** (skompresowane bezstratnie) i **jpeg** (skompresowane stratnie).

W tym laboratorium będziemy zajmowali się obrazami drugiego typu, to znaczy obrazami w odcieniach szarości. Kolorowe obrazy indeksowane nie

mogą być sensownie przetwarzane przy pomocy stosowanych przez nas metod, ponieważ numery kolorów w paletce nie mają związku z własnościami „kolorystycznymi”. Innymi słowy, to, że dwa kolory mają bliskie sobie indeksy nie oznacza, że są do siebie podobne. Oczywiście kolorowy obrazek w formacie indeksowanym można najpierw przekształcić do formatu *true color* lub odcieni szarości, i wtedy obrabiać. Obrazki *true color* można przetwarzać traktując każdy kolor składowy (tak zwany kanał) jako oddzielny obrazek. Lepiej jest najpierw przekształcić kanały z tak zwanej przestrzeni RGB (kanały to składowe czerwone, zielone i niebieskie) do przestrzeni YC_bC_r (kanały to mieszanki kolorów składowych, przy czym Y to kanał luminancji, czyli najważniejsza składowa zawierająca odcienie szarości, C_b i C_r to kanały chrominancji, odpowiednio niebieski i czerwony). Przekształcenia pomiędzy przestrzeniami RGB i C_bC_r implementuje się przy pomocy konkretnej, odwracalnej macierzy 3×3 . Być może starczy nam czasu na zajęcie się obrazkami kolorowymi, wtedy poznamy więcej szczegółów na temat zarządzania barwami.

Do wczytania obrazka służy funkcja `imread`. Jeżeli chcemy wczytać obraz `Lena.bmp`, który jest w bieżącym katalogu, wydajemy instrukcję

```
A=imread('Lena.bmp','bmp');
```

Powstaje w ten sposób macierz pikseli o wartościach od 0 do 255. Będziemy korzystali z formatów `gif` i `bmp`. Są to tak zwane formaty indeksowane, to znaczy wartość piksela jest numerem koloru w paletce. W obrazach z których będziemy korzystali paleta jest zawsze taka sama, standardowa. Składa się z 256 odcieni szarości, zmieniających się jednostajnie od czerni (pierwszy kolor palety, odpowiada mu indeks 0) do bieli (256 kolor, o indeksie 255). Jeżeli użyjemy innej palety to nasz obraz będzie miał zmienione kolory. Wczytując obrazek możemy wczytać też jego paletę, zakodowaną wraz z obrazkiem w pliku (dotyczy to, oczywiście formatów indeksowanych, takich jak `gif` lub `bmp`), wywołując funkcję `imread` w następujący sposób:

```
[A,MAP]=imread('goldhill.gif','gif');
```

Paleta obrazka składa się z trójek bajtów, reprezentujących intensywność kolorów czerwonego, zielonego i niebieskiego w skali od 0 do 255. Przy wczytywaniu do tablicy `MAP` (`MAP` jest tablicą 256 na 3, o wartościach typu `double`) te wartości są przeskalowane do zakresu $[0,1]$. Taki jest format używanej przez Matlab palety (tak zwanej *colormap*): dowolna ilość wierszy, 3 kolumny, i wartości typu `double` z przedziału $[0,1]$.

Otrzymana macierz pikseli `A` zawiera dane typu `uint8`. Niektórych obliczeń w Matlabie nie można wykonywać na liczbach typu `uint8` i najpierw trzeba je przekształcić do typu `double`.

```
B=double(A);
```

Podstawowe funkcje pakietu RWT, z którego będziemy korzystać nie mogą operować na danych typu `uint8`, więc nasze dane typu `uint8` zawsze będziemy przed obliczeniami przekształcać na `double`. To nic nie kosztuje. Zmieniany jest tylko typ danych, wartości pozostają bez zmian.

Do zapisu macierzy do pliku graficznego służy funkcja `imwrite`:

```
imwrite(A,'lena.bmp','bmp');
```

Macierz `A` powinna być typu `uint8`. Można ją do tego typu przekształcić z dowolnego innego (na przykład `double`) używając funkcji `uint8`

```
A=uint8(B);
```

