



BIOFIZYKA RADIACYJNA

wykład #7

Krzysztof Wojciech Fornalski

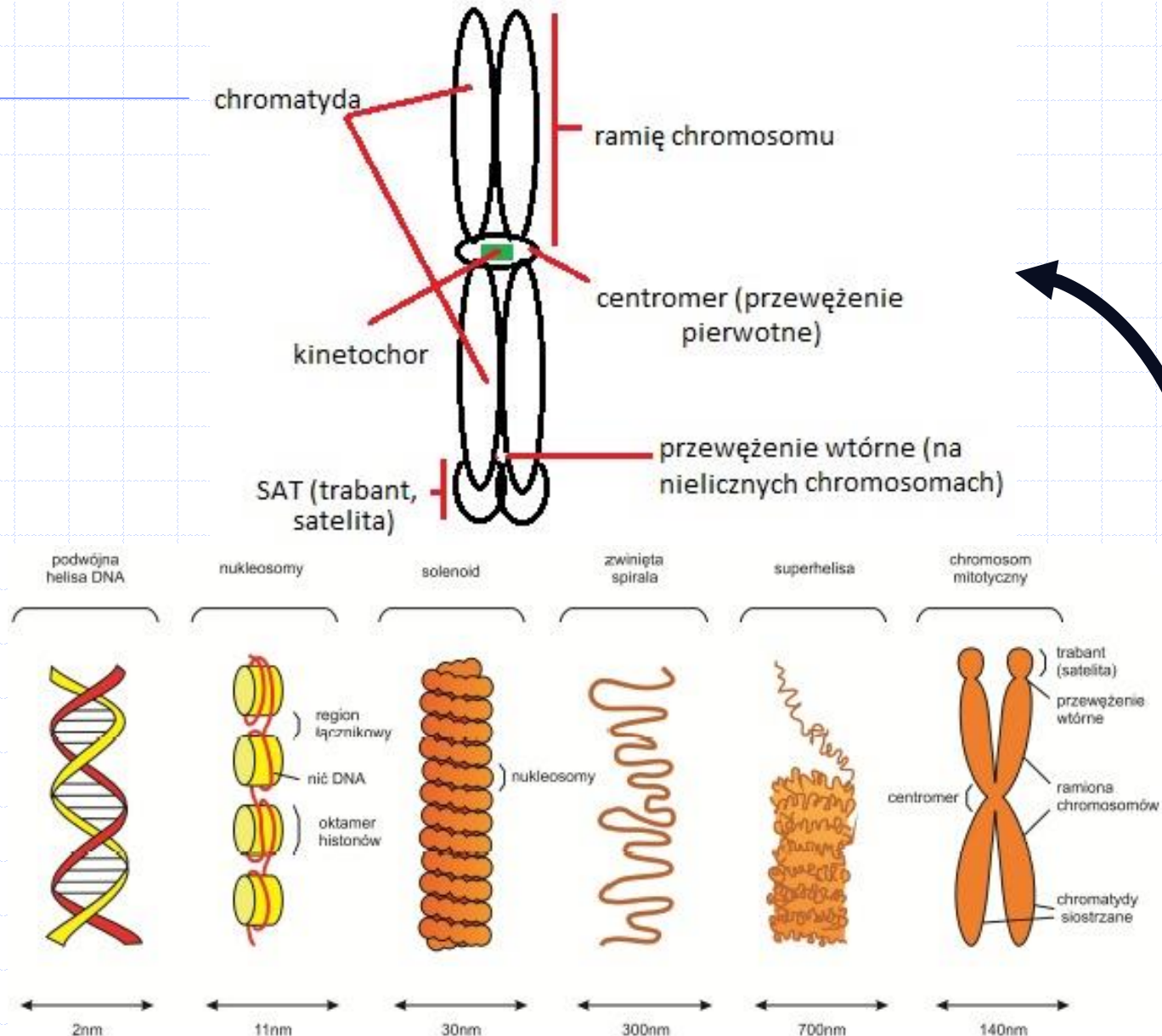
Wydział Fizyki, Politechnika Warszawska



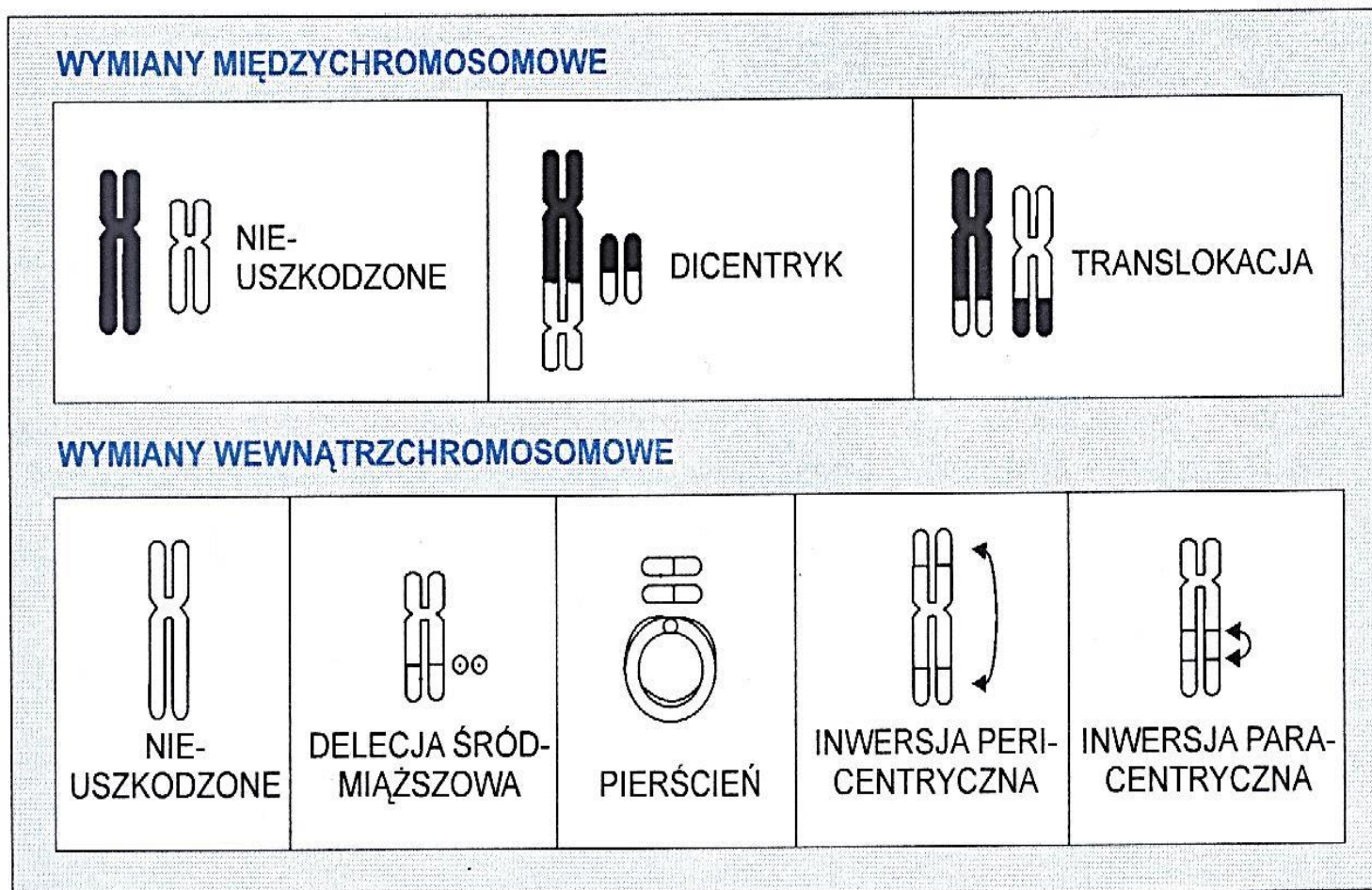


PRZYPOMNIENIE Z POPRZEDNIEGO WYKŁADU

Budowa chromosomu



Aberracje chromosomowe



Rys. 2. Schemat wymian wewnątrz- i międzychromosomowych obserwowanych w komórkach narażonych na działanie promieniowania jonizującego

BIODOZYMETRIA (DOZYMETRIA BIOLOGICZNA)

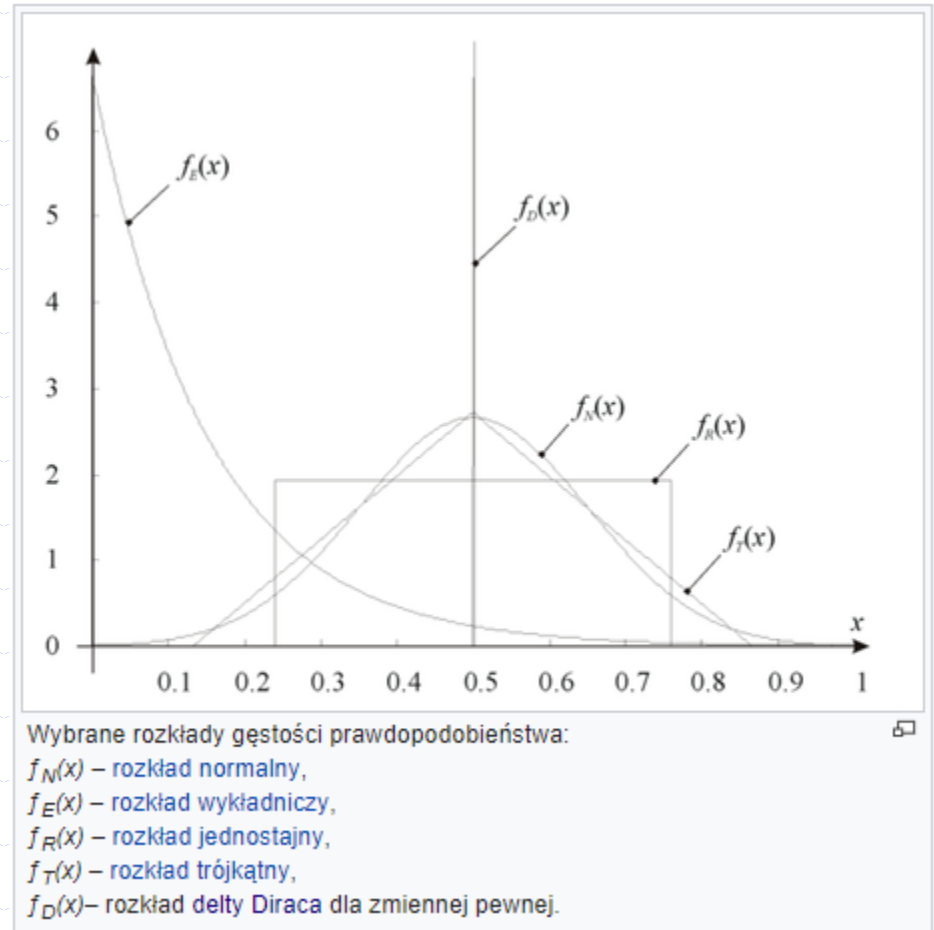
- BIODOZYMETRIA
CYTOGENETYCZNA
DICENTRYKÓW**
- ANALIZA BAYESOWSKA**

DYGRESJA: ZAAWANSOWANE METODY STATYSTYCZNE

- ROZKŁAD
PRAWDOPODOBIENSTWA**
- FUNKCJA WIARYGODNOŚCI**
- BAYESOWSKA ANALIZA
DANYCH**

Rozkład (funkcja) prawdopodobieństwa, PDF

- ◆ Nie definiujemy prawdopodobieństwa klasycznie (częstościowo), lecz poprzez funkcję prawdopodobieństwa
- ◆ Termin „funkcja prawdopodobieństwa” jest formalnie zarejestrowany do rozkładów dyskretnych
- ◆ Dla rozkładów ciągłych – formalnie „funkcja gęstości rozkładu prawdopodobieństwa” (tzw. PDF – probability density function)



Funkcja wiarygodności (likelihood, L)

wprowadzimy najpierw tzw. **funkcję wiarygodności**. Niech będzie dana próbka losowa x_i o liczebności n z rozkładu $f(x; \theta)$, gdzie x jest zmienną losową określoną, dla ustalenia uwagi, na całej osi rzeczywiste, a θ parametrem określającym rozkład. Funkcją $\mathcal{L}$ wiarygodności dla próbki x_i nazywamy wielkość:

$$\mathcal{L}(\bar{x}; \theta) = \prod_{i=1}^n f(x_i; \theta)$$

Podobnie, dla losowej próbki k_i ($i = 1, 2, \dots, n$) zmiennych dyskretnych z rozkładu $P_k(\theta)$, zakładając dla ustalenia uwagi, że $k = 0, 1, 2, \dots$, funkcję wiarygodności definiujemy jako:

$$\mathcal{L}(\bar{k}; \theta) = \prod_{i=1}^n P_{k_i}(\theta).$$

Należy zwrócić uwagę na to, że formalnie funkcja wiarygodności wygląda jak łączna funkcja gęstości rozkładu. I taką łączną funkcją gęstości jest ona tak długo, jak długo wielkości x_i oraz k_i w obu wyrażeniach są *zmiennymi* losowymi. Wielokrotnie, w dalszej części wykładu, spotkamy się z sytuacjami, gdy wielkości x_i oraz k_i to faktycznie wyniki pomiaru, a więc ściśle określone liczby, a nie zmienne. Wtedy wielkość $\mathcal{L}$ nie jest funkcją gęstości zmiennych losowych – jest to zwykła, matematyczna funkcja, zależna tylko i wyłącznie od parametru θ (jednego lub wielu). W literaturze statystycznej utarła się i bardzo głęboko zakorzeniła się tradycja wymiennego stosowania terminu *funkcja wiarygodności* dla obu tych sytuacji.

W praktyce częściej korzystamy z logarytmu z L, bo zarówno L jak i $\ln L$ osiągają ekstremum dla tych samych wartości, a obliczenia są prostsze (logarytm z iloczynu to suma logarytmów)

Funkcja wiarygodności (likelihood, L)

wprowadziliśmy funkcję wiarygodności:

$$\mathcal{L}(\bar{\mathbf{x}}; \theta) = \prod_{i=1}^n f(x_i; \theta)$$

jako łączną funkcję gęstości, w której w miejsce zmiennych losowych x_i podstawiamy wartości uzyskane w wyniku pomiaru, a więc w wyniku pobierania próbki prostej z rozkładu $f(x; \theta)$. Zakładamy, że postać funkcji gęstości $f(x; \theta)$ jest znana, nie znamy jednak wartości parametru θ . W takiej sytuacji, funkcja wiarygodności staje się funkcją tego i tylko tego parametru. Zgodnie z **zasadą (metodą) największej wiarygodności**, za wartość nieznanego parametru θ winniśmy wybrać taką liczbę $\hat{\theta}$, dla której funkcja wiarygodności osiąga maksimum:

$$\boxed{\mathcal{L}(\bar{\mathbf{x}}; \hat{\theta}) = \max.}$$

Żądanie maksymalnej wartości funkcji wiarygodności sprowadza się do znalezienia pierwiastka $\hat{\theta}$ równania:

$$\frac{\partial}{\partial \theta} \mathcal{L}(\bar{\mathbf{x}}; \theta) = 0 \quad \text{przy warunku:} \quad \frac{\partial^2}{\partial \theta^2} \mathcal{L}(\bar{\mathbf{x}}; \theta) \Big|_{\theta=\hat{\theta}} < 0.$$

Zazwyczaj funkcja największej wiarygodności ma jedno maksimum. Gdy nie jest to prawda, musimy poszukać dodatkowych informacji, pozwalających dokonać jednoznacznego wyboru.

