

Generatory liczb pseudolosowych

Plan wykładu:

- liczby losowe
- liczby pseudolosowe
- generatory o rozkładzie równomiernym
- generatory o zadanym rozkładzie prawdopodobieństwa:
 - a) metoda eliminacji
 - b) metoda odwracania dystrybuanty
- generatory wielowymiarowe

literatura:

„Random number generation and Monte Carlo methods” – James E. Gentle
„Handbook of Monte Carlo methods” – D.P. Kroese, T. Taimre, Z.I. Botev
„Computational Physics - An Introduction” – F. J. Vesely

Generatory liczb pseudolosowych

Generatory możemy podzielić na: fizyczne i komputerowe

Generatory fizyczne:

- wykorzystują zjawiska fizyczne jak szum aparatury elektronicznej, rozpad jądrowy itp.
- zalety: generują rzeczywiste ciągi liczb losowych (oczekiwany przez nas brak przewidywalności wyniku)
- wady:
 - a) wymagają ciągłej obsługi/kalibracji (wpływ temperatury, wilgotności, czas – zanik radionuklidu)
 - b) konieczne jest przechowywanie dużej liczby stabilizowanych wartości
 - c) generują zmienne losowe tylko dla jednego typu rozkładu

Generatory komputerowe:

- działanie oparte na wykorzystaniu algorytmów numerycznych
- zalety: łatwość obsługi, dostępność generatorów dla dowolnego rozkładu prawdopodobieństwa, bardzo wydajne/ekonomiczne w użyciu, powtarzalność generowanych ciągów (ważne w testowaniu oprogramowania)
- wady:
 - a) generowane liczby to liczby pseudolosowe, a nie losowe – znając algorytm możemy przewidzieć kolejne elementy ciągu, brak losowości
 - b) generowany ciąg liczb ma rozkład dyskretny, a nie ciągły
 - c) inne – które można zminimalizować stosując inny algorytm
- wymienione wady generatorów komputerowych są mało istotne z praktycznego punktu widzenia, natomiast korzyści są nie do przecenienia.

Arytmetyka modulo

Składnikiem generatora o dowolnym rozkładzie prawdopodobieństwa jest zawsze generator o rozkładzie jednorodnym, ten z kolei wykorzystując arytmetykę **dzielenia modulo**

Liczba a jest kongruentna do b modulo m jeśli

$$a \equiv b \pmod{m}, \quad 0 \leq a < m$$

co odpowiada relacji $\frac{b-a}{m} = k$ (liczba całkowita)

własności (relacje równoważności)

- symetria $a \equiv b \pmod{m} \Rightarrow b \equiv a \pmod{m}$
- zwrotność $a \equiv a \pmod{m}$
- przechodność $a \equiv b \pmod{m} \wedge b \equiv c \pmod{m} \Rightarrow a \equiv c \pmod{m}$

Przykład liczb kongruentnych:

$a=5, b=14$ modulo 3 ($k=3$)

ale także

$a=14, b=5$ modulo 3 ($k=-3$)

Inny sposób wyrażenia kongruencji (wyznaczania liczby kongruentnej modulo m)

$$a \equiv b \pmod{m} = b - \left\lfloor \frac{b}{m} \right\rfloor m \quad \lfloor arg \rfloor - \text{największa liczba całkowita} \leq arg$$

wówczas zakres liczby a zostaje ograniczony przez m

$$a \leq b \wedge a < m$$

tę własność wykorzystujemy w generatorach liczb pseudolosowych

Generatory kongruencyjne dzielą się na: **liniowe** i **nieliniowe**.

Generatory nieliniowe – mają dobre własności statystyczne (np. liczby nie układają się na hiperpowierzchniach), ale wykorzystują mechanizm wyznaczania liczby odwrotnej modulo m co jest mało wydajne (np.: generator Eichenauera-Hermanna). W typowych zastosowaniach raczej rzadko wykorzystywane.

Generatory liniowe – przede wszystkim szybkie, dobierając odpowiednie parametry generatora można też uzyskać dobre własności statystyczne generowanych ciągów liczb pseudolosowych.

Uwaga: wszystkie generatory produkują liczby pseudolosowe, a nie losowe – mimo iż tak się może wydawać, po wykonaniu testów statystycznych. Uzasadnienie jest proste: znając algorytm i ziarno generatora możemy przewidzieć kolejne elementy ciągu – zatem brak jest tu losowości.

Definicja generatora liniowego

$$x_i = (a_1x_{i-1} + a_2x_{i-2} + \dots + a_nx_n + c) \bmod m, \quad m, x_i, a_i \in \mathbb{Z}^+$$

$\{a_1, a_2, \dots, a_n, c, m\}$ - parametry definiujące generator (**liczby całkowite**)

m – duża liczba całkowita określająca zakres generowanych liczb

$\{x_1, x_2, \dots, x_n\}$ - ciąg inicjujący tzw. ziarno generatora
(może być ustalone, dostarczane przez użytkownika,
lub generowane z zegara systemowego etc.)

$\{x_{n+1}, x_{n+2}, \dots\}$ - ciąg generowanych liczb pseudolosowych

Generator dostarcza liczby całkowite z zakresu

$$x_i \in [1, m - 1]$$

należy je znormalizować (unormować)

$$X_i = \frac{x_i}{m} \quad \rightarrow \quad X_i \in (0, 1) \quad \rightarrow \quad X_i \in U(0, 1)$$

Ponieważ rozkład liczb X_i jest jednorodny (ang. *uniform*) w przedziale $(0,1)$ dlatego oznaczany będzie jako **U(0,1)**. Jeżeli rozkład ma obejmować inny zakres $[A,B]$ to dokonujemy prostej transformacji

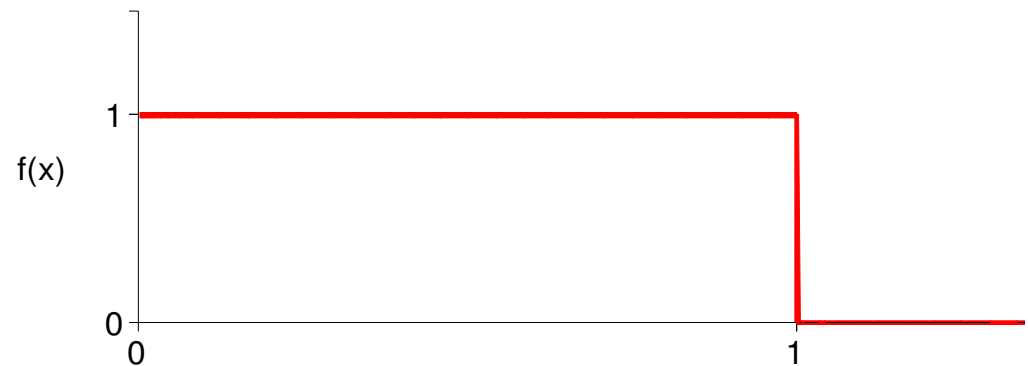
$$Y_i = A + X_i(B - A) \quad \rightarrow \quad Y_i \in U(A, B)$$

Końcowy ciąg liczb pseudolosowych ma rozkład jednorodny

funkcja gęstości prawdopodobieństwa

$$f_U(x) = \begin{cases} 1 & x \in (0, 1) \\ 0 & x \notin (0, 1) \end{cases}$$

rozkład jednorodny oznacza, że prawdopodobieństwo wylosowania dowolnej liczby jest takie samo (dla rozkładu ciągłego)



Własności dowolnego generatora liczb pseudolosowych:

- okres generatora - określa maksymalną długość łańcucha liczb pseudolosowych bez powtórzeń (okres generatora jest ograniczony liczbą m)
- ziarnistość ciągu – liczby pseudolosowe mają rozkład dyskretny, odległość między sąsiednimi liczbami jest odwrotnie proporcjonalna do okresu generatora
- własności statystyczne:
 - wartość oczekiwana x ,
 - odchylenie standardowe (wariancja) σ_x ,
 - współczynnik autokorelacji ciągu,
 - spełnienie testów statystycznych (zgodności rozkładu itp.)

Okres generatora

$$\underbrace{X_1, X_2, \dots, X_n}_{\text{ziarno}} \underbrace{X_{n+1}, X_{n+2}, \dots, X_{n+k}}_{\text{zakres aperiodyczności}} \underbrace{X_{n+k+1}, X_{n+k+2}, \dots, X_{n+k+T}}_{\text{okres generatora}} \underbrace{X_{n+k+T+1}, X_{n+k+T+2}, \dots, X_{n+k+2T}}_{\text{powtórzony ciąg}}$$

Dlaczego pożądane są generatory o dużych okresach?

Rozważmy proste całkowanie Monte Carlo (podstawy na następnym wykładzie)

$$C = \int_a^b g(x) f(x) dx \approx \frac{1}{T} \sum_{i=1}^T g(x_i)$$

$g(x)$ – funkcja podcałkowa
 $f(x)$ – rozkład prawdopodobieństwa (inaczej funkcja wagowa)
 T – okres generatora
 x_i – punkty wylosowane przy użyciu generatora o rozkładzie $f(x)$

Dokładność przybliżenia zależy od ilości wylosowanych punktów
 – zwiększmy więc dokładność biorąc 2 większą liczbę punktów.

$$C_2 = \frac{1}{2T} \sum_{i=1}^{2T} g(x_i) = \frac{1}{2T} \left[\sum_{i=1}^T g(x_i) + \underbrace{\sum_{i=T+1}^{2T} g(x_i)}_{\text{powtarzamy punkty}} \right] = \frac{1}{2T} 2 \sum_{i=1}^T g(x_i) = C$$

wynik identyczny jak dla pojedynczej sumy!

Pierwsza suma obejmuje niepowtarzający się ciąg liczb, druga suma korzysta dokładnie z tego samego ciągu, powiela go (całkowity determinizm) – całkowita zależność drugiej sumy od pierwszej nie powoduje jakichkolwiek zmian w wyniku.

Wykonaliśmy 2 razy więcej obliczeń a wynik się nie zmienił.

Ziarnistość generatora U(0,1)

- liczby pseudolosowe to liczby całkowite, które są normalizowane
- rozkład liczb jest dyskretny i jest wynikiem działania użytego algorytmu
- ponieważ generator liniowy ma rozkład jednorodny spodziewamy się, że **średnia odległość między najbliższymi sąsiadami** będzie zbliżona w dowolnym podprzedziale

$$\delta = \langle \min_{i < j} (x_i, x_j) \rangle$$

$$\delta \approx \frac{\Delta}{T} = \frac{1}{T}$$

$\Delta=1$ – szerokość znormalizowanego przedziału

Czy dyskretny rozkład liczb pseudolosowych jest wadą skoro odległości między nimi mogą być bardzo małe np. rzędu 10^{-9} lub mniejsze?

Ziarnistość generatora oznacza, że w metodzie Monte Carlo próbkujemy przestrzeń dopuszczalnych rozwiązań tylko w wybranych punktach. Omijając punkty, których wkład do rozwiązania jest duży (np. w pobliżu biegunów funkcji) popełniamy błąd.

Pozostaje tylko kwestia – jak duży jest ten błąd?

Im lepszy generator (duży okres) oraz im większa ilość wylosowanych liczb tym będzie on mniejszy.