Funkcja `uint8` zaokrągla liczby `double` do najbliższej całkowitej, p czym obcina zakres do przedziału $[0,255]$. Jeżeli funkcję `imwrite` zastosujemy do tablicy o wartościach typu `double` to wartości zostaną najpierw pomnożone przez 255, a następnie przekształcone przy pomocy funkcji `uint8`. Wartości poniżej 0 przejdą na 0 a powyżej 1 na 255. Cały używany przez macierz `A` zakres wartości powinien więc się zawierać w przedziale $[0,1]$. Dlatego przed zapisem do pliku warto ręcznie obrobić wartości macierzy `A` tak, aby zawierały się w przedziale $[0,1]$. Na przykład, niech

$$M = \max_{i,j} A(i,j), \quad N = \min_{i,j} A(i,j),$$

i

$$A'(i,j) = \frac{A(i,j) - N}{M - N}.$$

Po takim przekształceniu wartości macierzy `A'` są w przedziale $[0,1]$, a ponieważ przekształcenie wartości jest liniowe więc obraz „graficznie” nie uległ zmianie. Wraz z macierzą do pliku zapisywana jest także standardowa paleta, składająca się z 256 równomiernie rozłożonych odcieni szarości (dotyczy to formatów indeksowanych, takich jak `bmp`). Jeżeli chcemy zapisać inną paletę możemy użyć instrukcji

```
imwrite(A,MAP,'lena.bmp','bmp');
```

gdzie `MAP` jest dowolną colormapą Matlaba. Zostanie ona przeskalowana do formatu palety odpowiedniego pliku graficznego, ucięta do 256 wierszy jeżeli jest dłuższa, i uzupełniona wierszami zer, jeżeli jest krótsza niż 256 kolorów. Matlab nie zapisuje plików w formacie `gif`.

Macierz `A` można wyświetlić jako obraz nie zapisując jej do pliku graficznego. Służy do tego instrukcja `image`

```
image(A);
```

Wartości wyświetlanej macierzy *A* są traktowane jako indeksy do bieżącej colormapy Matlaba. Jeżeli są typu `uint8` to są bezpośrednio traktowane jako indeksy kolorów. Jeżeli są typu `double` to najpierw odejmowana jest od nich 1, następnie są zaokrąglane do najbliższej liczby całkowitej, a następnie obcinane na poziomie 0 i 255. Tak powstałe wartości są traktowane jako indeksy kolorów. Wygodną instrukcją jest też `imagesc`. Wartości tablicy *A* są najpierw przeskalowane liniowo, tak, żeby najmniejsza wartość (może być ujemna) przyjęła wartość 0 a największa 255. Tak otrzymane wartości są zaokrąglane w dół do liczb całkowitych, i traktowane jako indeksy kolorów bieżącej colormapy. Jeżeli chcemy przeskalować inaczej, możemy użyć składni

```
imagesc(A,[a,b]);
```

wtedy zakres `[a,b]` zostanie rozciągnięty liniowo do zakresu `[0,255]`.

Do podglądu naszych obrazów będzie nam potrzebna standardowa colormap, którą musimy sami wygenerować:

```
szara=zeros(256,3);  
for i=1:256  
    szara(i,1)=(i-1)/255;  
    szara(i,2)=szara(i,1);  
    szara(i,3)=szara(i,1);  
end;
```

Następnie ustalamy bieżącą colormapę instrukcją

```
colormap(szara);
```

Colormap pozostaje ustalona dla danego okna obrazka aż do jego zamknięcia. Przy następnym otwarciu colormapę trzeba ustawić ponownie. Można najpierw załadować obrazek, a następnie wydać instrukcję `colormap`. Dopóki nie zamkniemy okna obrazka ta sama colormap będzie stosowana do wszystkich kolejno ładowanych obrazków. Naszą colormapę *szara* możemy zachować w pliku instrukcją `save`, i na następnych zajęciach załadować z pliku instrukcją `load`. Matlab ma pewną ilość predefiniowanych colormap, jedną z nich jest colormap domyślna. Zawiera ona tylko 64 kolory. W przypadku gdy ustalona colormap ma mniej kolorów niż zakres wartości wyświetlanej macierzy *A*, to brakujące kolory są zastępowane czarnym. Obrazki w domyślnej colormapie są bardzo niewyraźne. Inne standardowe colormapy to `gray`, `hot`, `cool`, `copper` czy `pink`. Standardowe colormapy są z reguły krótsze niż 256

kolorów, ale każdej z powyższych nazw można użyć do wygenerowania colormapy dowolnej długości. Na przykład zamiast podanej powyżej procedury generowania colormapy „szara” można w Matlabie użyć instrukcji

```
szara=gray(256);
```

albo do ustawienia bieżącej colormapy instrukcji

```
colormap(gray(256));
```

Podstawowe operacje na obrazkach

Wypróbujemy niektóre typowe przekształcenia na obrazkach. Niektóre z operacji występujących w przykładowych skryptach można w Matlabie wykonać prościej, wykorzystując wbudowane funkcje. Na przykład w Matlabie są specjalne funkcje zwracające największą i najmniejszą wartość współczynników tablicy. W prezentowanych skryptach raczej wszystko wykonywane jest „na piechotę”.

