

**Wybrane metody matematyczne w
zastosowaniach na przykładzie projektu
„Domestic – domowy asystent osób
starszych i chorych”**

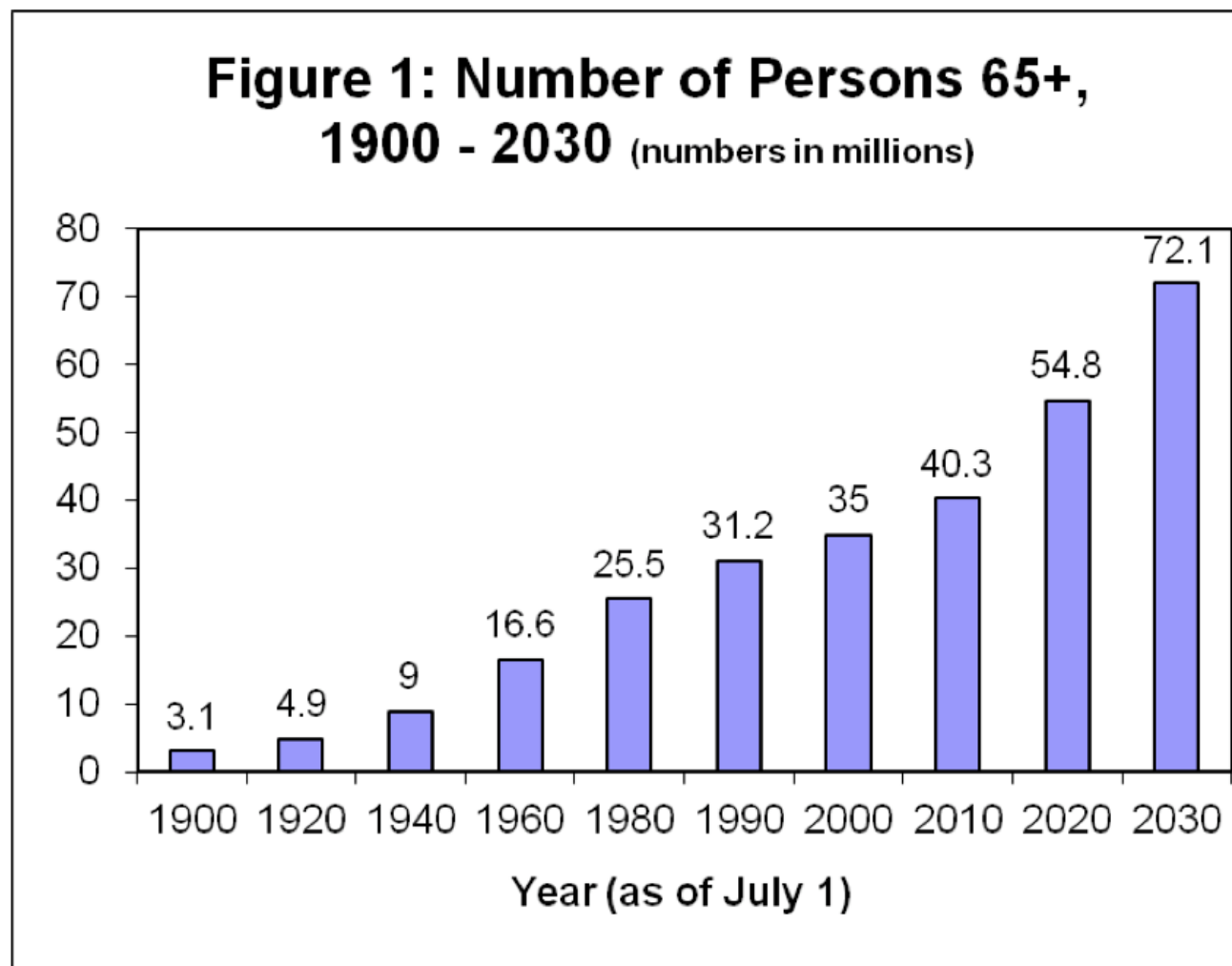
dr hab. inż. Jerzy Wtorek, prof. PG

dr inż. Artur Poliński

Katedra Inżynierii Biomedycznej, Politechnika
Gdańska

Przewiduje się, że zarówno w USA jak i Unii Europejskiej, pomijając problemy możliwe problemy ekonomiczne, zabraknie osób opiekujących się starszymi, chorymi i inwalidami.

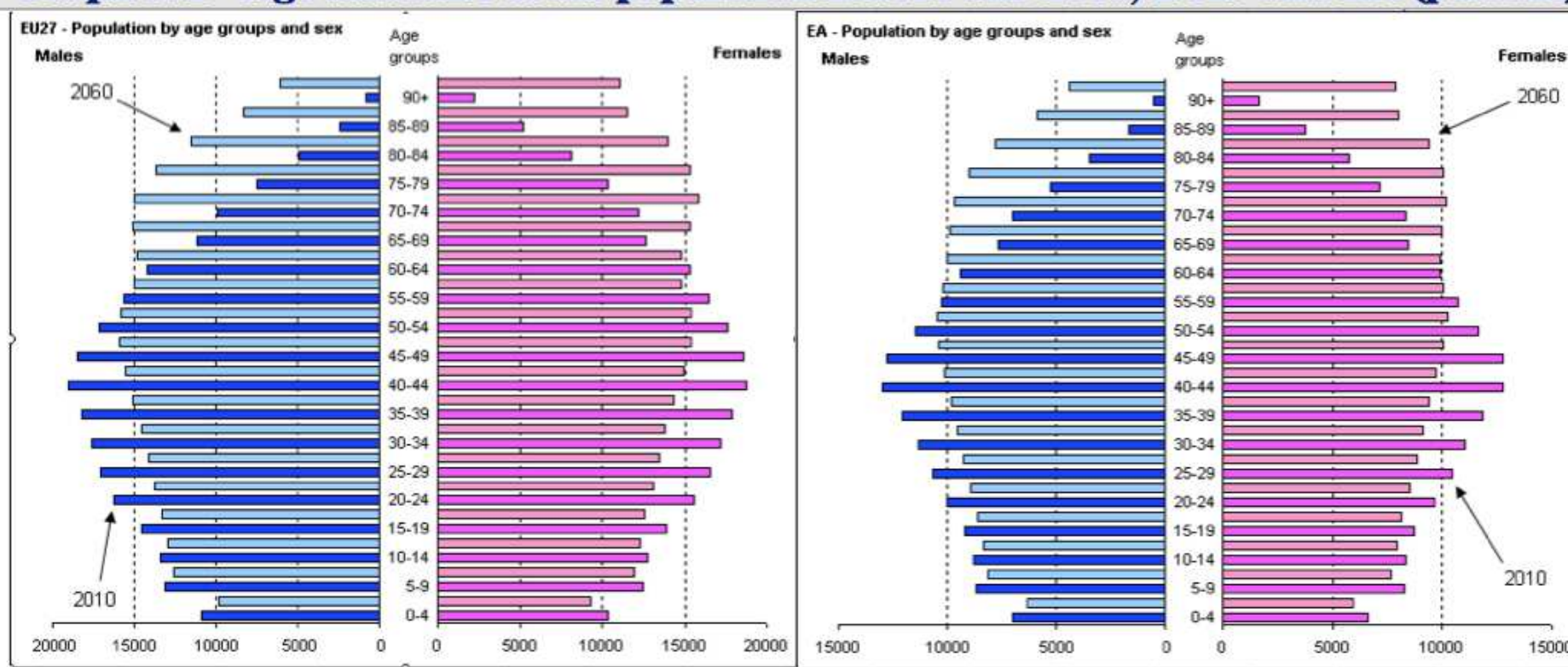
USA



Źródło: Administration on Aging, A Profile of Older Americans: 2011. Administration on Aging, U.S. Department of Health & Human Services, (2011).

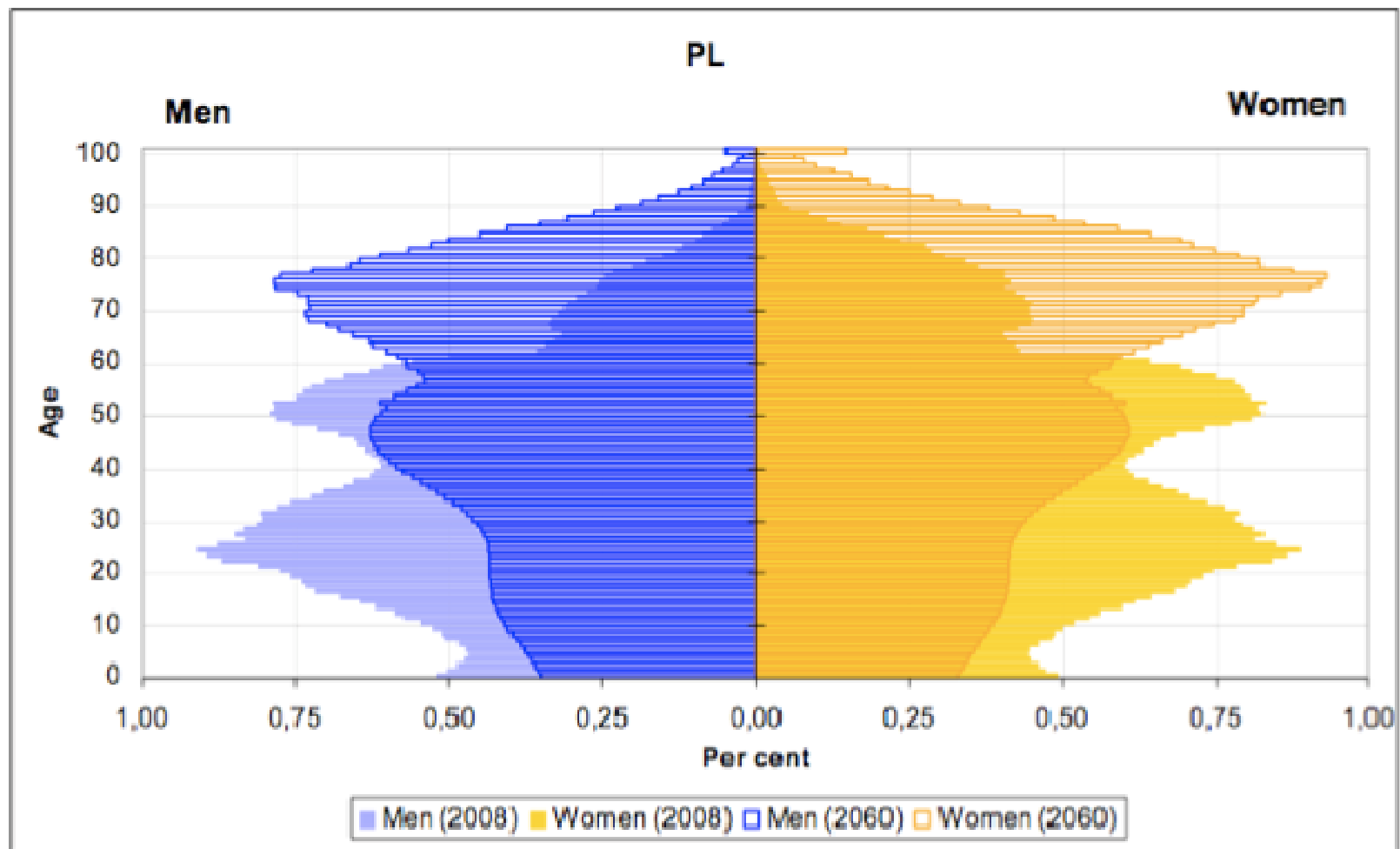
EU

Graph 0. 2 - Age structure of the population in 2010 and 2060, EU27 and EA (persons)

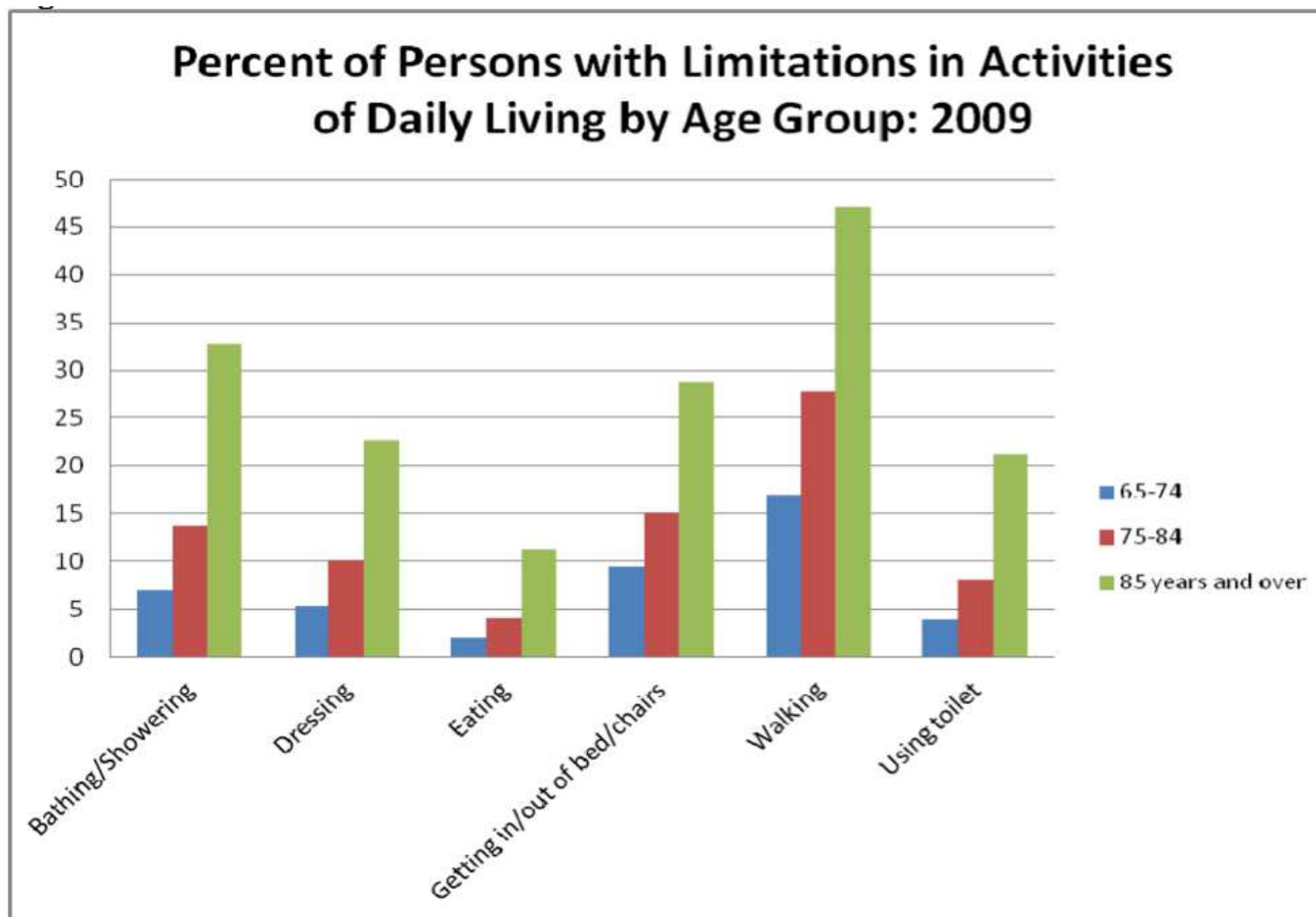


Źródło: European Union, The 2012 Ageing Report: Underlying Assumptions and Projection Methodologies. EUROPEAN ECONOMY 4, European Union, (2011).

PL



Źródło: Eurostat, EUROPOP2008, 2008



Źródło: Administration on Aging, A Profile of Older Americans: 2011. Administration on Aging, U.S. Department of Health & Human Services, (2011).

Znacznie dłuższy czas życia kreuje nowe problemy, a przynajmniej zmienia się skala niektórych, np. przewlekłe choroby, ograniczenie niektórych funkcji, itp.

Jednocześnie wzrastają koszty leczenia i zapotrzebowanie na profesjonalnych opiekunów.

Wstępne wnioski:

- 🏠 populacja 65+: około 30% całej populacji w EU/USA w 2060
- 🏠 Populacja 50+: istotny wzrost wydatków na opiekę zdrowotną
- 🏠 Problemy ekonomiczne (koszty opieki, koszty opieki zdrowotnej)
- 🏠 Zasoby ludzkie (ograniczona liczba w wieku produkcyjnym, ograniczona liczba opiekunów formalnych i nieformalnych)

Każde rozwiązanie redukujące konieczność udziału opiekuna i poprawiające jakość życia jest

OCZEKIWANE

Szczególnie oczekiwane są rozwiązania

INTELIGENTNE

Technologie wspomagające to te, które pomagają:

- Tobie i mi,
- Osobom z upośledzeniami,
- Osobom chorym,
- Osobom starszym

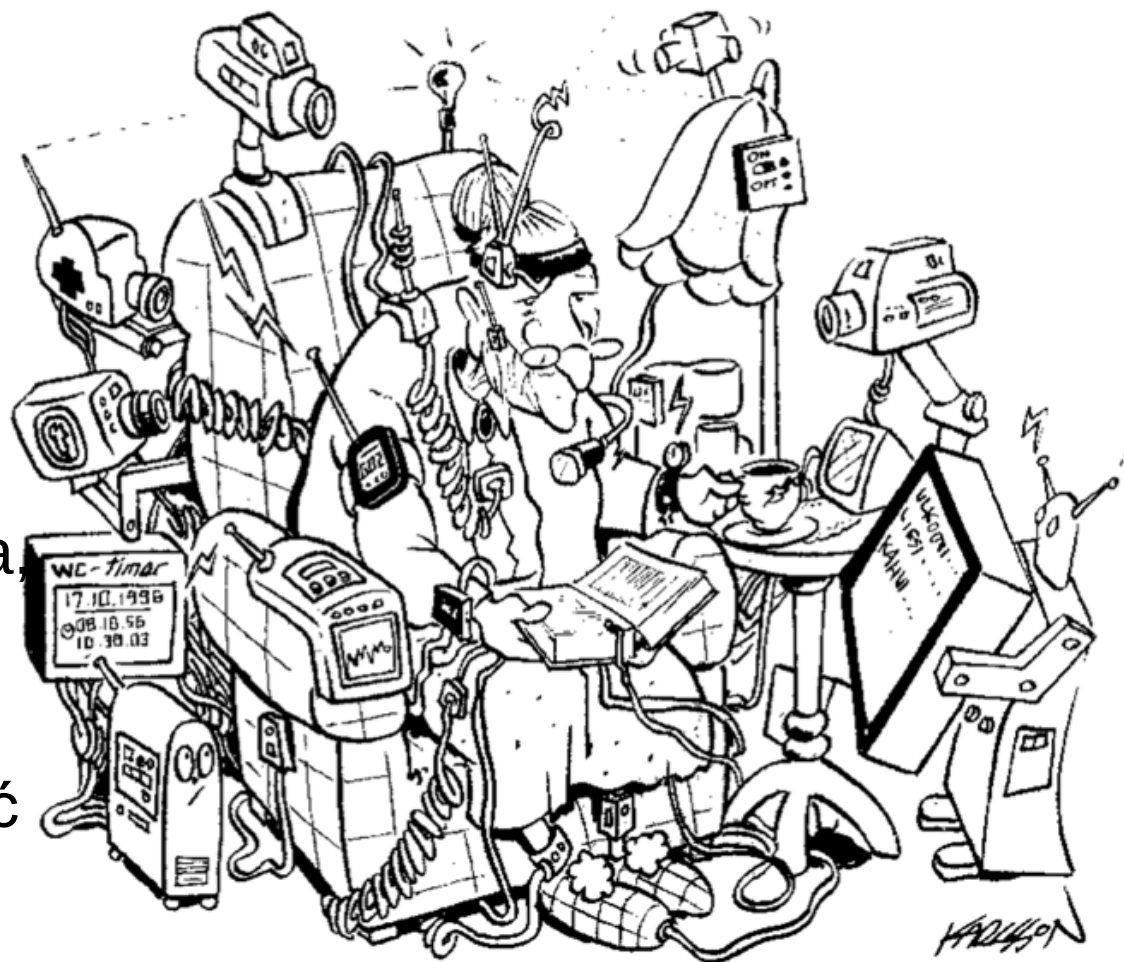
wykonywać czynności, które są trudne lub niemożliwe do wykonania



Źródło: IEEE Spectrum, September, 2012

Wymagania:

- Przyjazne w użyciu,
- Łatwe do zrozumienia,
- Ukryte,
- Tanie,
- O niskich kosztach użycia,
- Niezawodne,
- Bezobsługowe,
- Zapewniające prywatność
-



Spersonalizowana opieka medyczna

Rozwiązania ICT wspomagające opiekę zdrowotną:

- Typowe przyrządy medyczne (e.g. pulsoksymetry, ciśnieniomierze, glukometry, itd.)
- Dobry samopoczucie/stan fizyczny (np. wagi, mierniki aktywności, przyrządy ćwiczeń, inne)
- Specjalizowane przyrządy (np. zdalne rejestratory EKG)
- Aplikacje do zarządzania Indywidualnymi Kartami Pacjenta (Personal Health Records - PHR)
- Aplikacje do ćwiczeń mentalnych (demencja, Alzheimer, asysta w przemieszczaniu, układy przypominające, itp.)
- Zarządzanie procesem leczenia,
- inne.