Ponieważ funkcja $\mathcal{L}$ jak i jej logarytm osiągają maksimum dla tej samej wartości argumentu, estymator największej wiarygodności można często znaleźć z bardziej praktycznego równania:

$$\frac{\partial}{\partial \theta} \ln \mathcal{L}(\bar{\mathbf{x}}; \theta) = \sum_{i=1}^n \frac{\partial}{\partial \theta} \ln f(x_i; \theta) = 0, \quad \text{przy warunku:} \quad \sum_{i=1}^n \frac{\partial^2}{\partial \theta^2} \ln f(x_i; \theta) \Big|_{\theta=\hat{\theta}} < 0.$$

Dla ogólnego przypadku z k nieznanymi parametrami θ_i , musimy rozwiązać układ k równań:

$$\frac{\partial}{\partial \theta_j} \ln \mathcal{L}(\bar{\mathbf{x}}; \bar{\theta}) = \sum_{i=1}^n \frac{\partial}{\partial \theta_j} \ln f(x_i; \bar{\theta}) = 0, \quad \text{dla:} \quad j = 1, 2, \dots, k$$

Szacowanie niepewności parametru opisywanego funkcją wiarygodności

W języku funkcji L , a więc łącznej funkcji gęstości, ograniczenie od dołu na wartość wariancji nieobciążonego estymatora zadana jest przez nierówność określoną **twierdzeniem Rao–Cramera**:

$$V(\hat{\theta}) \geq V_{\min}(\hat{\theta}) = \frac{1}{\int L(\vec{x}; \theta) \left(\frac{\partial}{\partial \theta} \ln L(\vec{x}; \theta) \right)^2 d\vec{x}},$$

i podobnie dla zmiennej dyskretnej:

$$V(\hat{\theta}) \geq V_{\min}(\hat{\theta}) = \frac{1}{\sum_{\vec{k}} L(\vec{k}; \theta) \left(\frac{\partial}{\partial \theta} \ln L(\vec{k}; \theta) \right)^2},$$

Możemy przyjąć, że wariancja V to kwadrat niepewności σ^2

Twierdzenie Bayesa

Podstawy rozumowania bayesowskiego sięgają podstaw teorii prawdopodobieństwa, a mianowicie prawdopodobieństwa warunkowego zdarzenia X jeśli zachodzi zdarzenie Y .

Prawdopodobieństwo to dane jest znanym wzorem:

$P(X|Y) = P(X \cap Y) / P(Y)$. Z racji tego, iż prawdopodobieństwa iloczynów zdarzeń są sobie równe, $P(X \cap Y) = P(Y \cap X)$, zapisać można:

$$P(X | Y, I) = \frac{P(Y | X, I) \times P(X | I)}{P(Y | I)}$$

gdzie I oznacza posiadaną wcześniej informację o obiekcie (jeśli ona istnieje).

Powyższy wzór stanowi treść twierdzenia Thomasa Bayesa (1702-1761) i wiąże prawdopodobieństwa warunkowe $P(X|Y,I)$ oraz $P(Y|X,I)$ przy posiadaniu dodatkowej informacji I . W ogólności parametr I określa wszelkie dane, które są niezależne od X oraz Y , a które mówią o przedmiocie badań, np. wyniku poprzedniego eksperymentu.

Twierdzenie Bayesa

Jeśli przez X i Y rozumiemy zmienne fizyczne, które możemy wielokrotnie zmierzyć, sens prawdopodobieństwa może zostać sprowadzony do prawdopodobieństwa otrzymania takich właśnie wartości w nieskończonej liczbie eksperymentów powtarzanych w identycznych warunkach. Jednakże, jeśli któraś z tych wielkości miałaby oznaczać np. hipotezę, takie podejście, zwane „częstotliwościowym”, traci sens. Wówczas dobrym przykładem zastosowania twierdzenia Bayesa jest użycie równania w tzw. interpretacji „subiektywnej”. Przyjmując za X pewną teorię (hipotezę) fizyczną T , a za Y dane eksperymentalne E , określić można stan naszej wiedzy na temat prawdziwości danej teorii w świetle posiadanych wyników doświadczalnych i wcześniejszej wiedzy. Wówczas czynnik $P(E/T, I)$ oznacza funkcję wiarygodności (L) mówiącą o tym, jak dobrze dane eksperymentalne E odpowiadają teorii T , natomiast człon $P(T/I)$ określa stopień zaufania dla danej teorii, w skali ciągłej od 0 do 1, zanim przeprowadzono eksperyment. Tę wiedzę aprioryczną określać będziemy słowem *prior* jak w oryginalnej literaturze angielskojęzycznej. Człon z mianownika, $P(E/I)$, jest w praktyce członem normalizacyjnym.

Twierdzenie Bayesa

Ostatecznie otrzymujemy:

$$P(T | E, I) = \frac{P(E | T, I) \times P(T | I)}{P(E | I)} \propto P(E | T, I) \times P(T | I)$$

Funkcja wiarygodności może przyjąć postać konkretnego rozkładu (np. Gaussa), a *prior* być funkcją z góry nakładającą pewne obostrzenia np. na parametry teorii, która powinna opisać dany eksperyment. Równanie powyższe pokazuje ponadto, w jaki sposób zmienia się zaufanie do teorii T w wyniku przeprowadzenia eksperymentu i otrzymania wyników E .

Rozumowanie bayesowskie

- ◆ Na podstawie twierdzenia Bayesa zapisać więc można ogólne równanie na prawdopodobieństwo wynikowe (posterior)

$$\text{Posterior} \sim L \cdot \text{prior}$$

- ◆ gdzie L to likelihood (funkcja wiarygodności), zaś prior to funkcja aprioryczna
- ◆ Zarówno L jak i prior są rozkładami prawdopodobieństwa (PDF)

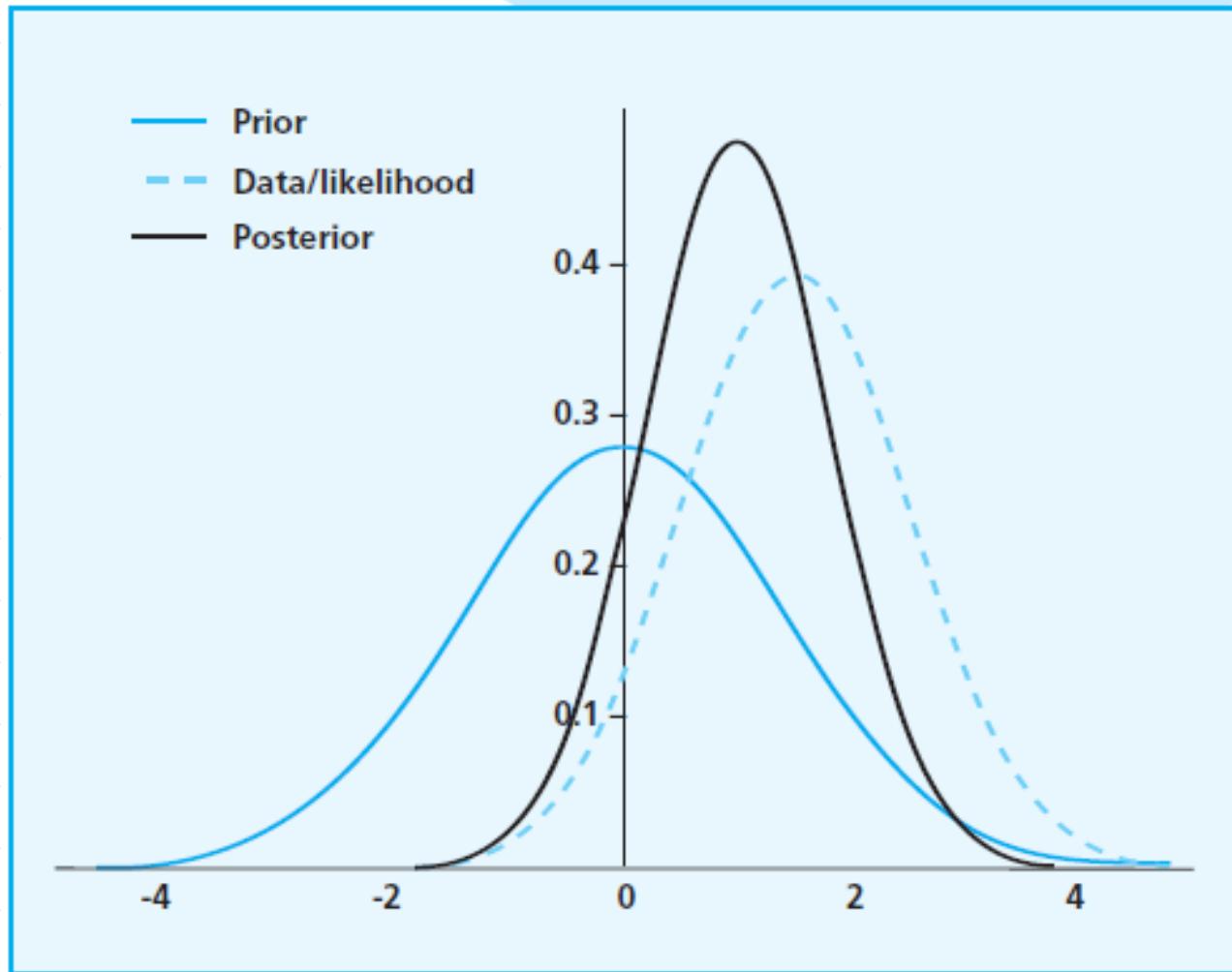
Rozumowanie bayesowskie

$$\text{Posterior} \sim \text{Likelihood} \times \text{Prior}$$

Powyższe równanie należy interpretować (i używać!) następująco: szukamy rozkładu p-twa jakiegoś parametru, np. wyniku w rzucie kostką. Rozkład ten jest naszym wynikiem (funkcja posterior). Wynik ten otrzymujemy wymnażając dwie inne funkcje: 1) wiarygodności (likelihood), która mówi o obiektywnych kryteriach (danych) wynikających np. z eksperymentu (np. p-two wyrzucenia każdego z 6 wyników na kostce jest równe), 2) prior, która określa dodatkową wiedzę (lub wyniki innych eksperymentów), np. że naroża kostki zostały spiłowane.

POSTERIOR PROB. = LIKELIHOOD PROB. $\times$ PRIOR PROB.

$$\mathbf{P(\Theta) = L(\Theta) \cdot p(\Theta)}$$



Mamy posterior PDF – i co dalej?

- ◆ Nasz wynik – posterior (P) – czyli funkcja rozkładu prawdopodobieństwa, posiada zazwyczaj jedno maksimum, które odpowiada najbardziej prawdopodobnej wartości szukanego parametru (Θ)
- ◆ Szukamy tego parametru szukając maksimum rozkładu P:

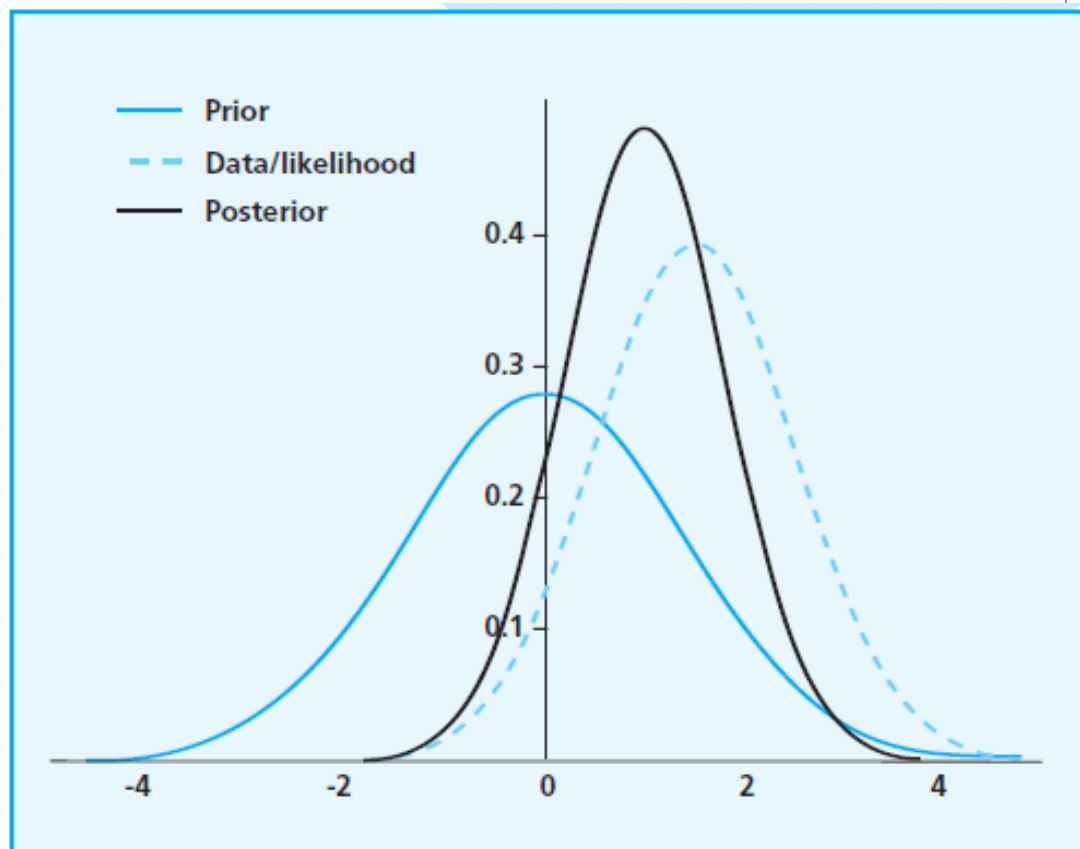
$$\frac{dP(\theta)}{d\theta} = \frac{d [L(\theta) \cdot p(\theta)]}{d\theta} = 0$$

Różnica między metodą bayesowską a metodą największej wiarygodności

- ◆ Kilka slajdów wcześniej opisywaliśmy analogiczne działanie dla samej funkcji wiarygodności (szukaliśmy jej maksimum aby znaleźć najbardziej prawdopodobny parametr)
- ◆ W metodzie bayesowskiej mamy dodatkową funkcję prior, którą mnożymy z f-cją wiarygodności
- ◆ Pojęcie funkcji prior, jej prawidłowe sformułowanie, jest tu krytycznym punktem

Przykład: uczenie się

- ◆ Wynik kolejnego eksperymentu/danej jest funkcją L
- ◆ Zaś wyniki wszystkich poprzednich: priorem
- ◆ I tak posterior i -tego działania staje się priorem dla $i+1$ działania