Rozjaśnij

Zwiększamy wartość pikseli. Wartości powyżej 255 będą ucięte

```
A=imread('Lena.bmp','bmp');  
B=double(A);  
B=1.2*B;  
A=uint8(B);  
imwrite(A,'Lena2a.bmp','bmp');
```

Przyciemnij

Zmniejszamy wartość pikseli:

```
A=imread('Lena.bmp','bmp');  
B=double(A);  
B=B/1.2;  
A=uint8(B);  
imwrite(A,'Lena2b.bmp','bmp');
```



Rysunek 6.1: Lekkie rozjaśnienie.



Rysunek 6.2: Lekkie przyciemnienie.

Zwiększ kontrast

Znajdziemy największą i najmniejszą wartość pikseli. Następnie przekształcimy, liniowo, obraz tak, żeby najmniejsza wartość wynosiła 0 a największa 255. W ten sposób maksymalizujemy kontrast. Jeżeli kontrast był już maksymalny to operacja nie daje żadnego efektu.

```
A=imread('Lena.bmp','bmp');
B=double(A);
min=B(1,1);
max=B(1,1);
for i=1:512
```

```

for j=1:512
    if B(i,j)>max
        max=b(i,j);
    end;
    if B(i,j)<min
        min=B(i,j);
    end;
end;
end;
c=255/(max-min);
for i=1:512
    for j=1:512
        B(i,j)=(B(i,j)-min)*c;
    end;
end;
A=uint8(B);
imwrite(A,'Lena2c.bmp','bmp');

```



Rysunek 6.3: Zwiększony kontrast.

Częstą operacją jest korekta kontrastu tylko w pewnym zakresie jasności. Na przykład przekształcenie $B(i, j) \mapsto B(i, j)^\gamma$ zwiększa kontrast w zakresie czerni jeżeli $\gamma < 1$ i w zakresie jasnym, jeżeli $\gamma > 1$ (obraz musi być wcześniej znormalizowany, tak aby $0 \leq B(i, j) \leq 1$).

Zmniejsz kontrast

Wartości pikseli pomnożymy przez 0.8 i dodamy do nich 25:

```

A=imread('Lena.bmp','bmp');
B=double(A);
for i=1:512
    for j=1:512
        B(i,j)=25+0.8*B(i,j);
    end;
end;
A=uint8(B);
imwrite(A,'Lena2d.bmp','bmp');

```



Rysunek 6.4: Zmniejszony kontrast.

Binaryzacja

Każdemu pikselowi przyporządkowujemy wartość 0 jeżeli jego wartość jest poniżej progu i 255 jeżeli powyżej.

Powiel obraz, używając odbicia

Rozszerzymy obraz w poziomie, uzupełniając prawą połówkę lustrzanym odbiciem lewej:

```

B=zeros(512,1024);
A=imread('Lena.bmp','bmp');
for i=1:512
    for j=1:512
        B(i,j)=A(i,j);

```



Rysunek 6.5: Binaryzacja na poziomach 50, 100, 150 i 200.

```

        B(i,1025-j)=A(i,j);
    end;
end;
A=uint8(B);
imwrite(A,'Lena2e.bmp','bmp');

```

Wycięcie fragmentu

Wytniemy kwadratowy fragment obrazka, zastępując resztę białym tłem:

```

A=imread('Lena.bmp','bmp');
for i=1:128
    for j=1:512
        A(i,j)=255;
    end
end

```



Rysunek 6.6: Odbicie poziome.

```

A(513-i,j)=255;r
A(j,i)=255;
A(j,513-i)=255;
end;
end;
imwrite(A,'Lena2f.bmp','bmp');

```



Rysunek 6.7: Wycięcie fragmentu.

Wycięcie gładkie

Wytniemy kwadratowy kawałek obrazka, ale tym razem przejście do białego tła będzie płynne:

```
A=imread('Lena.bmp','bmp');
B=double(A);
mask=zeros(512,1);
for i=1:256
    if (i>108)&(i<129)
        mask(i)=(i-108)/21;
        mask(513-i)=mask(i);
    end;
    if i>128
        mask(i)=1;
        mask(513-i)=1;
    end;
end;
for i=1:512
    for j=1:512
        B(i,j)=255-mask(i)*mask(j)*(255-B(i,j));
    end;
end;
A=uint8(B);
imwrite(A,'Lena2g.bmp','bmp');
```



Rysunek 6.8: Wycięcie fragmentu z gładkim przejściem do tła.