Podstawowe parametry statystyczne: pierwszy i drugi moment, współczynnik autokorelacji ciągu

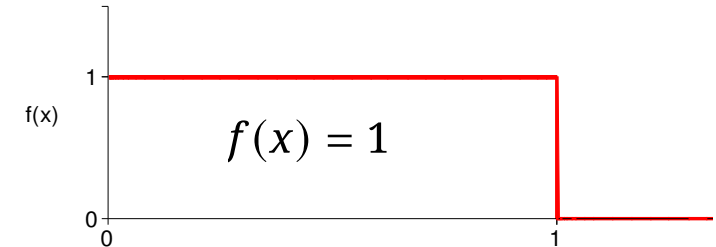
Dla rozkładu jednorodnego możemy wyznaczyć pierwszy i drugi moment analitycznie

- wartość oczekiwana x (pierwszy moment)

$$\langle x \rangle = \int_0^1 x f(x) dx = \left. \frac{x^2}{2} \right|_0^1 = \frac{1}{2}$$

- estymator

$$\bar{x} = \frac{1}{N} \sum_{i=1}^N x_i$$



- wariancja (drugi moment)

$$\begin{aligned} \sigma^2 &= \int_0^1 (x - \langle x \rangle)^2 f(x) dx = \int_0^1 (x^2 - 2\langle x \rangle x + \langle x \rangle^2) f(x) dx \\ &= \int_0^1 x^2 f(x) dx - 2\langle x \rangle \int_0^1 x f(x) dx + \langle x \rangle^2 \int_0^1 f(x) dx \\ &= \langle x^2 \rangle - \langle x \rangle^2 \end{aligned}$$

$$\sigma^2 = \langle x^2 \rangle - \langle x \rangle^2 = \frac{1}{12}$$

- estymator

$$\sigma^2 = \langle x - \langle x \rangle \rangle \approx \frac{1}{N} \sum_{i=1}^N (x_i - \bar{x})^2$$

Funkcja autokorelacji (unormowana) definiowana dla procesów stochastycznych (ciągłych lub dyskretnych)

$$R_r = \frac{E[(X_t - \bar{x})(X_s - \bar{x})]}{\sigma^2}, \quad r = t - s$$

t,s – numerują kolejność w ciągu

r – odległość między elementami ciągu

$$-1 \leq R_r \leq 1$$

określa zależność pomiędzy przesuniętymi o r elementami generowanego ciągu
($R_r=0$ oznacza brak zależności)

- estymatorem jest współczynnik autokorelacji

$$R_r = \frac{\frac{1}{(N-r)} \sum_{i=1}^{N-r} (x_i - \mu)(x_{i+r} - \mu)}{\sigma^2}$$

- generatory bada się ze względu na występowanie korelacji (zależności) w ciągu liczb pseudolosowych.
- testy wykonuje się dla niewielkich przesunięć

$$r = s - t = 1, 2, 3, 4, 5, 6$$

- **od generatora oczekujemy, że współczynnik R_r będzie bliski 0 - jak najmniejsza korelacja elementów**

Liniowy generator multiplikatywny

- najprostszy generator ma postać

$$x_i \equiv ax_{i-1} \pmod{m}, \quad 0 < x_i < m$$

przykładowe ciągi liczb dla $X_0 = 1, m = 11$

	a=1	2	3	4	5	6	7	8	9	10
i=0	1	1	1	1	1	1	1	1	1	1
1	1	2	3	4	5	6	7	8	9	10
2		4	9	5	3	3	5	9	4	1
3		8	5	9	4	7	2	6	3	
4		5	4	3	9	9	3	4	5	
5		10	1	1	1	10	10	10	1	
6		9				5	4	3		
7		7				8	6	2		
8		3				4	9	5		
9		6				2	8	7		
10		1				1	1	1		

- dla ustalonego m , zmiana parametru a wpływa na okres generatora
- maksymalny okres nie przekracza wartości $m-1$ (wykluczamy 0 oraz m)

Okres generatora determinuje najmniejsza dodatnia liczba k spełniająca relację

$$a^k \equiv 1 \pmod{m}$$

Na tym prostym przypadku widać że dobór parametrów generatora jest kluczowy dla jego okresu. Dość często m jest **liczbą pierwszą Mersenne'a**

$$m = 2^p - 1, \quad p - \text{liczba całkowita}$$

- częstym wyborem jest $a = 7^5 = 16807$, $m = 2^{31}$
- **inny wybór – generator RANDU** $a = 65539 = 2^{16} + 3$ $m = 2^{31}$

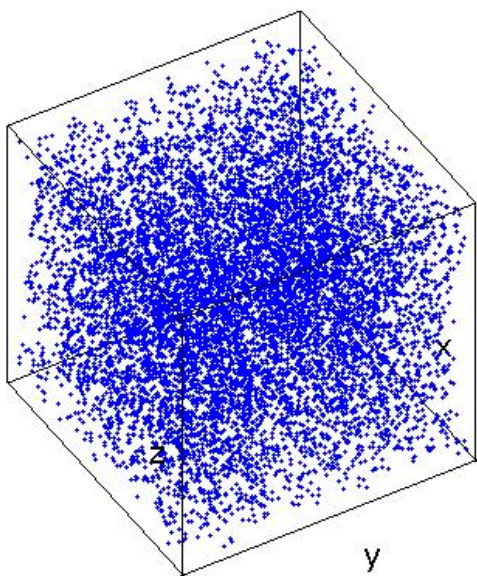
$$x_{i+1} = ax_i \bmod m = (2^{16} + 3)x_i \bmod m$$

$$x_{i+2} = ax_{i+1} \bmod m = (2^{16} + 3)x_{i+1} \bmod m = (2^{16} + 3)^2 x_i \bmod m$$

$$x_{i+2} = 6x_{i+1} - 9x_i \quad \text{- silna zależność (korelacja) trzech kolejnych liczb}$$

Silna korelacja elementów ciągu powoduje że układają się one na hiperpowierzchniach
- generowane punkty próbują w sposób wybiórczy przestrzeń.

przypadkowy rzut wygląda dobrze

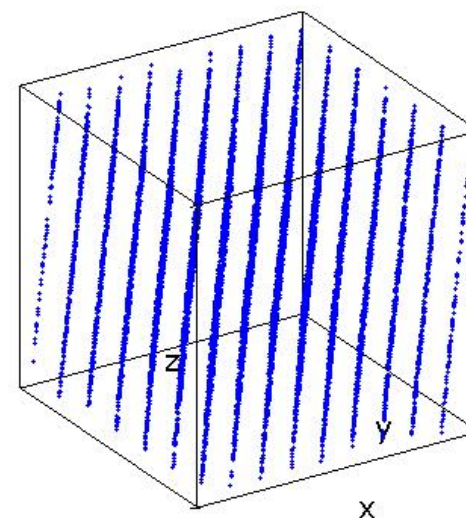


**10 tysięcy punktów
wygenerowanych
przez RANDU**

rysujemy punkty

$$[X_i, X_{i+1}, X_{i+2}]$$

... ale po obrocie odkrywamy, że wyniki nie są
takimi jakich oczekujemy – widać hiperpłaszczyzny



Generatory liniowe (multiple recursive generators - MRG)

Generatory liniowe typu

$$x_i = (a_1x_{i-1} + a_2x_{i-2} + \dots + a_kx_{i-k}) \bmod m$$

mają maksymalny okres równy

$$T = m^k - 1$$

pod warunkiem, że m jest liczbą pierwszą

- przykłady generatorów o dobrych parametrach statystycznych

$$x_i = (1176x_{i-1} + 1476x_{i-2} + 1776x_{i-3}) \bmod (2^{32} - 5)$$

$$x_i = 2^{13}(x_{i-1} + x_{i-2} + x_{i-3}) \bmod (2^{32} - 5)$$

$$x_i = (1995x_{i-1} + 1998x_{i-2} + 2001x_{i-3}) \bmod (2^{35} - 849)$$

Ich okresy są maksymalne $T = m^k - 1 = m^3 - 1$

Generatory na rejestrach przesuwnych (Tausworthe)

- operację dzielenia modulo można wykorzystać do tworzenia ciągów bitów, działania wykonujemy w arytmetyce binarnej więc **m=2**

$$b_i = (a_1 b_{i-1} + a_2 b_{i-2} + \dots + a_p b_{i-p}) \mod 2$$

$$a_1, a_2, \dots, a_p, b_i \in \{0, 1\}$$

$$b_1, b_2, \dots, b_k \in \{0, 1\} \quad - \text{potrzebny jest także ciąg inicjujący}$$

- ponieważ **m=2** jest liczbą pierwszą, więc odpowiednio dobrany parametr **p** pozwala uzyskać maksymalny okres

$$T_{\max} = 2^p - 1$$

- z reprezentacji usuwamy zero
(bo z niego nic nie wygenerujemy)

- ze względów prakycznych (wydajność) do generowania nowego bitu wykorzystuje się tylko 2 bity z ciągu

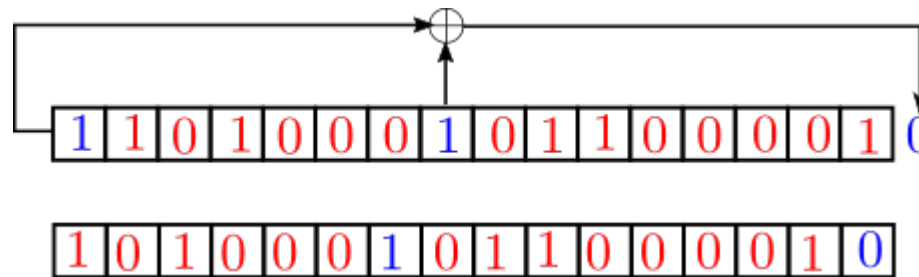
$$b_i = (b_{i-p} + b_{i-p+q}) \mod 2$$

$$a_{i-p}, a_{i-p+q} = 1, \quad (\text{pozostałe } a_i = 0)$$

powyższe wyrażenie odpowiada operacji XOR

$$b_i = b_{i-p} \oplus b_{i-p+q}$$

- nowy bit można generować wykorzystując rejestr, który sam jest reprezentacją liczby zmiennopozycyjnej (wygoda)



Bity są przesuwane w lewo, a na skrają prawą pozycję umieszczany bit z operacji **XOR**

- z ciągu L bitów możemy utworzyć liczbę zmiennopozycyjną o rozkładzie jednorodnym $U(0,1)$

$$x_i = \sum_{j=1}^L 2^{-j} b_{is+j}, \quad i = 1, 2, 3, \dots$$

$s < L$ – liczby U_i oraz U_{i+1} mają wspólne bity
 $s = L$ – liczby U_i oraz U_{i+1} tworzone są z rozłącznych fragmentów ciągu bitów

Dla liczb x_i zapisanych w reprezentacji binarnej nadal obowiązuje relacja

$$x_i = x_{i-p} \oplus x_{i-p+q}$$

- w generatorze Fushimiego (**biblioteka IMSL – International Statistical and Mathematical Library**)

$$p = 521, \quad a_{521} = a_{32} = 1 \quad T = 2^{521} - 1$$

nie spełnia on jednak niektórych testów - lepsze własności statystyczne uzyskuje się dla 3 operacji XOR

$$b_i = b_{i-a} \oplus b_{i-b} \oplus b_{i-c} \oplus b_{i-s}$$

- poprawienie własności statystycznych generatora bitowego można uzyskać stosując obrót (**twisting bits**) - metoda Matsumoto-Kurita

Definiuje się macierz A wypełnioną 0 i 1 – odpowiada ona za zamianę bitów miejscami (mieszanie)

$$x_i = x_{i-p} \oplus Ax_{i-p+q}$$

Na podstawie tej modyfikacji stworzyli generator nazwany **Marsenne Twister (MT19937)**

okres generatora jest bardzo duży

$$T = 2^{19937} - 1$$

Generator ten ma dość skomplikowany proces inicjalizacji, ale generuje rozkład

$$U^{623}(0, 1)$$

Generatory o dowolnym rozkładzie prawdopodobieństwa

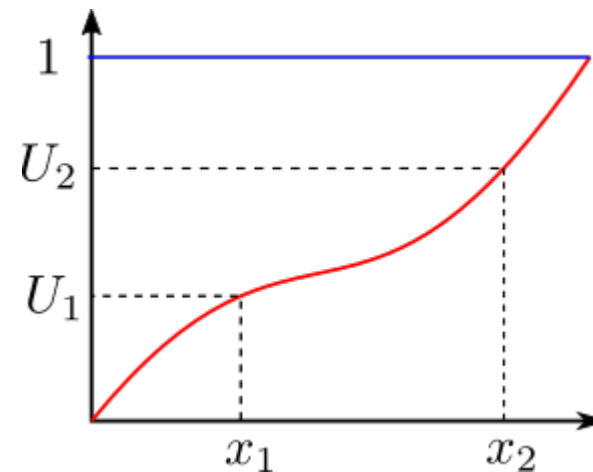
- generatory o rozkładzie $U(0,1)$ są podstawą konstrukcji generatorów o dowolnym rozkładzie
- nie ma jednej dobrej (wydajnej) metody, która pozwalałaby generować dowolny rozkład, najczęściej stosujemy jedną z poniższych metod:
 - **metoda odwracania dystrybucyj**
 - w zasadzie najlepsza pod względem statystycznym i wydajności, o ile potrafimy znaleźć funkcję odwrotną funkcji gęstości prawdopodobieństwa
 - **metodę eliminacji von Neumanna**
 - ma dobre własności statystyczne, ale jest mało wydajna
 - **metodę błędzenia przypadkowego (algorytm Metropolis)**
 - jest wydajna, ale w generowanym ciągu istnieje silna korelacja jego elementów
- w metodzie Monte Carlo wybór generatora musi być poprzedzony analizą jakościową, aspekty które należy wziąć pod uwagę to:
 - wydajność
 - własności statystyczne bo może to wpłynąć na wynik
 - ewentualny nakład pracy związany z jego obsługą (ręczny podział na podobszary w metodzie eliminacji, wielowątkowość, etc.)