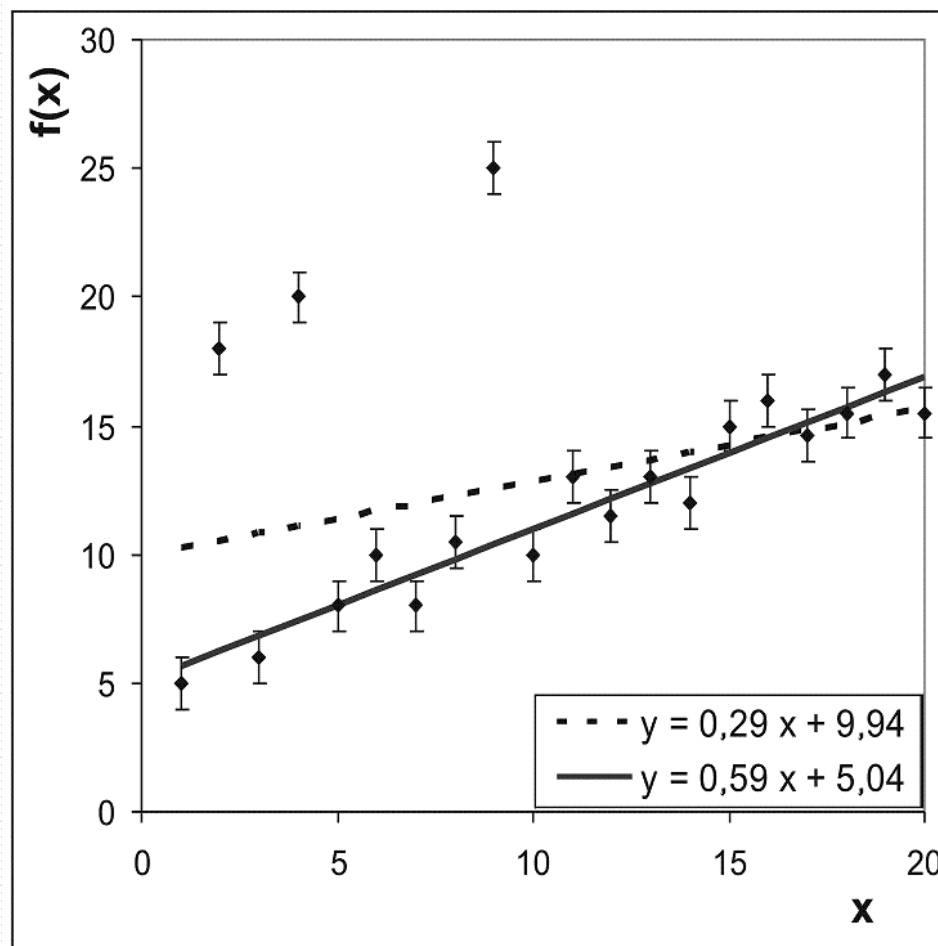


Przykład – regresja liniowa: porównanie metody najmniejszych kwadratów a dopasowania bayesowskiego (odpornego na punkty wybite)

- W dopasowaniu bayesowskim funkcja wiarygodności to rozkład Gaussa (jak w metodzie najmniejszych kwadratów), zaś prior kwestionuje poprawność istniejących niepewności i rozciąga je powodując, że ich wkład do finalnego posterioru (iloczynu) jest nieznaczący

$$P = \frac{1}{\sqrt{2\pi}\sigma} \exp\left\{-\frac{(T - E)^2}{2\sigma^2}\right\} p(\sigma|I)$$

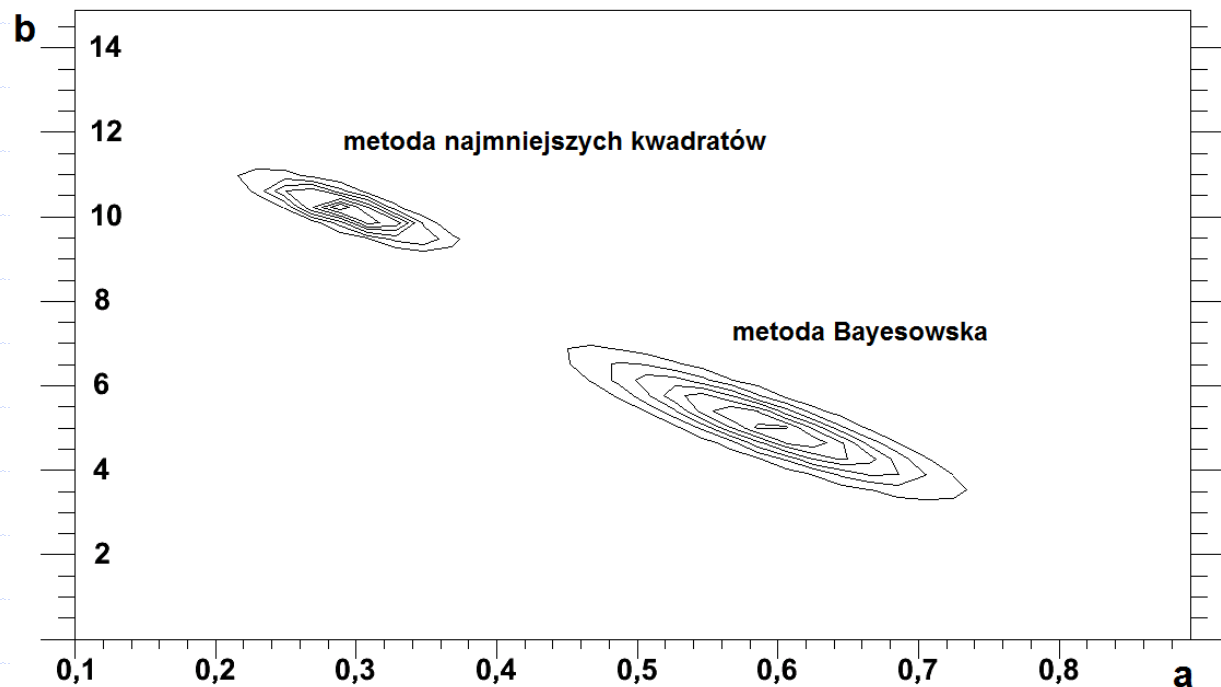
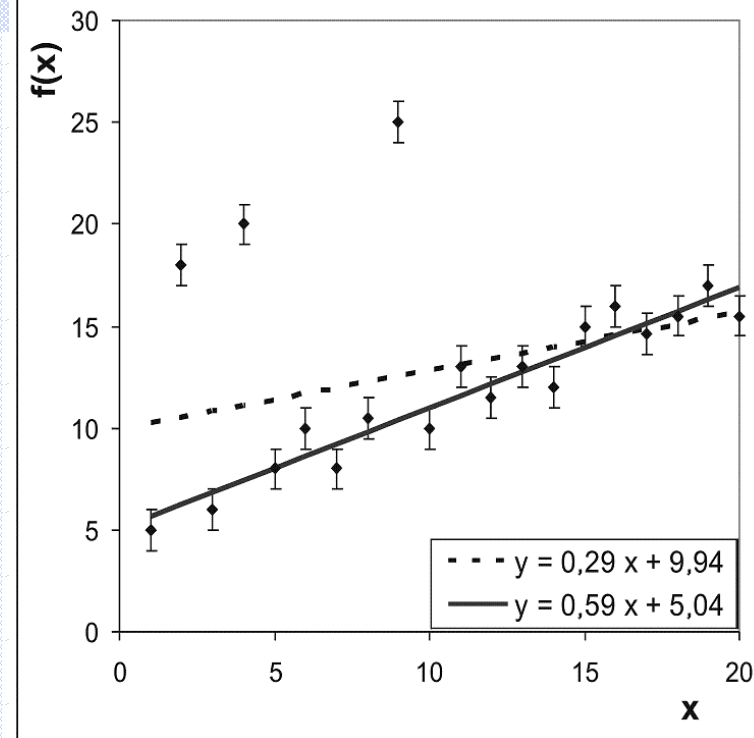
gdzie prior: $p(\sigma) = \frac{\sigma_0}{\sigma^2}$



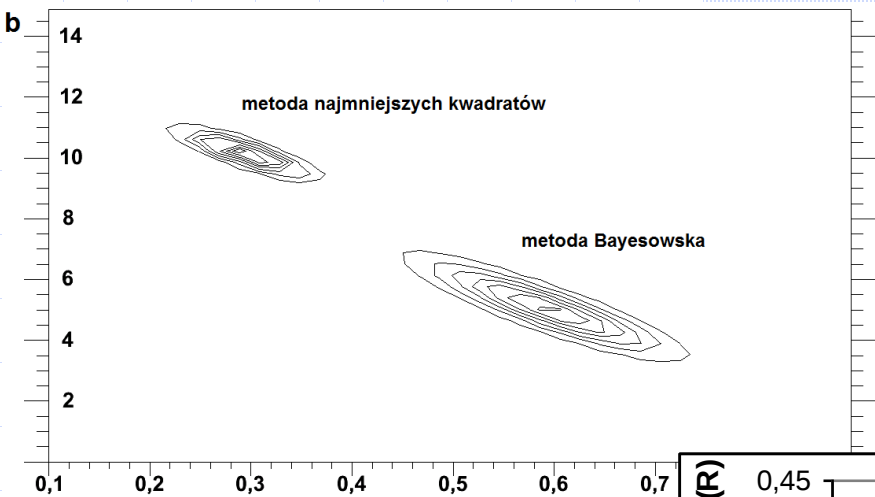
Regresja liniowa

– c.d.

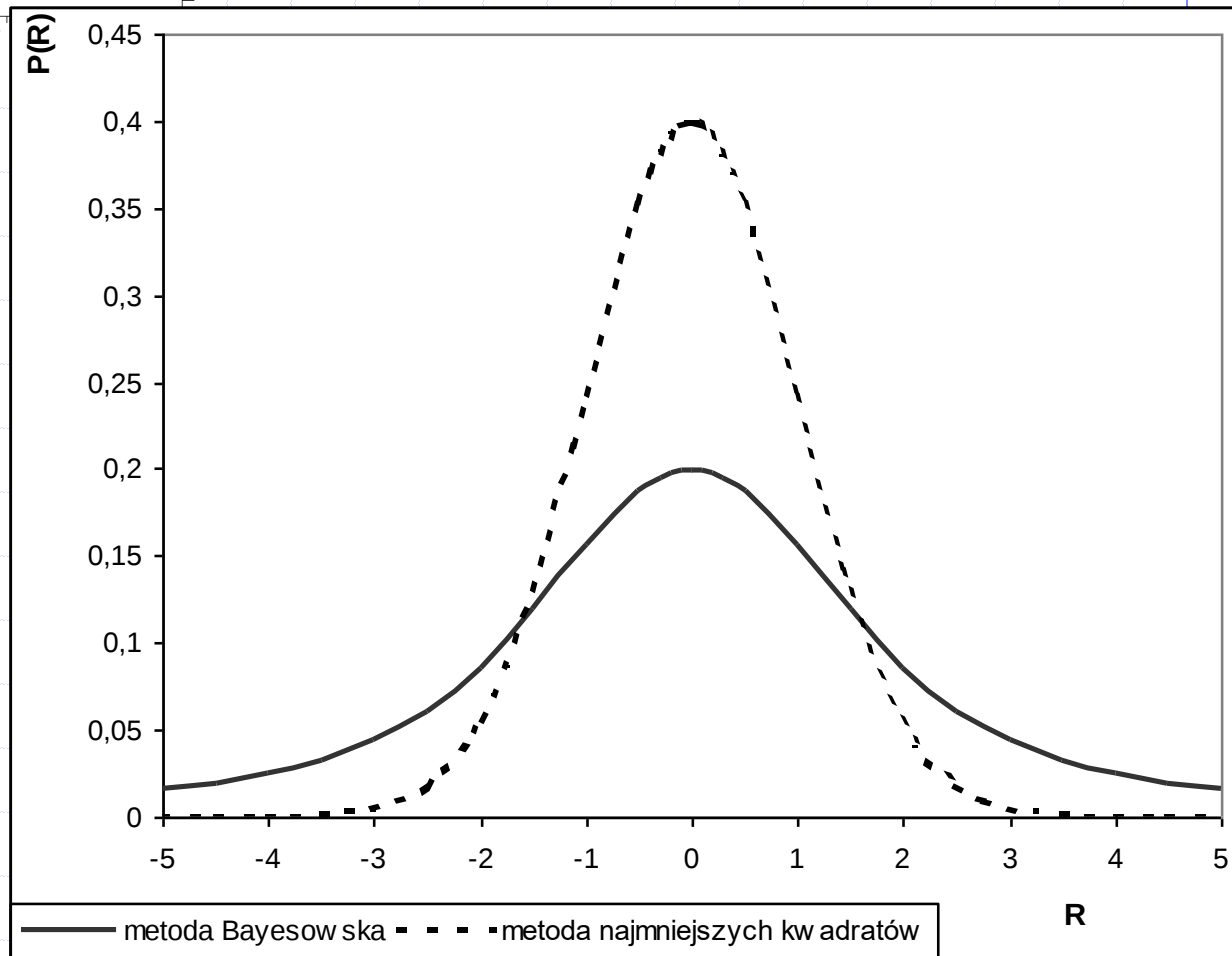
- ◆ Przedstawiona metoda jest przykładem tzw. robust regression analysis, czyli metody nieczułej na punkty wybite
- ◆ Oczywiście skutkuje to większymi niepewnościami otrzymanych wyników (parametrów dopasowania prostej)



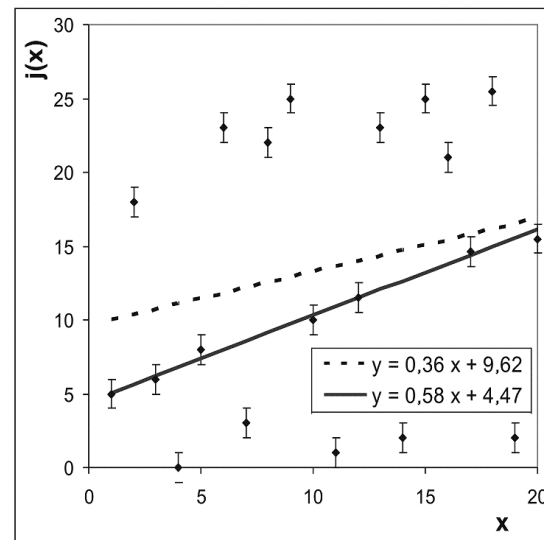
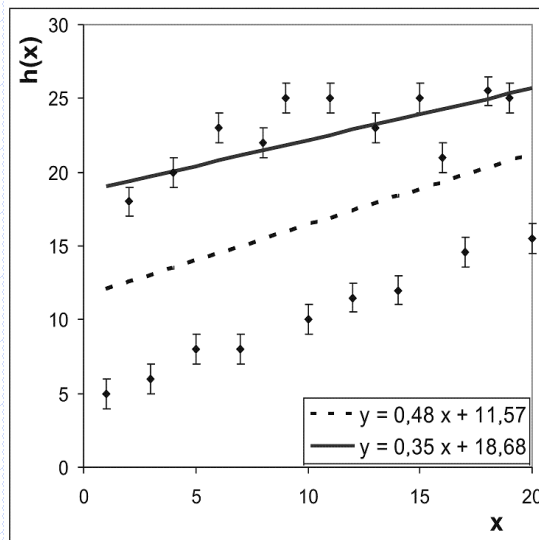
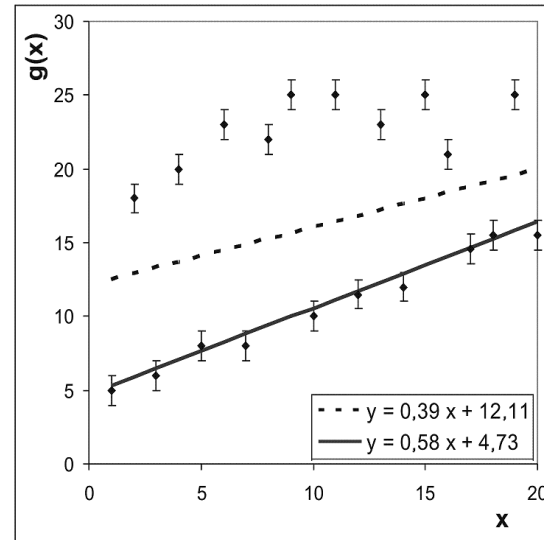
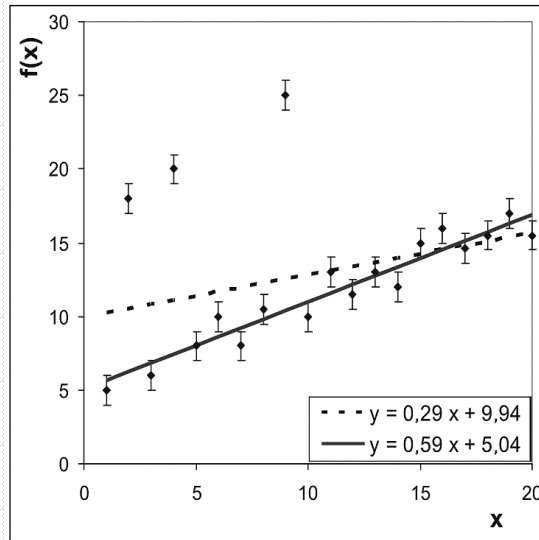
b