Zaszumienie

Dodamy do obrazka tak zwany „biały szum”, to znaczy do każdego piksela niezależnie dodamy liczbę losową o rozkładzie normalnym o średniej 0 i jakimś odchyleniu. W Matlabie jest funkcja `randn` która generuje liczbę pseudolosową o rozkładzie normalnym (gaussowskim), o średniej 0 i odchyleniu standardowym 1. Można też użyć składni `randn(512)`, która od razu generuje tablicę 512×512 liczb pseudolosowych, niezależnych o tym samym rozkładzie. W naszym przykładzie tak wygenerowane liczby mnożymy przez 20, w ten sposób rozkład generowanych zmiennych ma odchylenie standardowe 20. Wielkość odchylenia jest związana ze stopniem zaszumienia. Takie zaszumienie symuluje zniekształcenie obrazka często występujące w praktyce, przy przesyłaniu sygnałów, lub przy rejestracji bardzo słabych sygnałów.

```
A=imread('Lena.bmp','bmp');
B=double(A);
for i=1:512
    for j=1:512
        B(i,j)=B(i,j)+20*randn;
    end;
end;
A=uint8(B);
imwrite(A,'Lena2h.bmp','bmp');
```



Rysunek 6.9: Zaszumienie obrazka, $\sigma^2 = 20$.

Rozmycie

Obrazek zostaje rozmyty, pozbawiony ostrości. Jest to zabieg stosowany na przykład jako wstępna obróbka przed algorytmem wykrywania krawędzi. Efektem rozmycia jest osłabienie niektórych rodzajów szumu (na przykład szumu białego). Krawędzie też zostają osłabione (rozmyte), ale z reguły w mniejszym stopniu niż szum.

```
A=imread('Lena.bmp','bmp');
B=double(A);
C=zeros(512);
mask=[2 4 5 4 2;4 9 12 9 4;5 12 15 12 5;4 9 12 9 4;2 4 5 4 2];
mask=mask/159;
for i=1:512
    for j=1:512
        suma=0;
        for k=-2:2
            for l=-2:2
                if (i-k<513)& (j-l<513)& (i-k>0)& (j-l>0)
                    prod=B(i-k,j-l);
                else
                    prod=0;
                end;
                suma=suma+mask(k+3,l+3)*prod;
            end;
        end;
        C(i,j)=suma;
    end;
end;
A=uint8(C);
imwrite(A,'Lena2i.bmp','bmp');
```

Dwa często stosowane filtry Gaussa (maski), 3×3 i 5×5 :

$$\begin{pmatrix} 1 & 2 & 1 \\ 2 & 4 & 2 \\ 1 & 2 & 1 \end{pmatrix}, \quad \begin{pmatrix} 2 & 4 & 5 & 4 & 2 \\ 4 & 9 & 12 & 9 & 4 \\ 5 & 12 & 15 & 12 & 5 \\ 4 & 9 & 12 & 9 & 4 \\ 2 & 4 & 5 & 4 & 2 \end{pmatrix},$$

filtry te należy podzielić przez sumę współczynników, odpowiednio 16 i 159.



Rysunek 6.10: Rozmycie filtrem Gaussa 5x5.



Rysunek 6.11: Rozmycie filtrem Gaussa 3x3.