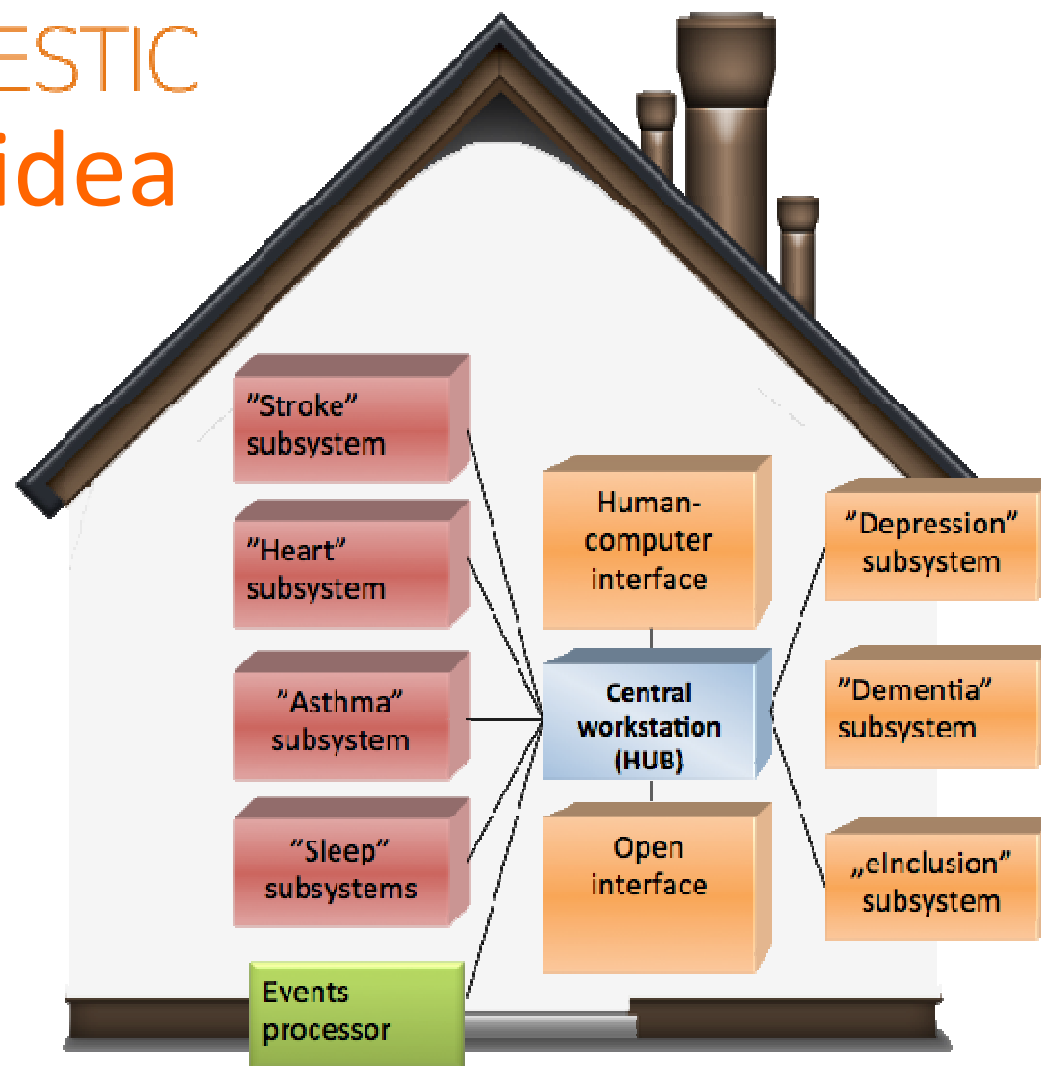


Zaprojektowane w KIB, PG

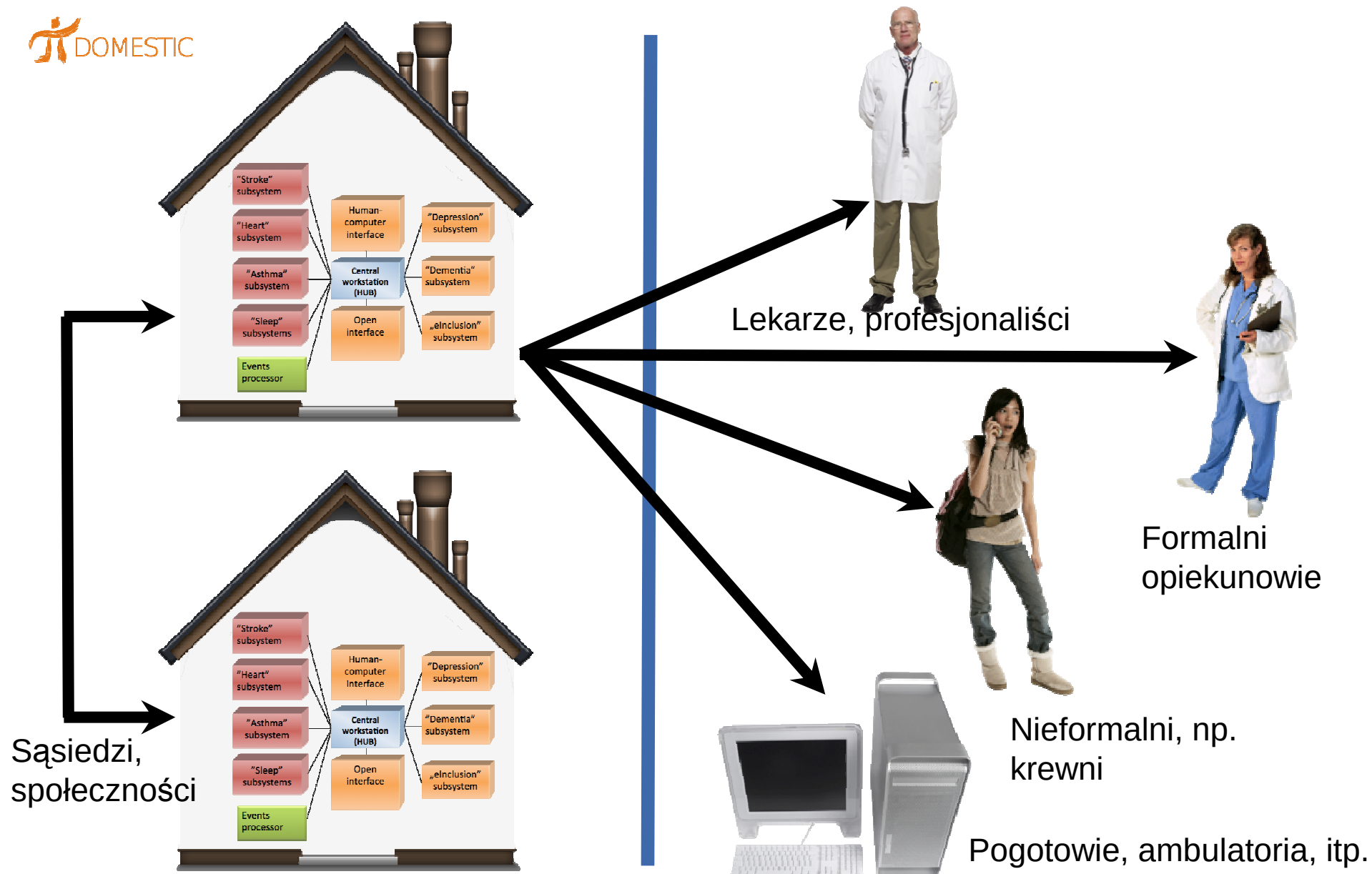


**Praca współfinansowana przez Europejski
Funduszu Rozwoju Regionalnego, projekt:
UDA-POIG.01.03.01-22-139/09-00 -“Domowy
asystent osób starszych i chorych –
DOMESTIC”, Innowacyjna Gospodarka 2007-
2013, Narodowa Strategia Spójności**

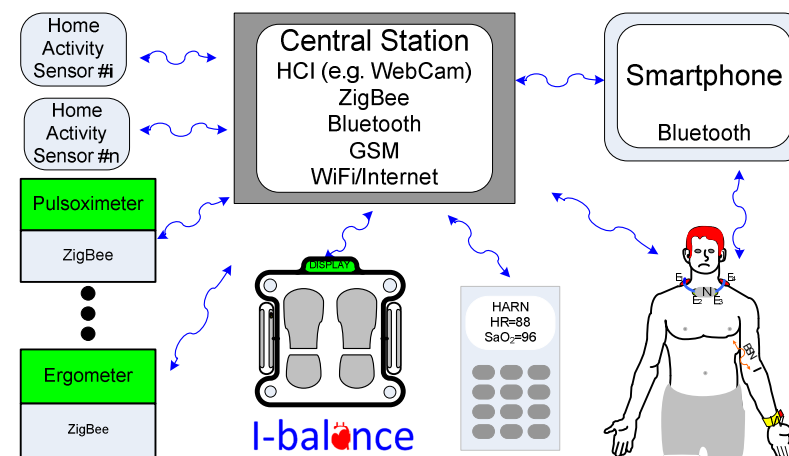
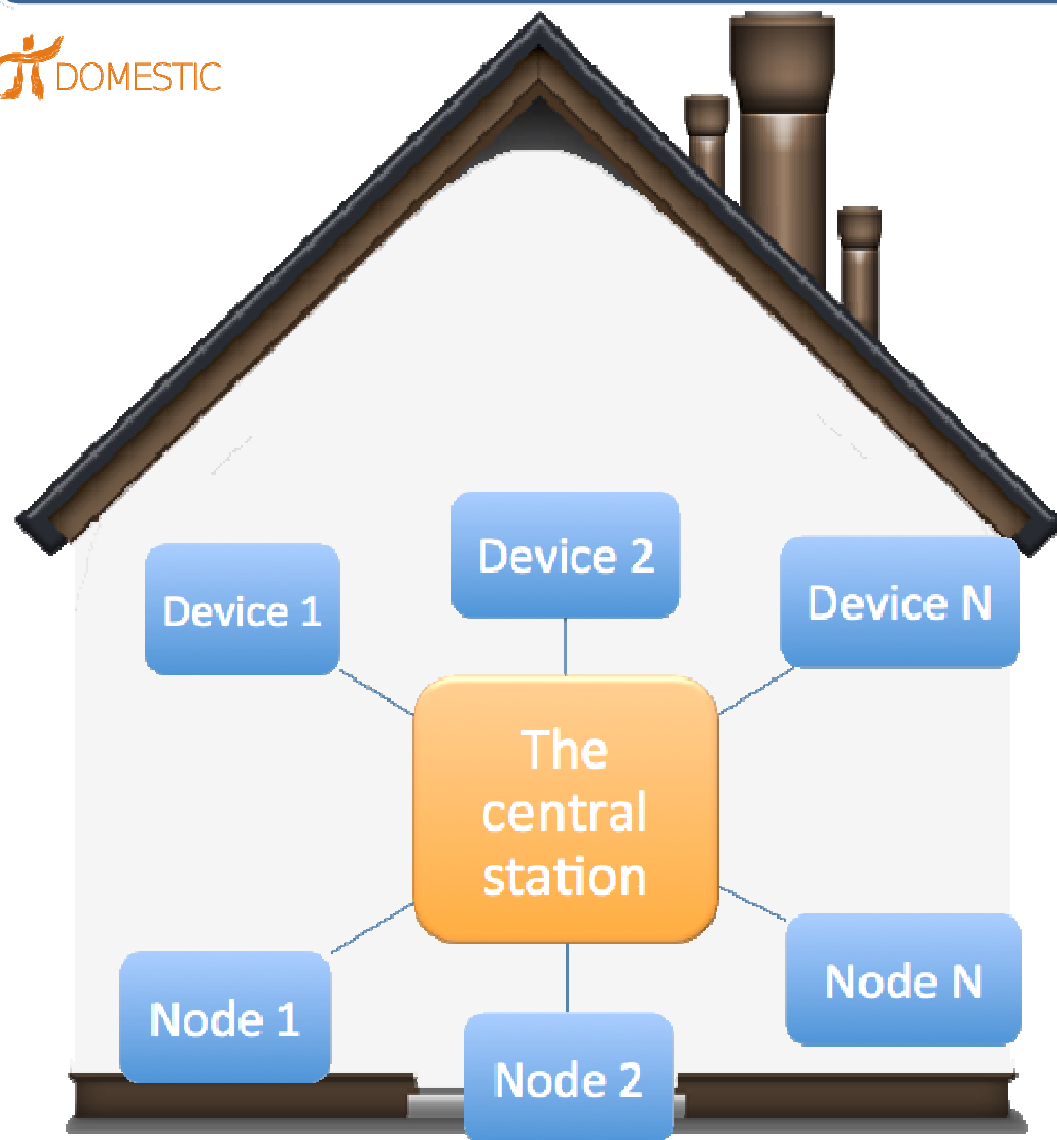
DOMESTIC idea



Spersonalizowana opieka medyczna

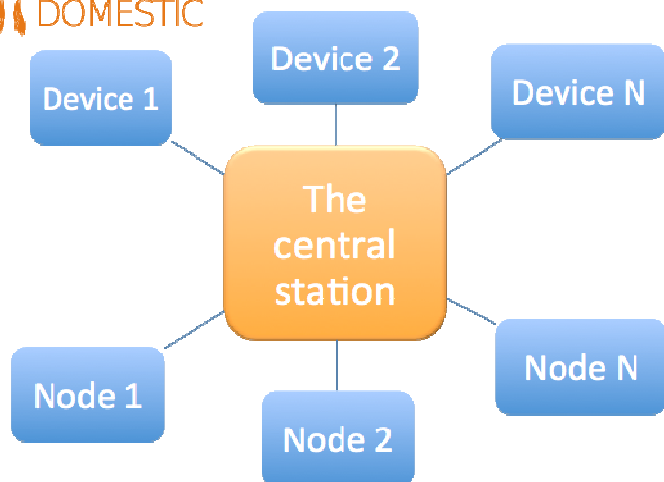


Spersonalizowana opieka medyczna -> Stacja centralna



Krótkookresowe i długookresowe zbierania danych i analiza: dane medyczne, aktywność, bezpieczeństwo, zdarzenia, itp.

Spersonalizowana opieka medyczna -> KIB PG

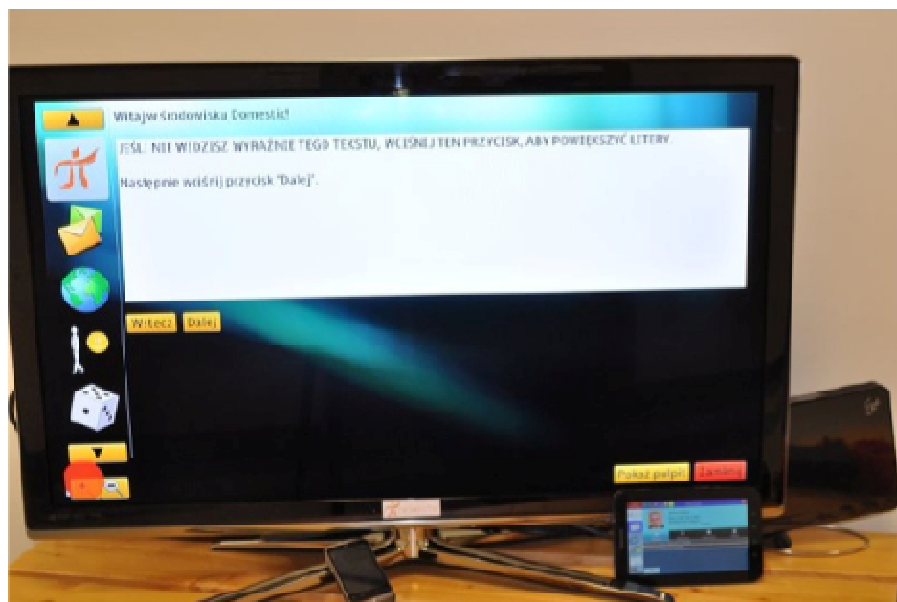


Interfejsy:

Bluetooth, Zigbee, WiFi, GSM, inne

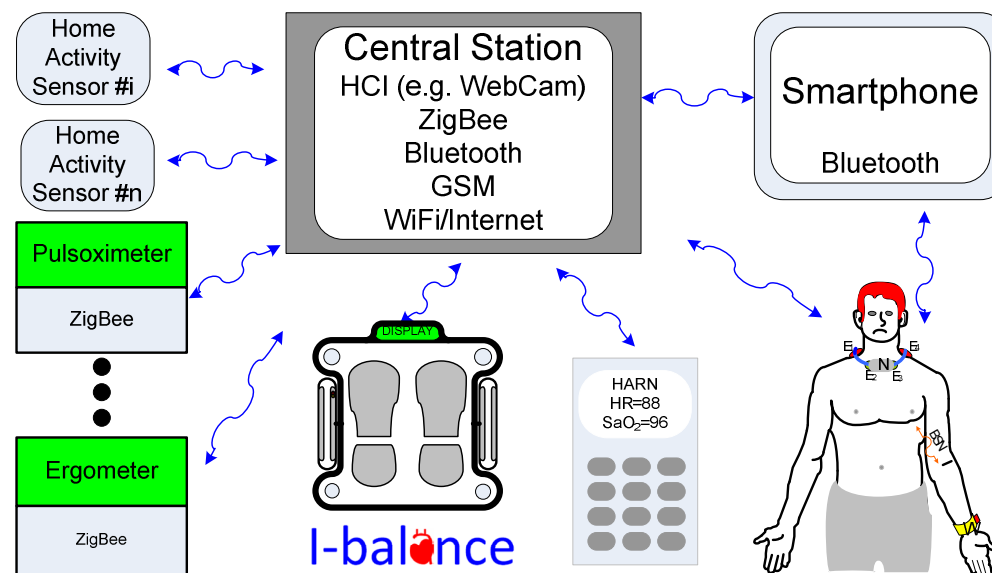
Funkcje:

Dane: zbieranie, przetwarzanie,
odkrywanie wiedzy, wizualizacja,
informacja (alerty), itp.





- Stacja centralna
- Waga
- Krzesło (narzuta)
- Pilot
- Wanna
- SleAP
- PathMon
- Dmuchawka
- Systemy wykorzystujące Kinect
- Inne (inteligentne gniazdka, czujniki aktywności, ...)





Urządzenie składa się z układu pomiaru masy za pomocą czujników tensometrycznych oraz rozbudowane jest o układ umożliwiający dodatkowe pomiary: Body Mass Index i sześciu odprowadzeń elektrokardiograficznych optyczny detektor dla pomiaru saturacji krwi SaO_2 oraz o układ transmisji bezprzewodowej do komputera-laptopa



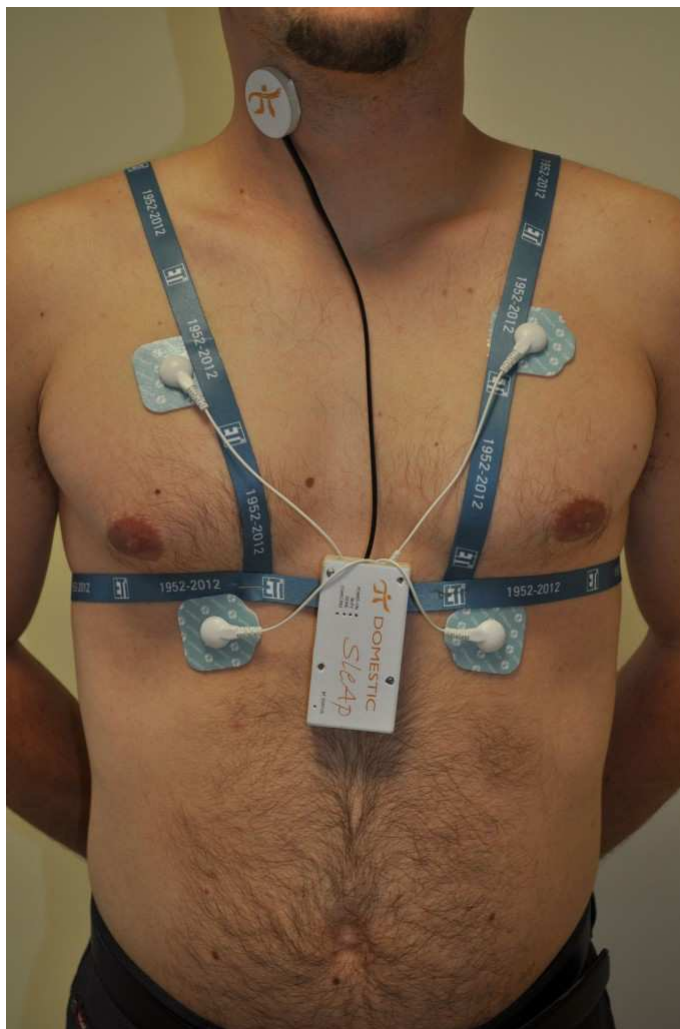
Poza przetwarzaniem EKG narzuta wyposażona jest w dwusektorowy, ciśnieniowy czujnik położenia osoby siedzącej. Wartości pomiarowe wskazują na aktywność fizyczną osoby siedzącej i są raportowane do stacji nadzorującej.