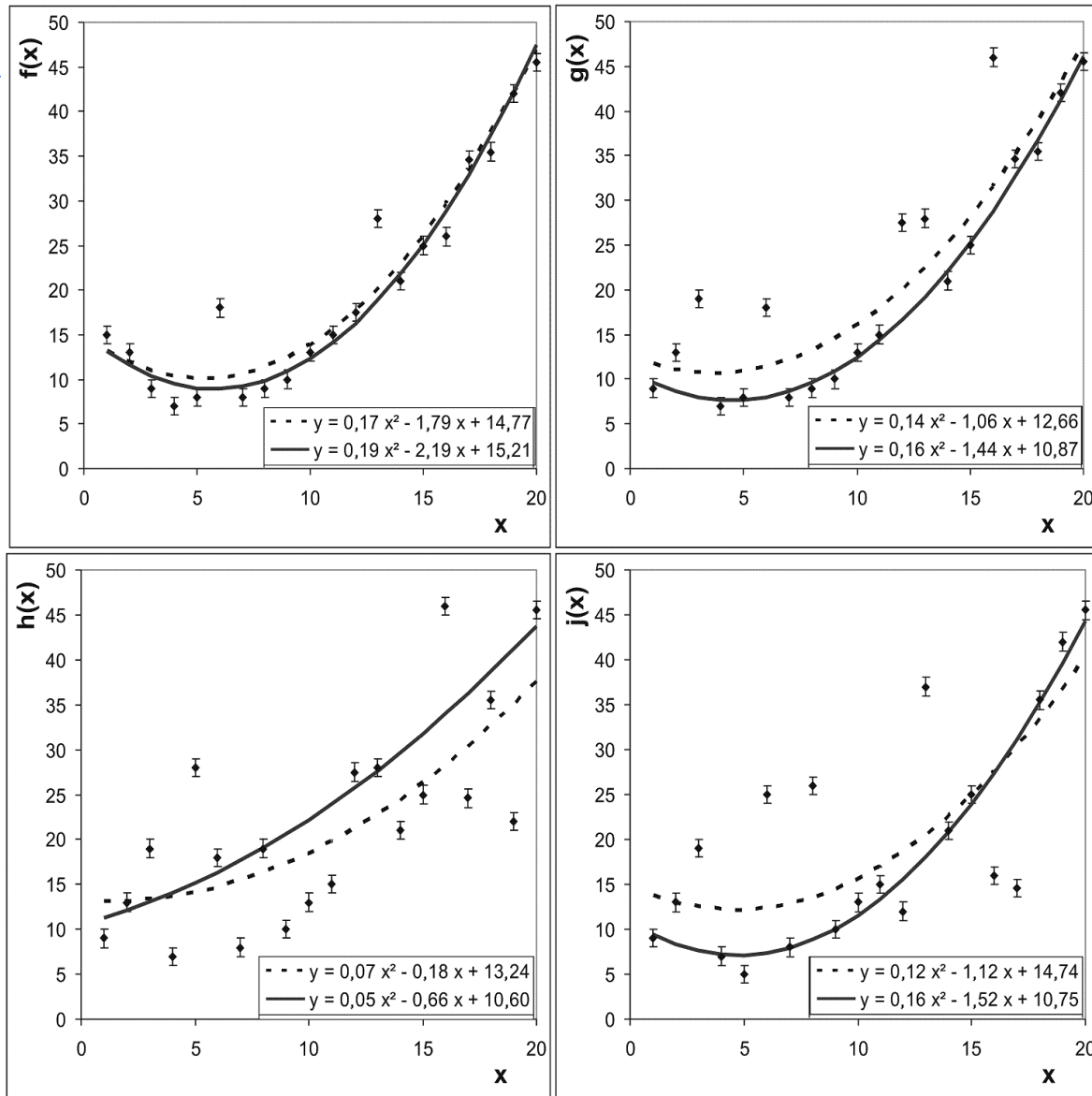
◆ Rozkład p-twa dla jednego punktu – porównanie metody bayesowskiej i najmniejszych kwadratów; prior rozmywa nam ogony



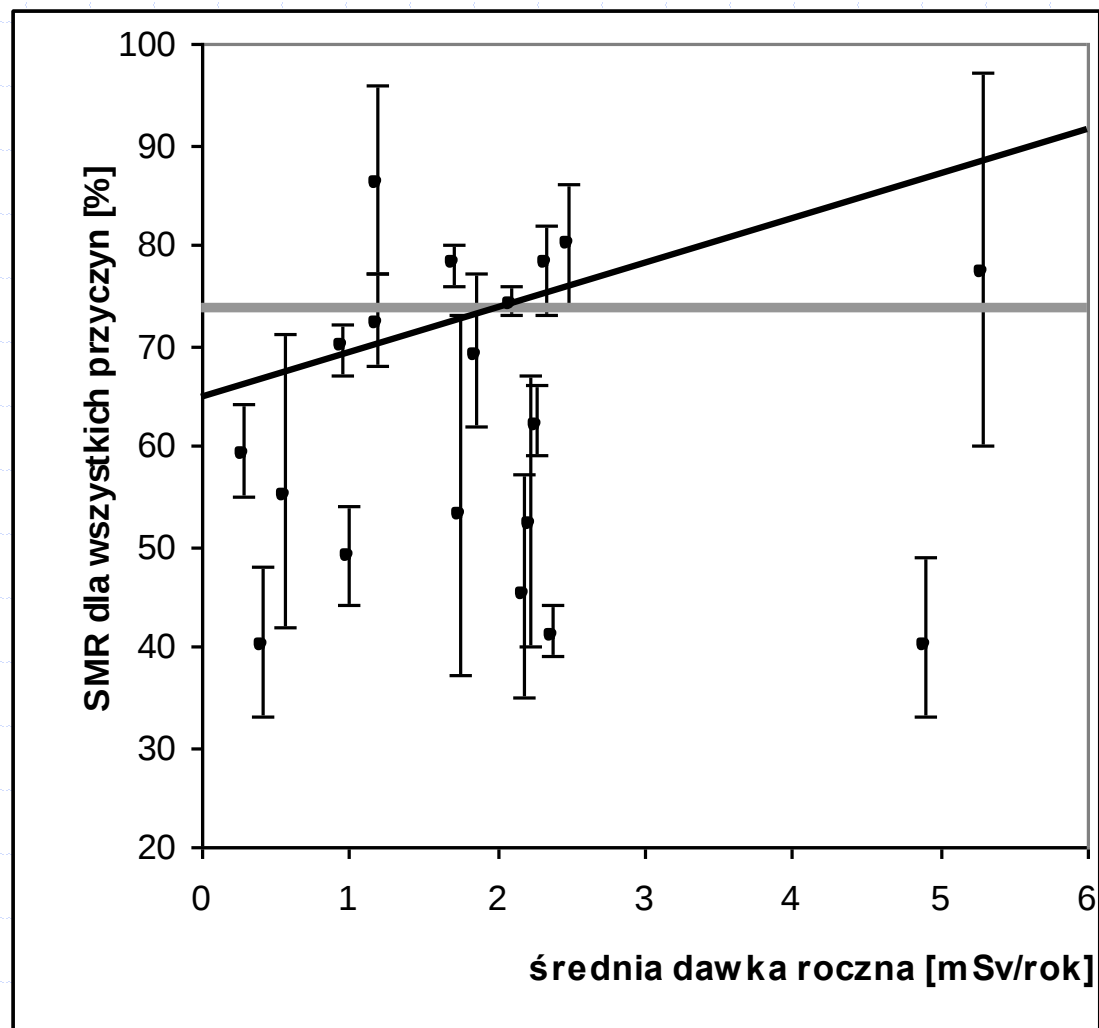
Dajmy jeszcze więcej punktów wybitych (tzw. outlier'ów)



Regresja dla wielomianu



Dla dużego rozrzutu danych



Bayesowska procedura wyboru najbardziej wiarygodnego modelu

Współczynnik wyboru modelu: $W_m = \frac{P(A | D, I)}{P(B | D, I)}$

gdzie A oraz B oznaczają dwa porównywane modele, D dane do których się one odnoszą, a I wszelkie wcześniejsze informacje. Jeśli wartość W_m będzie większa od jedności, wówczas model A jest bardziej prawdopodobny.

Prawdopodobieństwo słuszności modelu A dla danych D oraz modelu B dla danych D można zapisać w myśl twierdzenia Bayesa jako:

$$\begin{cases} P(A|D,I) = P(D|A,I) \times P(A|I) / P(D|I) \\ P(B|D,I) = P(D|B,I) \times P(B|I) / P(D|I) \end{cases}$$

Wybór modelu – c.d.

W sytuacji, jeśli żaden z modeli A i B nie jest z góry preferowany, człony $P(A/I)$ oraz $P(B/I)$ są sobie równe i po podstawieniu ulegają skróceniu. Analogicznie skracają się identyczne człony $P(D/I)$.

Dla modelu B , zawierającego jeden parametr dopasowania λ , używając procedury marginalizacji względem tego parametru, człon $P(D/B,I)$ zapisać można bayesowsko jako

$$P(D | B, I) = \int P(D, \lambda | B, I) d\lambda = \int P(D | \lambda, B, I) \cdot P(\lambda | B, I) d\lambda$$

Obierając za prior rozkład jednostajny ciągły otrzymujemy:

$$W_m = \frac{P(A | D, I)}{P(B | D, I)} = \frac{P(A | I)}{P(B | I)} \cdot \frac{P(D | A, I)}{P(D | \lambda_0, B, I)} \cdot \frac{\lambda_{\max} - \lambda_{\min}}{\delta\lambda\sqrt{2\pi}}$$

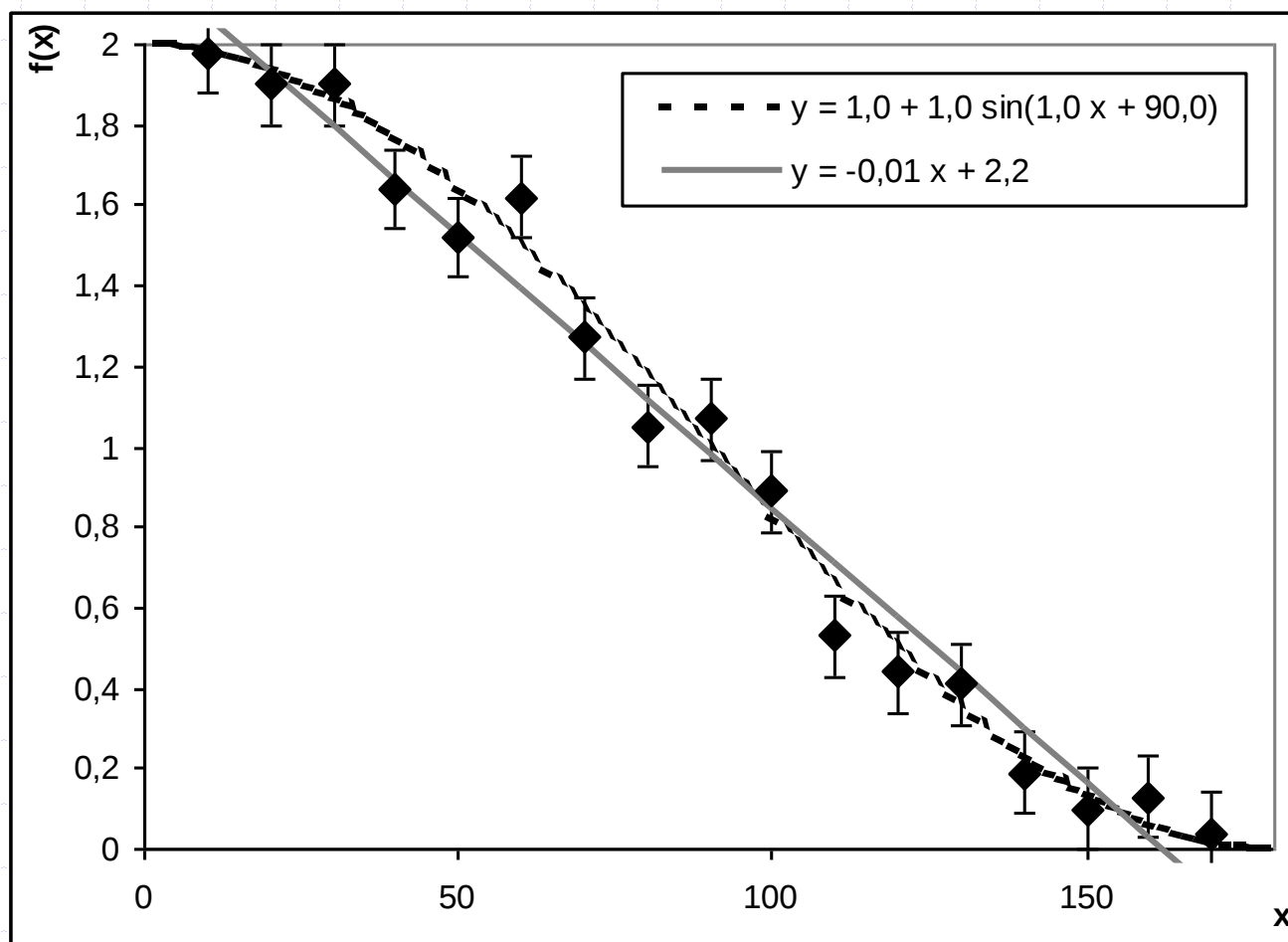
$$W_m \approx \frac{P(D | A, I)}{P(D | \lambda_0, B, I)} \cdot \frac{\lambda_{\max} - \lambda_{\min}}{\delta\lambda\sqrt{2\pi}}$$

Wybór modelu – c.d.