Metoda odwracania dystrybuanty - transformacja zmiennych losowych

- wykorzystujemy własności dystrybuanty
 - nieujemna
 - niemalejąca
 - ograniczona $F(x) \in [0, 1]$
- losując zmienną U_i z rozkładu $U(0,1)$ określamy losową wartość dystrybuanty i jej argumentu x_i

$$U_i = F(x_i) \in (0, 1)$$

$$x_i = F^{-1}(U_i), \quad x \in (-\infty, \infty)$$



Metoda w zasadzie sprowadza się do znalezienia funkcji odwrotnej:

- jeśli ją znamy to problem jest częściowo rozwiązany (przykład: rozkład normalny $\rightarrow$ funkcja błędu)
- funkcję odwrotną można przybliżać/aproksymować – wówczas rozwiązanie przybliżone

Sposób generowania zmiennych losowych x z rozkładu o dystrybuancie F sprowadza się do wylosowania ciągu liczb

$$U_1, U_2, \dots, U_n \in (0, 1)$$

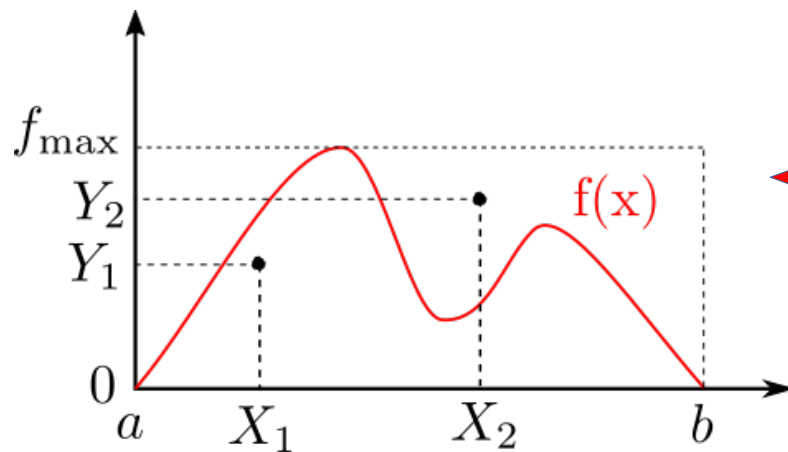
i ich transformacji w elementy ciągu o docelowym rozkładzie

$$x_1, x_2, \dots, x_n \in (-\infty, \infty)$$

Metodę odwracania dystrybuanty można stosować do rozkładów ciągłych i dyskretnych

Metoda eliminacji von Neumanna (acceptance-rejection method)

- jedna z pierwszych metod losowania zmiennych o zadanym rozkładzie
- mało wydajna, używana gdy liczba wymiarów jest niewielka ($d < 5$) i gdy nie można zastosować wydajniejszych metod (odwracania dystrybucyj) lub błędzenia przypadkowego (korelacja w ciągu)



Metoda podstawowa eliminacji

losujemy 2 liczby:

$$X \in U(a, b)$$

$$Y \in U(0, f_{\max})$$

$$\begin{cases} Y \leq f(X) & \text{akceptujemy punkt } x_i = X \\ Y > f(X) & \text{odrzucaamy punkt} \end{cases}$$

- algorytm stosujemy dopóki nie uzyskamy złożonej liczby zmiennych o zadanym rozkładzie, trudno przewidzieć ile należy wykonać losowań
- wadą podstawowej metody jest mała wydajność, część wylosowanych zmiennych odrzucamy, prawdopodobieństwo akceptacji zmiennej łatwo oszacować

$$p_{\text{accept}} = \frac{S_f}{S_{\max}} = \frac{\int_a^b f(x) dx}{(b-a)f_{\max}}$$

Wniosek: im większa różnica powierzchni obejmowanej przez funkcję $f(x)$ oraz obszar regularny, tym mniejsze prawdopodobieństwo akceptacji (większa strata w wydajności)

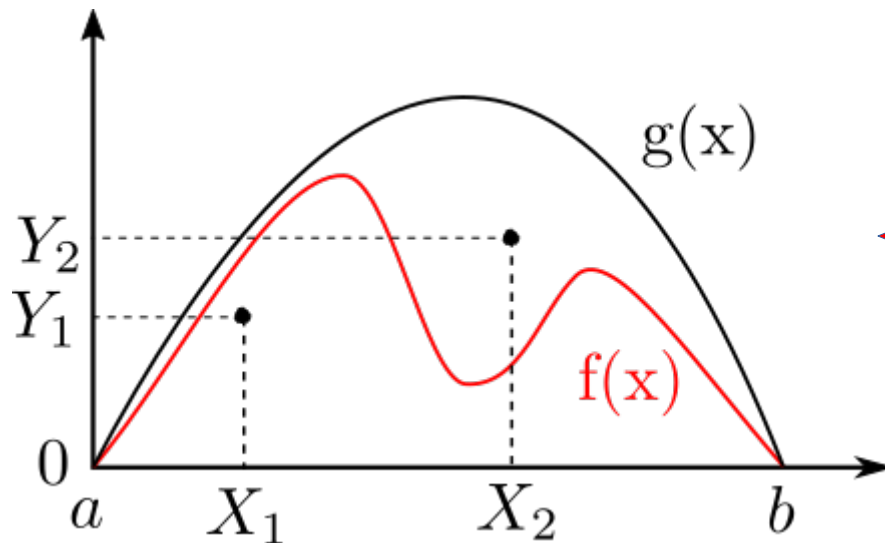
Modyfikacje metody – kombinacja metody eliminacji i odwracania dystrybuanty

Skoro duża różnica między powierzchniami jest przyczyną małej wydajności, to trzeba ją zmniejszyć. Zamiast ograniczać $f(x)$ od góry funkcją stałą, można poszukać funkcji $g(x)$ o zbliżonym kształcie. Szukamy, tzn. dobieramy, funkcję gęstości prawdopodobieństwa $g(x)$ o własnościach

$$g(x) : \quad g(x) \geq f(x), \quad |g(x) - f(x)| = \min, \quad x \in [a, b]$$

dla funkcji $g(x)$ powinniśmy dysponować dobrym (wydajnym) generatorem liczb pseudolosowych (np. wykorzystującego metodę odwracania dystrybuanty)

$G(a,b)$ – generator liczb pseudolosowych o rozkładzie $g(x)$



losujemy 1 liczbę z rozkładu $g(x)$:

$$Y \sim G(a, b)$$

$$\begin{cases} Y \leq f(X) & \text{akceptujemy punkt } x_i = X \\ Y > f(X) & \text{odrzucaamy punkt} \end{cases}$$

- problem polega na znalezieniu „dobrej” funkcji $g(x)$ o zbliżonym kształcie do $f(x)$ w całym zakresie.
- kolejna modyfikacja polega na przybliżaniu $f(x)$ w podprzedziałach, można to zrobić przy nawet użyciu funkcji stałej, konieczne jest jednak ustalenie liczby zmiennych, które należy wylosować w każdym przedziale – określamy ją na podstawie stosunku powierzchni $f(x)$ w danym podprzedziale do całkowitej objętości

$$n_i = N \frac{\int_{a+\Delta(i-1)}^{a+\Delta i} f(x) dx}{\int_a^b f(x) dx}$$

Kolejny problem: czy $f(x)$ jest analitycznie całkowalna? jeśli nie to można próbować numerycznie

Losowanie wartości $\sin(\theta)$ i $\cos(\theta)$ z wykorzystaniem metody eliminacji

(problem istotny w zagadnieniach transportu / rozpraszania cząstek)

- najprostszy sposób losowania to wylosować kąt (etap szybki)

$$\theta = 2\pi \cdot U_1, \quad U_1 \sim U(0, 1)$$

i obliczyć wartość funkcji (etap wolny – sumujemy elementy szeregu)

- możemy też określić losowy kąt w pierwszej ćwiartce układu, metodą eliminacji (rozkład kąta losowego w ćwiartce jest jednorodny)

$$f(\theta)d\theta = \frac{2}{\pi}d\theta, \quad \theta \in (0, \pi/2)$$

1) losujemy dwie liczby w obszarze $[0,1] \times [0,1]$

$$X = U_1, \quad Y = U_2, \quad U_i \sim U(0, 1)$$

2) sprawdzamy czy punkt leży w kole

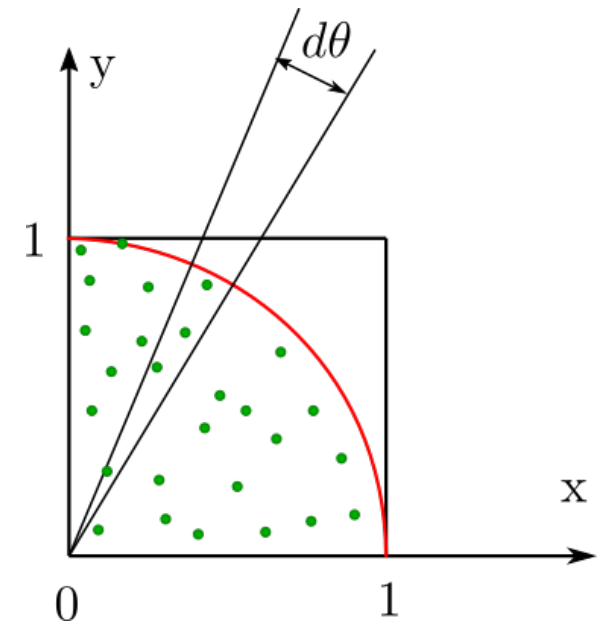
$$r^2 = X^2 + Y^2 = \begin{cases} r \leq 1 & \text{X,Y akceptujemy} \\ r > 1 & \text{X,Y odrzucamy i wracmy do pkt(1)} \end{cases}$$

3) wykorzystujemy relację

$$\theta = \arctan\left(\frac{Y}{X}\right), \quad \theta \in \left(0, \frac{\pi}{2}\right)$$

$$\cos(\theta) = \frac{X}{(X^2 + Y^2)^{1/2}} \quad \sin(\theta) = \frac{Y}{(X^2 + Y^2)^{1/2}}$$

- sumowanie elementów szeregu (sin/cos) zastąpiliśmy alternatywnymi formułami, ale część wyników losowań odrzucamy