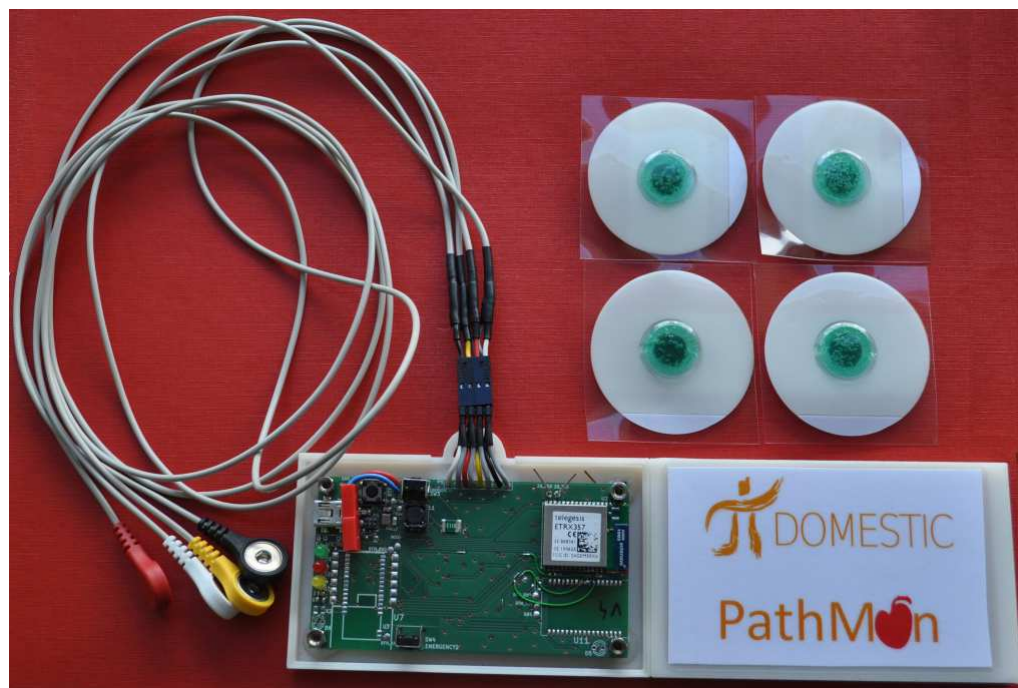
**Mobilny monitor pracy
serca, aktywności
oddechowej, ruchowej,
temperatury i postawy.**



- Kontroluje proces napełniania wanny – stabilizując temperaturę wody i jej poziom,
- Wykryje osobę w wannie oraz w przybliżeniu jej pozycję,
- Zna aktualny poziom i temperaturę wody w wannie,
- Zasugeruje moment opuszczenia kąpiel,
- Sprawdza czy osoba w wannie aktywnie się porusza,
- Zmierzy częstość oddechu osoby w wannie,
- Zmierzy częstość uderzeń serca,
- Wykryje moment wchodzenia i wychodzenia z wanny,
- Poinformuje opiekuna o wykrytym stanie zagrożenia.



Detekcja
bezdechu
za pomocą
pomiaru EKG i
impedancji
klatki piersiowej



Monitoring sygnału EKG oraz mechanicznej pracy serca (za pomocą pomiaru zmian impedancji klatki piersiowej)

Predykcja sygnałów

Klasyfikacja

Analiza EKG

Detekcja bezdechu

Modelowanie pomiarów impedancyjnych

W przypadku określenia liczby próbek branych do predykcji pożyteczna może być znajomość występowania trendów okresowych zmian ciśnienia, jak i innych sygnałów.

W celu przebadania, czy takie trendy występują przeprowadzono analizę widmową Fouriera.

Przyjmijmy, że funkcja jednej zmiennej $f(x)$ jest transformowalna (nie wnikając, jakie założenia muszą być spełnione, aby tak było).

Transformata Fouriera (prosta)

$$F(\omega) = \int_{-\infty}^{+\infty} f(x) e^{-i\omega x} dx$$

Transformata odwrotna

$$f(x) = \frac{1}{2\pi} \int_{-\infty}^{+\infty} F(\omega) e^{i\omega x} d\omega$$

Niech będzie dany sygnał y_m określony w równoodległych punktach t_m na przedziale $[0, 2\pi]$

Transformata prosta

$$Y_k = \frac{1}{M} \sum_{m=0}^{M-1} y_m e^{-ikt_m} = \frac{1}{M} \sum_{m=0}^{M-1} y_m e^{-ik \frac{2\pi}{M} m} \quad k = -M/2, \dots, M/2$$

$$t_m = 2\pi m / M, \quad m = 0, \dots, M-1$$

Dla nierównomiernych

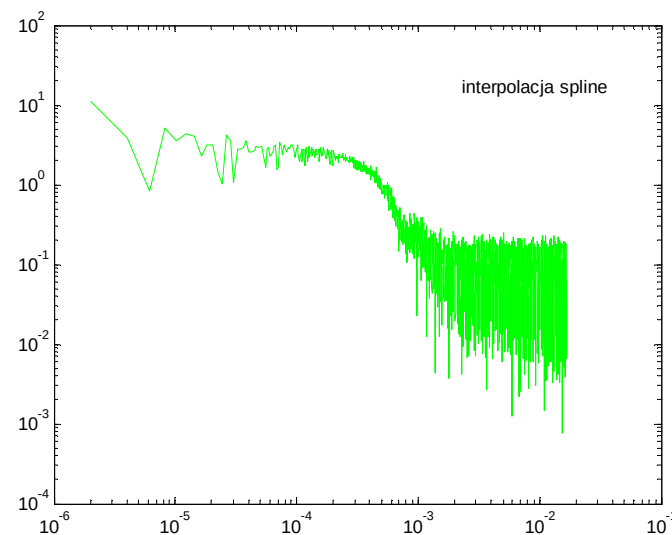
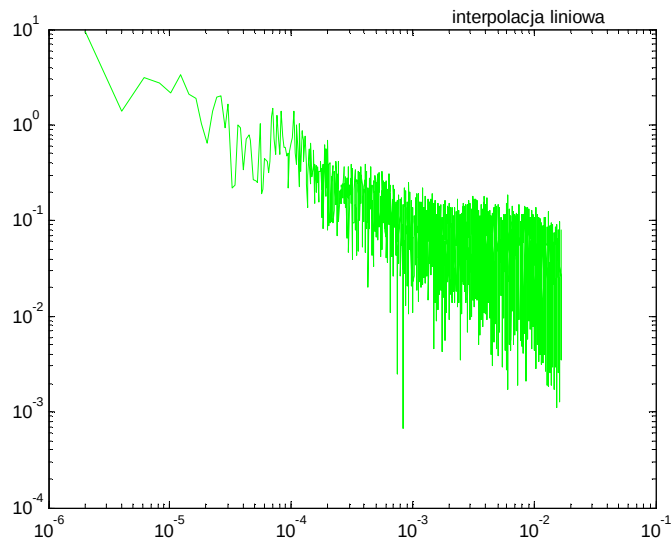
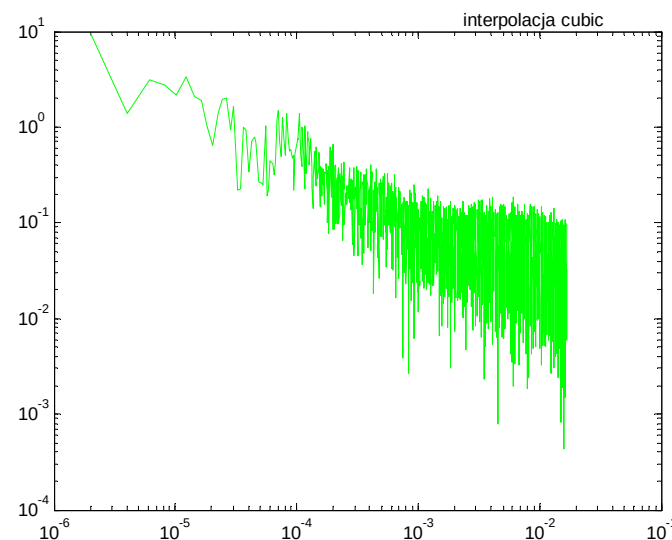
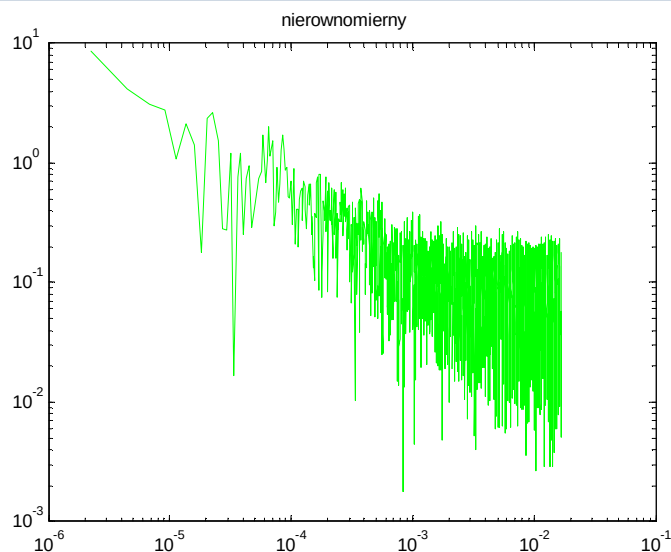
$$Y_k = \frac{1}{N} \sum_{n=0}^{N-1} y_n e^{-ikt_n} \quad k = -N/2, \dots, N/2$$

$$0 \leq t_n \leq 2\pi, \quad n = 0, \dots, N-1$$

Przykład – widmo sygnału ciśnienia

- analiza sygnału dla próbkowania nierównomiernego
- interpolacja i następnie analiza
 - liniowa,
 - spline – funkcje sklejące 3 stopnia,
 - cubic – interpolacja Hermita wielomianem stopnia 3.

Analiza Fouriera



Przykładowe wyniki dla ciśnienia skurczowego mierzonego inwazyjnie. Na osi poziomej częstotliwość w Hz.

- Wykorzystanie filtru adaptacyjnego
- Wykorzystanie rozwinięcia w szereg Taylora
- Ekstrapolacja
- Metody łączone
- Łańcuchy Markowa

- **Model**

$$p(n | P_{n-1}) = \sum_{k=1}^N a_k p(n-k)$$

- **Minimalizujemy**

$$E[(p(n) - p(n | P_{n-1}))^2]$$

$$E[p(n-1)p(n)] = \begin{bmatrix} E[p(n-1)p(n)] \\ E[p(n-2)p(n)] \\ \dots \\ E[p(n-N-1)p(n)] \\ E[p(n-N)p(n)] \end{bmatrix} = \begin{bmatrix} R_{pp}(1) \\ R_{pp}(2) \\ \dots \\ R_{pp}(N-1) \\ R_{pp}(N) \end{bmatrix} = r$$

$$R = E[p(n-1)p^T(n-1)] = \begin{bmatrix} R_{pp}(0) & R_{pp}(1) & \dots & R_{pp}(N-2) & R_{pp}(N-1) \\ R_{pp}(1) & R_{pp}(0) & \dots & R_{pp}(N-3) & R_{pp}(N-2) \\ \dots & \dots & \dots & \dots & \dots \\ R_{pp}(N-2) & R_{pp}(N-3) & \dots & R_{pp}(0) & R_{pp}(1) \\ R_{pp}(N-1) & R_{pp}(N-2) & \dots & R_{pp}(1) & R_{pp}(0) \end{bmatrix}$$

$$Ra = r$$

$$R_{pp}(j-k) = R_{pp}(k-j) = E[p(n-k+1)p(n-j+1)]$$

- **Wada**
 - próbkowanie równomierne

- **Problem**
 - długość filtru
 - estymacja funkcji autokorelacji

Estymacja funkcji autokorelacji

– Możemy wybrać estymator obciążony

$$\hat{R}_{pp}(k) = \frac{1}{M} \sum_{n=k+1}^M p(n)p(n-k), \quad k = 0, 1, \dots, N$$

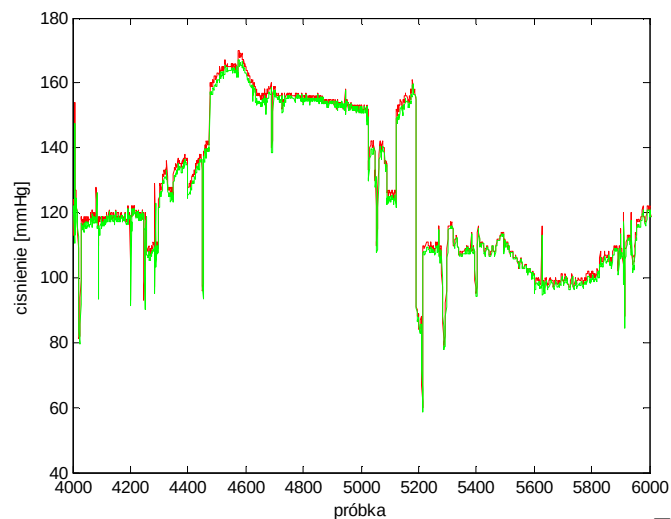
– Analogicznie możemy wybrać estymator nieobciążony postaci

$$\hat{R}_{pp}(k) = \frac{1}{M-k} \sum_{n=k+1}^M p(n)p(n-k), \quad k = 0, 1, \dots, N$$

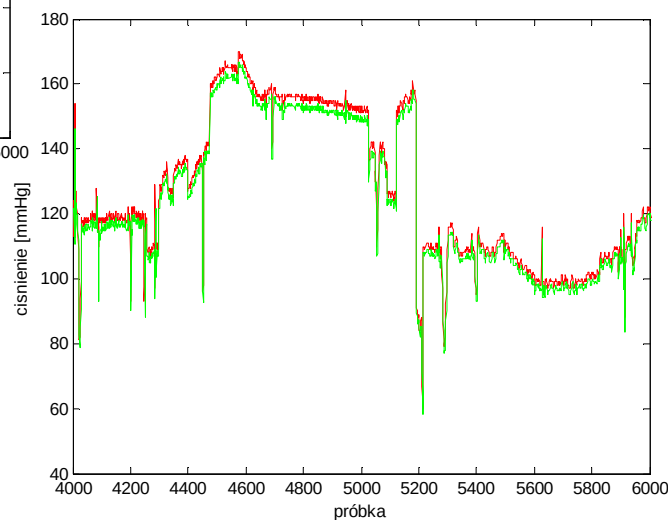
Zapominanie danych historycznych

$$\hat{R}_{pp}(k) = \frac{1}{M} \sum_{n=k+1}^M \lambda^{M-n} p(n) p(n-k), \quad k = 0, 1, \dots, N$$

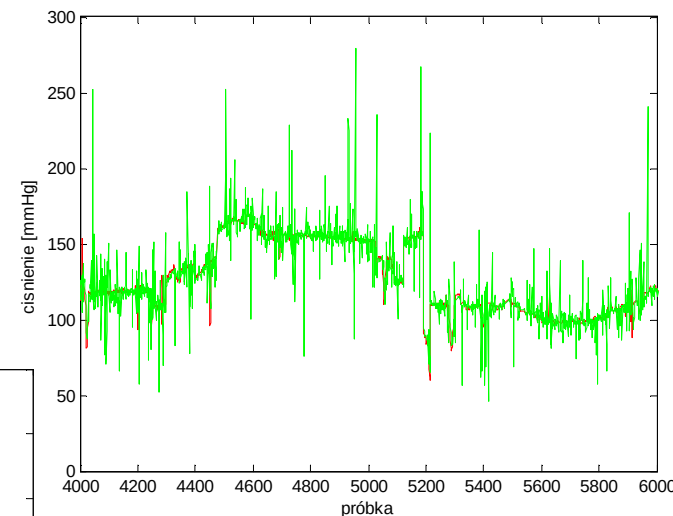
Filtr adaptacyjny



a)



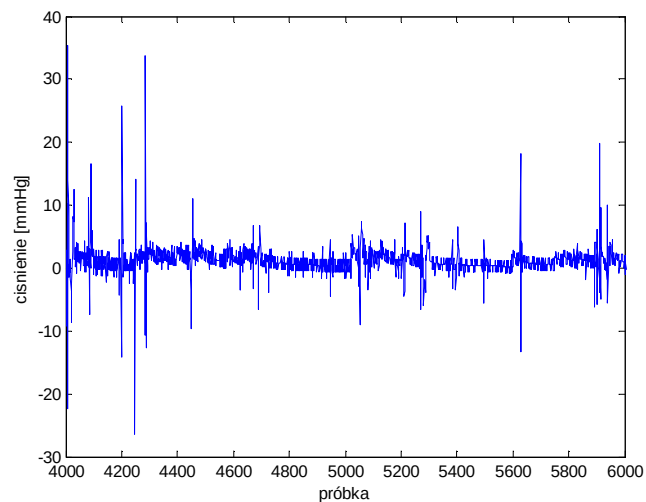
b)



c)

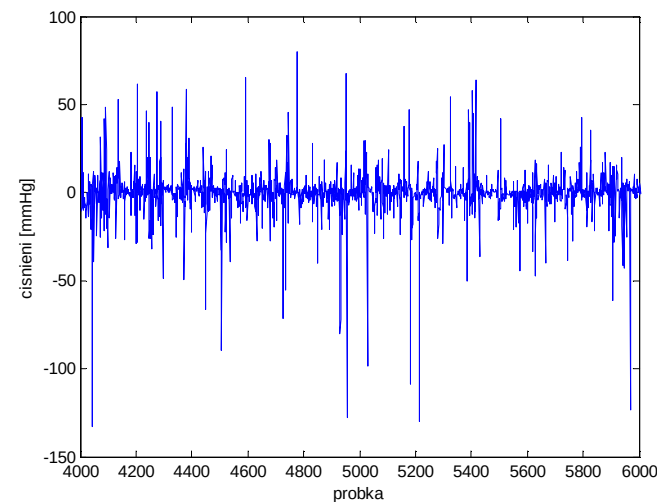
Wyniki dla predykcji ciśnienia tętniczego dla filtra długości 12 próbek. Rysunek a) pełna informacja, b) model z obcięciem o długości 5 krotność długości filtra, c) model z zapominaniem, współczynnik zapominania równy 0.9. Kolorem czerwonym oznaczono pomiary, a zielonym przewidywane wartości.

Filtr adaptacyjny

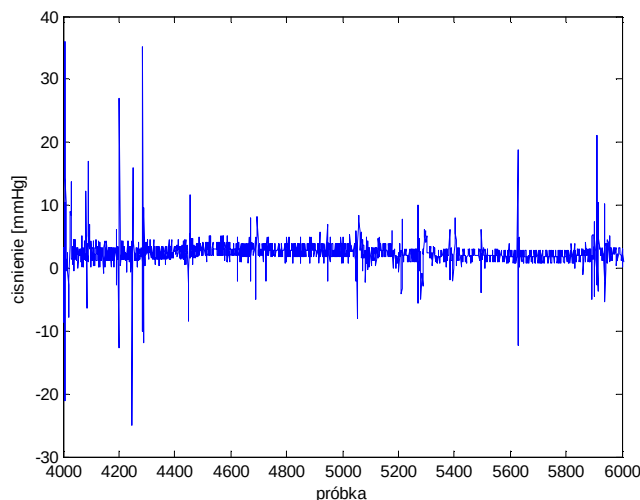


a)

b)



c)



Błędy dla predykcji ciśnienia tętniczego dla filtra długości 12 próbek. Rysunek a) pełna informacja, b) model z obcięciem o długości 5 krotność długości filtra, c) model z zapominaniem, współczynnik zapominania równy 0.9.