Dla modeli A i B zadanych wielomianami (postaci T_A i T_B o k i m parametrach) dopasowanymi do punktów eksperymentalnych E_i z niepewnościami σ_i otrzymujemy ostatecznie

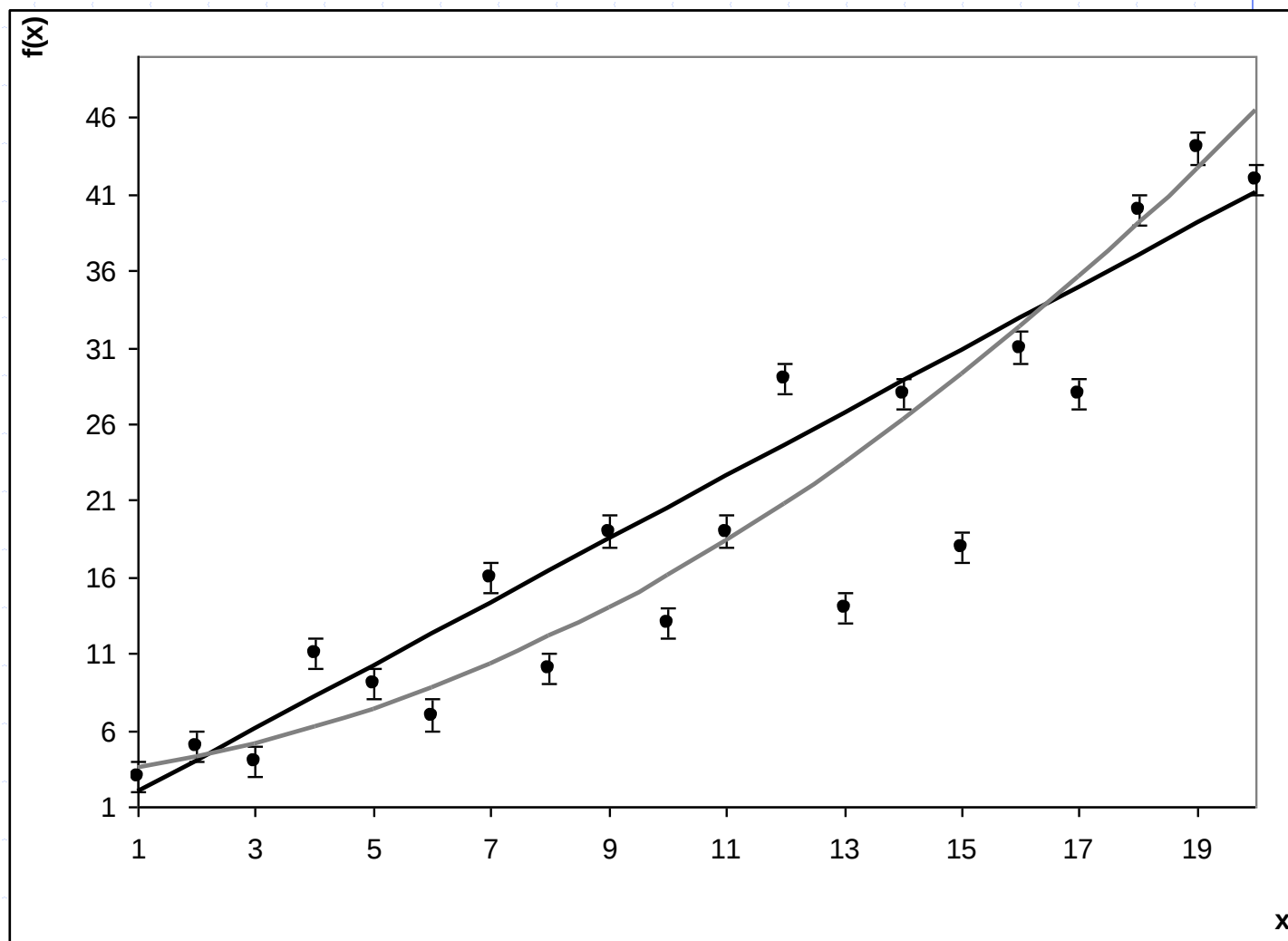
$$W_m = \frac{\sum_{i=1}^N \frac{1}{(T_{Ai} - E_i)^2} \left[1 - \exp\left(-\frac{(T_{Ai} - E_i)^2}{2\sigma_{0i}^2}\right) \right]}{\sum_{i=1}^N \frac{1}{(T_{Bi} - E_i)^2} \left[1 - \exp\left(-\frac{(T_{Bi} - E_i)^2}{2\sigma_{0i}^2}\right) \right]} \times \frac{\prod_{b=1}^m \frac{b_{\max}^{(B)} - b_{\min}^{(B)}}{\sigma_b^{(B)} \sqrt{2\pi}}}{\prod_{a=1}^k \frac{a_{\max}^{(A)} - a_{\min}^{(A)}}{\sigma_a^{(A)} \sqrt{2\pi}}}$$

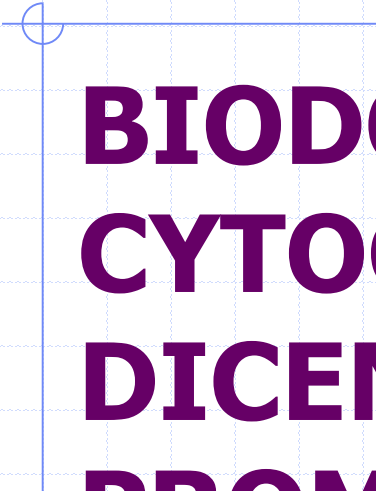
Przykład zastosowania #1



Przykład zastosowania #2

◆ Także i tu model liniowy okazał się bardziej prawdopodobny (brzytwa Ockham'a)





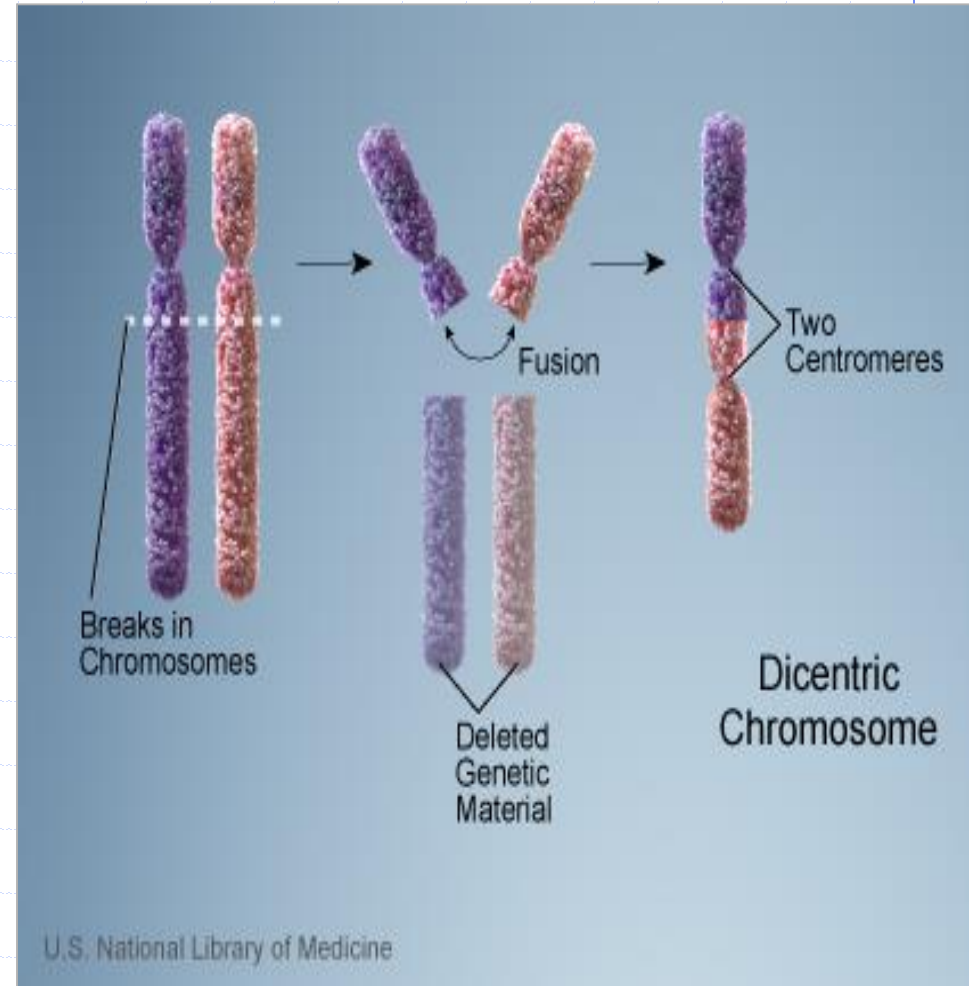
BIODOZYMETRIA CYTOGENETYCZNA DICENTRYKÓW DLA PROMIENIOWANIA MIESZANEGO

Biodozymetria cytogenetyczna dicentryków dla promieniowania mieszanego

- ◆ Biodozymetria cytogenetyczna – zespół metod/technik szacowania dawki na podstawie uszkodzeń DNA
- ◆ Dicentryki – aberracje chromosomowe dicentryczne (tj. posiadające 2 centromery)
- ◆ Promieniowanie mieszane - chodzi o napromienienie różnymi typami promieniowania jonizującego jednocześnie – w naszym przypadku będzie to promieniowanie gamma oraz neutronowe (występuje w awariach jądrowych!)

Dicentryczne aberracje chromosomowe (dicentryki)

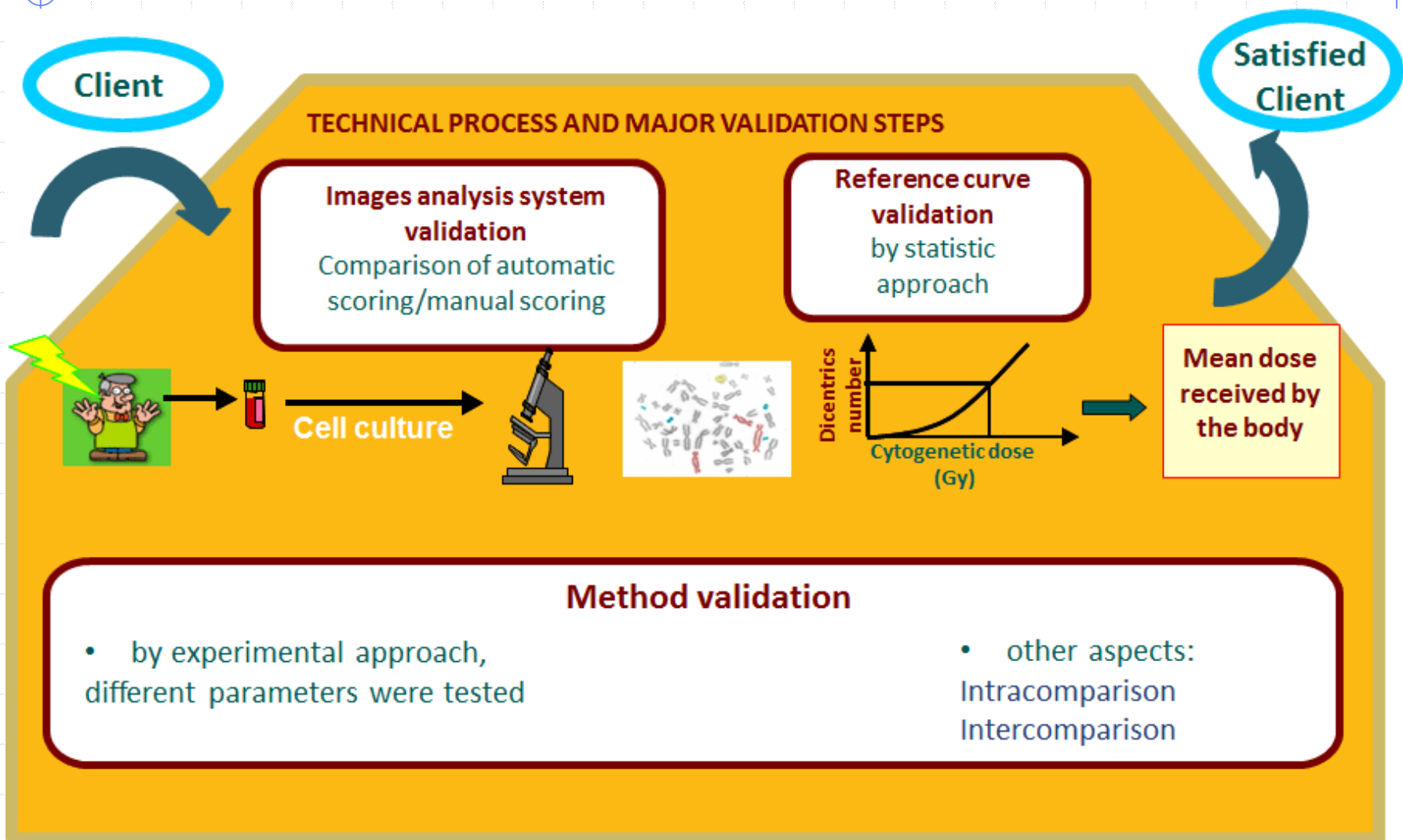
- Jedna z najgroźniejszych aberracji – chromosom dicentryczny i fragment acentryczny
- Powstaje w wyniku połączenia się dwóch chromosomów
- Aberracja nietrwała
- Prowadzi do śmierci komórki
- Fragment acentryczny, który nie zawiera centromeru w czasie mitozy nie jest segregowany do jądra komórkowego. Fragment taki obserwuje się najczęściej jako mikrojądro



Biodozymetria na dicentrykach – c.d.

- ◆ Jest to w chwili obecnej najpowszechniej używana technika w Polsce
- ◆ CLOR posiada akredytację na przeprowadzanie tego rodzaju badań
- ◆ W praktyce jest wiele metod obliczeniowych: iteracyjna, analityczna, bayesowska, Monte Carlo

Procedura w uproszczeniu



Zalety limfocytów krwi obwodowej jako modelu komórkowego do analizy dicentryków

- Łatwość pobierania
- Możliwość przechowywania i transportowania
- Zdolność do dzielenia się w hodowli po stymulacji mitogenem
- Cyrkulacja pomiędzy krwią obwodową i tkankami
- Łatwe do rozpoznania
- Powtarzalność i zgodność wyników in vivo i in vitro

Cytogenetyczna rekonstrukcja dawki

Częstość występowania dicentryków w limfocytach krwi obwodowej osoby narażonej na działanie promieniowania przelicza się na wartość dawki pochłoniętej za pomocą współczynników opracowanej in vitro **krzywej wzorcowej (kalibracyjnej)**.

Fizyczna jakość promieniowania użytego do opracowania krzywej wzorcowej powinna być jak najbardziej zbliżona do promieniowania, które jest przedmiotem oceny.