Filtr medianowy

Czasem chcielibyśmy osłabić szum w obrazku, ale bez ogólnego „zmiękczenia”. Można wtedy zastosować tak zwany filtr medianowy. Odczytujemy wartość piksela i jego sąsiadów (na przykład najbliższych sąsiadów). Tak otrzymane wartości sortujemy. Pikselowi przypisujemy wartość znajdującą się w środku posortowanej listy (jeżeli mamy parzystą liczbę wartości, to pikselowi przyporządkowujemy średnią arytmetyczną wartości w środku posortowanej listy). Filtr medianowy dobrze usuwa niektóre rodzaje szumu, na przykład tak zwany „speckle noise”, bez ogólnego „zmiękczenia” obrazka.

```
A=imread('Lena.bmp','bmp');
```

```

B=double(A);
C=zeros(512);
lista=zeros(9,1);
for i=2:511
    for j=2:511
        m=1;
        for k=-1:1
            for l=-1:1
                lista(m)=B(i+k,j+l);
                m=m+1;
            end;
        end;
        for k=1:8
            for l=1:9-k
                if lista(l)>lista(l+1)
                    t=lista(l);
                    lista(l)=lista(l+1);
                    lista(l+1)=t;
                end;
            end;
        end;
        C(i,j)=lista(5);
    end;
end;
A=uint8(C);
imwrite(A,'Lena2j.bmp','bmp');

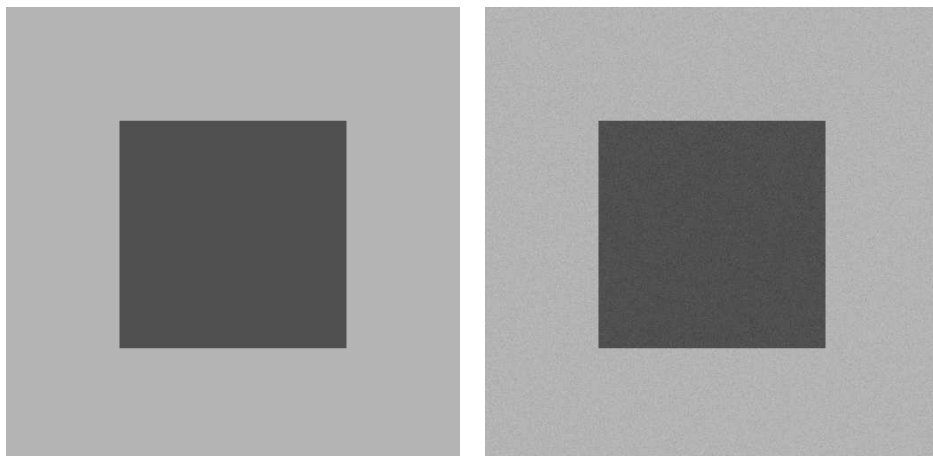
```

Zastosowanie filtru medianowego, podobnie jak filtru Gaussa daje efekt osłabienia szumu białego. Efekt ten najlepiej widać na obrazku kontrastowym. Obrazek `test` składa się z pikseli o wartości 80 i 180 (dosyć ciemne i dosyć jasne), i ma wyraźne krawędzie. Do obrazka dodajemy nieco szumu (odchylenie standardowe 5), a następnie stosujemy filtr Gaussa z maską 5×5 , oraz filtr medianowy, też o głębokości 5×5 . Na obrazkach możemy porównać efekty.

Oba filtry osłabiają szum, ale filtr Gaussa zmiękcza krawędzie, natomiast filtr medianowy pozostawia krawędzie nienaruszone, jedynie „obgryza” narożniki. Łatwo sobie wytłumaczyć jak to się dzieje.



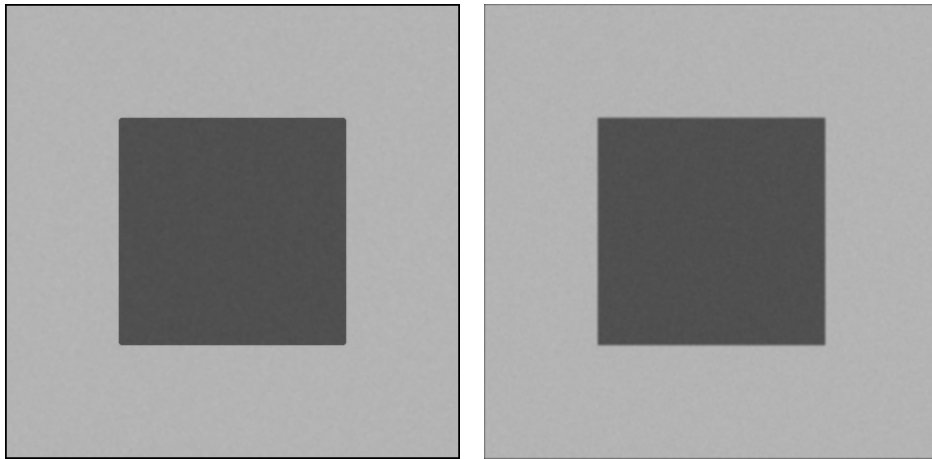
Rysunek 6.12: Filtr medianowy, obejmujący najbliższych sąsiadów.



Rysunek 6.13: Obraz testowy, bez szumu i z lekkim szumem, $\sigma^2 = 5$.