- możemy rozszerzyć interesujący nas zakres

$$\theta \in \left(0, \frac{\pi}{2}\right) \rightarrow 2\theta \in (0, \pi)$$

$$\cos 2\theta = \cos^2 \theta - \sin^2 \theta = \frac{X^2 - Y^2}{X^2 + Y^2}, \quad \cos(2\theta) \in (-1, 1)$$

$$\sin 2\theta = 2 \sin \theta \cos \theta = \frac{2XY}{X^2 + Y^2}, \quad \sin(2\theta) \in (0, 1)$$

- oraz dla pełnego zakresu

$$\sin 2\theta = \pm \frac{2XY}{X^2 + Y^2}, \quad \theta \in (0, 2\pi), \quad \sin(2\theta) \in (-1, 1)$$

ale musimy określić znak funkcji **sin()** w sposób losowy (**cos()** jest parzysta)

$$U_3 \sim U(0, 1) \rightarrow \begin{cases} U_3 \leq \frac{1}{2} & \text{sign} = "+" \\ U_3 > \frac{1}{2} & \text{sign} = "-" \end{cases}$$

Metoda odwracania dystrybuanty dla rozkładu dyskretnego

- w rozkładzie dyskretnym definiujemy prawdopodobieństwa wylosowania danej liczby/własności/cechy etc., której nadajemy indeks k , przy czym zakładamy że rozkład może być dowolny (nierównomierny)

$$P\{X_k\} = p_k, \quad k = 0, 1, 2, 3, \dots, K$$

i narzucamy warunek normalizacji

$$\sum_{k=0}^K p_k = 1$$

- algorytm generowania pojedynczej liczby o zadanym rozkładzie dyskretnym

1) ustalamy parametry startowe: $k = 0, S = p_0$

2) losujemy zmienną: $U_i \in (0, 1)$

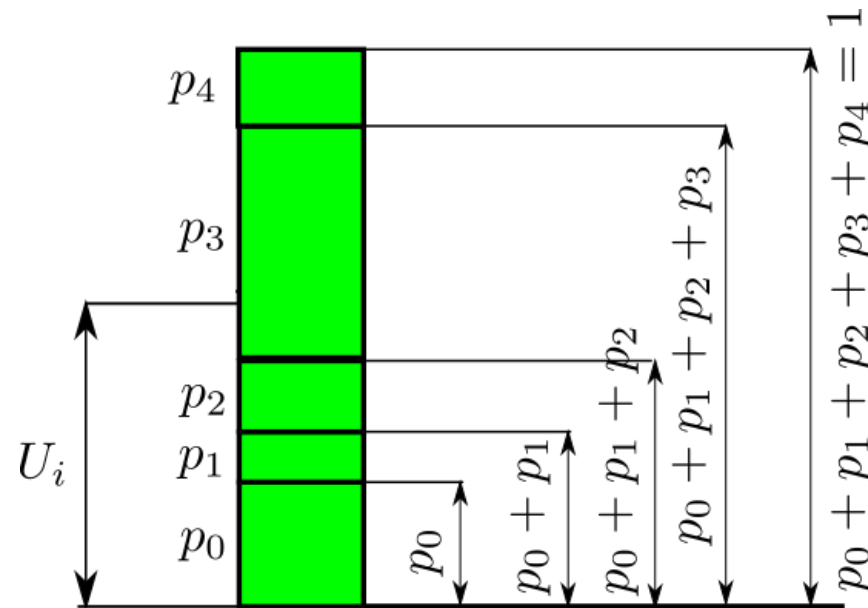
3) sprawdzamy warunek: $U > S?$

TAK – zmieniamy parametry k i S : $k = k + 1, S = S + p_k$

NIE – akceptujemy aktualną wartość $k \rightarrow X_k$

kroki 1-3 wykonujemy aż do uzyskania wymaganej ilości liczb $\{X_1, X_2, X_3, \dots, X_n\}$

- generowanie liczb o rozkładzie dyskretnym – algorytm podstawowy



START $X = 0, \quad S = p_0, \quad U > S(?) \quad \text{TAK}$

$k = 0 + 1 = 1, \quad S = p_0 + p_1, \quad U > S(?) \quad \text{TAK}$

$k = 1 + 1 = 2, \quad S = p_0 + p_1 + p_2, \quad U > S(?) \quad \text{TAK}$

$k = 2 + 1 = 3, \quad S = p_0 + p_1 + p_2 + p_3, \quad U > S(?) \quad \text{NIE} \rightarrow \text{STOP}$

nowa zmienna: $k = 3 \rightarrow X_3$

- algorytm jest prosty i skuteczny, ale jego konstrukcja wymaga sprawdzania warunku dla kolejnych podprzedziałów – idąc z dołu do góry wydajność algorytmu maleje w przypadku, gdy ciąg dyskretnych liczb jest duży np. kilkadziesiąt (losowanie kart w talii) !!!!!
- potrzebna jest modyfikacja algorytmu

Rozkład dyskretny – metoda równomiernego podziału

- zakładamy jak poprzednio $P\{X_k\} = p_k, \quad k = 0, 1, 2, 3, \dots, K$ $\sum_{k=0}^K p_k = 1$
- przedział $[0,1]$ dzielimy na podprzedziały o identycznej długości

$$\Delta = \frac{1}{K+1} \quad [(i-1) \cdot \Delta, i \cdot \Delta], \quad i = 1, 2, \dots, K+1$$

- zmienna wygenerowana z rozkładu jednorodnego

$$U \in U(0, 1)$$

wpada do podprzedziału o indeksie (transformacja liniowa + obcięcie części ułamkowej)

$$Y = \lfloor (K+1)U + 1 \rfloor \quad - \text{lokalizujemy się gdzieś w środku rozkładu}$$

- konstruujemy dwa ciągi liczb, których używamy do generowania wszystkich liczb o zadanym rozkładzie

$$\{q_i\} = \{q_0, q_1, \dots, q_K\}$$

$$q_i = \sum_{j=0}^i p_j, \quad i = 0, 1, 2, \dots, K$$

$$\{g_i\} = \{g_1, g_2, g_3, \dots, g_{K+1}\}$$

$$g_i = \max \{j : q_j \leq i \cdot \Delta\}, \quad i = 1, 2, \dots, K+1$$

- dysponując Y , $\{q\}$ oraz $\{g\}$ wyznaczamy zmienną losową X z rozkładu dyskretnego wykonując „tylko” 2 sprawdzenia - ale jeśli p_k znacznie się różnią to więcej

Algorytm metody równomiernego podziału

1) losujemy zmienną z rozkładu jednorodnego

$$U \sim U(0, 1)$$

2) określamy indeks wylosowanego przedziału Y i przybliżoną wartość X

$$Y = \lfloor (K + 1)U + 1 \rfloor$$

$$k = g_Y + 1$$

3) iteracyjnie sprawdzamy warunek

$$q_{k-1} < U(?)$$

TAK: akceptujemy wynik (STOP)

$$X = X_k$$

NIE: zmieniamy k i sprawdzamy nowy wynik (jeśli warunek spełniony: STOP)

$$k \leftarrow k - 1, \quad X = X_{k-1}$$

W ten sposób dokonujemy sprawdzenia warunku conajwyżej kilka (2-3) razy zamiast 10, 50 czy 100.

Przykład działania metody równomiernego podziału

$$X = 0, 1, 2, 3 \quad p_k = 0.1, 0.2, 0.3, 0.4 \quad K = 3 \quad \Delta = \frac{1}{4} = 0.25$$

Wyznaczamy elementy obu ciągów

$$\{q_i\} = \{q_0, q_1, q_2, q_3\} = \{0.1, 0.3, 0.6, 1\}$$

$$q_i = \sum_{j=0}^i p_j, \quad i = 0, 1, 2, \dots, K$$

$$\{g_i\} = \{g_1, g_2, g_3, g_4\} = \left\{ \underbrace{0}_{q_0 < 1/4}, \underbrace{1}_{q_1 < 2/4}, \underbrace{2}_{q_2 < 3/4}, \underbrace{3}_{q_3 < 4/4} \right\}$$

$$g_i = \max \{j : q_j \leq i \cdot \Delta\} \\ i = 1, 2, \dots, K + 1$$

$$U = 0.05 \rightarrow Y = \lfloor (K + 1)U + 1 \rfloor = \lfloor 4 \cdot 0.05 + 1 \rfloor = 1 \rightarrow k = g_Y + 1 = 0 + 1 = 1$$

sprawdzamy: $q_{k-1} = q_{1-1} = q_0 = 0.1 \stackrel{?}{<} U = 0.05$ **NIE**

redukujemy k: $k \leftarrow k - 1 = 1 - 1 = 0$ **$X = x_{k-1} = 0$**

drugi raz już nie sprawdzamy bo nie ma niższego indeksu k

$$U = 0.11 \rightarrow Y = \lfloor (K + 1)U + 1 \rfloor = \lfloor 4 \cdot 0.11 + 1 \rfloor = 1 \rightarrow k = g_Y + 1 = 0 + 1 = 1$$

sprawdzamy: $q_{k-1} = q_{1-1} = q_0 = 0.1 \stackrel{?}{<} U = 0.11$ **TAK** **$X = x_k = 1$**

$$U = 0.59 \rightarrow Y = \lfloor (K + 1)U + 1 \rfloor = \lfloor 4 \cdot 0.59 + 1 \rfloor = 3 \rightarrow k = g_Y + 1 = 2 + 1 = 3$$

sprawdzamy: $q_{k-1} = q_{3-1} = q_2 = 0.6 \stackrel{?}{<} U = 0.59$ **NIE**

redukujemy k: $k \leftarrow k - 1 = 3 - 1 = 2$

sprawdzamy: $q_{k-1} = q_{2-1} = q_1 = 0.3 \stackrel{?}{<} U = 0.59$ **TAK** **$X = x_k = 2$**

dwukrotnie sprawdziliśmy warunek

Algorytm Metropolisa-Hastingsa – metoda błędzenia przypadkowego

Generując liczby pseudolosowe wykonujemy pewien proces stochastyczny w którym elementy generowanego ciągu określają aktualny stan układu czyli generatora.

Ponieważ proces jest losowy i odbywa się według zadanego algorytmu możemy założyć, że prawdopodobieństwo pojawienia się w ciągu sekwencji dwóch stanów x i y

$\dots, x, y, \dots$

$\dots, y, x, \dots$

powinno być identyczne

W każdym przypadku, nowy stan (liczba) musi zależeć od stanu go poprzedzającego. Możemy określić **prawdopodobieństwo przejścia** pomiędzy tymi stanami.