Prognozujemy $p(t+a)$ wykorzystując dane historyczne $p(t-b_k)$

$$p(t + a) = \sum_{k=0}^n \alpha_k p(t - b_k)$$

Rozwijamy $p(t+a)$ i $p(t-b_k)$ w szereg Taylora

$$p(t + a) = \sum_{i=0}^n \frac{p^{(i)}(t)}{i!} a^i$$

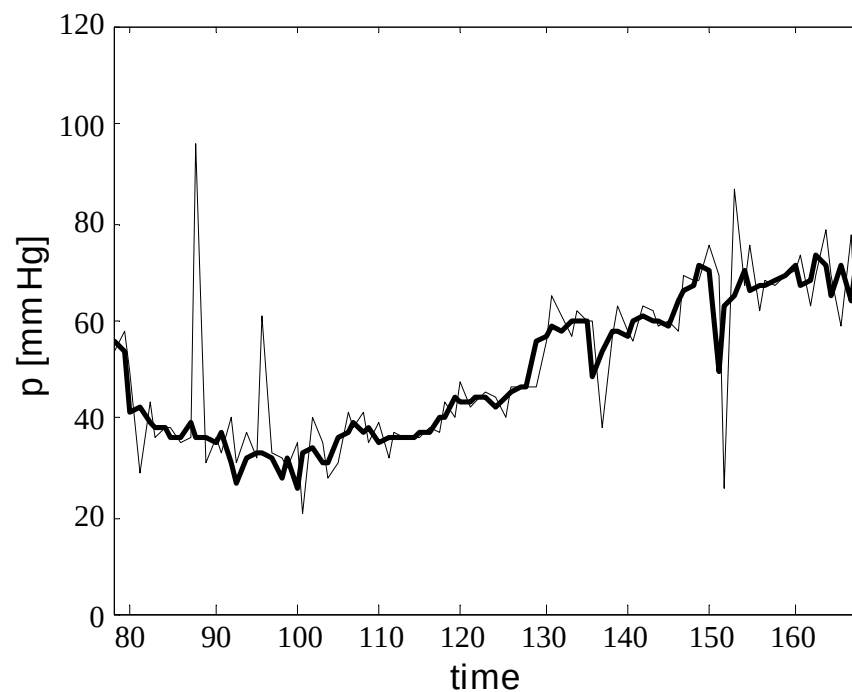
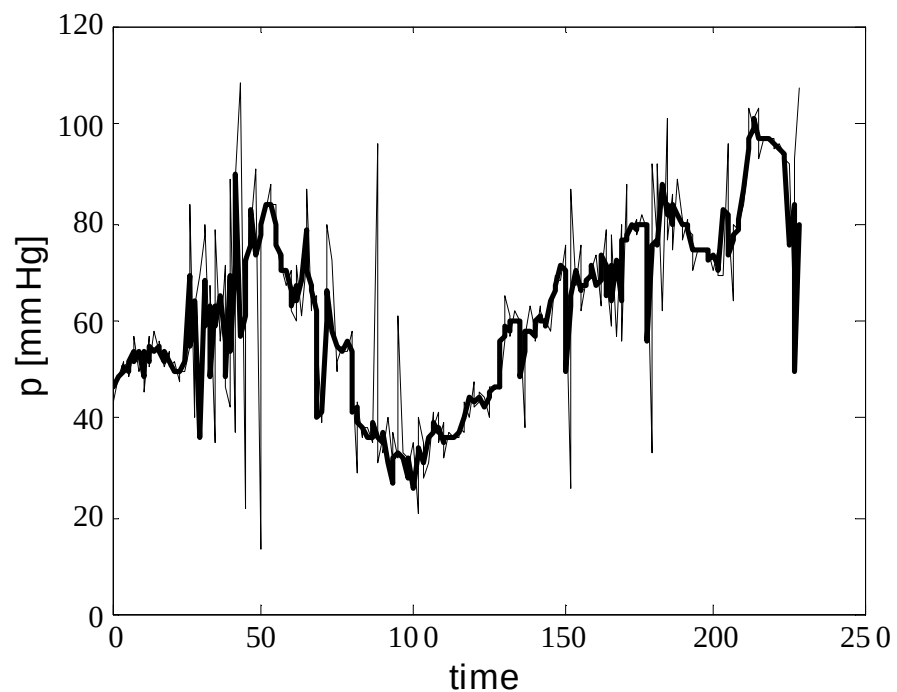
$$p(t - b_k) = \sum_{i=0}^n \frac{p^{(i)}(t)}{i!} (-1)^i b_k^i$$

Stąd

$$\begin{bmatrix} 1 & 1 & \dots & 1 & 1 \\ -b_1 & -b_2 & \dots & -b_{n-1} & -b_n \\ \dots & \dots & \dots & \dots & \dots \\ (-b_1)^{n-1} & (-b_2)^{n-1} & \dots & (-b_{n-1})^{n-1} & (-b_n)^{n-1} \\ (-b_1)^n & (-b_2)^n & \dots & (-b_{n-1})^n & (-b_n)^n \end{bmatrix} \begin{bmatrix} \alpha_1 \\ \alpha_2 \\ \dots \\ \alpha_{n-1} \\ \alpha_n \end{bmatrix} = \begin{bmatrix} a^0 \\ a^1 \\ \dots \\ a^{n-1} \\ a^n \end{bmatrix}$$

- **Zaleta**
 - nie zależy od momentu czasu
 - próbkowanie nierównomierne
- **Wady**
 - dobre dla sygnałów aproksymowanych wielomianem (możliwie niskiego stopnia)
 - wrażliwe na szumy

Przykład predykcji



Dane

f – dany sygnał o skończonej energii, o widmie zawartym w obszarze Ω

g - znany fragment sygnału f

Zadanie

odtworzyć f znając g

Metoda iteracyjna

$$y_0(x_j) = 0$$

$$y_{m+1}(x_j) = y_m(x_j) + \sum_{x_k \in B} \sin c_{\Omega}(x_j - x_k) (g(x_k) - y_m(x_k))$$

Sanz i Huang "Discrete and continuous band-limited Signac extrapolation"
IEEE Trans. Acoustics, Speech, and signal processing vol. assp-31, no.5, 1983,
pp 1276-1285

- **Stabilność dla funkcji *sinc***
- **Inne jądra h i regularyzacja λ**
- **Próbkowanie równomierne i nierównomierne**

**Podobne podejście jest przedstawione w pracy
T. Strohmer, On discrete band-limited signal
extrapolation, Contemp. Math. 190:323-337, 1995.**

**Zamiast iterować równanie splotowe rozwiązuje się
metodą sprzężonych gradientów problem**

$$p_0 = \lambda S p,$$

$$p_{m+1} = p_m + \lambda S(p - p_m)$$

**S jest operatorem, D jądrem Dirichleta, a r liczbą próbek
sygnału ($0 \leq t_1 < t_2 < \dots < t_r < 1$)**

$$2M+1 \leq r$$

$$Sp(t) = \sum_{j=1}^r p(t_j) D_M(t - t_j)$$

$$D_M(t) = \sum_{k=-M}^M e^{2\pi i k t} = \frac{\sin(2M+1)\pi t}{\sin \pi t}$$

Inne podejście jest w pracy

**H.G. Feichtinger, K.Gröchenig, and T.Strohmer,
Efficient Numerical Methods in Non-uniform Sampling
Theory, Numerische Mathematik, 69:423-440, 1995.**

Aproksymują poszukiwaną funkcję wzorem

$$p(t) = \sum_{k=-M}^M a_k e^{2\pi i k t}$$

**gdzie a_k są to nieznane współczynniki aproksymacji
(trzeba je wyznaczyć)**

wtedy

$$Sp(t) = \sum_{j=1}^r p(t_j) D_M(t - t_j) = \sum_{j=1}^r \sum_{k=-M}^M a_k e^{2\pi i k t_j} \sum_{l=-M}^M e^{2\pi i l t} e^{-2\pi i l t_j}$$

Przestawiając sumy dostajemy

$$= \sum_{l=-M}^M \left[\sum_{k=-M}^M \left(\sum_{j=1}^r e^{-2\pi i (l-k) t_j} \right) a_k \right] e^{2\pi i l t}$$

Zdefiniujmy macierz

$$T_{lk} = \sum_{j=1}^r e^{-2\pi i(l-k)t_j} \quad |l|, |k| \leq M$$

Korzystając z tego mamy

$$Sp(t) = \sum_{l=-M}^M \left[\sum_{k=-M}^M T_{lk} a_k \right] e^{2\pi i l t}$$

Z drugiej strony

$$Sp(t) = \sum_{j=1}^r p(t_j) D_M(t - t_j) = \sum_{k=-M}^M \left(\sum_{j=1}^r p(t_j) e^{-2\pi i k t_j} \right) e^{2\pi i k t}$$

Definiując

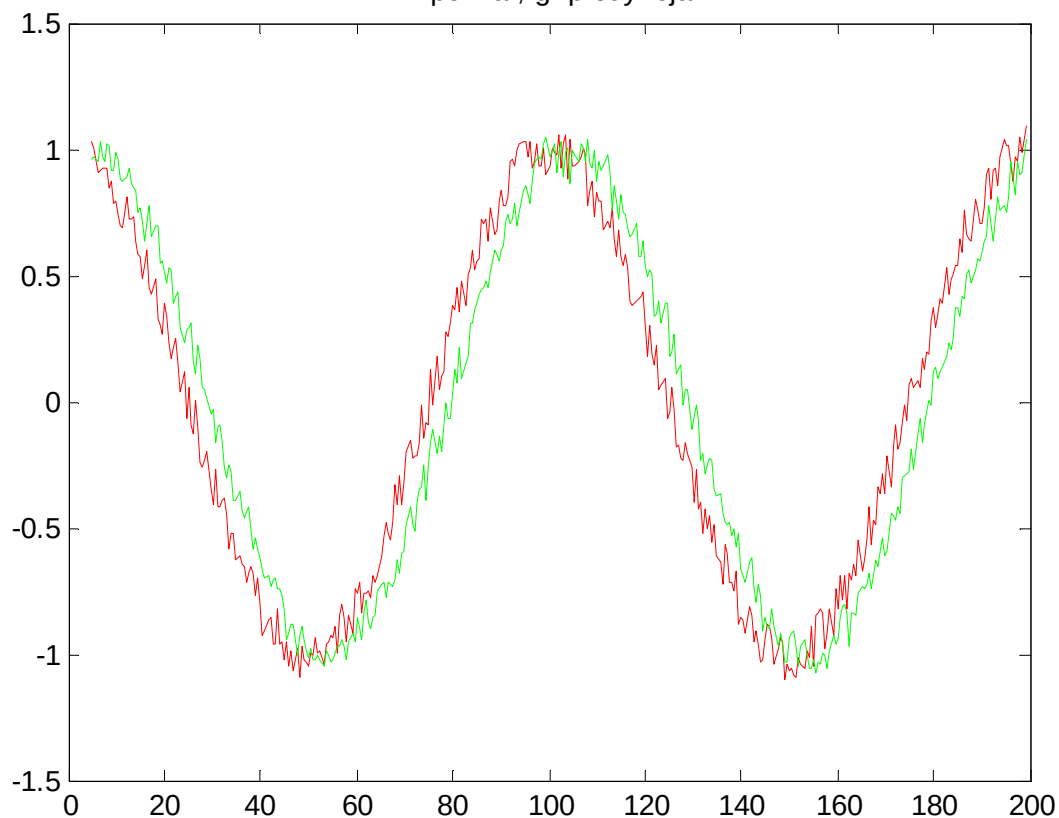
$$b_k = \sum_{j=1}^r p(t_j) e^{-2\pi i k t_j} \quad |k| \leq M$$

Dostajemy układ $Ta=b$, z którego wyznaczamy a , a to daje aproksymację p .

To jest prostsza wersja, którą można rozszerzyć na próbkowanie nierównomierne (algorytm jest podany w cytowanej pracy)

Przykładowy wyniki

r- pomiar, g- predykcja



czerwony - oryginalny sygnał, zielony – predykcja

Rozważamy dwa modele

$$s(t + a) = \sum_{k=0}^n \alpha_k s(t - b_k)$$

$$s(t + a) = \sum_{k=0}^n \alpha_k s(t - b_k) + \alpha_{n+1}$$

$$a, b_k > 0, b_i \neq b_j \text{ dla } i \neq j$$

Rozwinięcie w szereg Taylora - modyfikacja

$$\begin{bmatrix} 1 & 1 & \dots & 1 & 1 \\ -b_0 & -b_1 & \dots & -b_{n-1} & -b_n \\ \dots & \dots & \dots & \dots & \dots \\ (-b_0)^{m-1} & (-b_1)^{m-1} & \dots & (-b_{n-1})^{m-1} & (-b_n)^{m-1} \\ (-b_0)^m & (-b_1)^m & \dots & (-b_{n-1})^m & (-b_n)^m \end{bmatrix} \cdot \begin{bmatrix} \alpha_0 \\ \alpha_1 \\ \dots \\ \alpha_{n-1} \\ \alpha_n \end{bmatrix} \approx \begin{bmatrix} a^0 \\ a^1 \\ \dots \\ a^{m-1} \\ a^m \end{bmatrix}$$

$$\left\| s(t + \Delta t) - \sum_{k=0}^n \alpha_k s(t - k\Delta t) \right\|_2^2 =$$

$$\sum_{i=0}^m \left(s(t_i + \Delta t) - \sum_{k=0}^n \alpha_k s(t_i - k\Delta t) \right)^2$$

$$\left\| s(t + \Delta t) - \sum_{k=0}^n \alpha_k s(t - k\Delta t) - \alpha_{n+1} \right\|_2^2 =$$

$$\sum_{i=0}^m \left(s(t_i + \Delta t) - \sum_{k=0}^n \alpha_k s(t_i - k\Delta t) - \alpha_{n+1} \right)^2$$

Problem - jak wybrać m i n ?

- Zmiana parametrów pacjenta w czasie – jak szybka?
- Dokładność pomiarów

Przyjmujemy

- $n=1, 2, \dots, 4$
- $m=n, 2(n+1)-1, 3(n+1)-1, \dots, 5(n+1)-1$

Algorytm konkurencyjny

$$\hat{s}(t) = \begin{cases} \hat{s}_1(t) & \text{dla } |\hat{s}_1(t - \Delta t) - s(t - \Delta t)| \leq |\hat{s}_2(t - \Delta t) - s(t - \Delta t)| \\ \hat{s}_2(t) & \text{dla } |\hat{s}_2(t - \Delta t) - s(t - \Delta t)| < |\hat{s}_1(t - \Delta t) - s(t - \Delta t)| \end{cases}$$

Można wykorzystać więcej metod

Kombinacja liniowa predyktorów

$$\hat{s}(t) = w_1(t)\hat{s}_1(t) + w_2(t)\hat{s}_2(t)$$

$$w_1(t) = \frac{\text{Var}((\hat{s}_2(t) - s(t)) / s(t))}{\text{Var}((\hat{s}_1(t) - s(t)) / s(t)) + \text{Var}((\hat{s}_2(t) - s(t)) / s(t))}$$

$$w_2(t) = \frac{\text{Var}((\hat{s}_1(t) - s(t)) / s(t))}{\text{Var}((\hat{s}_1(t) - s(t)) / s(t)) + \text{Var}((\hat{s}_2(t) - s(t)) / s(t))}$$

$$w_2(t) + w_1(t) = 1$$

Można wykorzystać więcej metod

Jak porównać wyniki dla danych różnej długości?

$$q = \sqrt{\frac{\sum_{i=1}^N (\hat{s}(t_i) - s(t_i))^2}{N}}$$

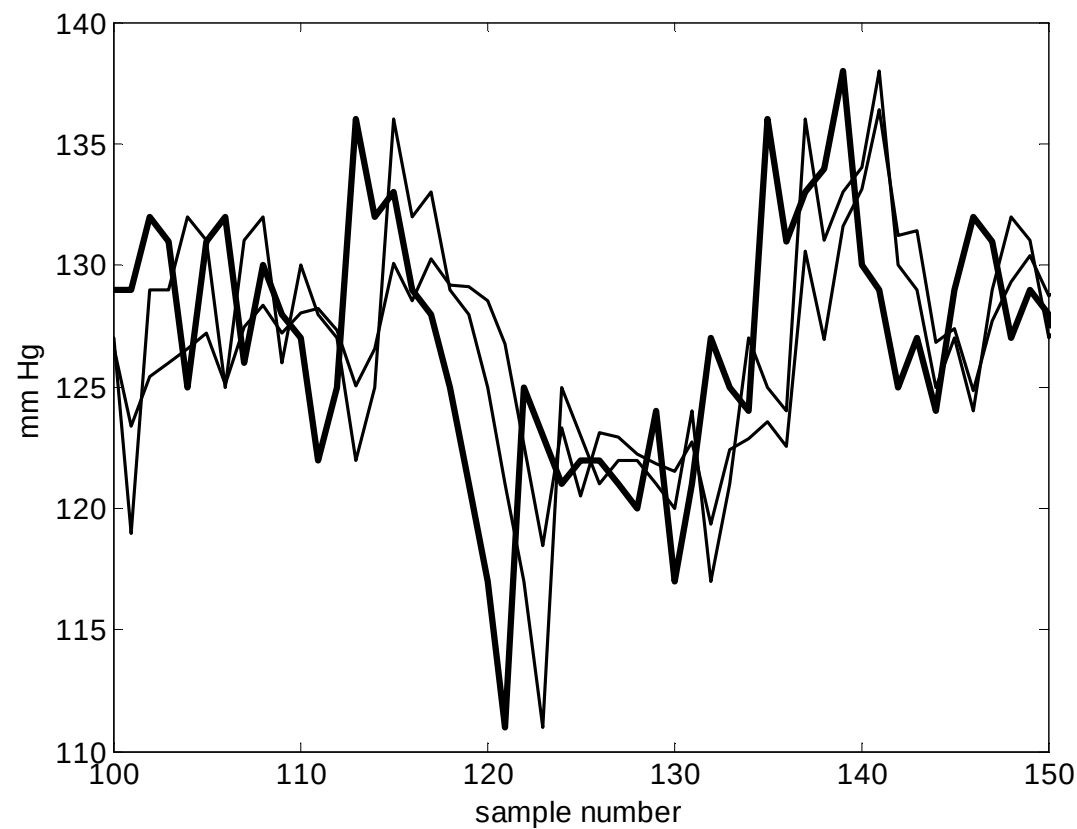
Dane pomiarowe:

**Trzy kobiety 47, 52 i 70 lat oraz jedno
mężczyzna w wieku 86 lat.**

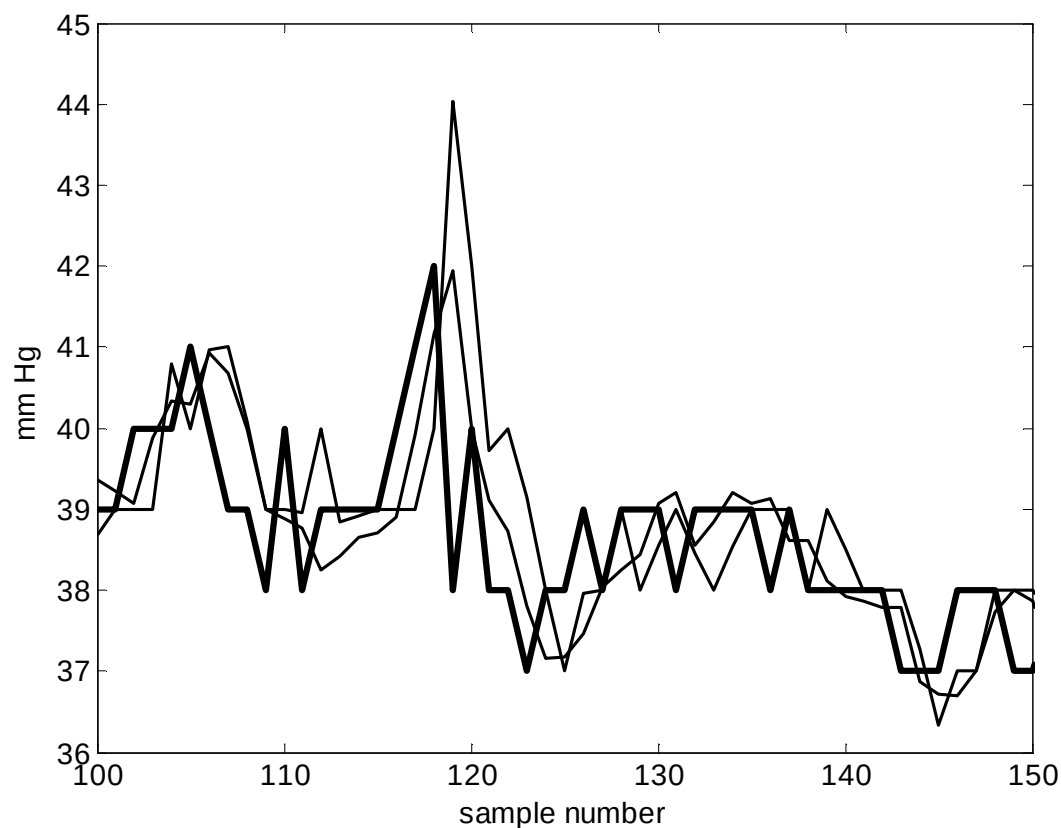
**Ciśnienie skurczowe i rozkurczowe z tętnicy
promieniowej.**

**Wybrane po 2 najdłuższe nieprzerwanie
mierzone wartości ciśnienia (od 257 do 1177
próbek co 30 sekund)**

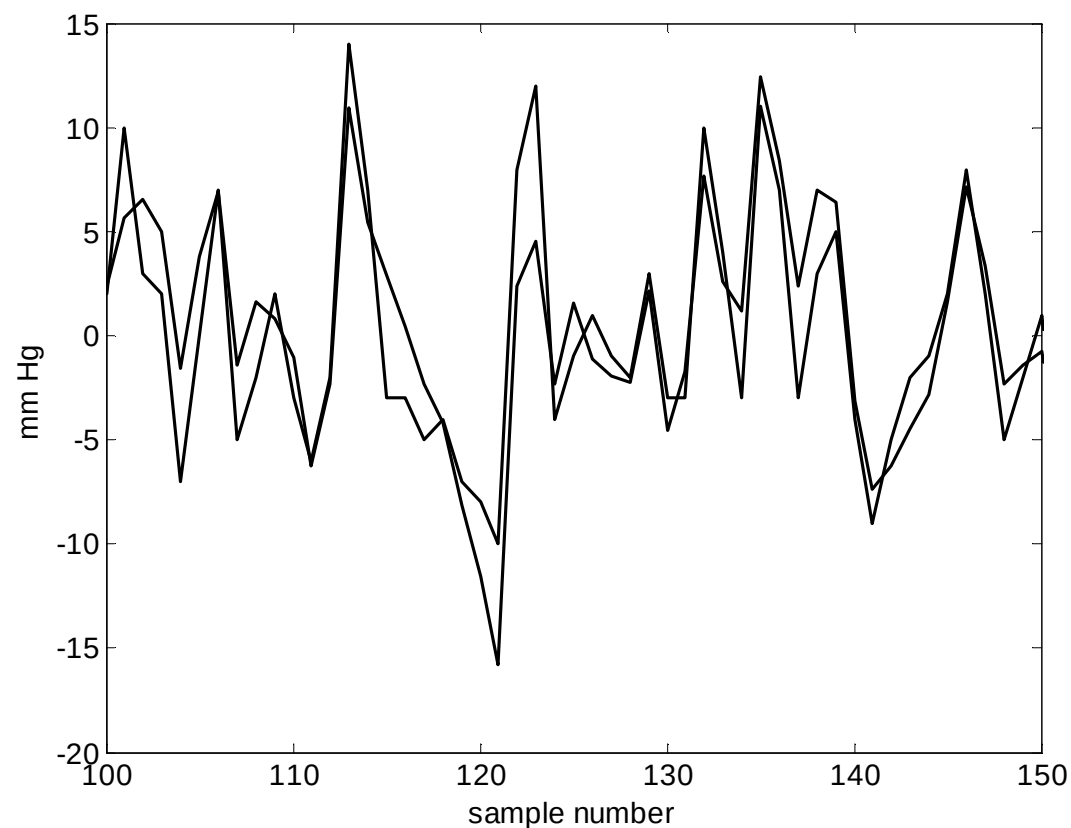
Przykład słabej predykcji – najlepsza i najgorsza