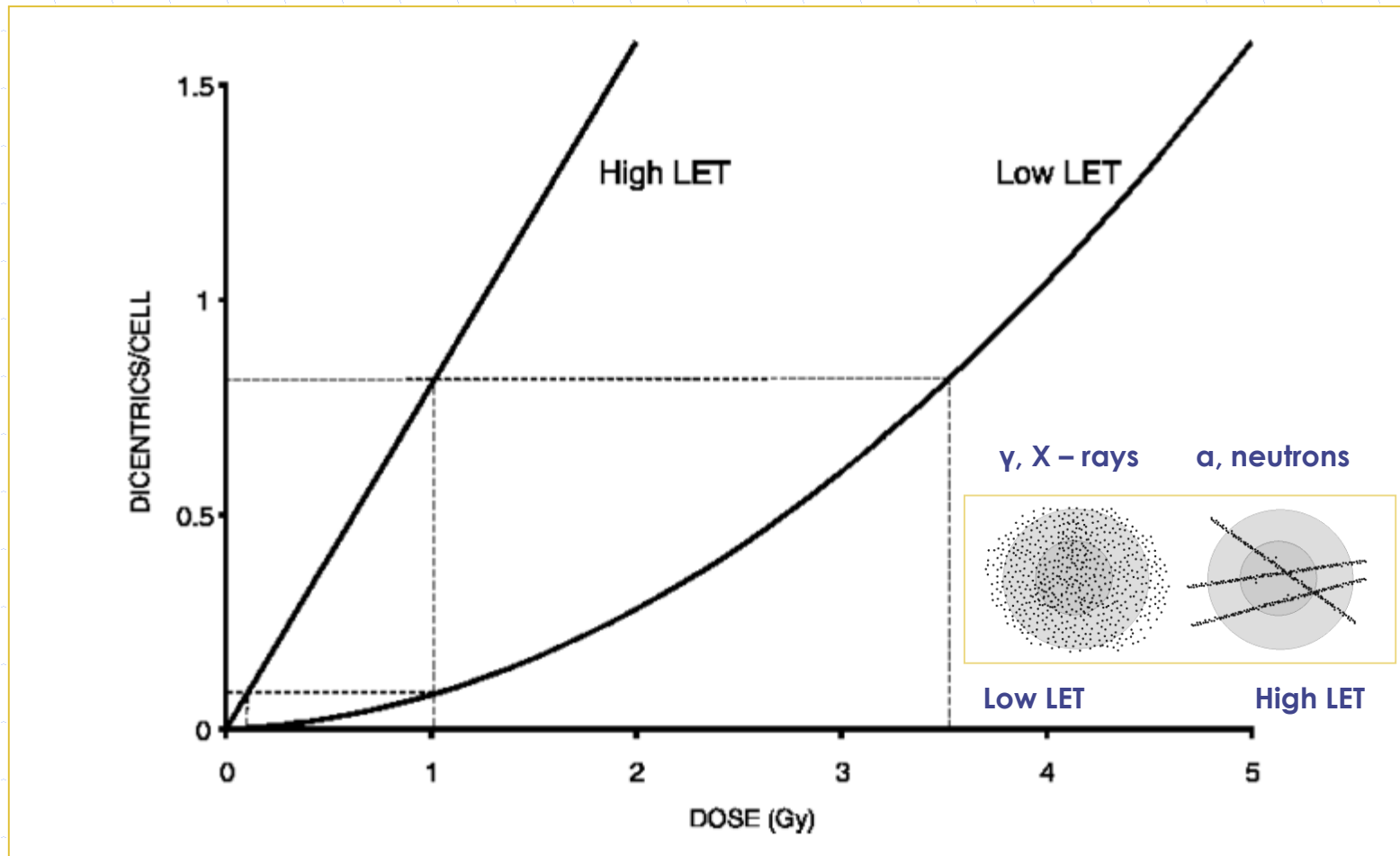
Innymi słowy: w warunkach laboratoryjnych zadajemy znaną dawkę znanym typem promieniowania i patrzymy jaka liczba dicentryków jej odpowiada (in-vitro). To nam wyznacza krzywą kalibracyjną, dzięki której później szacujemy nieznaną dawkę po zliczeniu częstości występowania dicentryków w rzeczywistym przypadku napromienionego człowieka.

- Mała częstość spontaniczna, ok. 1-5 dicentryków na 1000 komórek

Cytogenetyczna rekonstrukcja dawki

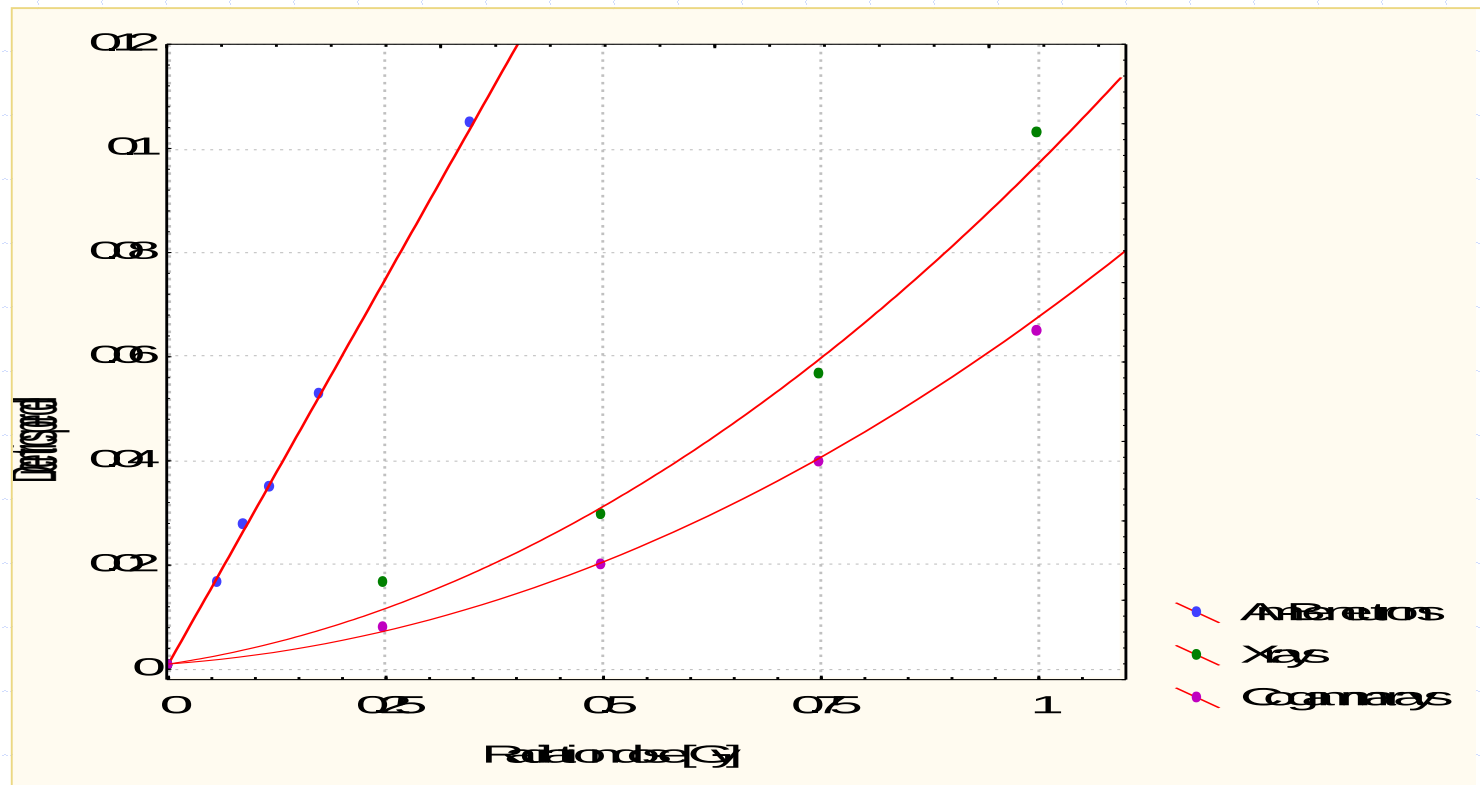


Krzywa kalibracyjna dla dicentryków



$$Y = C + \alpha D + \beta D^2$$

$$Y = C + \alpha D$$



Y is the yield of dicentrics, D is the dose, C is the control (background frequency), α is the linear coefficient, β is the dose squared coefficient.

Krzywe kalibracyjne w CLOR dla prom. mieszanego

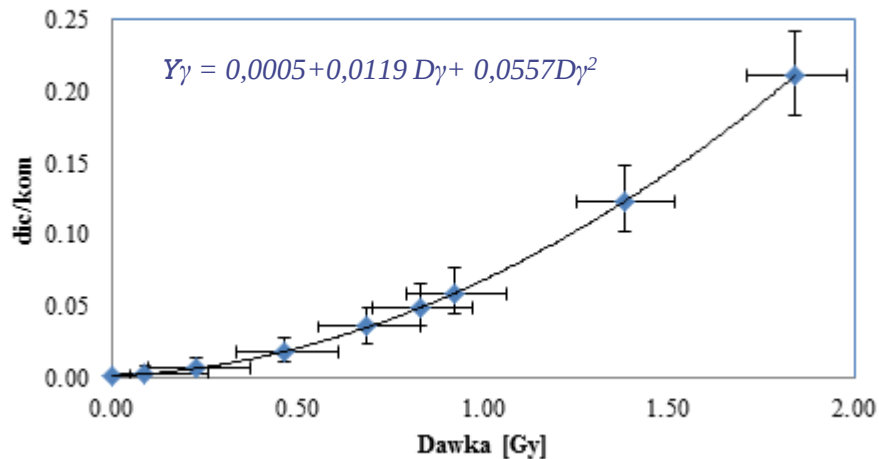
prom. γ

$$Y = C + \alpha D + \beta D^2$$

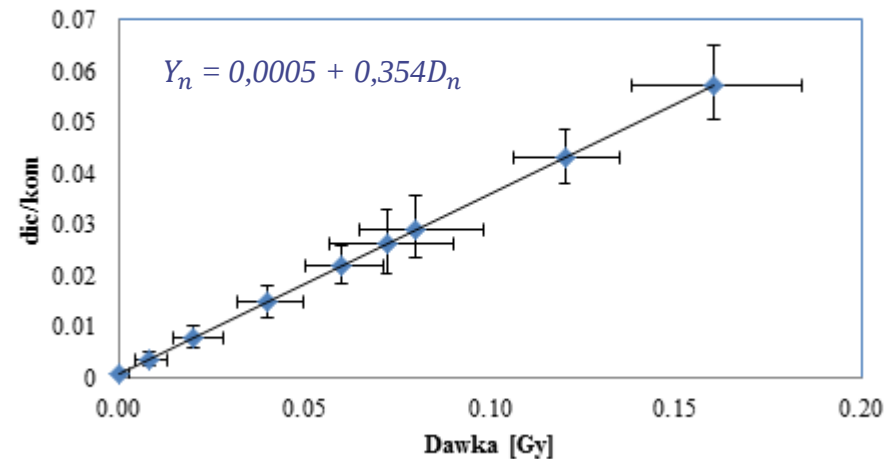
prom. neutronowe

$$Y = C + \alpha D$$

Krzywa dawka - skutek od promieniowania gamma

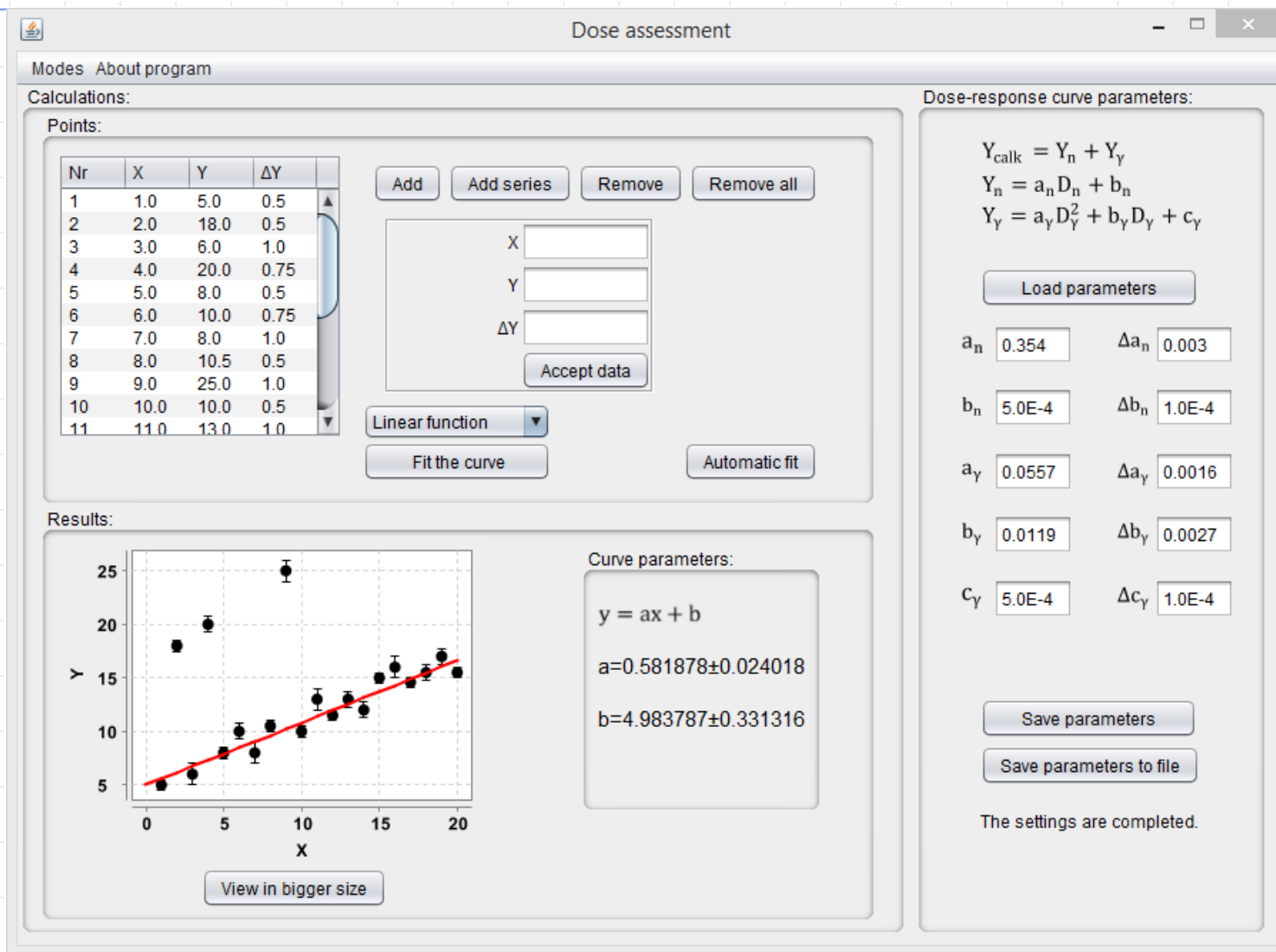


Krzywa dawka - skutek dla neutronów



Y – częstość dicentryków,
 D – dawka,
 C – parametr kontrolny,
 α – współczynnik liniowy,
 β – współczynnik kwadratowy.

Krzywe można wyznaczać metodami klasycznymi lub bayesowskimi (gdy są outlier'y)



Jak szacować dawkę?