- przejście $x \rightarrow y$

$$p(y|x) = \alpha(y|x)q(y|x)f(x)$$

- analogicznie możemy dla przejścia $y \rightarrow x$

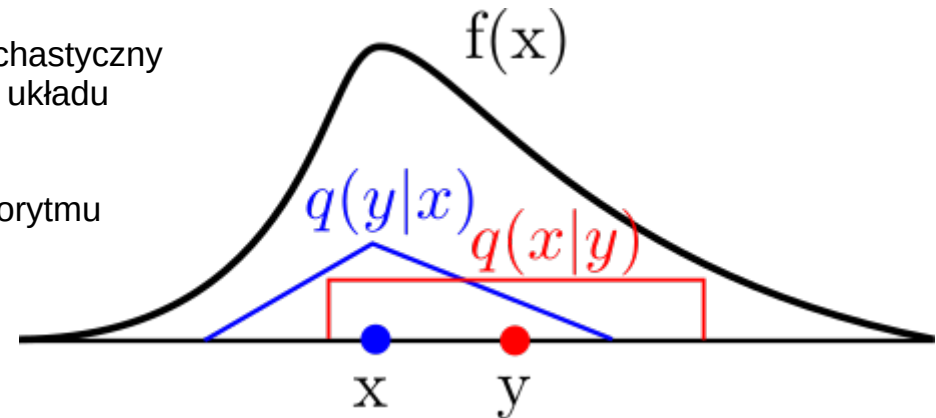
$$p(x|y) = \alpha(x|y)q(x|y)f(y)$$

- zgodnie z postawionym warunkiem wymagamy

$$p(y|x) = p(x|y)$$

- tzw. warunek „detailed balance”

$$\alpha(y|x)q(y|x)f(x) = \alpha(x|y)q(x|y)f(y)$$



- $f(x)$ – określa prawdopodobieństwo znalezienia się w danym punkcie (proces jest stacjonarny - nie zmienia fgp)
- $q(y|x)$ – prawdopodobieństwo warunkowe wylosowania zmiennej y dla ustalonej zmiennej x (**to my je definiujemy**)
- $\alpha(x|y)$ – określa prawdopodobieństwo akceptacji przejścia

Przekształcając relację możemy określić prawdopodobieństwo akceptacji

$$\alpha(y|x) = \alpha(x|y) \frac{q(x|y)f(y)}{q(y|x)f(x)}$$

← $\alpha(x|y)$ nieznane, musimy coś założyć

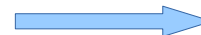
$$\frac{q(x|y)f(y)}{q(y|x)f(x)} = \begin{cases} > 1 & \text{przejście niemal pewne} (\alpha(y|x) \approx 1) \\ << 1 & \text{przejście mało prawdopodobne} (\alpha(y|x) << 1) \end{cases}$$

Metropolis zaproponował, aby przyjąć następujący sposób wyznaczania α

$$\alpha(y|x) = \min \left\{ \frac{q(x|y)f(y)}{q(y|x)f(x)}, 1 \right\}$$

Jeśli przyjmiemy, że prawdopodobieństwo warunkowe jest symetryczne

$$q(x|y) = q(y|x)$$



zazwyczaj przyjmuje się, że $q(x|y)$ ma rozkład jednorodny w lokalnym otoczeniu x/y

$$U_1 \sim U(-\delta, \delta)$$

$$y = x + U_1$$

to dostaniemy uproszczoną postać dla akceptacji nowego stanu, które bazuje wyłącznie na funkcji gęstości prawdopodobieństwa

$$\alpha(y|x) = \min \left\{ \frac{f(y)}{f(x)}, 1 \right\}$$

funkcja $f(x)$ nie musi być nawet unormowana (w pewnych przypadkach możemy nie znać stałej normalizacyjnej)

Algorytm Metropolisa-Hastingsa dla dowolnego rozkładu prawdopodobieństwa

- wybieramy stan początkowy (x_0), dla $i=1,2,3,\dots$ generujemy liczby:

- dla aktualnego stanu x_i generujemy nowy **próbny** stan (liczbę)

$$Y \in q(y|x_i)$$

- sprawdzamy akceptację nowego stanu

$$U_i \sim U(0, 1)$$

$$x_{i+1} = \begin{cases} Y, & \text{jeśli } U \leq \alpha(Y|x_i) \\ x_i, & \text{jeśli } U > \alpha(Y|x_i) \end{cases}$$

W ten sposób generujemy ciąg zwany **łańcuchem Markowa**

$$\{x_0, x_1, x_2, \dots, x_n\}$$

Uwagi:

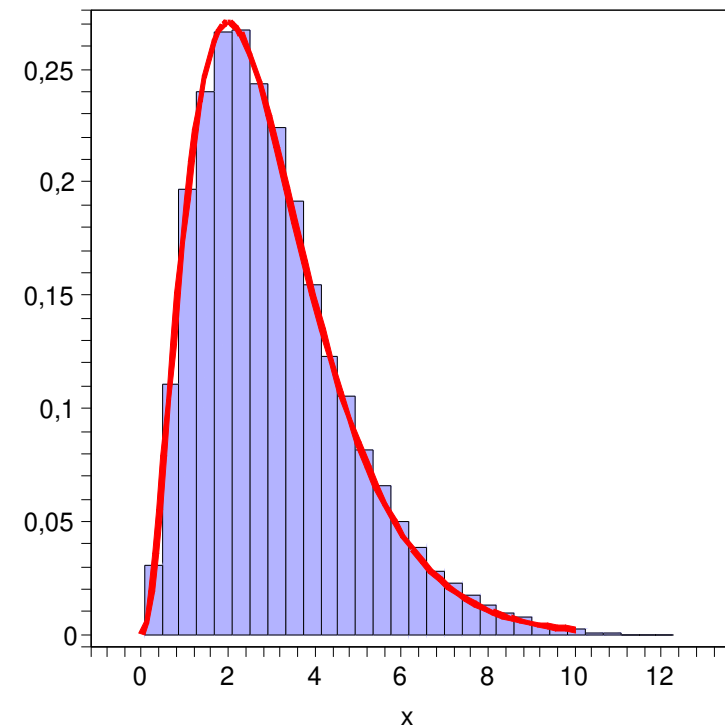
- jeśli warunek akceptacji nie jest spełniony to kolejnym elementem ciągu jest aktualny stan (liczba)
- to powoduje występowanie silnej korelacji elementów ciągu
- korelacja rośnie również dlatego, że kolejny stan losowany jest w lokalnym otoczeniu stanu aktualnego.

Przykład

$$f(x) = \frac{1}{2}x^2e^{-x}, \quad x \in [0, \infty)$$

$$q(y|x) = q(x|y) = U(-1, 1)$$

Porównanie histogramu dla $N=50000$ punktów z rozkładem $f(x)$.



Uzyskany rozkład wygląda poprawnie
- więc nasz generator działa

Rozkład jednomianowy

Nieunormowana funkcja gęstości prawdopodobieństwa

$$f(x) = Cx^n, \quad x \in (0, 1), \quad n = 1, 2, 3, \dots$$

znajdujemy stałą normalizacji

$$C \int_0^1 dx x^n = C \frac{x^{n+1}}{n+1} \Big|_0^1 = \frac{C}{n+1} = 1$$

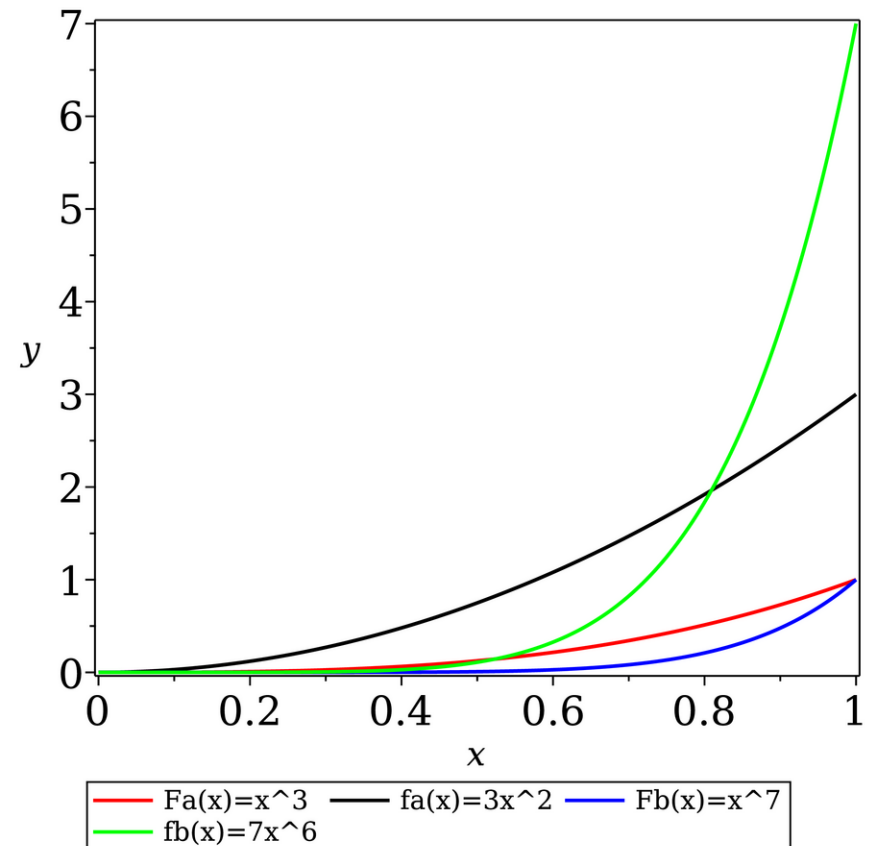
$$C = n + 1$$

unormowana fgp rozkładu jednomianowego

$$f(x) = (n + 1) x^n$$

Dystrybuanta rozkładu

$$F(x) = (n + 1) \int_0^x dx' (x')^n = x^{n+1} = P(x' \leq x)$$



Uwaga: im większy wykładnik „n” tym rozkład jest bardziej zakumulowany w pobliżu prawej granicy (możemy zatem manipulować wykładnikiem w zależności od tego jak bardzo chcemy zagęścić liczby pseudolosowe w pobliżu $x=1$)

Do generowania liczb pseudolosowych wykorzystamy **metodę odwracania dystrybuanty**

$$F(x) = y \rightarrow F^{-1}(y) = x$$

$$F(x) = (n+1) \int_0^x dx' (x')^n = x^{n+1} = P(x' \leq x) \in [0, 1]$$

Losujemy liczbę z rozkładu równomiernego $U(0,1)$

$$U_1 \in U(0, 1)$$

określa ona **losową** wartość dystrybuanty

$$F(x) = x^{n+1} = U_1$$

$$x = U_1^{\frac{1}{n+1}}, \quad x \in (0, 1)$$



znaleźliśmy funkcję odwrotną dystrybuanty

$$F^{-1}(y) = y^{\frac{1}{n+1}} = x$$

Zmienna losowa x ma rozkład jednomianowy

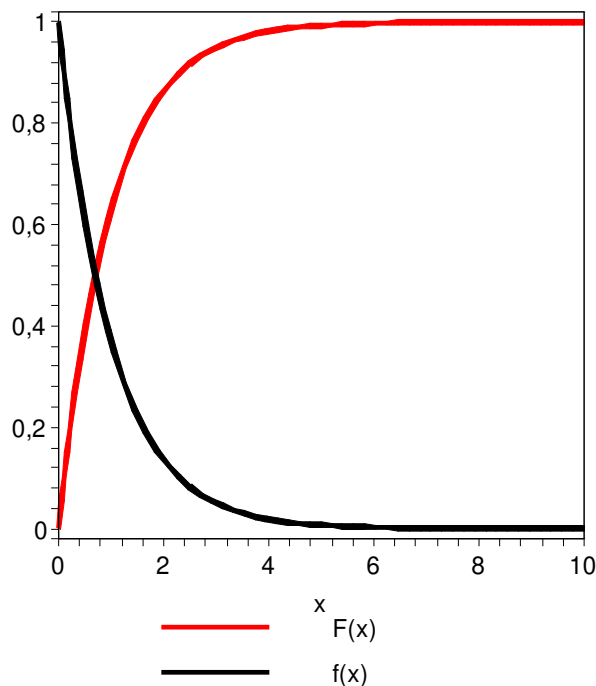
Rozkład eksponencjalny

- funkcja gęstości prawdopodobieństwa

$$f(x) = e^{-x}, \quad x \in [0, \infty)$$

- dystribuanta rozkładu

$$F(x) = \int_0^x dx' e^{-x'} = 1 - e^{-x}$$



Do generowania liczb pseudolosowych o rozkładzie eksponencjalnym wykorzystujemy metodę odwracania dystrybuanty

$$\int_0^x f(x') dx' = P(x' \leq x) = U_1, \quad U_1 \sim U(0, 1)$$

$$1 - e^{-x} = U_1$$

$$x = -\ln(1 - U_1), \quad x \in (0, \infty)$$

Rozkład używany np. do opisu rozpadu jądrowego

$$\frac{dN}{dt} = -\lambda N \quad N(t) = N_0 e^{-\lambda t} \quad T_{1/2} = \frac{\ln(2)}{\lambda}$$

λ – stała rozpadu, $T_{1/2}$ – czas półrozpadu

funkcja gęstości prawdopodobieństwa dla rozpadu

$$f(t, \lambda) = \lambda e^{-\lambda t}$$

Rozkład normalny (Gaussa) $N(0,1)$

Rozkład używany często do opisu procesów stochastycznych (proces Wienera), szumu itp.