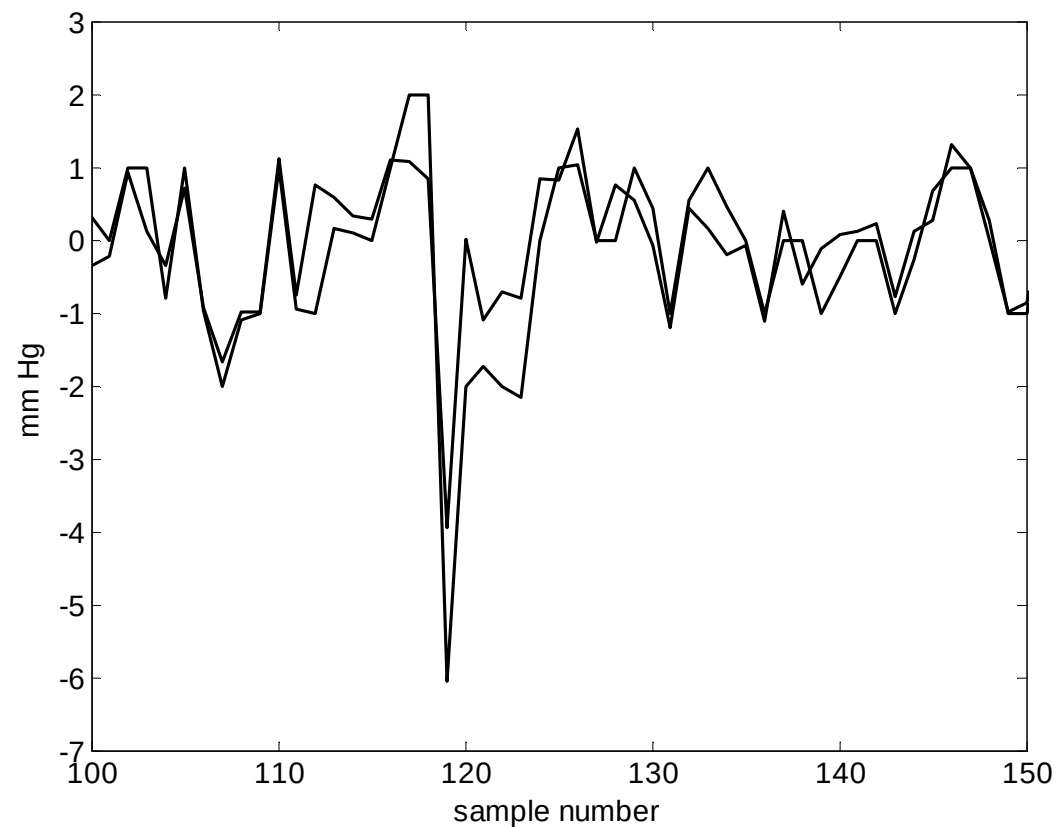
Przykład dobrej predykcji – najlepsza i najgorsza



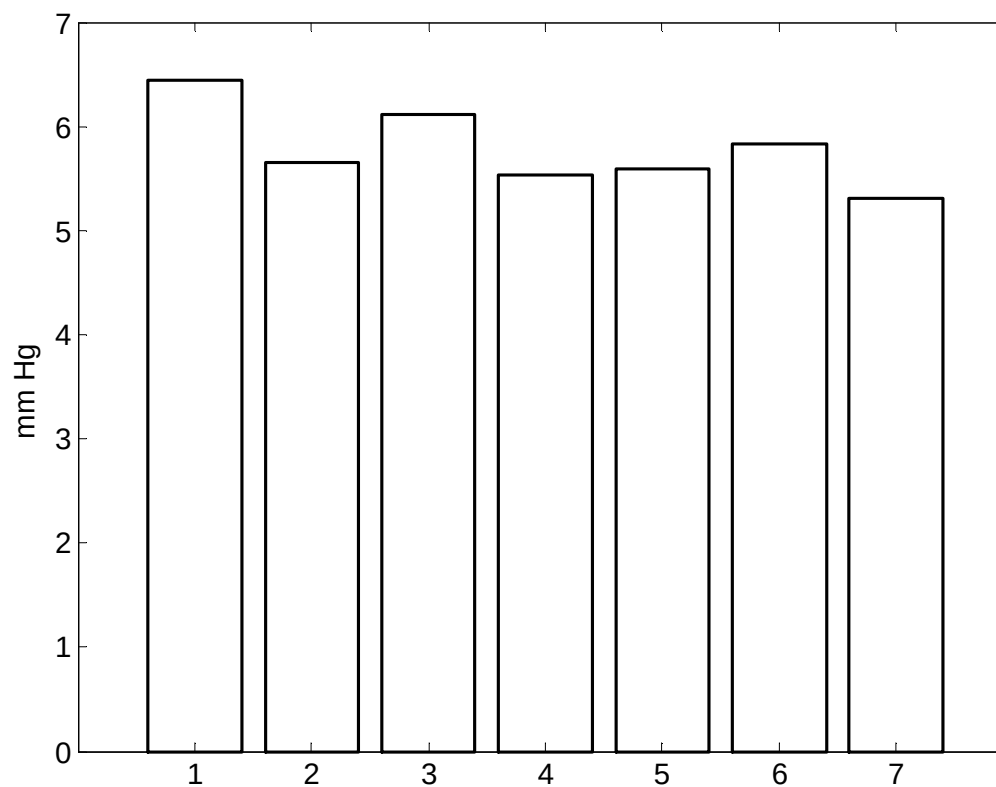
Przykład błędu dla słabej predykcji – najlepsza i najgorsza



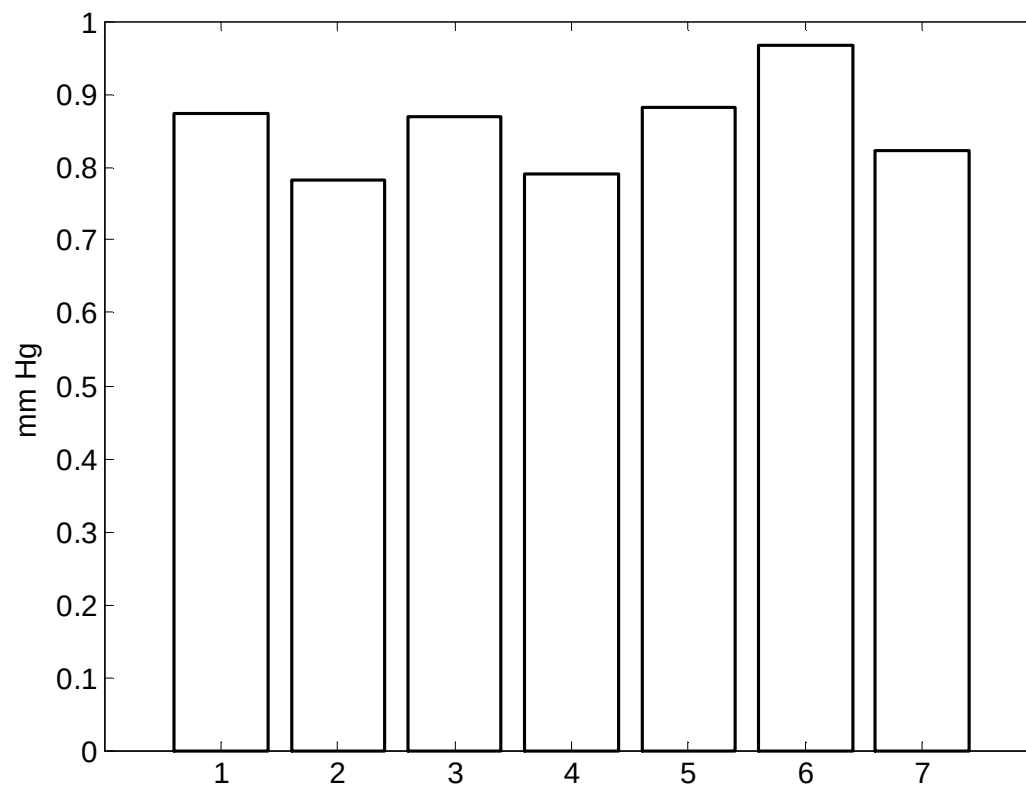
Przykład błędu dla dobrej predykcji – najlepsza i najgorsza



Błąd predykcji dla różnych metod dla słabej predykcji



Błąd predykcji dla różnych metod dla dobrej predykcji



- **Wiele czynników wpływających na mierzone parametry (np. ciśnienie)**
- **Czynniki zmienne w czasie – także zależność od leków**
- **Zależność wyników od stabilności pacjenta**
- **Nie wszystkie parametry wpływające na dany sygnał fizjologiczny są możliwe do monitorowania**
- **Błędy pomiarowe**

- **Brak najlepszej metody**
- **Zadowalająca dokładność predykcji w większości przypadków**

Cel

Nieinwazyjny ciągły monitoring ciśnienia

Czy można wykorzystać np. częstość akcji serca?

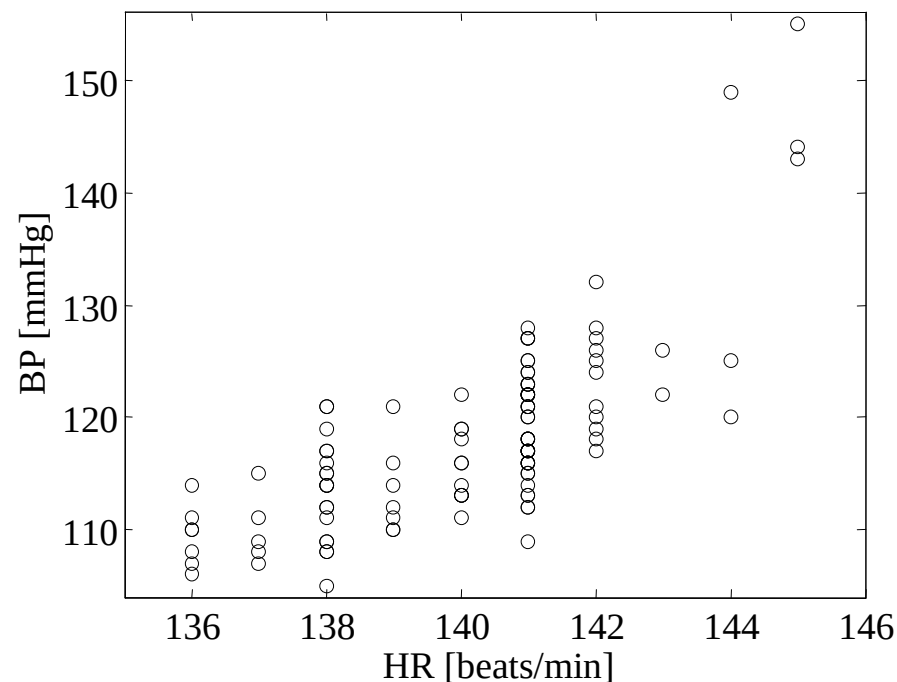
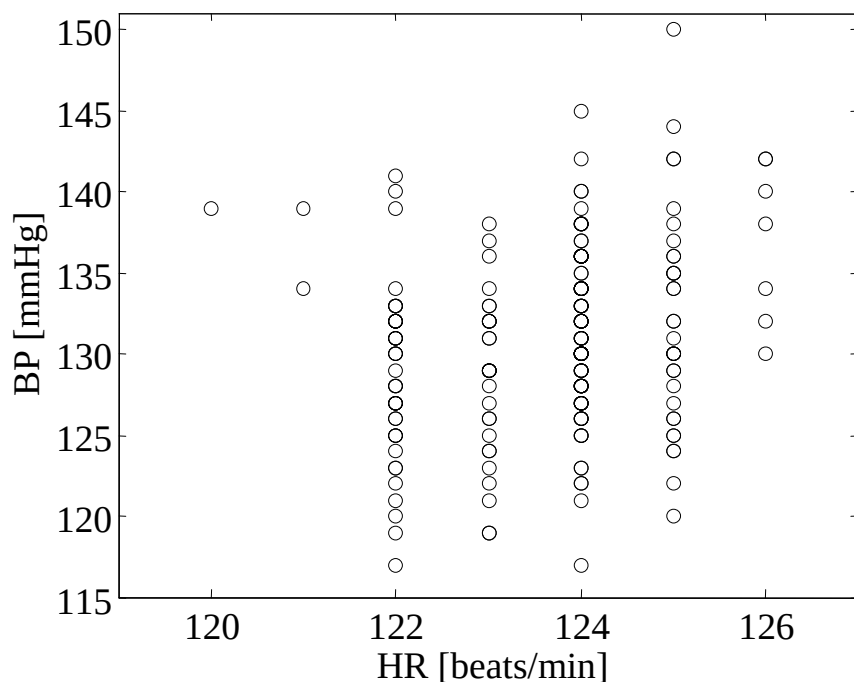
Korelacja

- surowe dane
- uśrednione dane (MA(3),MA(4),MA(5)) -
odszumianie
- współczynnik nachylenia prostej ($y=ax+b$)
 - trendy

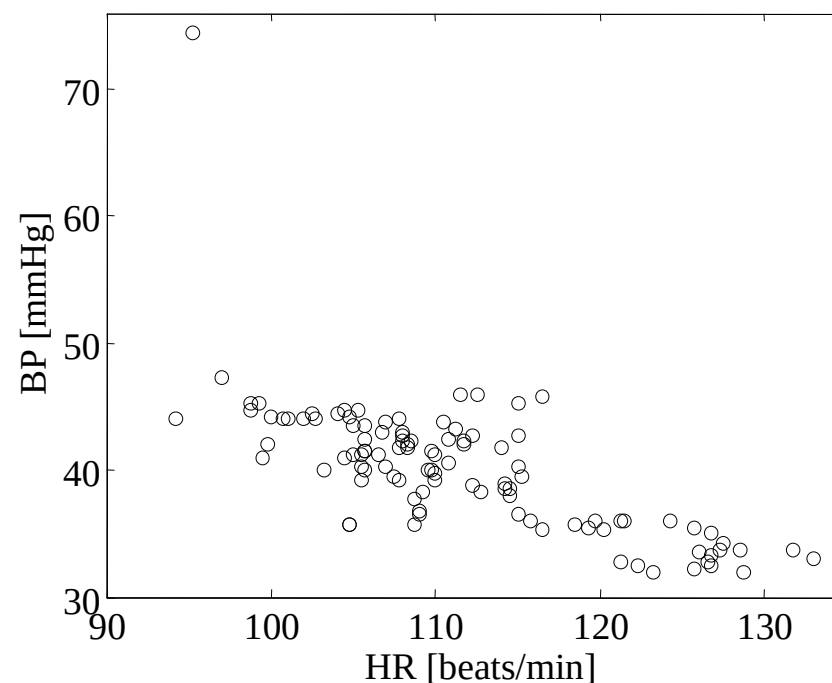
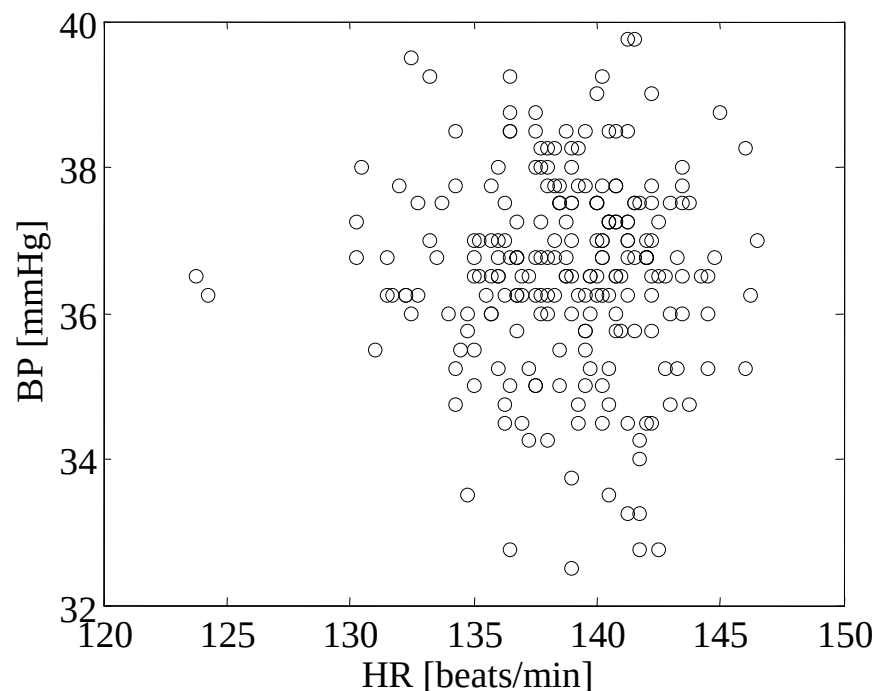
Analiza ciśnienia i częstości akcji serca ciągi min długości 100 próbek

SBP vs HR – ten sam pacjent, różne momenty czasu. Współczynnik korelacji 0.229 ($p < 0.001$) i 0.734 ($p < 0.001$).

Surowe dane

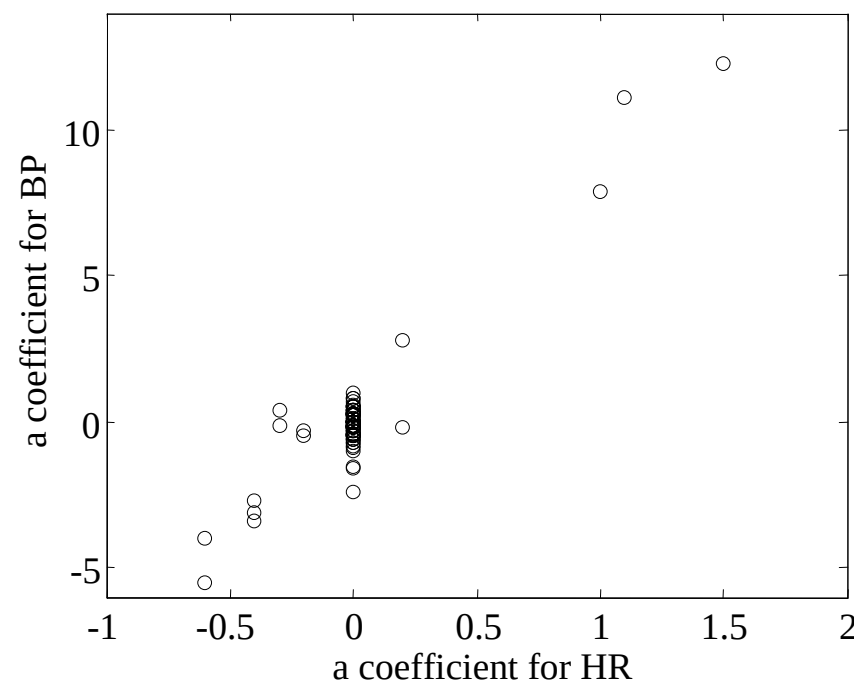
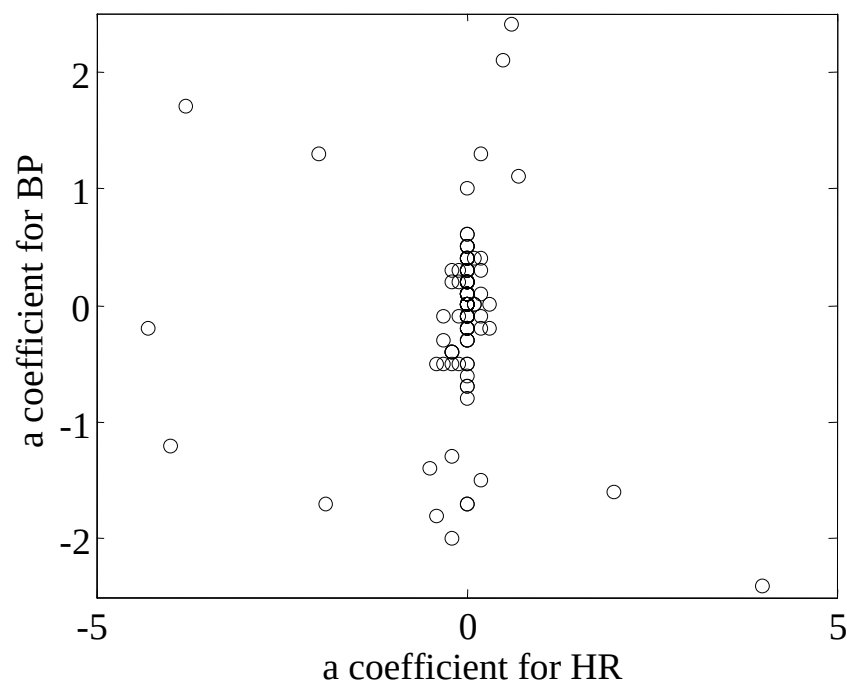


**DBP vs HR – ten sam pacjent, różne momenty czasu. Współczynnik korelacji -0.029 ($p > 0.05$) i -0.708 ($p < 0.001$).
Dane po filtracji**