- ◆ Mając już wcześniej wyznaczoną in-vitro krzywą kalibracyjną i mikroskopowo zliczone dicentryki możemy rozpocząć szacowanie dawki
- ◆ Problem: dicentryki po fotonach niczym się nie różnią od dicentryków po neutronach
- ◆ Zazwyczaj znamy stosunek strumienia gamma/neutrony → metody klasyczne (iteracyjna i analityczna)
- ◆ Ale gdy nie znamy, to pojawia nam się dodatkowy problem → zakładamy pewien rozkład prawdopodobieństwa stosunku gamma/neutrony jako nasz prior → metody bayesowskie, metoda Monte Carlo

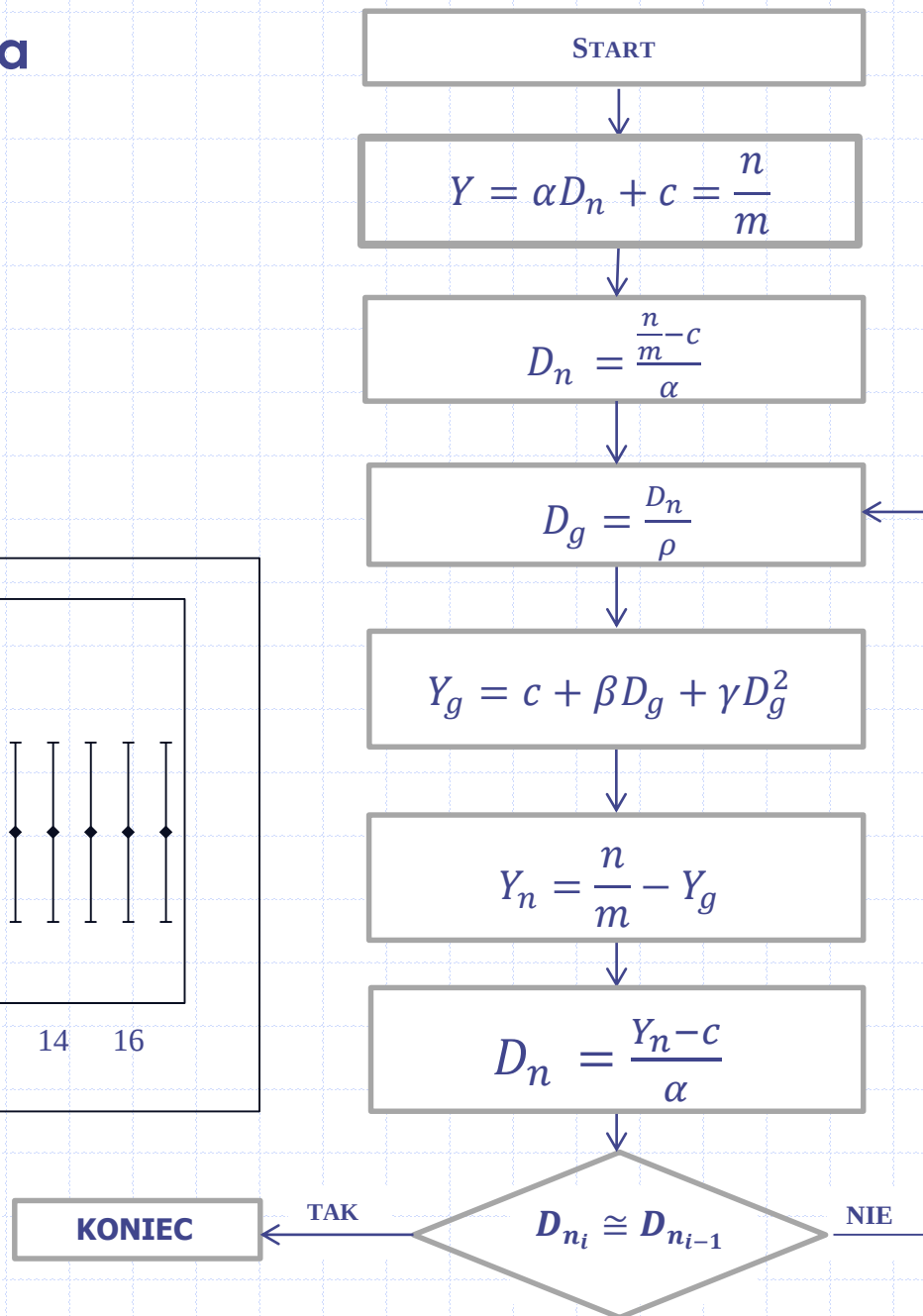
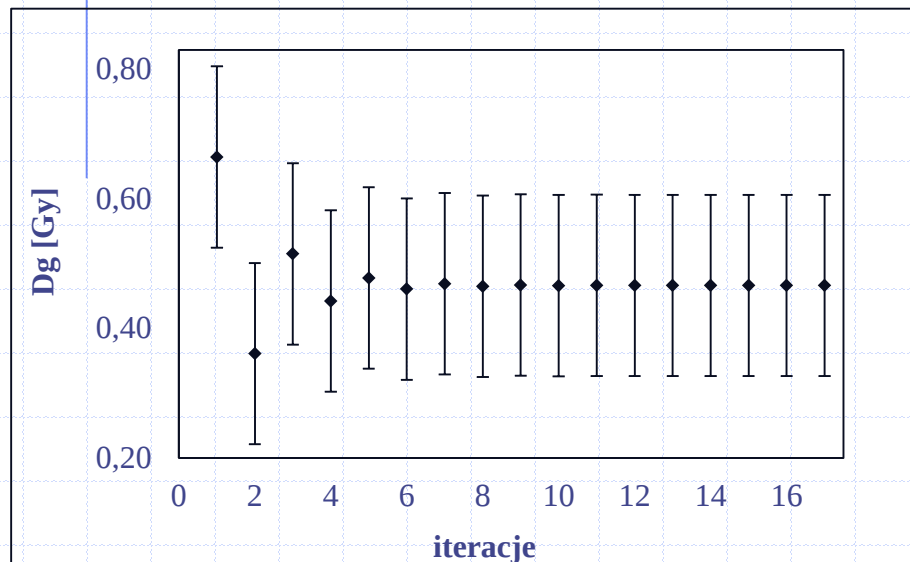
Metoda iteracyjna

$$Y_{n+\gamma} = Y_n + Y_\gamma$$

$$\rho = D_n / D_\gamma$$

$$Y_\gamma = c + \alpha D + \beta D^2$$

$$Y_n = c + \alpha D$$



Metoda iteracyjna i analityczna

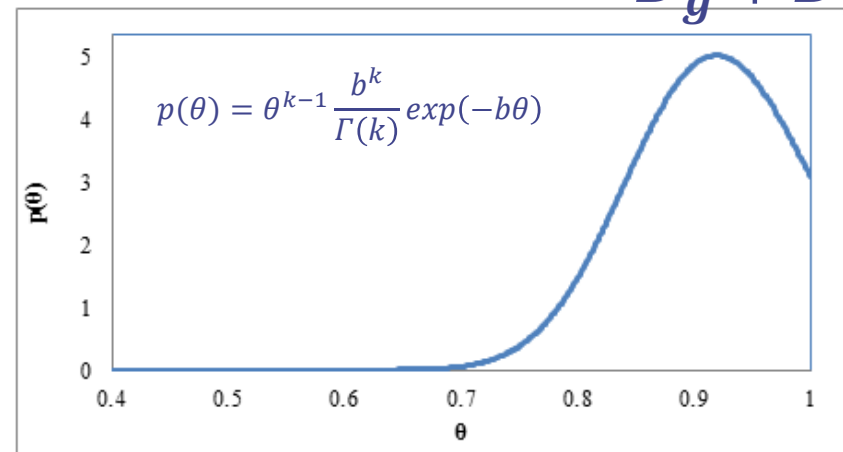
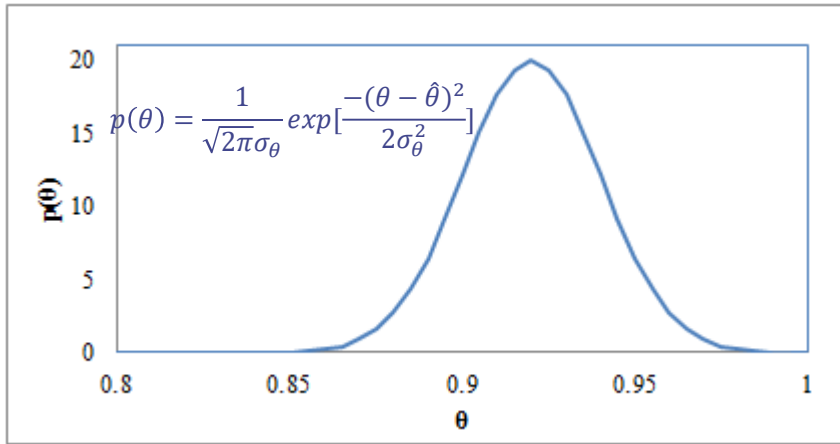
- ◆ Metoda iteracyjna jest prosta, ale poprzez efekt motyla przy kolejnych przybliżeniach zwiększa się nam niepewność
- ◆ De facto metoda ta bazuje na 2 równaniach z 2 niewiadomymi, więc wystarczy je rozwiązać (→ metoda analityczna)

Metody bayesowskie

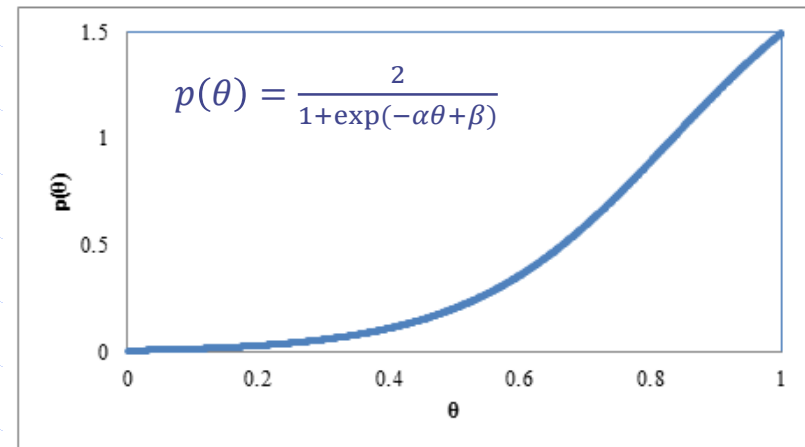
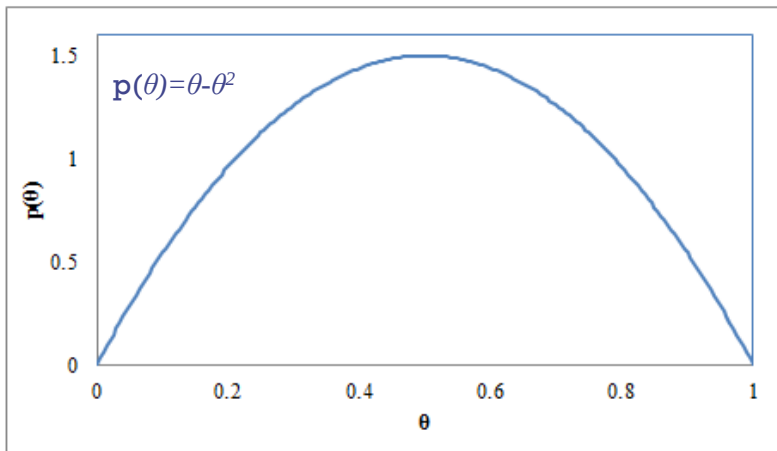
- ◆ Metoda iteracyjna i analityczna zakładają, że znamy stosunek strumienia fotonów do neutronów (dzięki klasycznej dozymetrii fizycznej)
- ◆ Ale dla przypadkowego narażenia osoby nieposiadającej dozymetru nie mamy tej informacji
- ◆ Wówczas stosunek ten opisujemy za pomocą prioru

- Informatywny prior $p(\theta)$:

$$\theta = \frac{D_g}{D_g + D_n}$$



- Nieinformatywny prior $p(\theta)$:

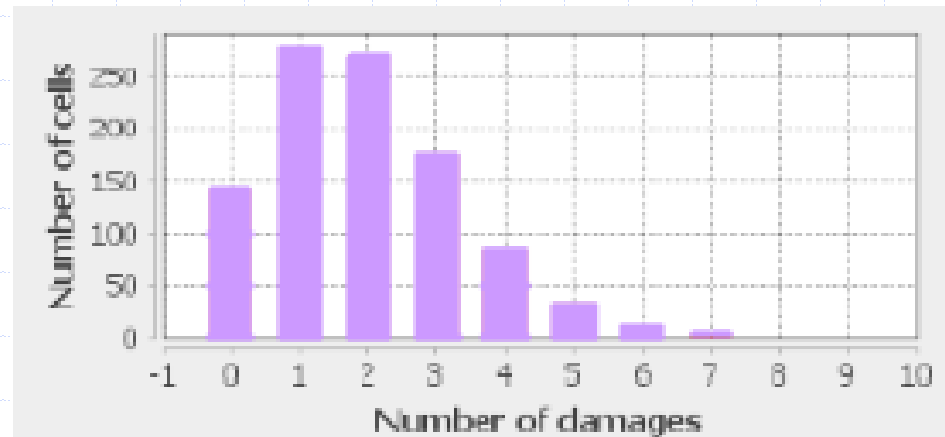


Bayesowska funkcja wiarygodności (statystyka Poissona)

$$L = \frac{(my_f)^n e^{-my_f}}{n!}$$

$$\rho = \frac{D_n}{D_g}$$

$$\theta = \frac{D_g}{D_g + D_n} = \frac{1}{\rho + 1}$$

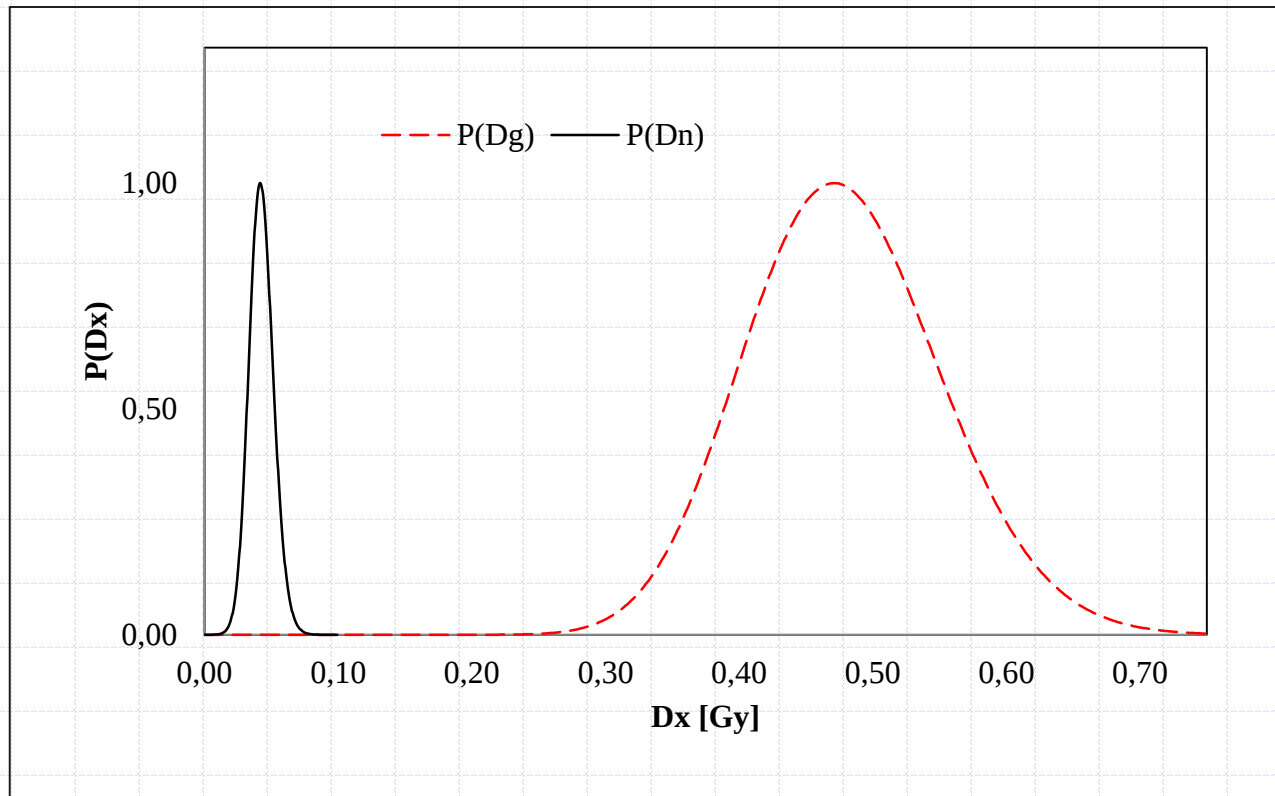


$$L(D_g|\theta) = \frac{[m(c + \alpha \frac{1-\theta}{\theta} D_g + \beta D_g + \gamma D_g^2)]^n}{n!} \cdot e^{-m(c + \alpha \frac{1-\theta}{\theta} D_g + \beta D_g + \gamma D_g^2)}$$

$$L(D_n|\theta) = \frac{[m(c + \alpha D_n + \beta \frac{\theta}{1-\theta} D_n + \gamma (\frac{\theta}{1-\theta} D_n)^2)]^n}{n!} \cdot e^{-m(c + \alpha D_n + \beta \frac{\theta}{1-\theta} D_n + \gamma (\frac{\theta}{1-\theta} D_n)^2)}$$

Dzięki temu mamy wynikowe prawdopodobieństwo a-posteriori

$$P(D_x|\theta) = \int_0^1 L(D_x|\theta) p(\theta) d\theta$$





Modes About program

Iterative method

Analytical method

Bayesian method

Calculations:

Prior:

Gaussian

i

$$p(\theta) = \frac{1}{\sqrt{2\pi}\sigma_\theta} \exp\left[-\frac{(\theta - \hat{\theta})^2}{2\sigma_\theta^2}\right]$$

 $\hat{\theta}$ 0.92

i

 σ_θ 0.92

Draw prior

Sample parameters:

Number of dicentrics:

i

Number of cells:

Calculate

Results:

 D_n [Gy] D_Y [Gy] Y_n [dic/cell] Y_Y [dic/cell]

Dose-response curve parameters:

$$Y_{\text{calc}} = Y_n + Y_Y$$

$$Y_n = a_n D_n + b_n$$

$$Y_Y = a_Y D_Y^2 + b_Y D_Y + c_Y$$

Load parameters

 a_n 0.354 Δa_n 0.003 b_n 5.0E-4 Δb_n 1.0E-4 a_Y 0.0557 Δa_Y 0.0016 b_Y 0.0119 Δb_Y 0.0027 c_Y 5.0E-4 Δc_Y 1.0E-4

Save parameters

Save parameters to file

The settings are completed.