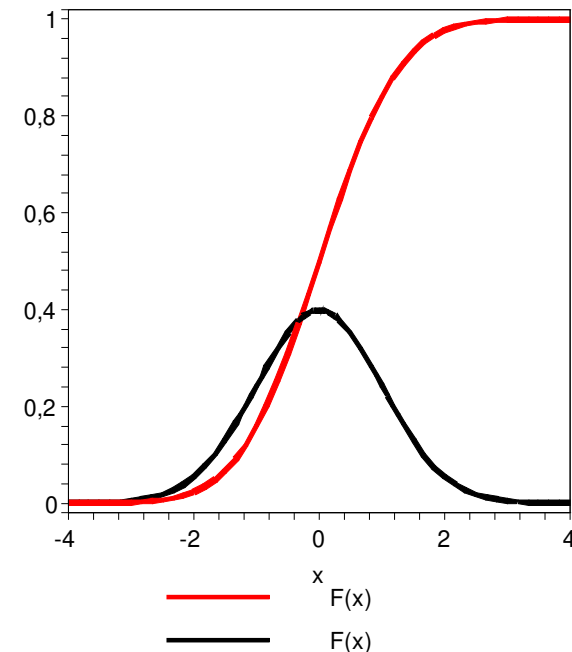
- funkcja gęstości prawdopodobieństwa

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

- dystribuanta

$$\begin{aligned} F(x) &= \int_{-\infty}^x f(x') dx' = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^x e^{-\frac{x'^2}{2}} dx' \\ &= \frac{1}{2} \left[1 + \operatorname{erf} \left(\frac{x}{\sqrt{2}} \right) \right] = P(x' \leq x) \end{aligned}$$

$\operatorname{erf}(x)$ – funkcja błędu



Metoda Box-Muller'a dla rozkładu normalnego

Zapiszmy fgp rozkładu Gaussa w 2D - to będzie iloczyn fgp dla zmiennych x i y

$$f(x, y) = f(x)f(y) = \frac{1}{2\pi} e^{-\frac{x^2+y^2}{2}} \quad \iint_{-\infty}^{\infty} dx dy f(x, y) = 1$$

Dokonujemy transformacji

$$(x, y) \rightarrow w(r, \phi) \quad x = r \cos(\phi) \quad y = r \sin(\phi)$$

co prowadzi do zmiany sposobu liczenia prawdopodobieństwa

$$dP = f(x, y) \underbrace{dx dy}_{d\Omega} = f(r, \phi) \underbrace{r dr d\phi}_{d\Omega}$$

$$\int_0^{2\pi} d\phi \int_0^{\infty} dr r f(r, \phi) = \int_0^{2\pi} d\phi \underbrace{\frac{1}{2\pi}}_{f_\phi(\phi)} \cdot \int_0^{\infty} dr r \underbrace{e^{-r^2/2}}_{f_r(r)} = 1$$

zmienne: r i ϕ są niezależne, a ich fgp odczytujemy z warunku unormowania

drugą całkę możemy jeszcze przekształcić tj. uprościć

$$z = \frac{1}{2}r^2 \rightarrow dz = r dr \rightarrow \int_0^{\infty} dr r e^{-r^2/2} = \int_0^{\infty} dz e^{-z}$$

$$f_z(z) = e^{-z}$$

- to fgp rozkładu eksponencjalnego
(z tym sobie łatwo poradzimy)

Teraz wystarczy posłużyć się rozkładem eksponencjalnym

$$\int_0^Z e^{-z} dz = 1 - e^{-Z} = P(z \leq Z) = U_1 \sim U(0, 1)$$

losujemy liczbę z rozkładu równomiernego $U_1 \in U(0, 1)$

$$1 - e^{-Z} = U_1$$

$$Z = -\ln(1 - U_1), \quad Z \in (0, \infty)$$

i zamieniamy na zmienną radialną $Z = \frac{r^2}{2}$

$$r = \sqrt{-2\ln(1 - U_1)}, \quad r \in (0, \infty)$$

losujemy też drugą zmienną niezależną - kątową

$$U_2 \sim U(0, 1)$$

$$\phi = 2\pi \cdot U_2$$

I wracamy do współrzędnych kartezjańskich

$$x = r \cos(\phi) = \sqrt{-2\ln(1 - U_1)} \cos(2\pi U_2)$$

$$y = r \sin(\phi) = \sqrt{-2\ln(1 - U_1)} \sin(2\pi U_2)$$

x,y – to para zmiennych **niezależnych**
o rozkładzie N(0,1)

Metoda **Boxa-Mullera** generuje liczby z rozkładu $N(0,1)$, aby uzyskać rozkład $N(\mu, \sigma)$ dokonujemy prostej transformacji

$$y = \frac{x - \mu}{\sigma} \rightarrow f_y(y)dy = f_x(x)dx$$

$$f_x(x; \mu, \sigma) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{(x-\mu)^2}{2\sigma^2}}$$

$$Y = \mu + X\sigma, \quad X \in N(0, 1), \quad Y \in N(\mu, \sigma)$$

W ten sposób manipulujemy położeniem i szerokością rozkładu.

Rozkład Cauchy

Stosowany często do opisu kształtu pików rezonansowych, jego szerokość jest skorelowana z czasem życia stanu rezonansowego (wzór **Breita-Wignera** w fizyce jądrowej, w fizyce kwantowej **funkcja Lorentza** („lorencjan”) wykorzystywany m.in. w analizie spektralnej widm (identyfikacja położień pików rezonansowych i ich czasów życia)

ogólna postać funkcji gęstości prawdopodobieństwa

$$g(E, E_0, \gamma) = \frac{1}{\pi\gamma \left[1 + \left(\frac{E-E_0}{\gamma} \right)^2 \right]}$$

$$E \in (-\infty, \infty)$$

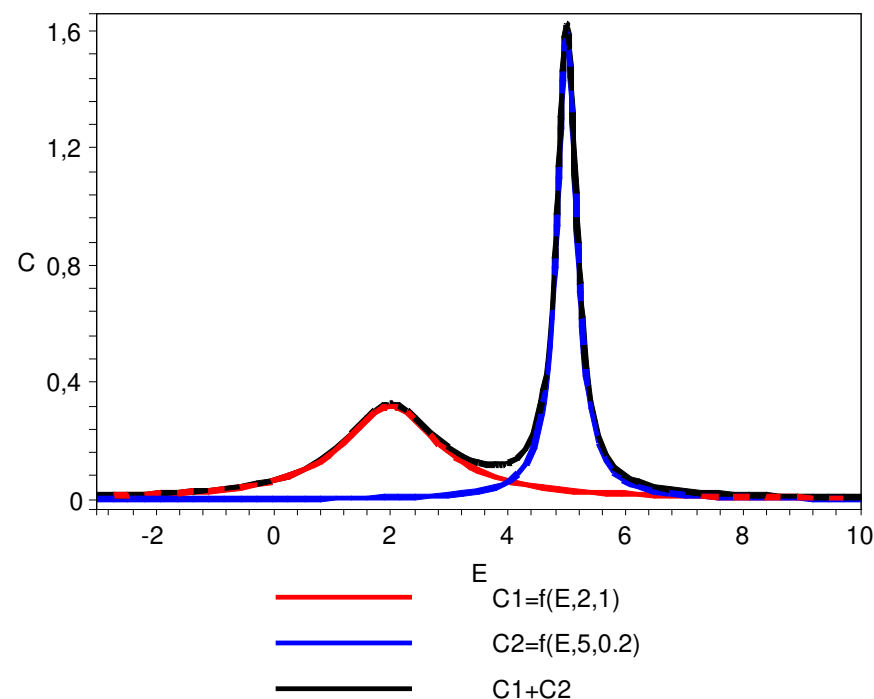
E_0 – położenie maksimum rozkładu,
 λ – szerokość połówkowa rozkładu
 (szerokość linii spektralnej)

Uwaga:

- w 1D rozkład Cauchy jest unormowany
- w 2D i 3D nie można go unormować (rozkład „nienormowalny”)
- dla

$$\gamma \rightarrow 0$$

rozkład Cauchy dąży do delty Diraca



Na rysunku wioczne są dwa piki rezonansowe:
 1-pik, szeroki (długożyjący), niski (mała intensywność w widmie)
 2-pik, wąski (krótkożyjący), wysoki (duża intensywność w widmie)

Skonstruujmy generator dla **standardowego rozkładu Cauchyego** metodą odwracania dystrybuanty

Najpierw dokonujemy podstawienia i dokonujemy transformacji fgp

$$x = \frac{E - E_0}{\gamma}$$

$$g(E, E_0, \gamma) dE = f(x) dx$$

$$f(x) = g(E(x), E_0, \gamma) \frac{dE}{dx} = \frac{1}{\pi(1+x^2)}$$

liczymy dystrybuantę

$$F(x) = \frac{1}{\pi} \int_{-\infty}^x \frac{1}{1+x^2} dx = \frac{1}{\pi} \tan^{-1}(x) + \frac{1}{2} = y$$

$$\pi \left(y - \frac{1}{2} \right) = \tan^{-1}(x)$$

$$x = \tan \left[\pi \left(y - \frac{1}{2} \right) \right], \quad x \in (-\infty, \infty), \quad y \in (0, 1)$$

Generowanie liczb pseudolosowych dla **standardowego rozkładu Cauchyego**

$$U_1 \sim U(0, 1)$$

$$x = \tan \left[\pi \left(U_1 - \frac{1}{2} \right) \right]$$

i dla rozkładu pierwotnego/ogólnego

$$x = \frac{E - E_0}{\gamma} \rightarrow E = \gamma x + E_0$$

Jeszcze o splocie rozkładów i konsekwencjach eksperymentalnych.

W spektroskopii wyznacza się tzw. profile Voigta, są to linie spektralne które uległy poszerzeniu powodowanym np. efektem Dopplera.

Rozkład spektralny linii opisuje **trójparametrowy rozkład Voigta**.

Jest splotem **rozkładu Cauchyego** charakteryzującego stany rezonansowe oraz **rozkładu Gaussa** odpowiadającego za poszerzenie linii spektralnych (**efekt Dopplera**)

$$V_{\sigma,\gamma,E_0}(E) = \int_{-\infty}^{\infty} G_{\sigma}(E') L_{\gamma,E_0}(E - E') dE'$$

$$G_{\sigma}(E) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{E^2}{2\sigma^2}}$$

$$L_{\gamma,E_0}(E) = \frac{1}{\pi\gamma\left(1 + \left(\frac{E-E_0}{\gamma}\right)^2\right)}$$

Poszerzenie dopplerowskie jest wynikiem oddziaływania cząstek z otoczeniem (np. fononami) i jest funkcją temperatury

$$\sigma = \sqrt{\frac{kT}{m_{eff}c^2}} E_0$$

Jak wygenerować zmienną losową o takim rozkładzie?