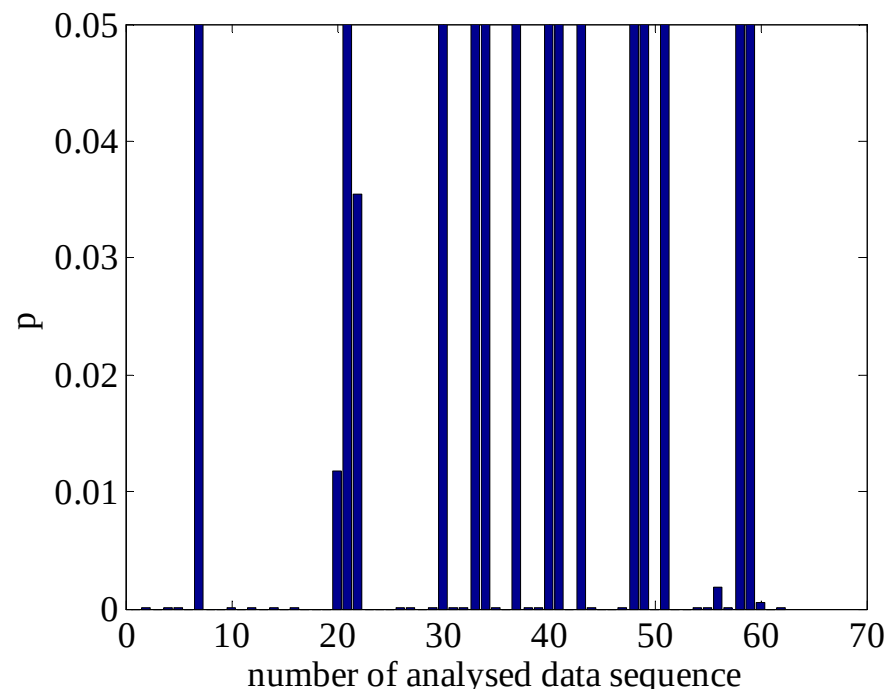
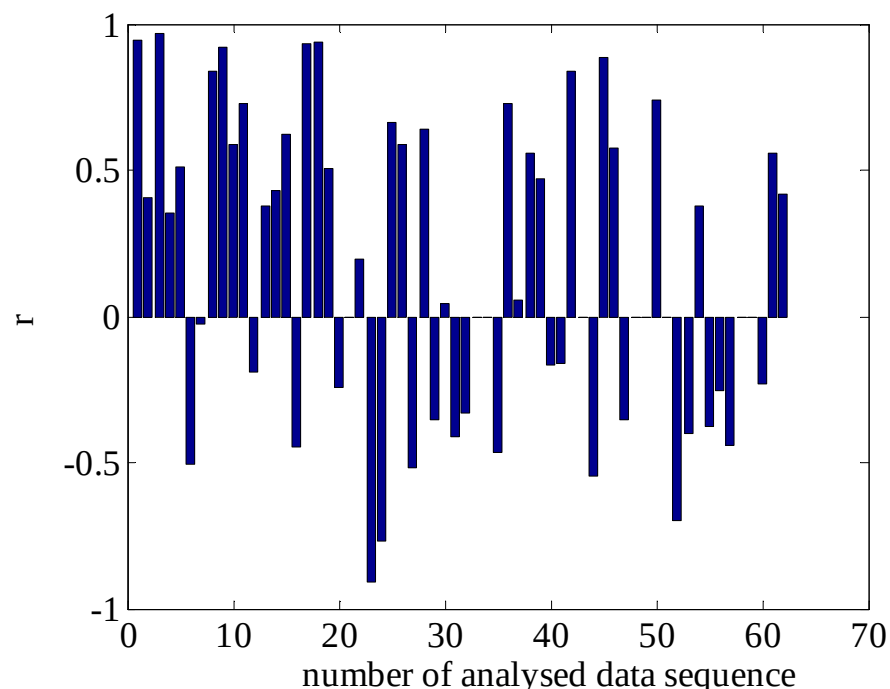


SBP vs HR – ten sam pacjent, różne momenty czasu. Współczynnik korelacji -0.161 ($p > 0.05$) i 0.939 ($p < 0.001$).

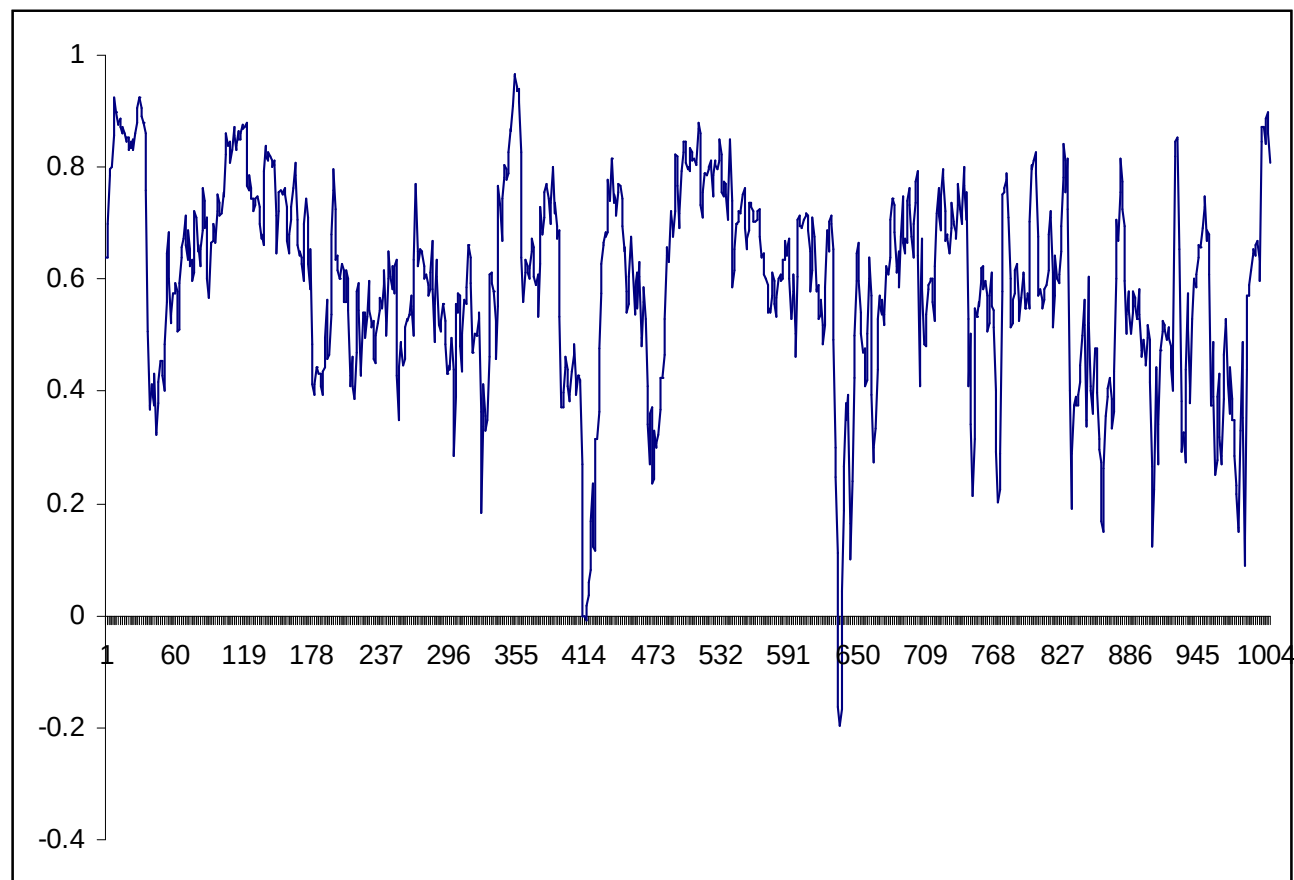
Trendy



Przykładowe wyniki zmian współczynnika korelacji pomiędzy SBP i HR w czasie (sygnały filtrowane)



Krótkookresowa korelacja



Duża zmienność

Zależność liniowa lub jej brak

Idea zaczerpnięta z pracy

**Baier V., Baumert M., Caminal P., Vallverdú
M., Faber R., Voss A.,
Hidden Markov Models Based on Symbolic
Dynamics for Statistical Modeling of
Cardiovascular Control in Hypertensive
Pregnancy Disorders,
IEEE Transactions on biomedical
engineering: vol. 53, nr 1 2006.**

$S = \{s_0, s_1, s_2, \dots, s_r\}$ – zbiór stanów

**Jeżeli łańcuch jest aktualnie w stanie s_i
i przechodzi do stanu s_j z prawdopodobieństwem p_{ij}
oraz to prawdopodobieństwo nie zależy od stanów
w jakich łańcuch znajdował się poprzednio, to
mówimy, że proces jest łańcuchem Markowa.**

Ciąg zmiennych losowych $(X_n)_{n=0}^{\infty}$ o wartościach w przeliczalnym zbiorze S (przestrzeni stanów) nazywamy łańcuchem Markowa wtedy i tylko wtedy, gdy dla każdego $n \in N$ i każdego ciągu $s_0, s_1, s_2, \dots, s_n \in S$ mamy

$$\begin{aligned} P(X_n = s_n | X_{n-1} = s_{n-1}, \dots, X_1 = s_1, X_0 = s_0) \\ = P(X_n = s_n | X_{n-1} = s_{n-1}), \end{aligned}$$

jeżeli

$$P(X_{n-1} = s_{n-1}, \dots, X_1 = s_1, X_0 = s_0) > 0.$$

Macierz $P = (p_{ij})_{s_i, s_j \in S}$ nazywamy macierzą przejścia na S , gdy wszystkie jej wyrazy są nieujemne, a ponadto suma każdego wiersza wynosi 1, czyli dla każdego

$$s_i \in S$$

$$\sum_{s_j \in S} p_{ij} = 1$$

Jeżeli $(X_n)_{n=0}^{\infty}$ jest łańcuchem Markowa, to rozkład zmiennej losowej X_0 nazywamy rozkładem początkowym. Łańcuch Markowa nazywamy jednorodnym (jednorodnym w czasie), gdy istnieje macierz

$$\hat{P} = (\hat{p}_{ij})_{s_i, s_j \in S}$$

będąca dla każdego n jego macierzą przejścia w n -tym kroku.

Wektorem stanu w chwili t nazywamy wektor postaci

$$y(t) = (y_1(t), y_2(t), \dots, y_r(t)),$$

**gdzie $y_i(t)$ dla $i \in \{1, 2, \dots, r\}$ jest
prawdopodobieństwem znalezienia się w
stanie $s_i \in S$ w chwili t .**

Prognoza

$$y_p(T+1) = y(T)\hat{P}$$

$$y_p(T+1) = y(T)\hat{P}(T+1)$$

Histogram

Wybór szerokości klasy h

$$f_n(x) = \frac{N \leftarrow x_i \in x}{nh}$$

$$N \leftarrow x_i \in x$$

liczba wartości x_i należących do x ,

n – liczba przedziałów

Ocena dopasowania estymatora

$$R(f_n) = E \int_{-\infty}^{\infty} [f_n(x) - f(x)]^2 dx$$

Jakie h jest optymalne?

$$h_0 = \frac{c}{\sqrt[3]{n}}$$

Dla dużych wartości n

$$c = \sqrt[3]{6 \int [f'(x)]^2 dx}$$

Jeżeli rozkład jest normalny to

$$c = 2\sqrt[3]{3^6\sqrt{\pi}\sigma} \approx 3.486\sigma$$

$$S_n = \sqrt{\frac{1}{n} \sum_{i=1}^n (X_i - \bar{X}_n)^2}$$

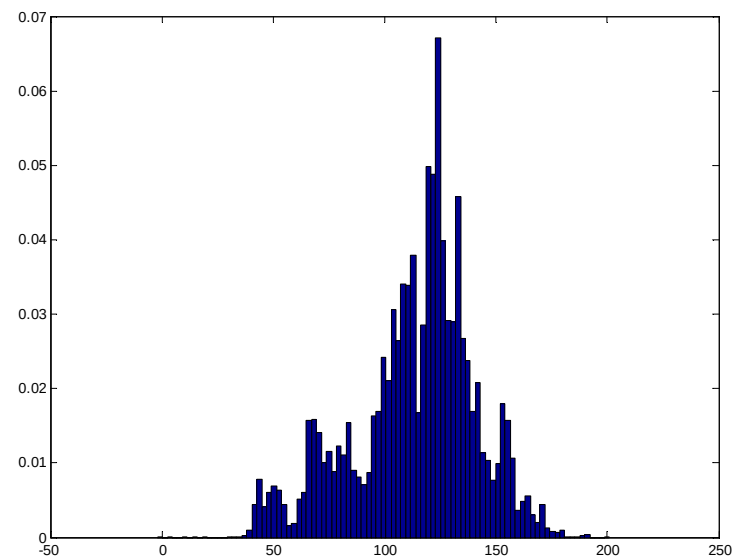
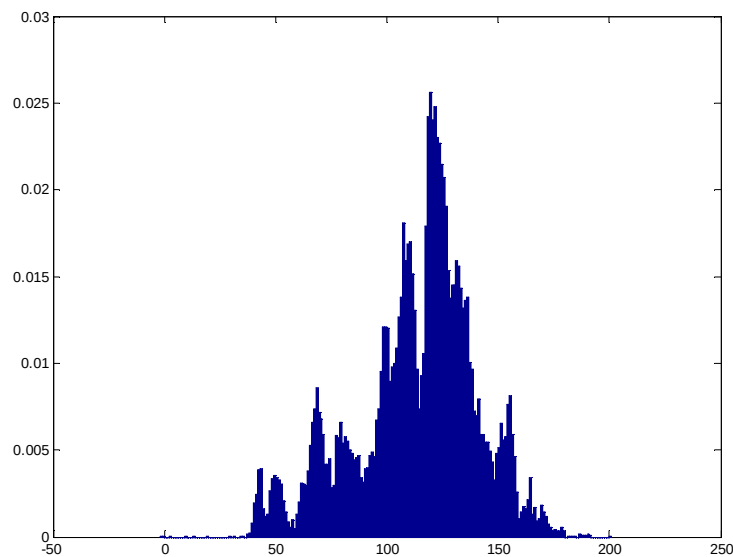
Inny dobór

**David Freedman and Persi Diaconis,
On the Histogram as a Density
Estimator: L_2 Theory,
Z. Wahrscheinlichkeitstheorie verw.
Gebiete 57, 453-476 (1981)**

Brak znajomości c , a więc

**Choose the cell width as twice the
interquartile range of the data,
divided by the cube root of the sample
size.**

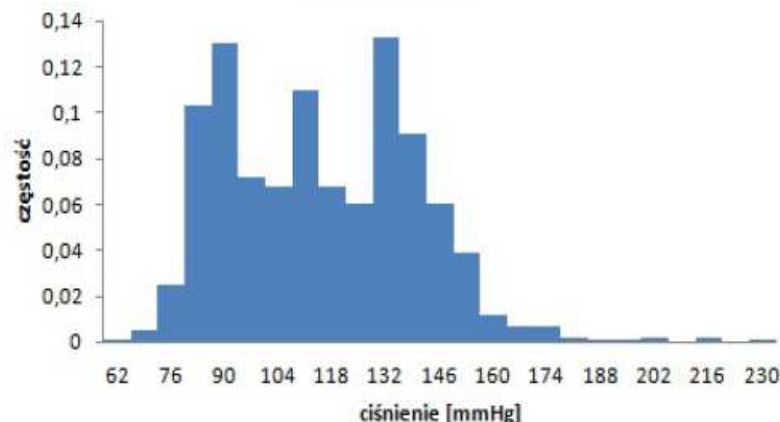
Histogramy – te same dane, różna ilość podprzedziałów



Przykładowe wyniki

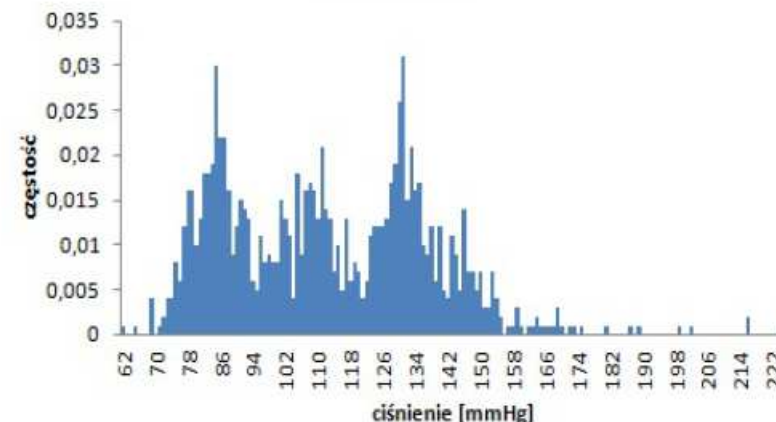
$$h_* = 8,7 \approx 9$$

Próbka 1.



$$h = 1,3 \approx 1$$

Próbka 1.



* oznacza h wyznaczone z założeniem rozkładu normalnego,
bez * - estymacja z wykorzystaniem przybliżenia pochodnych

Przykładowa estymowana macierz przejścia

[illegible]

Problemy

- wyznaczeniem rozkładu
- weryfikacją jednorodności łańcuchów Markowa
 - testy wykazały jej brak, nawet przy analizie uwzględniającej rytm dobowy

**Elementami zbioru treningowego X_T są
instancje (x_i, k_i) będące listą wartości
atrybutów i wybranego
atrybutu decyzyjnego - klasy**

$$X_T = \{(x_1, k_1), (x_2, k_2), \dots, (x_M, k_M)\}, x_i \in X \subset R^N, k_i \in \{k_1, k_2, \dots, k_K\}$$

Przykładowy wektor

$$x_i = (a_{i1}, a_{i2}, \dots, a_{iN})$$

Klasyfikator przypisuje nieznanym obserwacjom określoną klasę. Przypisanie to dokonywane jest na podstawie funkcji $f(x)$, za której zbudowanie odpowiada klasyfikator.

Najprostszy przykład – klasyfikator binarny

Zakładamy dwie klasy – stan pacjenta stabilny oraz stan pacjenta niestabilny

Jeśli instancja jest przypisana do klasy niestabilnej i klasyfikator przewiduje dla niej taką klasę, to zalicza się ją do *true positive* (TP). Jeśli przypisuje ją do klasy stabilnej – to liczy się ją jako *false negative* (FN). Jeśli instancja w rzeczywistości posiada klasę stabilną i zostaje sklasyfikowana tak samo, to zalicza się ją do *true negative* (TN). W przeciwnym przypadku liczy się ją jako *false positive* (FP)

Zakładamy, że mamy dane prawdopodobieństwa przynależności do klas $p(k_1), p(k_2), \dots, p(k_N)$ oraz prawdopodobieństwa warunkowe wystąpienia wektora obserwacji x w poszczególnych klasach $p(x|k_1), p(x|k_2), \dots, p(x|k_N)$.

Korzystając z twierdzenia Bayesa można obliczyć prawdopodobieństwo wystąpienia klasy k pod warunkiem zaobserwowania wektora x :

$$p(k_i | x) = \frac{p(x | k_i) p(k_i)}{p(x)}$$

$$p(x) = \sum_{i=1}^N p(x | k_i)$$

Wyszukując maksymalne prawdopodobieństwo warunkowe możemy określić do jakiej klasy należą zaobserwowane dane. Problemem jest znajomość rozkładu prawdopodobieństwa.

$$f(x) = \arg \max_{k_i} p(k_i | x)$$

Przy konstrukcji klasyfikatora można wziąć pod uwagę także różne koszty klasyfikacji (poprawne i błędne, różne dla różnych klas). Można to zrobić dodając macierz kosztów

Ferdinand van der Heijden, Image Based measurement Systems – object recognition and parameter estimation, John Wiley & Sons, 1994

Szukamy minimum funkcji kosztu

$$\begin{aligned} f(x) &= \arg \min_{\substack{k_i \\ 1 \leq i \leq N}} \left\{ \sum_{j=1}^N C(k_i, k_j) p(k_j | x) \right\} \\ &= \arg \min_{\substack{k_i \\ 1 \leq i \leq N}} \left\{ \sum_{j=1}^N C(k_i, k_j) \frac{p(x | k_j) p(k_j)}{p(x)} \right\} \\ &= \arg \min_{\substack{k_i \\ 1 \leq i \leq N}} \left\{ \sum_{j=1}^N C(k_i, k_j) p(x | k_j) p(k_j) \right\} \end{aligned}$$

W przypadku macierzy kosztów na diagonalu możemy mieć 0 (poprawne przypisanie do klasy), a poza nią wartości dodatnie.

W przypadku, gdy

$$C(k_i, k_j) = 1 - \delta(i, j)$$

(δ - delta Kroneckera) to nasz problem sprowadza się do klasycznego przypadku znalezienia maksimum funkcji

$$f(x) = \arg \max_{k_i} p(k_i | x)$$

W przypadku zastosowania tzw. „rejection option”
dodajemy klasę k_0

$$k = k_0 \quad C_{rej} < e(x)$$

$$k = \arg \max_{\substack{k_i \\ 1 \leq i \leq N}} p(k_i | x) \quad C_{rej} \geq e(x)$$

gdzie

$$e(x) = 1 - \max_{\substack{k_i \\ 1 \leq i \leq N}} \{p(k_i | x)\}$$

C_{rej} jest to koszt odrzucenia

Klasyfikator naiwny – zakładamy wzajemną niezależność atrybutów

$$p(x_i | K_j) = \prod_{k=1}^N p(a_{ik} | k_j) = p(a_{i1} | k_j) \cdot p(a_{i2} | k_j) \cdot \dots \cdot p(a_{iN} | k_j)$$

Zaleta – redukcja wymiaru

Przekleństwo wymiarowości (ang. *curse of dimensionality*) - im więcej wymiarów tym bardziej unikalne stają się posiadane dane.