Transformata Fouriera

Transformata Fouriera to jest nasze główne narzędzie na wykładzie. W praktycznej obróbce obrazków nie będziemy stosowali transformaty Fouriera. W praktyce lepsza jest transformata falkowa. Spróbujmy obliczyć transformatę Fouriera obrazka. W Matlabie są funkcje `fft` i `fft2` obliczające 1- i 2-wymiarową transformatę Fouriera, używające algorytmu szybkiej transformaty. Transformata ma wartości zespolone, nawet dla sygnałów, które mają wartości tylko rzeczywiste. Żeby zwizualizować transformatę osobno wyświetlimy tablicę modułów wartości, i osobno tablicę argumentów (faz) wartości. Do obliczania modułu używamy funkcji `abs(X)`, którą można stosować do



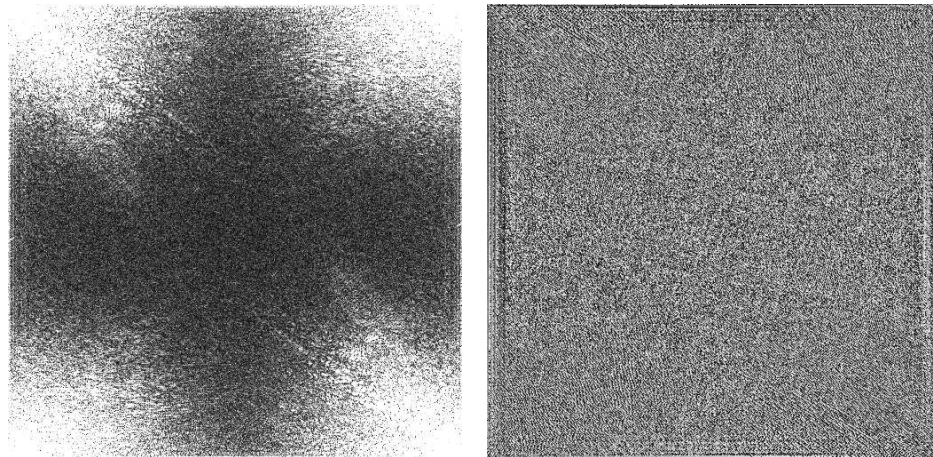
Rysunek 6.14: Z lewej zastosowano filtr medianowy, a z prawej filtr Gaussa, oba o wielkości 5×5 .

zmiennych typu rzeczywistego i zespolonego, a do obliczania argumentu używamy funkcji `angle`, która zwraca argument jako kąt w zakresie $(-\pi, \pi]$.

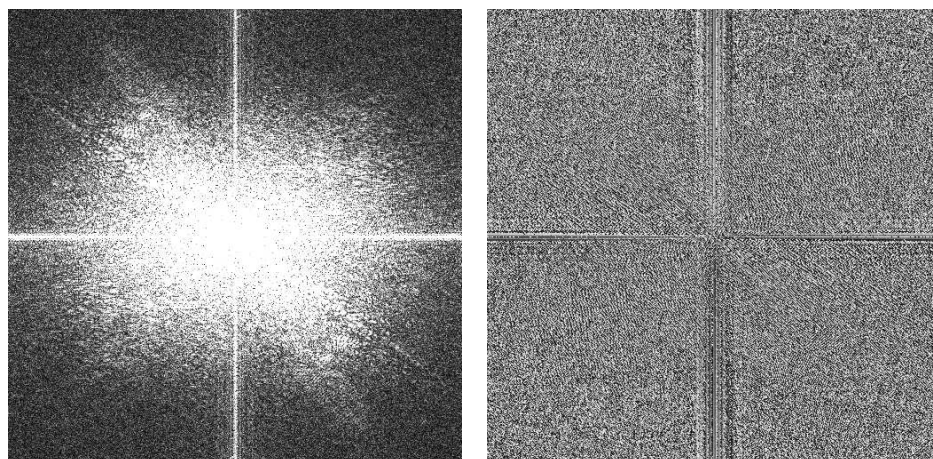
```
A=imread('Lena.bmp','bmp');
B=double(A);
C=fft2(B);
A=uint8(abs(C)/20);
imwrite(A,'Lena3a.bmp','bmp');
A=uint8((pi+angle(C))*128/pi);
imwrite(A,'Lena3b.bmp','bmp');
```

Na obrazku 3.1 współczynniki odpowiadające niskim częstotliwościom rozmieszczone są w rogach obrazka. Czasem obraz częstotliwościowy przedstawia się ze współczynnikami niskich częstotliwości umieszczonymi w centrum. W Matlabie jest specjalna funkcja do takiego przesunięcia obrazka `fftshift(C)`. Na obrazku 3.2 przedstawiona jest tak przesunięta transformata Fouriera.

Transformata Fouriera nie jest lokalna. To znaczy, że nawet jeżeli dwa obrazki różnią się tylko na jakimś małym obszarze, to ich transformaty różnią się wszędzie. Na obrazku 3.3 widzimy obraz Leny, z niewielką modyfikacją w okolicach środka. Obok przedstawiony jest obraz modułu różnicy transformaty Fouriera Leny zwykłej i zmodyfikowanej.