Initial data

Measured chromosome aberration frequency: Y_f

Calibration curves:

$$Y_n = \alpha D_n + c$$

 α β γ c

$$Y_g = \beta D_g + \gamma D_g^2 + c$$

 $\Delta\alpha$ $\Delta\beta$ $\Delta\gamma$ Δc

Open from file

Save to file

Clear

☒ Use a prior probability distribution 

$$p(\theta) \propto \frac{1}{\sqrt{2\pi}\sigma_\theta} \exp\left[-\frac{(\theta - \hat{\theta})^2}{2\sigma_\theta^2}\right]$$

Gaussian

Gaussian

Scaled Gaussian

Beta

Simplified Beta

Sigmoidal

Avrami

Uniform

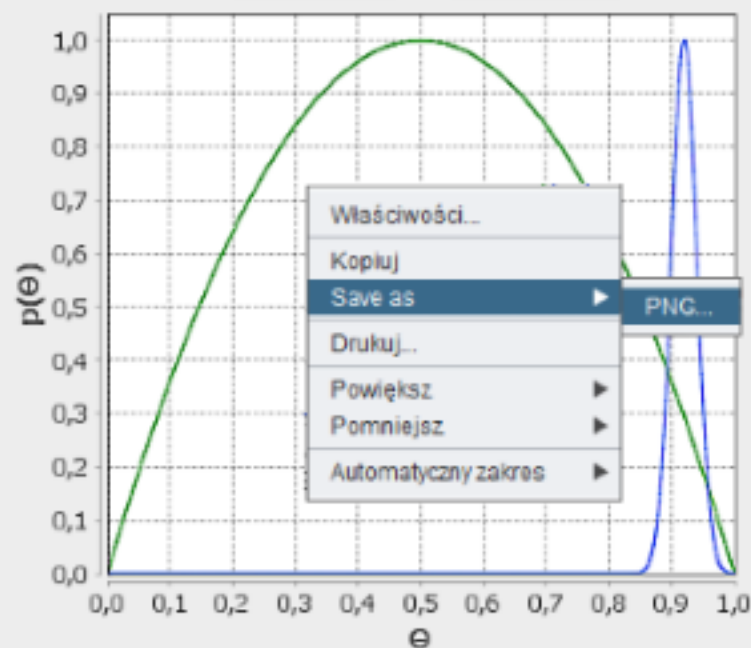
Constant

 σ_θ σ_θ p σ_p

Calculate

Draw prior

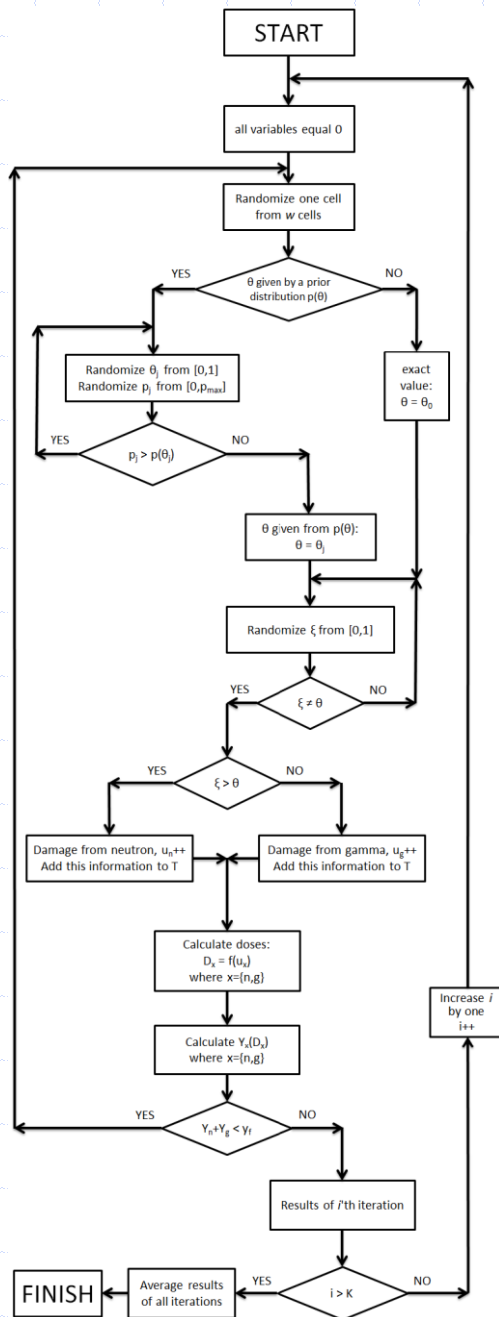
Clear graph



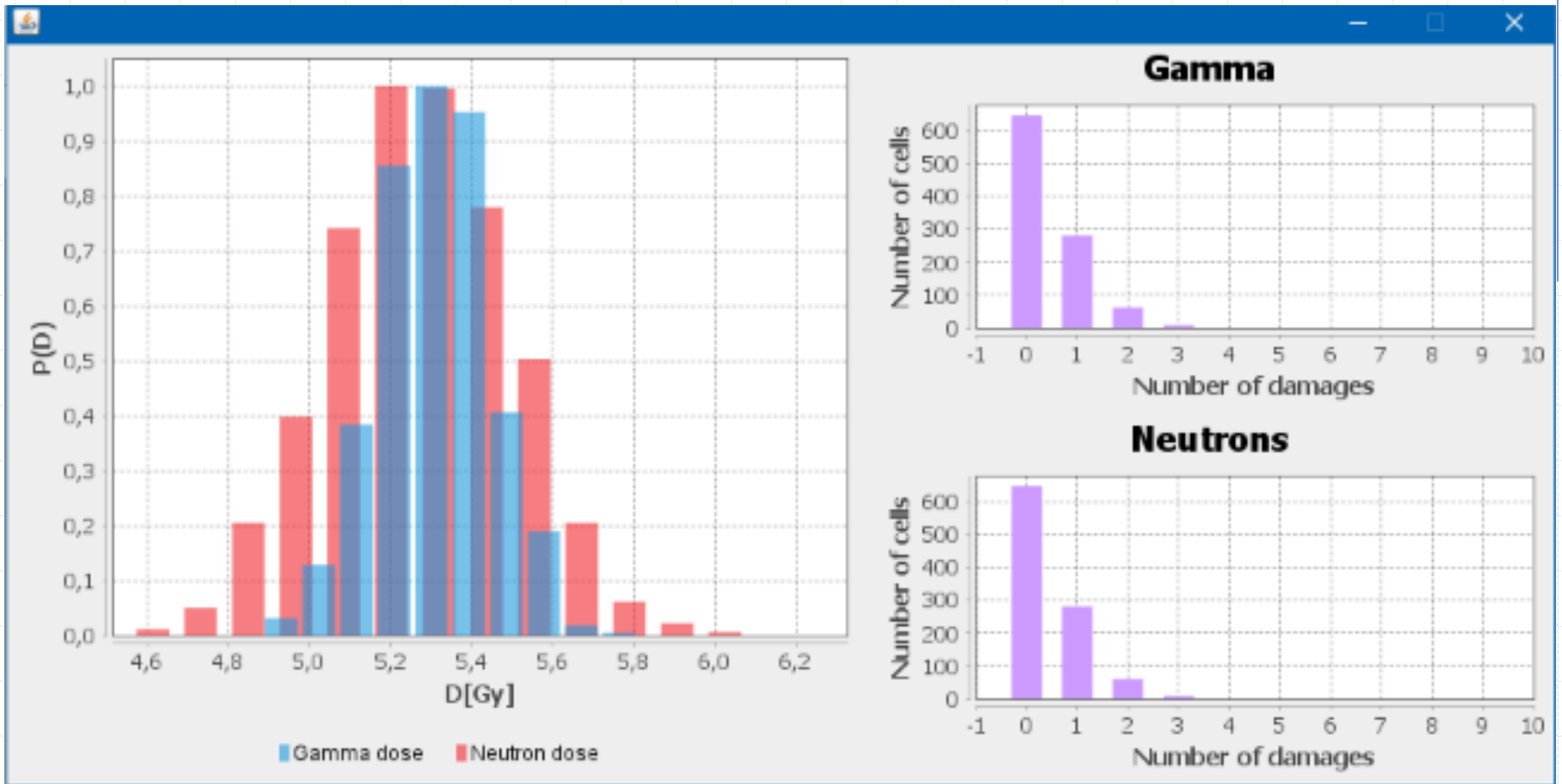
Results

 D_n [Gy] D_g [Gy] D_{n+g} [Gy] Y_n [%/cell] Y_g [%/cell] Y_{n+g} [%/cell]

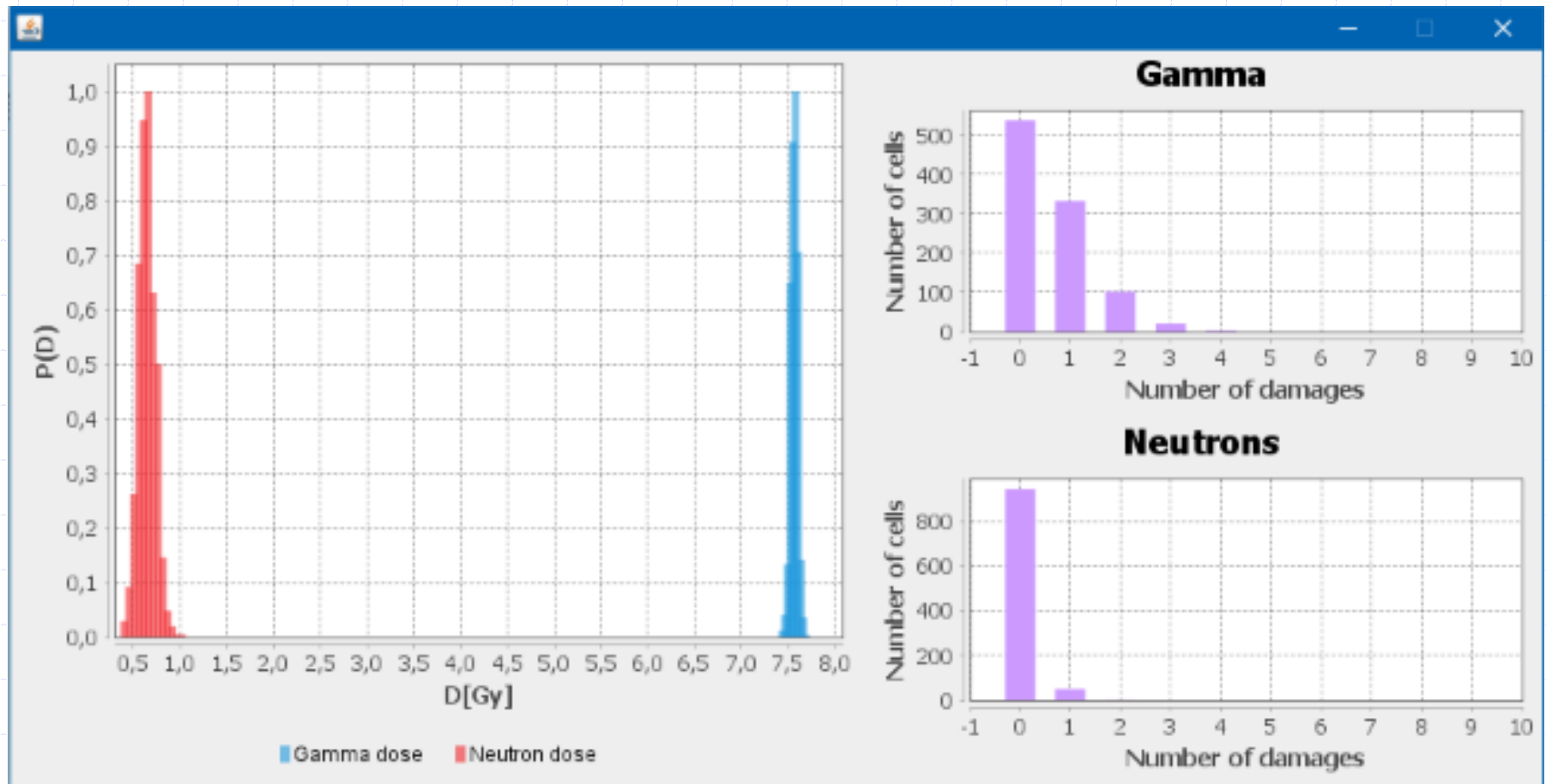
Metoda Monte Carlo



- podanie wartości θ lub priora
- losowanie komórek z próbki wirtualnej
- program wykonuje kolejne iteracje, aż do momentu osiągnięcia w próbce podanej na wstępie częstości aberracji
- obliczenie dawek pochłoniętych od neutronów i od promieniowania gamma



Results for $y_f = 3.5$ [dic/cell], $\theta = 0.5$, $D_t = 10.59$ [Gy].



Results for $y_f = 3.5$ [dic/cell], $\theta = 0.92$, $D_t = 8.2$ [Gy].

Alternative statistical methods for cytogenetic radiation biological dosimetry

Krzysztof Wojciech Fornalski ^{1,2}

¹⁾ PGE EJ 1 Sp. z o.o., Technology and Operations Office, ul. Mysia 2, 00-496 Warszawa, Poland

²⁾ contract of specified task with Central Laboratory for Radiological Protection (CLOR), ul. Konwaliowa 7, 03-194 Warszawa, Poland
e-mail: krzysztof.fornalski@gmail.com, krzysztof.fornalski@gkpge.pl

ABSTRACT

The paper presents alternative statistical methods for biological dosimetry, such as the Bayesian and Monte Carlo method. The classical Gaussian and robust Bayesian fit algorithms for the linear, linear-quadratic as well as saturated and critical calibration curves are described. The Bayesian model selection algorithm for those curves is also presented. In