Można stosować klasyfikatory bayesowskie łączące dane zależne (pośrednie rozwiązanie)

Estymacja rozkładu

- **Histogram**
- **Estymator jądrowy**

Estymator jądrowy

Niech $(y_1, \dots, y_L)$ będzie próbką pobraną z rozkładu o nieznanej gęstości f , gdzie L to wielkość próby prostej pobranej z rozkładu. Wartość funkcji gęstości prawdopodobieństwa estymowanej funkcją gaussowską dla punktu wynosi

$$g(y, \mu, \sigma) = \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{(y-\mu)^2}{2\sigma^2}}$$

Typowo w naiwnym klasyfikatorze bayesowskim zakłada się, że każda klasa ma rozkład normalny

Odchylenie standardowe i wartość oczekiwaną estymuje się z próbki

$$\hat{f}(y) = \frac{1}{L} \sum_{i=1}^L g(y, y_i, \sigma_k)$$

Klasyfikator NS poszukuje najbliższego wektora ze zbioru wektorów uczących.

Wynik zależy od sposobu liczenia odległości.

Modyfikacja - kNS - wybiera się odległości do k najbliższych sąsiadów.

Podobnie działa klasyfikator najbliższego centrum (NC).

Przykładowe miary odległości

- odległość Czebyszewa
- odległość miejska
- odległość euklidesowa
- odległość Minkowskiego
- odległość Hamminga
- odległość Mahalanobisa

Odległość Czebyszewa (maksimum)

$$d = \max_{1 \leq i \leq n} |x_i - y_i|$$

Odległość miejska (taksówkowa, Manhattan distance, city-block)

$$d = \sum_{i=1}^n |x_i - y_i|$$

Odległość euklidesowa

$$d = \sqrt{\sum_{i=1}^n |x_i - y_i|^2}$$

Odległość Minkowskiego

$$d = \left(\sum_{i=1}^n |x_i - y_i|^q \right)^{1/q}$$

$q=1$, $q=2$, $q=\infty$ – szczególne przypadki

Odległość Hamminga

Określa na ilu pozycjach różnią się dwa wektory

Przykładowo

$$d(0011, 1100) = 4, d(122, 220) = 3.$$

Można także stosować wektory literowe itp.

Odległość Mahalanobisa

$$d = \sqrt{(x - y)^T S^{-1} (x - y)}$$

gdzie S jest macierzą kowariancji

Można warzyć współrzędne

$$d = \left(\sum_{i=1}^n w_i |x_i - y_i|^q \right)^{1/q}$$

Można skalować współrzędne

$$a' = (a - \min_A(a)) / (\max_A(a) - \min_A(a))$$

Detekcja anomalii czy też nowości (ang. *outlier detection, novelty detection*) pozwala wykryć dane w znacznym stopniu odmienne, wyjątkowe i niespójne z większością danych wejściowych.

Metody

- **Statystyczne**
- **Odległościowe**
- **Na podstawie gęstości**
- **inne**

Statystyczne wykrywanie anomalii

Bazują na rozkładzie prawdopodobieństwa konstruowanym w sposób parametryczny (zakładamy rozkład) bądź nieparametryczny

Trening – budowa modelu

Przykładem metody nieparametrycznej jest wykorzystanie histogramu do określenia prawdopodobieństwa wystąpienia danych

Problem przy większej liczbie wymiarów

Statystyczne wykrywanie anomalii

Dla danych składających się z wielu atrybutów można wyznaczyć *próg anomalii* (ang. *outlier score*).

Niskie prawdopodobieństwo występowania danej instancji oznacza wysoki *próg anomalii*.

$$próg = \frac{1}{\#F} \sum_{f \in F} w_f (1 - p_f)$$

gdzie w_f oznacza wagę nadaną każdemu atrybutowi, p_f oznacza prawdopodobieństwo a priori wystąpienia wartości atrybutu, a F zbiór atrybutów

Odległościowe wykrywanie anomalii

Można wykorzystać ideę najbliższego sąsiedztwa badanej instancji.

Brak założeń dotyczących rozkładu danych, na których operujemy.

Problem – dobór miary odległości.

Odległościowe wykrywanie anomalii

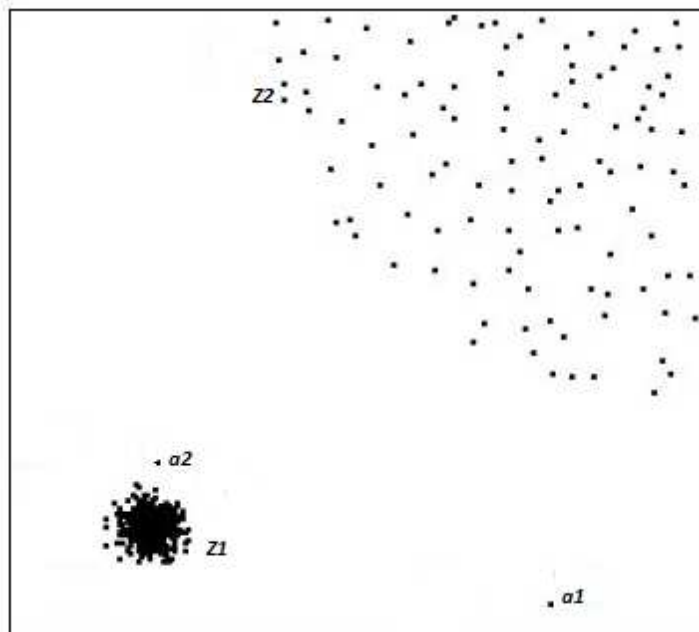
Przykładowa metoda – $DB(k, \lambda)$ – Outlier

Instancję x ze zbioru testowego zalicza się jako nietypową jeśli nie więcej niż k instancji treningowych jest w odległości λ lub mniejszej od x .

Modyfikacja – metoda $DB(pct, d_{min})$

Instancje zalicza jako nietypową, jeśli co najmniej $pct\%$ instancji w zbiorze treningowym jest w odległości większej niż d_{min} od tej instancji.

Klasyfikacja – wykrywanie anomalii



$Z1$ - 400 obiektów

$Z2$ - 100 obiektów

$a1, a2$ - anomalie

Problem - tylko $a1$ jest anomalią albo cały $Z2$ to anomalie

Nie da się dobrać współczynników tak, aby tylko $a1$ i $a2$ były anomaliami

Wykrywanie anomalii na podstawie gęstości

Przykład - metoda LOF (ang. *local outlier factor*)

Nadaje każdej instancji stopień, w jakim można ją uważać za anomalię. Metoda ta sprawdza jak bardzo odizolowana jest instancja od otaczającego ją sąsiedztwa.

Markus M. Breunig, Hans-Peter Kriegel, Raymond T. Ng, Jörg Sander, LOF: Identifying Density-Based Local Outliers, Proc. ACM SIGMOD 2000 Int. Conf. On Management of Data, Dalles, TX, 2000

LOF – local outlier factor

$$LOF_k(p) = \frac{\#N_k(p)}{\sum_{a \in N_k(p)} \frac{lrd_k(a)}{lrd_k(p)}}$$

gdzie:

- $lrd_k(a)$ - local reachability density punktu a -
określa zagęszczenie w k -sąsiedztwie obiektu a
- $lrd_k(p)$ - local reachability density punktu p -
określa zagęszczenie w k -sąsiedztwie obiektu p
- $\#N_k(p)$ - liczba elementów w k -sąsiedztwie punktu p

LOF – local outlier factor

$$lrd_k(p) = \frac{\# N_k(p)}{\sum_{o \in N_k(p)} reach-dist_k(p, o)}$$

$$reach-dist_k(p, o) = \max\{k - distance(o), dist(p, o)\}$$

LOF – local outlier factor

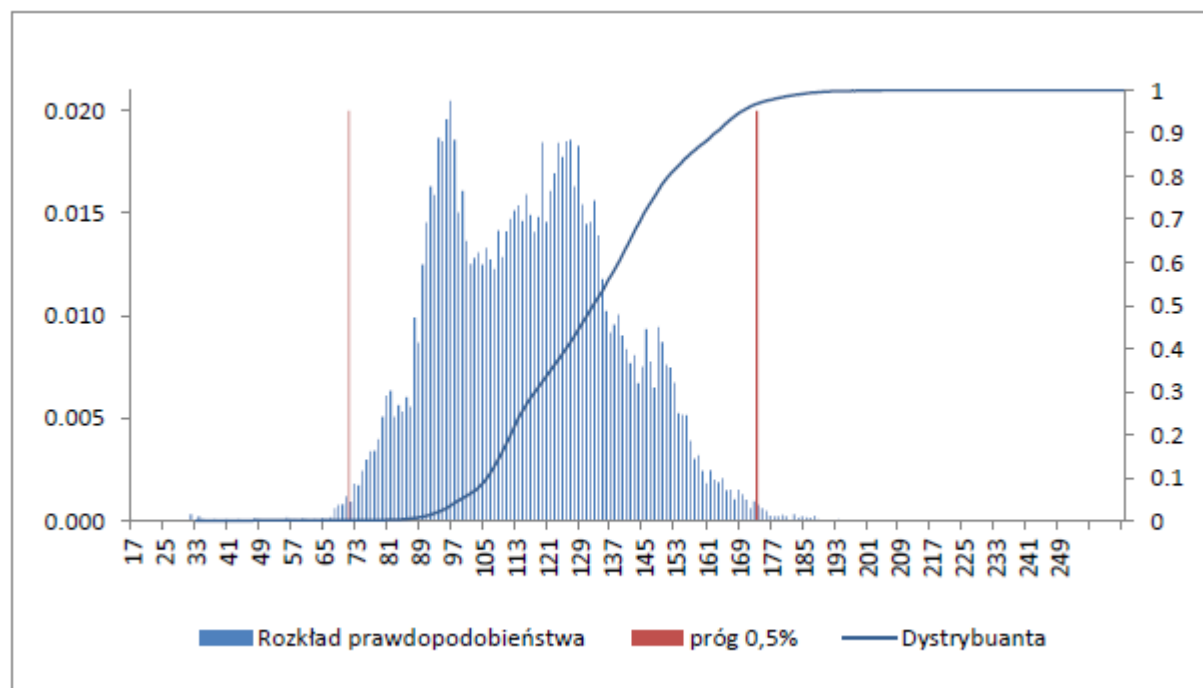
- metoda lokalna
- brak kategoryzacji
- jeżeli nowy pomiar jest "bliski" pomiarom treningowym, to $LOF \sim 1$
- ocena "jak mocno" dany obiekt jest "odbiegający"

Przykładowe wyniki

Dane typowe – $[param \cdot \sigma - \mu, param \cdot \sigma + \mu]$,
 $param = 1, 2, 3$

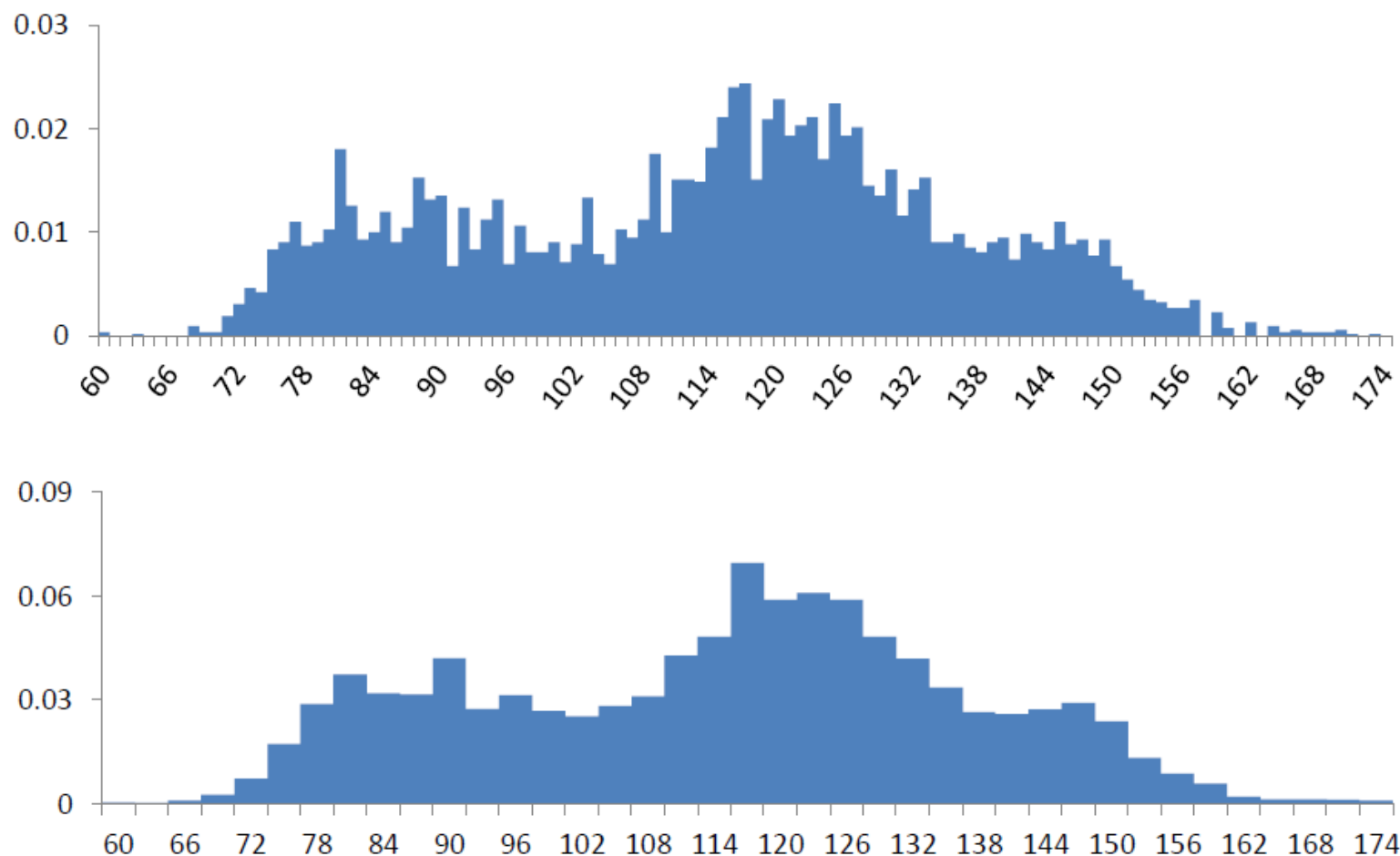
param	instancje stabilne (% ogółu)	instancje niestabilne (% ogółu)
1	21185 (82,3%)	4554 (17,7%)
2	25529 (99,2%)	211 (0,8%)
3	25724 (99,9%)	16 (0,1%)

Klasyfikacja – przykłady wyników



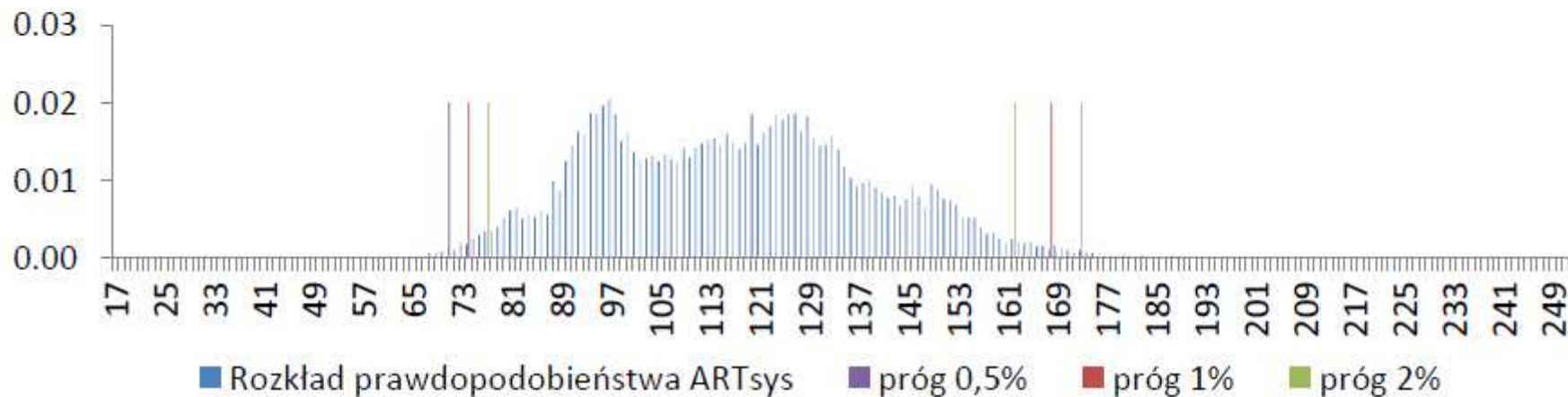
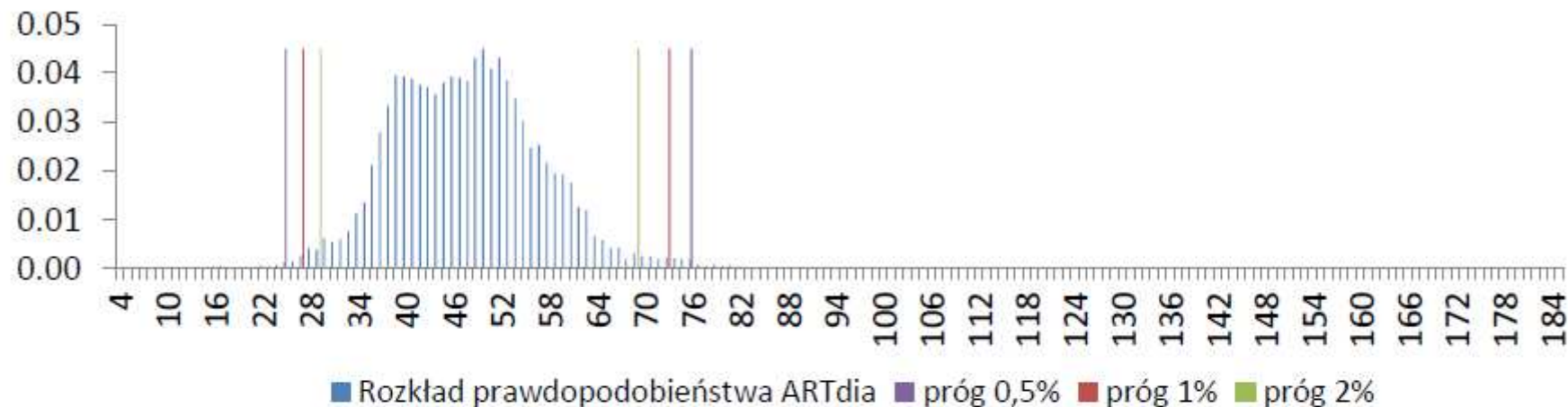
Rozkład prawdopodobieństwa tętniczego ciśnienia skurczowego (skala na lewej osi) wraz z dystrybuantą (skala na prawej osi) oraz progiem wyznaczonym na jej podstawie.