Rysunek 6.15: Moduł i argument transformaty Fouriera Leny.



Rysunek 6.16: Moduł i argument transformaty Fouriera, niskie częstotliwości w środku obrazka.

Kodowanie pasmowe

Typowym przekształceniem obrazka jest tak zwane kodowanie pasmowe. Zerujemy te współczynniki transformaty Fouriera, które leżą w wybranych obszarach. Tradycyjnie różne obszary geometryczne transformaty Fouriera nazywają się pasmami, i stąd nazwa. Typowym przykładem kodowania pasmowego jest filtr dolnoprzepustowy. Zerujemy wszystkie współczynniki transformaty Fouriera, które leżą poza pewnym, powiedzmy kwadratowym otoczeniem początku układu. Fig. 3.4 pokazuje maskę przykładowego filtra dolnoprzepustowego, oraz przefiltrowany obraz. Obraz jest zupełnie dobrej



Rysunek 6.17: Lena z niewielką modyfikacją (prawe oko). Obok moduł różnicy transformaty Fouriera Leny zwykłej i zmodyfikowanej.

jakości, jeżeli wziąć pod uwagę, że został odtworzony z 12544 współczynników (pozostałe zostały wyzerowane), co stanowi ok. 4,78% całości. Jedyne co optycznie przeszkadza, to okresowo powielone, „drgające” krawędzie.



Rysunek 6.18: Maska filtru dolnoprzepustowego, i przekształcona Lena. Stopień kompresji. wynosi ok. 21 (obraz po prawej ma w rezultacie ok. 0,38 bita na piksel (bpp))

```
clear;
mask=zeros(512);
for i=200:311
```

```

        for j=200:311
            mask(i,j)=1;
        end
    end;
end;
A=imread('Lena.bmp','bmp');
B=double(A);
C=fft2(B);
C=fftshift(C);
C=mask.*C;
C=ifftshift(C);
B=ifft2(C);
A=uint8(abs(B));
imwrite(A,'low.bmp','bmp');
mask=255*mask;
A=uint8(mask);
imwrite(A,'mask.bmp','bmp');

```

Rice Wavelet Toolbox

Naszym narzędziem do analizy falkowej obrazków jest darmowy toolbox falkowy (Rice wavelet toolbox) napisany przez grupę ludzi na uniwersytecie Rice. Można go znaleźć w sieci używając słowa kluczowego „rwt”. Będziemy używali następujących funkcji z tego pakietu: `mdwt` - transformata falkowa, `midwt` - odwrotna transformata falkowa oraz `daubcwf` - program generujący filtry falkowe Daubechies. Funkcje używamy następująco

```
h0=daubcwf(N);
```

lub

```
[h0,h1]=daubcwf(N);
```

N jest długością filtru, musi być liczbą parzystą. Dłuższe filtry powinny dawać lepsze rezultaty, ale działają wolniej, i generują większe błędy zaokrągleń. Typowe długości to 2,6,10. Będziemy porównywać nasze algorytmy dla różnych długości filtrów. Uzyskane przy pomocy funkcji `daubcwf` filtry stanowią parametr transformaty falkowej i transformaty odwrotnej. Podajemy tylko współczynniki filtru $h0$ (dolnoprzepustowego), Matlab sam wyliczy współczynniki pasującego filtru $h1$.

```
[B,L]=mdwt(A,h0,L);          [A,L]=midwt(B,h0,L);
```



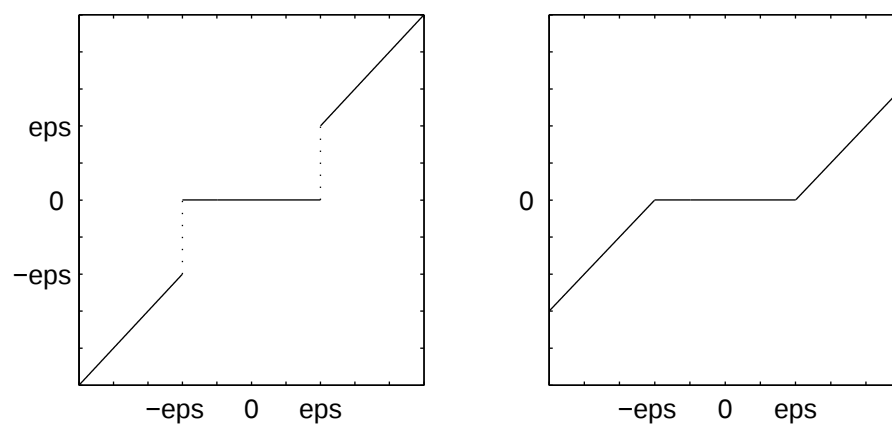
Rysunek 6.19: Transformata falkowa głębokości 1 i głębokości 9 (silnie rozjaśniona)