Rozkłady złożone

Rozkład zmiennej może być złożeniem dwóch rozkładów

$$F(x) = \int_{\mathcal{T}} H_t(x) dG(t), \quad dG(t) = g(t)dt, \quad t \in \mathcal{T}$$

interpretacja wyrażenia jest następująca

- zmienna losowa $T=t$ ma rozkład o dystrybuancie $G(t)$ [fgp: $g(t)$]
- rozkład zmiennej X_t określa prawdopodobieństwo warunkowe dane dystrybuantą $H_t(x)$, (t stanowi parametr rozkładu)
- rezultat: $X=X_t$ to zmienna o dystrybuancie $F(x)$

Algorytm generowania ciągu $\{X_1, X_2, \dots\}$ w naturalny sposób wynika z powyższej interpretacji:

1) wylosuj zmienną

$$T_i \sim \mathcal{G}$$

2) dla danej zmiennej T_i (teraz jest parametrem) wylosuj X_i

$$X_i \sim \mathcal{H}_t$$

3) utwórz ciąg zmiennych o rozkładzie $F(x)$

$$\{X_1, X_2, X_3, \dots\}$$

Rozkład sumy dwóch i więcej zmiennych

- dysponujemy fgp **dwóch zmiennych niezależnych**

$$x_1 \rightarrow f_1(x_1)$$

$$x_2 \rightarrow f_2(x_2)$$

- wykorzystajmy je do stworzenia dwóch **nowych zmiennych**

$$u_1 = x_1 + x_2 \rightarrow f_{u_1}(u_1) = ?$$

$$u_2 = x_2 \rightarrow f_{u_2}(u_2) = f_2(u_2)$$

ze względu na relacje pomiędzy starymi a nowymi zmiennymi spełniony jest warunek

$$\iint_{\Omega_x} f_x(x_1, x_2) dx_1 dx_2 = \iint_{\Omega_u} f_u(u_1, u_2) du_1 du_2$$

$$f_x(x_1, x_2) dx_1 dx_2 = f_u(u_1, u_2) du_1 du_2$$

wykorzystujemy niezależność starych zmiennych

$$f_x(x_1, x_2) = f_1(x_1) f_2(x_2)$$

oraz relację dla transformacji odwrotnej

$$x_1 = u_1 - u_2 \rightarrow dx_1 = d(u_1 - u_2)$$

$$x_2 = u_2 \rightarrow dx_2 = du_2$$

$$f_x(x_1, x_2) dx_1 dx_2 = f_1(u_1 - u_2) f_2(u_2) \underbrace{d(u_1 - u_2) du_2}_{=?}$$

Całkując prawą stronę po zmiennej u_1 zauważamy, że **$u_2 = \text{const}$**

$$\int du_2 f_2(u_2) \int f_1(u_1 - u_2) d(\underbrace{u_1}_{\text{const}}) = \int du_2 f_2(u_2) \int f_1(u_1 - u_2) d(u_1)$$

Dostaliśmy więc relację

$$f_x(x_1, x_2)dx_1dx_2 = f_1(u_1 - u_2)f_2(u_2)du_1du_2 = f_u(u_1, u_2)du_1du_2$$

$$f_u(u_1, u_2) = f_1(u_1 - u_2)f_2(u_2)$$

Wracamy do postawionego problemu, szukamy $f_{u_1}(u_1) = ?$

więc musimy dokonać redukcji (**marginalizacji**) zmiennej u_2 - całkujemy

$$f_{u_1}(u_1) = \int du_2 f_1(u_1 - u_2)f_2(u_2) = f_1 * f_2$$

Funkcja gęstości prawdopodobieństwa sumy dwóch zmiennych niezależnych jest splotem fgp obu zmiennych.

Uogólniając wynik dla większej liczby zmiennych dostajemy **unormowaną fgp**

$$f_{u_n}(u_n) = (((f_1 * f_2) * f_3) * \dots * f_n)$$

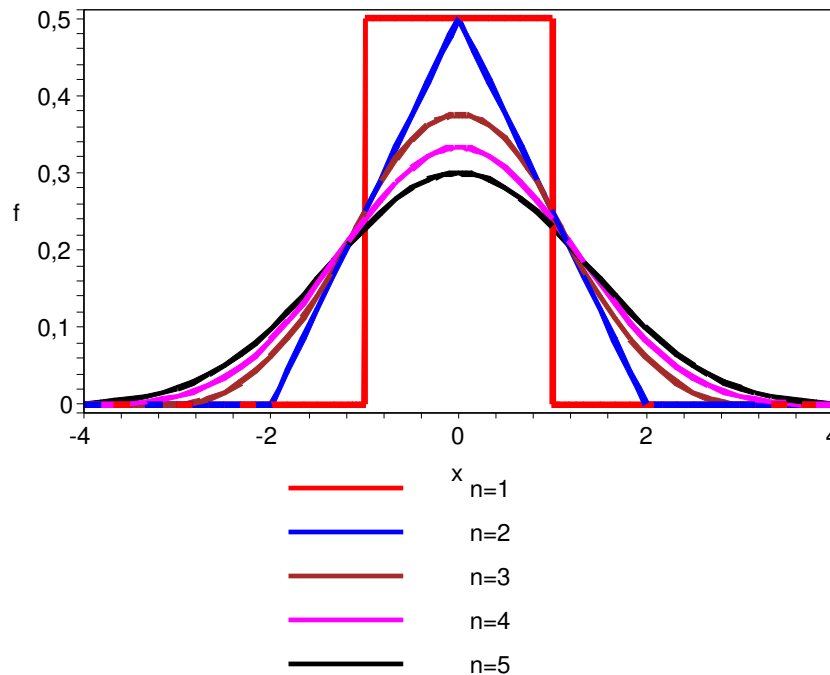
Zobaczmy jakie rozkłady uzyskamy dla sumy zmiennych o rozkładzie jednorodnym $U(-1,1)$

$$x_i \in U(-1, 1), \quad i = 1, 2, 3, \dots$$

fgp sumy kilku zmiennych ($n=1,2,3,4,5$),
każda zmienna ma rozkład jednorodny $U(-1,1)$

$$f_1(x) = \frac{1}{2} [H(x+1) - H(x-1)]$$

$H(x)$ -funkcja Heaviside'a



- $n=1$ oznacza tylko jedną zmienną o rozkładzie $U(-1,1)$
- $n=2$ – suma dwóch zmiennych daje rozkład trójkątny
- dla większych n rozkład staje się gładki, jego zasięg powiększa się, a maksimum obniża
- dla dużych n rozkład sumy zmiennych dąży do **rozkładu normalnego** zgodnie z centralnym twierdzeniem granicznym

Uwaga:

Dla dużego $n \sim 10-20$, fgp sumy tylu zmiennych tylko przybliża rozkład normalny. Dlaczego?
Zakres zmienności nowej zmiennej rozciąga się pomiędzy $[-n, n]$,
nie obejmuje więc tego samego zakresu co rozkład normalny (przedział obustronnie otwarty).
W pewnych szczególnych zastosowaniach, taki generator może mieć zastosowanie
(głównie ze względu na swoją szybkość – nie liczymy funkcji $\sin/\cos/\log$ co wymaga czasu)

Przykład: rozważamy rozkład złożony - dyskretny + ciągły

$$f(x) = \sum_{i=1}^n g_i h_i(x)$$

z warunkiem normalizacji

$$\sum_{i=1}^n g_i = 1, \quad g_i \in (0, 1) \quad (\text{dyskretny})$$

$$\int_{-\infty}^{\infty} h_i(x) dx = 1, \quad i = 1, 2, \dots, n \quad (\text{ciągły})$$

weźmy prostą kombinację rozkładów

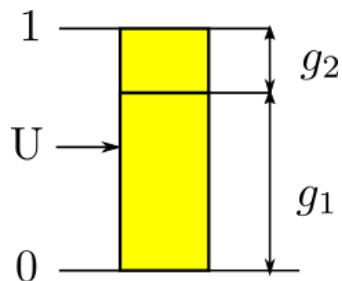
$$f(x) = g_1 h(x; \mu_1, \sigma_1) + g_2 h(x; \mu_2, \sigma_2)$$

$$h(x; \mu, \sigma) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{(x-\mu)^2}{2\sigma^2}} \quad - \text{rozkład normalny } N(\mu, \sigma)$$

sprawdzamy warunek unormowania

$$\int_{-\infty}^{\infty} f(x) dx = g_1 \int_{-\infty}^{\infty} h(x; \mu_1, \sigma_1) dx + g_2 \int_{-\infty}^{\infty} h(x; \mu_2, \sigma_2) dx = g_1 + g_2 = 1$$

g_1 i g_2 definiują dystrybuantę rozkładu dyskretnego



parametry rozkładu:

$$x_1 = 0, \quad \sigma_1 = 1$$

$$x_2 = 5, \quad \sigma_1 = 2$$

$$g_1 = \frac{3}{4}, \quad g_2 = \frac{1}{4}$$

użyty algorytm:

$$U_1 \sim U(0, 1)$$

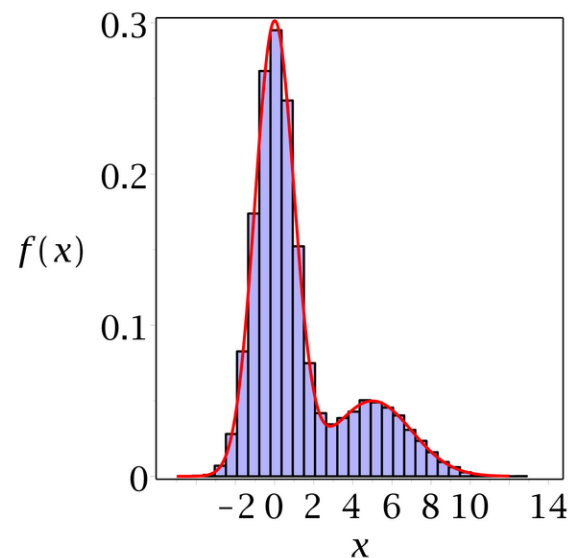
$$U_2 \sim N(0, 1)$$

$$\begin{cases} U_1 < g_1, & X_i = U_2\sigma_1 + \mu_1 \\ U_1 \geq g_1, & X_i = U_2\sigma_2 + \mu_2 \end{cases}$$

$$\{X\} = \{X_1, X_2, X_3, \dots\}$$

czerwony – rozkład zadany

histogram – rozkład uzyskany z generatora



$$f(x) = \sum_{i=1}^n g_i h_i(x)$$

Rozkłady wielowymiarowe

Efektywność metody MC względem innych (deterministycznych) metod rośnie wraz ze wzrostem wymiarowości rozwiązywanego problemu. Konieczne zatem staje się zastąpienie pojedynczej zmiennej losowej n-wymiarowym wektorem losowym, którego elementy stanowią zmienne losowe o zadanym rozkładzie i często są to zmienne zależne (skorelowane).