Klasyfikacja – przykłady wyników

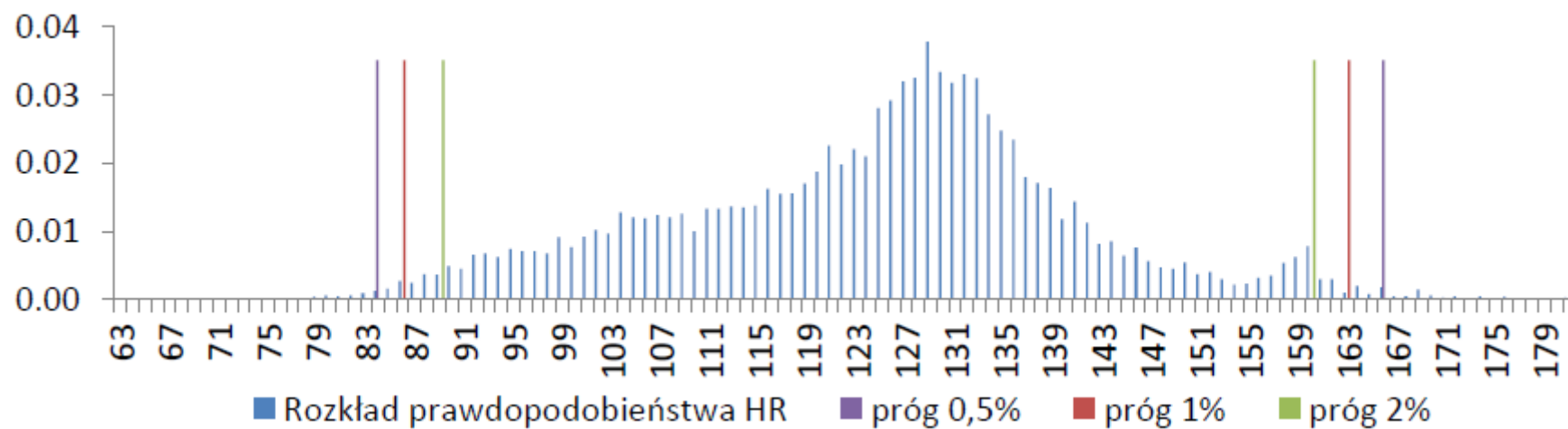


Przykład wpływu szerokości przedziału na estymowany rozkład prawdopodobieństwa dla tętniczego ciśnienie skurczowego

Klasyfikacja – przykłady wyników



Klasyfikacja – przykłady wyników



Przykładowe miary jakości klasyfikacji

$$\text{TP rate} = \text{TP} / (\text{TP} + \text{FN})$$

$$\text{TN rate} = \text{TN} / (\text{TN} + \text{FP})$$

Klasyfikacja – przykłady wyników

próg [%]	param	Dokładność klasyfikacji [%]	TN rate	TP rate
0.5	1	87%	90%	71%
	2	92%	92%	78%
	3	94%	94%	56%
		91%	92%	68%
1	1	87%	90%	73%
	2	91%	91%	75%
	3	93%	93%	38%
		90%	91%	62%
2	1	87%	90%	74%
	2	90%	90%	70%
	3	93%	93%	36%
		90%	91%	60%

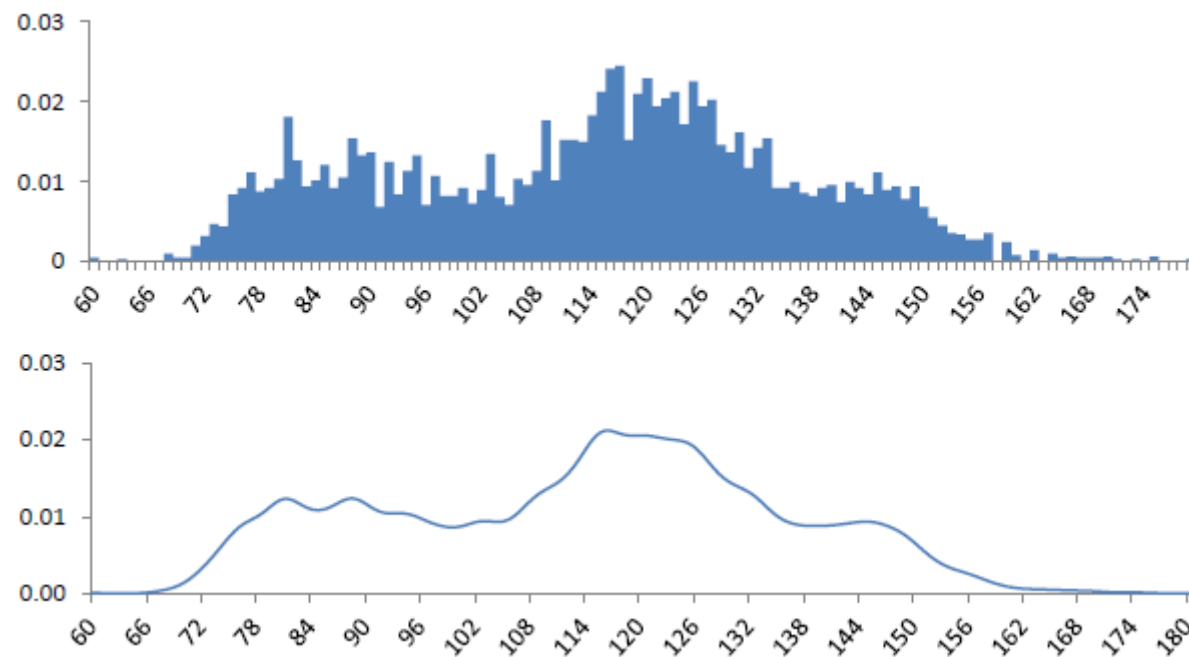
Wyniki działania naiwnego klasyfikatora bayesowskiego dla rozkładów prawdopodobieństwa z różnym umiejscowieniem skrajnych progów cięcia dla trzech zbiorów instancji

Klasyfikacja – przykłady wyników

	param					
	1		2		3	
Nazwa zestawu obowiązująca w programie	TP rate	TN rate	TP rate	TN rate	TP rate	TN rate
$h_{ARTdia} = 2, h_{ARTsys} = 2, h_{EtCO2} = 0.04, h_{HR} = 2, h_{RR} = 1$	72%	90%	85%	90%	46%	91%
"Optymalne szerokości"	72%	90%	85%	90%	43%	91%
$h_{ARTdia} = 2, h_{ARTsys} = 4, h_{EtCO2} = 0.05, h_{HR} = 3, h_{RR} = 1$	72%	90%	86%	91%	49%	91%
$h_{ARTdia} = 3, h_{ARTsys} = 5, h_{EtCO2} = 0.06, h_{HR} = 4, h_{RR} = 1$	72%	90%	86%	91%	49%	91%
$h_{ARTdia} = 4, h_{ARTsys} = 6, h_{EtCO2} = 0.07, h_{HR} = 5, h_{RR} = 1$	71%	90%	86%	92%	55%	91%
$h_{ARTdia} = 5, h_{ARTsys} = 7, h_{EtCO2} = 0.08, h_{HR} = 6, h_{RR} = 1$	72%	90%	85%	92%	47%	91%

**Wpływ estymacji rozkładu na wyniki klasyfikacji
naiwnym klasyfikatorem bayesowskim**

Klasyfikacja – przykłady wyników



Rozkład prawdopodobieństwa dla tętniczego ciśnienia skurczowego oraz gęstość prawdopodobieństwa tego rozkładu estymowana z wykorzystaniem estymatora jądrowego.

Klasyfikacja – przykłady wyników

param	Dokładność klasyfikacji	TN rate	TP rate
1	76%	75%	84%
2	87%	87%	56%
3	100%	100%	0%

Wyniki klasyfikacji naiwnym klasyfikatorem bayesowskim z wykorzystaniem rozkładów prawdopodobieństwa otrzymanych metodą estymacji jądrowej.

Klasyfikacja – przykłady wyników

Klasyfikator	Nazwa zestawu obowiązująca w programie	param	Dokładność	TN rate	TP rate
NBC	<i>"Optymalne szerokości"</i>	1	87.0%	90.2%	71.8%
		2	89.8%	89.8%	84.6%
		3	91.0%	91.2%	43.2%
NBC z połączonymi ARTdia i ARTsys	<i>"Optymalne szerokości"</i>	1	81.6%	89.2%	50.0%
		2	94.8%	95.0%	74.6%
		3	95.8%	95.8%	45.2%
	$h_{ARTdia} = 5, h_{ARTsys} = 7,$ $h_{EtCO2} = 0.08, h_{HR} = 6, h_{RR} = 1$	1	84.2%	90.0%	58.2%
		2	94.2%	94.2%	79.6%
		3	94.8%	94.8%	42.8%

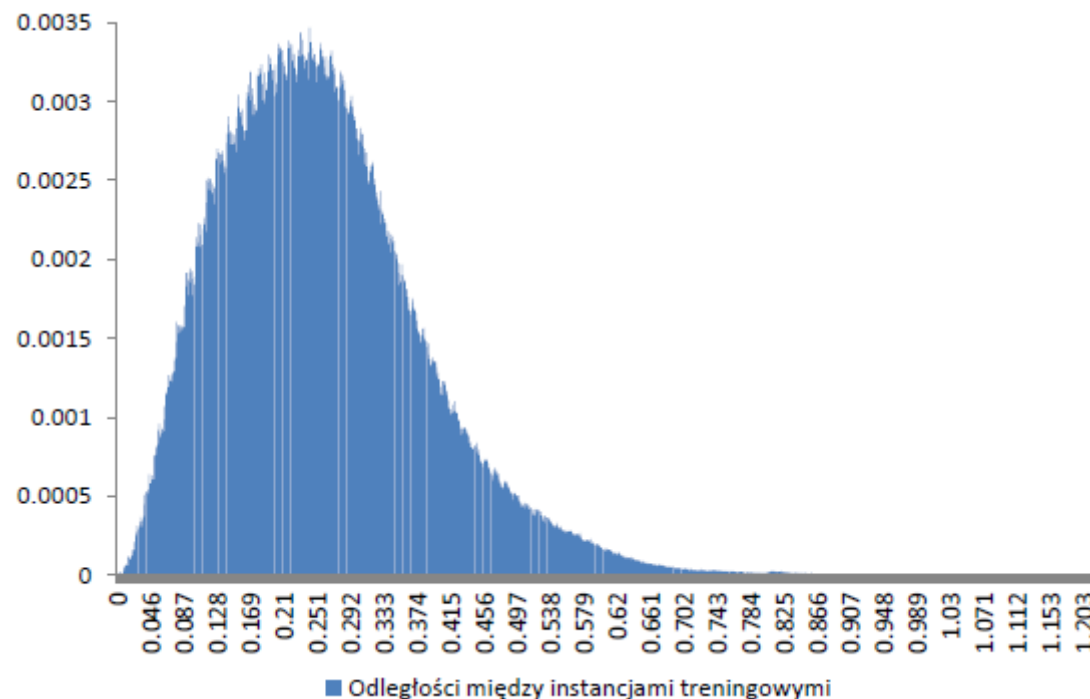
Wyniki klasyfikacji dla naiwnego klasyfikatora bayesowskiego oraz naiwnego klasyfikatora bayesowskiego wykorzystującego łączny rozkład prawdopodobieństwa dla ciśnienia rozkurczowego oraz skurczowego.

Klasyfikacja – przykłady wyników

Liczba najbliższych sąsiadów k	param	Dokładność klasyfikacji	TN rate	TP rate
1	1	86.4%	97.4%	37.0%
	2	99.2%	100.0%	9.8%
	3	100.0%	100.0%	4.4%
Średnia dla $k = 1$				
		95.2%	99.1%	17.1%
3	1	86.4%	98.4%	29.2%
	2	99.2%	100.0%	4.2%
	3	100.0%	100.0%	1.6%
Średnia dla $k = 3$				
		95.2%	99.5%	11.7%
5	1	86.2%	98.6%	28.0%
	2	99.2%	100.0%	3.8%
	3	100.0%	100.0%	0.0%
Średnia dla $k = 5$				
		95.1%	99.5%	10.6%

Dokładność klasyfikacji algorytmem k-NN dla różnej ilości najbliższych sąsiadów k .

Problem – liczba instancji niestabilnych ze wzrostem k



Częstość występowania euklidesowych odległości między instancjami w zbiorze treningowym.

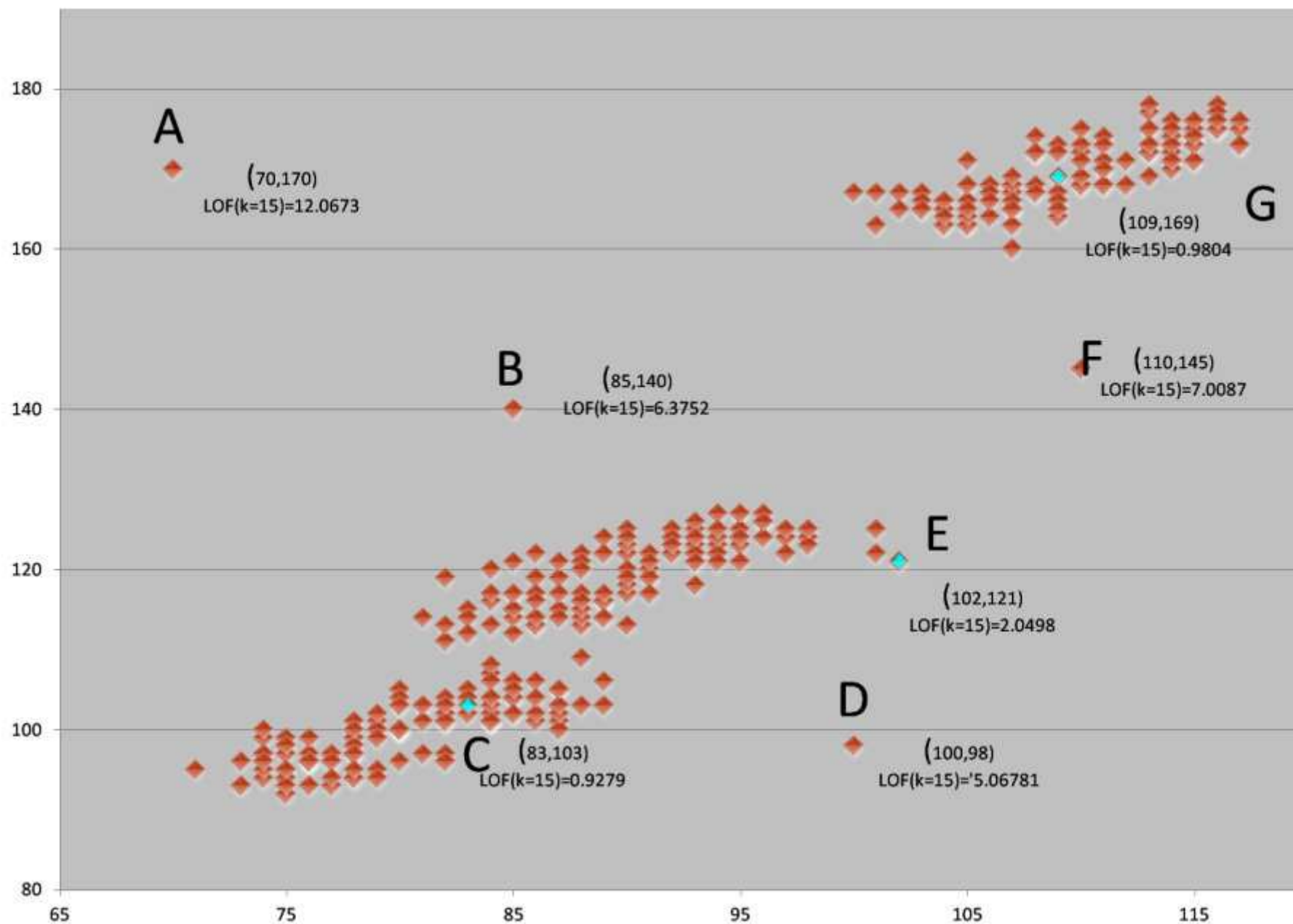
Próg Anomalii	Procent wykrytych anomalii w zbiorze testowym
0.98	7.93%
0.99	1.84%
0.995	0.52%

Procent anomalii wykrytych w zbiorze instancji testowych dla różnych wartości *progu anomalii*

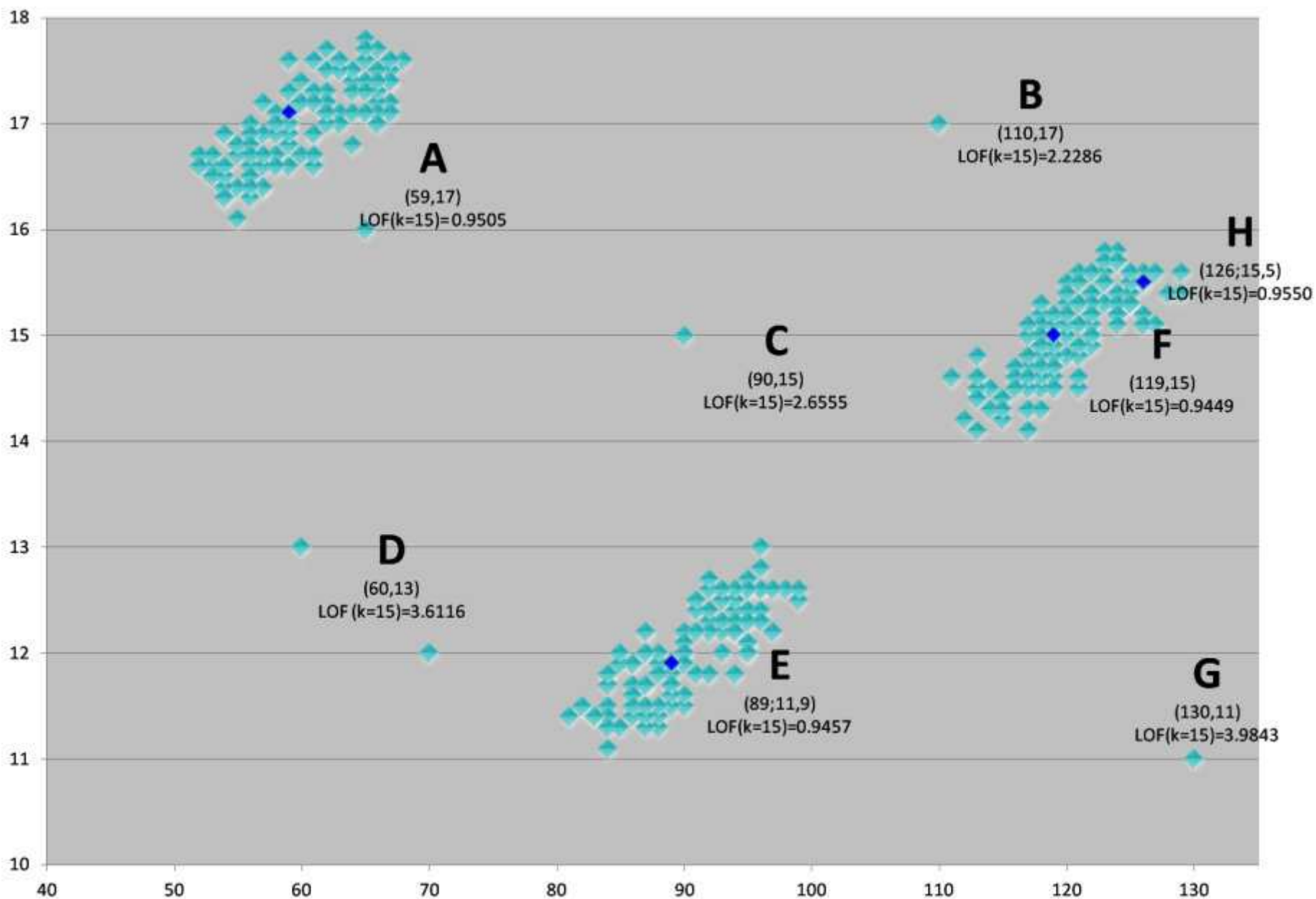
d_{Min}	pct	Procent wykrytych anomalii
0.8	98.0%	14.67%
	99.0%	21.90%
	99.5%	30.27%
1	98.0%	2.40%
	99.0%	9.65%
	99.5%	15.18%

Procent anomalii wykrytych w zbiorze instancji testowych dla różnych wartości d_{min} oraz pct

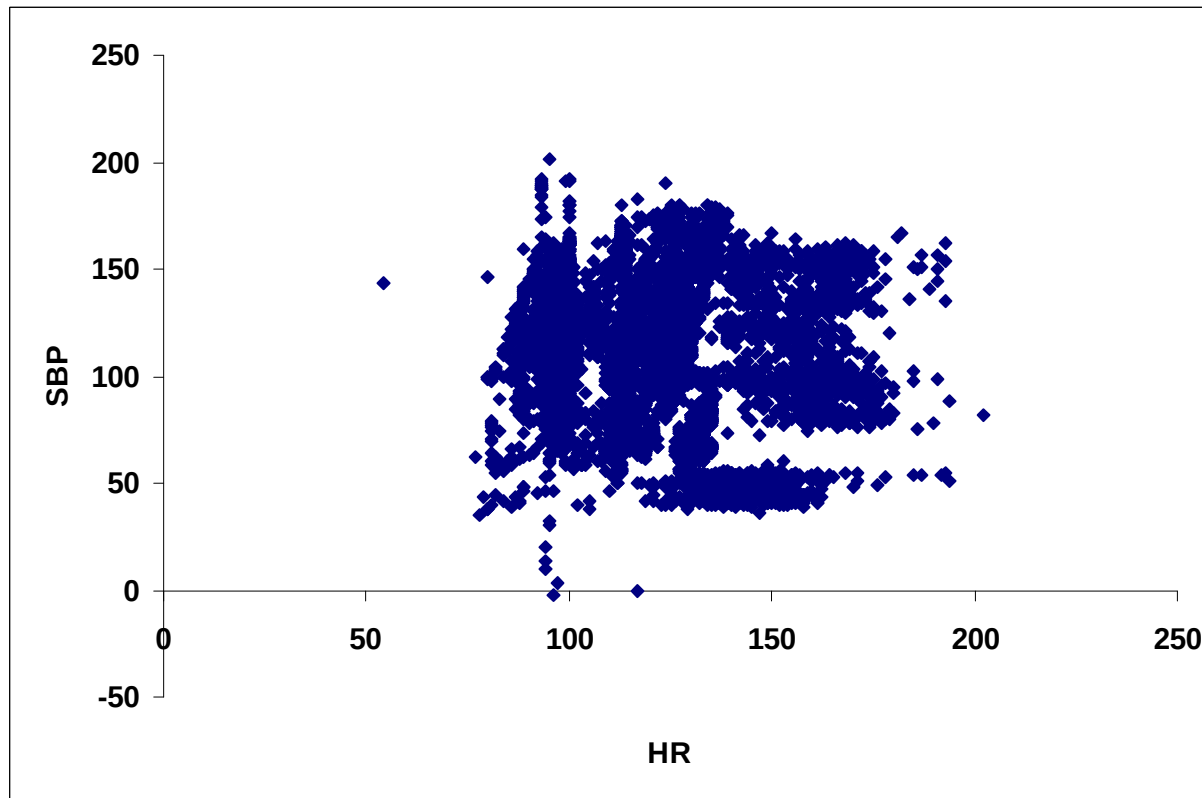
Klasyfikacja – anomalie - przykłady



Klasyfikacja – anomalie - przykłady



Klasyfikacja – anomalie - przykłady



Wnioski z symulacji LOF

- jeżeli pomiar jest „w normie”, to klasyfikator LOF~1
- duże wartości LOF (>3) wskazują na mocno odbiegający pomiar
- LOF w przedziale (2,3) – reguły decyzyjne

Eksperymenty dotyczące klasyfikatorów bayesowskich pokazały duże znaczenie sposobu estymacji rozkładu prawdopodobieństwa.

Problemem może też być liczba parametrów wykorzystywanych w klasyfikatorze.

Dla modelowanych danych dla klasyfikatora k-NN instancji o klasie niestabilnej było zbyt mało, aby skutecznie wykorzystywać go do identyfikacji stanów zagrożonych.

Algorytmy wykrywające anomalie nie muszą mieć sklasyfikowanych instancji, ale pojawia się problem ustawienia odpowiednich progów detekcji stanów nietypowych.

- nieregularne/niedokładne pomiary
- wpływ rytmu dobowego/tygodniowego/...
- nieregularne przyjmowanie leków
- zróżnicowana aktywność dzienna
- możliwa zmiana stanu pacjenta w czasie (np. zmiana leków)

**Dziękujemy
za uwagę**