A jest obrazkiem, który powinien być kwadratowy o boku będącym potęgą 2. L jest parametrem określającym głębokość transformaty. Wartość parametru L powinna być taka sama dla transformaty i transformaty odwrotnej — w przeciwnym razie obraz nie będzie prawidłowo zrekonstruowany. Wartość L może się zawierać w przedziale od 1 do $\log_2 N$, gdzie N jest długością (w pikslach) boku obrazu. Transformata `mdwt` i transformata odwrotna `midwt` wymagają danych typu `double`. Jeżeli zastosujemy ją do danych typu `uint8` (takich, jakie zwraca `imread`), to transformaty albo będą źle policzone, albo program się wysypie. Przetransformowany obraz możemy obejrzeć instrukcją `image`, pamiętając o ustawieniu odpowiedniej colormapy, i o ewentualnym przekształceniu zakresu wartości transformaty do przedziału $[0,255]$.

Kompresja obrazów

Nasze podejście do kompresji będzie bardzo proste. Pierwsza obserwacja jest taka, że transformata falkowa obrazu jest prawie cała czarna. Większość współczynników (dla realistycznego obrazu, takiego jak fotografia) jest bardzo mała. Ustalimy sobie próg ϵ , i wszystkie współczynniki o wartości bezwzględnej poniżej progu zmienimy na 0. W ten sposób w obrazku pozostanie niewiele niezerowych współczynników. Format obrazu z którym pracujemy (`.bmp`) nie kompresuje danych, więc efektów kompresji nie zauważymy w rozmiarze kompresowanego pliku. Dlatego będziemy obliczali stopień kompresji w sposób uproszczony. Każdy wyzerowany współczynnik policzymy, i na koniec podzielimy przez ilość wszystkich pikseli.

Zerowanie współczynników o wartościach poniżej ustalonego progu będziemy nazywać progowaniem (po angielsku „thresholding”). Będziemy rozważać dwa rodzaje progowania, tak zwane progowanie twarde i progowanie miękkie. Różnicę objaśniają wykresy funkcji progujących. Programowanie twarde jest koncepcyjnie prostsze, ale wprowadza do obrazka sztuczne nieciągłości. Z kolei progowanie miękkie obniża kontrast transformaty, i w przypadku wysokich progów należy kontrast wyrównać przed zastosowaniem transformaty odwrotnej.



Rysunek 6.20: Progowanie twarde i miękkie

Będziemy porównywać subiektywną, optyczną jakość obrazów skompresowanych filtrami Daubechies różnej długości, o różnym stopniu kompresji, i kompresowanych z użyciem obu metod progowania. Przykładowe procedury mogą więc być następujące.

```
A=imread('Lena.bmp','bmp');
A=double(A);
N=6;% 2,10 itp
h0=daubcwf(N);
L=9;
[B,L]=mdwt(A,h0,L);
eps=50;% 30, 100 itp
il=0;
for i=1:512
    for j=1:512
        if abs(B(i,j))<eps
            B(i,j)=0;
            il=il+1;
        end
    end
end
```

```

        end;
    end;
end;
A=midwt(B,h0,L);
image(A);
100*il/(512*512)

```

Progowanie miękkie możemy zaimplementować następująco

```

if B(i,j)>0
    B(i,j)=B(i,j)-eps;
    if B(i,j)<0
        B(i,j)=0;
        il=il+1;
    end;
else;
    B(i,j)=B(i,j)+eps;
    if B(i,j)>0
        B(i,j)=0;
        il=il+1;
    end;
end;

```

Jeżeli będziemy używać dużej wartości progu ϵ , to należy wyrównać poziomy. To znaczy, należy przed progowaniem znaleźć

$$M = \max_{i,j} |B(i,j)|,$$

a następnie, po progowaniu pomnożyć

$$B(i,j) = B(i,j) * M / (M - \epsilon).$$

Powinniśmy eksperymentować z ϵ tak, aby uzyskać typowe wartości stopnia kompresji: 90%, 95%, 98%. W przypadku trafienia we właściwy stopień kompresji obraz skompresowany należy zapisać, nadając mu nazwę umożliwiającą identyfikację stopnia kompresji, długości filtru i rodzaju progowania. postarajmy się wyciągnąć wnioski na temat roli długości filtru oraz stopnia progowania. Porównajmy wyniki dla kilku różnych obrazków. Wypróbujmy ten algorytm również na obrazkach typu grafika komputerowa.