Rozkład jednorodny w n-wymiarowej kuli $K^n(0,1)$ metodą eliminacji

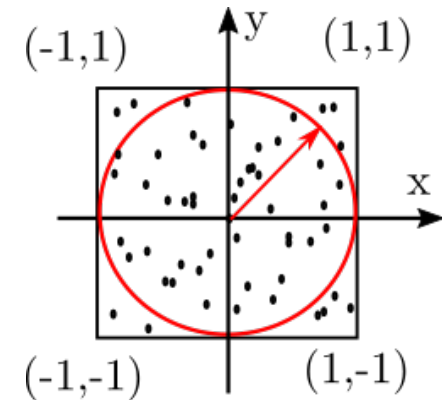
generujemy wektory w n-wymiarowej kostce

z warunkiem akceptacji

$$\vec{x} = [x_1, x_2, \dots, x_n], \quad x_i \in U(-1, 1)$$

$$\alpha = \begin{cases} 1, & \|x\|_2 \leq 1 \\ 0, & \|x\|_2 > 1 \end{cases}$$

Częstość („prawdopodobieństwo”) akceptacji generowanych wektorów



$$p_n = \frac{V_n}{2^n} \quad V_n = \frac{\pi^{n/2}}{\Gamma(\frac{n}{2} + 1)} r^n$$

$$N_n = \frac{1}{p_n}$$

n	p_n	N_n
2	0.785	1.27
5	0.164	6.08
10	$2.5 \cdot 10^{-3}$	401
20	$2.46 \cdot 10^{-8}$	$4.6 \cdot 10^7$

2^n – objętość n-wymiarowej kostki o boku 2
 V_k – objętość n-wymiarowej kuli zamkniętej w kostce
 N_n – liczba wygenerowanych wektorów na 1 akceptację

- dla $n \leq 5$ metoda akceptowalna
- dla $n > 5$ metoda eliminacji jest nieefektywna
- wektory należy losować inną metodą **(wydajniejszą)**

Rozkład jednorodny na sferze $S^n(0,1)$

Metody służące do generowania punktów wewnątrz sfery wykorzystują etap pośredni

- najpierw musimy rozmieścić losowo w sposób jednorodny punkty na sferze,
a w następnym kroku wprowadzić je do środka.

n-wymiarowy wektor losowy określa położenie na n-wymiarowej sferze jeśli

$$\|\vec{x}\|^2 = \sum_{i=1}^n x_i^2 = 1$$

Losowy wektor n-wymiarowy ma rozkład **sferycznie konturowany** jeśli jego fgp zależy jedynie od długości wektora wodzącego

$$f_{\vec{x}}(\vec{x}) = f(\|\vec{x}\|)$$

Taką własność ma n-wymiarowy rozkład normalny $N^n(0,1)$

$$f_n(\vec{x}) = f_1(x_1)f_1(x_2)\dots f_1(x_n) = \prod_{i=1}^n \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}x_i^2} = \frac{1}{(2\pi)^{n/2}} e^{-\frac{1}{2}\|\vec{x}\|^2} = f(\|\vec{x}\|)$$

Uwaga: każdy element wektora x musi być losowany z rozkładu o tych samych parametrach $N(0,1)$

Generowanie jednorodnego rozkładu punktów na sferze z rozkładu normalnego

1. losujemy wektor z rozkładu normalnego $\vec{x} = [x_1, x_2, \dots, x_n] \quad x_i \in N(0, 1)$

2. normalizujemy długość wektora
(rzutowanie na sferę) $\vec{z} = \frac{\vec{x}}{\|\vec{x}\|}$

Wylosowane punkty leżą na sferze o promieniu $R=1$, a rozkład punktów jest jednorodny.

Rozkład jednorodny w kuli $K^n(0,1)$ generowany z rozkładu sferycznego

Punkty rozłożone w sposób jednorodny na sferze wprowadzamy do środka odpowiednio skalując ich odległość od środka kuli.

W tym celu losujemy zmienną R z rozkładu o fgp

$$h(r) = nr^{n-1}, \quad r \in [0, 1), \quad n = 2, 3, 4$$

Dystrubuanta rozkładu

$$H(r) = \int_0^r dr' n(r'^{n-1}) = r^n$$

Liczbę losową R generujemy odwracając dystrybuantę

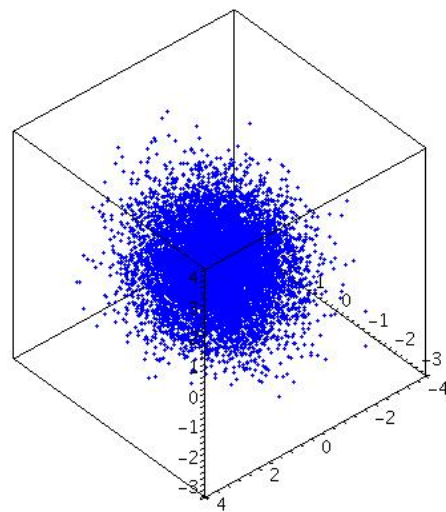
$$U_1 \sim U(0, 1)$$

$$r^n = U_1 \quad \longrightarrow \quad R = U_1^{\frac{1}{n}}, \quad r \in [0, 1)$$

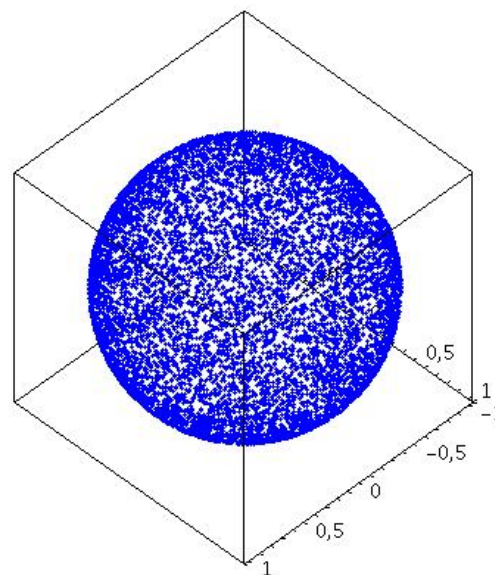
I skalujemy wektor losowy na sferze zamieniamy na losowy w kuli

$$\vec{X}_S = [X_1, X_2, \dots, X_n] \quad \longrightarrow \quad \vec{X}_K = [RX_1, RX_2, \dots, RX_n]$$

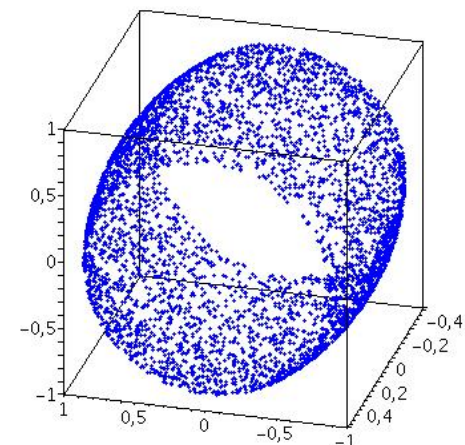
Przykład. generowanie rozkładu: sferycznie konturowanego, na sferze i w kuli 3D (10 tysięcy punktów)



rozkład normalny

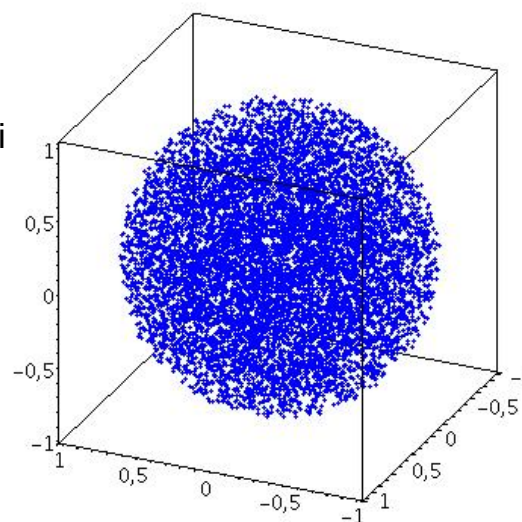


rozkład na sferze



rozkład na sferze - przekrój

rozkład w kuli



potwierdzenie, że rozkład w kuli jest
jednorodny uzyskamy dopiero po
zrobieniu histogramu

Losowanie punktów na sferze 3D i 4D metodą **Marsagli**
(wydajność 2 razy większa niż z wykorzystaniem rozkładu normalnego)

3D

- 1) wylosuj dwie zmienne
- 2) sprawdź warunek, jeśli niespełniony wykonaj punkt (1)
- 3) warunek (2) spełniony - wyznacz 3 współrzędne wektora

$$U_1, U_2 \in U(0, 1)$$

$$s_1 = U_1^2 + U_2^2 \leq 1$$

$$\begin{aligned} x_1 &= 2 U_1 \sqrt{1 - s_1} \\ x_2 &= 2 U_2 \sqrt{1 - s_1} \\ x_3 &= 1 - 2s_1 \end{aligned}$$

4D

- 1) wylosuj dwie zmienne $U_1, U_2 \in U(0, 1)$
- 2) sprawdź warunek $s_1 = U_1^2 + U_2^2 \leq 1$
- 3) wylosuj dwie zmienne $U_3, U_4 \in U(0, 1)$
- 4) sprawdź warunek $s_2 = U_3^2 + U_4^2 \leq 1$

- 5) wyznacz 4 współrzędne

$$\begin{aligned} x_1 &= 2 U_1 \sqrt{1 - s_1} & x_3 &= U_3 \sqrt{\frac{1 - s_1}{s_2}} \\ x_2 &= 2 U_2 \sqrt{1 - s_1} & x_4 &= U_4 \sqrt{\frac{1 - s_1}{s_2}} \end{aligned}$$

Zastosowanie: punkt na sferze (n-wymiarowej) wyznacza kierunek w przestrzeni,
w MC kierunek ten może zostać wykorzystany do przesunięcia cząstki
(np. w metodzie kinetycznej dla gazów, cieczy itp.)

Wielowymiarowy rozkład Gaussa, korelacja, macierz kowariancji

Założmy że dwie zmienne losowe x_1 i x_2 mają rozkład normalny ale o różnych parametrach

$$x_1 \in N(\mu_1, \sigma_1) \quad x_2 \in N(\mu_2, \sigma_2)$$

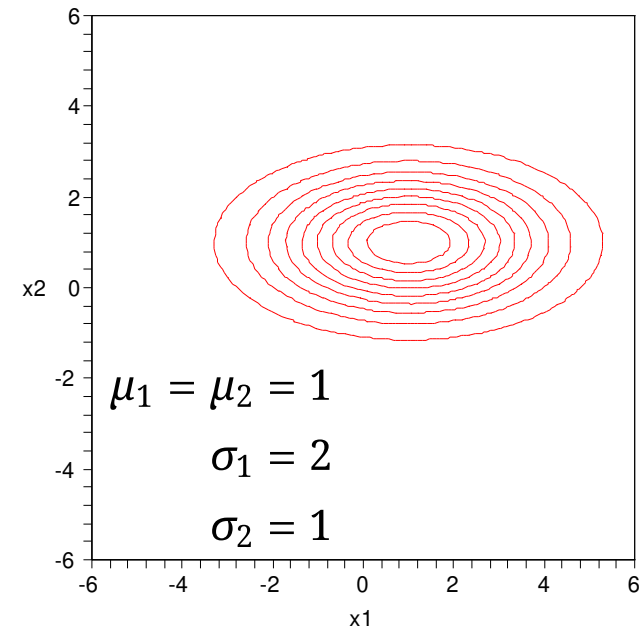
funkcja gęstości prawdopodobieństwa opisująca łączony rozkład obu zmiennych jest iloczynem rozkładów pierwotnych

$$f(x_1, x_2) = f_1(x_1)f_2(x_2) = \frac{1}{\sigma_{x_1}\sqrt{2\pi}}e^{-\frac{(x_1-\mu_1)^2}{2\sigma_1^2}} \frac{1}{\sigma_{x_2}\sqrt{2\pi}}e^{-\frac{(x_2-\mu_2)^2}{2\sigma_2^2}}$$

z warunkiem normalizacji

$$\int_{-\infty}^{\infty} dx_1 \int_{-\infty}^{\infty} dx_2 f(x_1, x_2) = 1$$

Rozkład $f(x_1, x_2)$ jest specyficzny, jego izolinie w płaszczyźnie x_1 - x_2 opisują elipsy o osiach pokrywających się z osiami układu kartezjańskiego



Zapiszmy fgp w ogólnej postaci: **wektorowo-macierzowej**

$$\vec{x} = [x_1, x_2]$$

$$\vec{\mu} = [\mu_1, \mu_2]$$

$$\mathbf{D} = \begin{bmatrix} \sigma_{x_1}^2 & 0 \\ 0 & \sigma_{x_2}^2 \end{bmatrix}$$

$$|\mathbf{D}| = \det(\mathbf{D}) = \sigma_{x_1}^2 \sigma_{x_2}^2$$

$$\mathbf{D}^{-1} = \begin{bmatrix} 1/\sigma_{x_1}^2 & 0 \\ 0 & 1/\sigma_{x_2}^2 \end{bmatrix}$$

$$f(\vec{x}) = \frac{1}{2\pi\sqrt{|\mathbf{D}|}} e^{-(\vec{x}-\vec{\mu})^T \mathbf{D}^{-1} (\vec{x}-\vec{\mu})/2}$$

Co się stanie jeśli obrócimy wektor $(\mathbf{x}'-\boldsymbol{\mu}')$ przy użyciu dowolnej macierzy obrotu $\mathbf{A}$? (**obróć wykonujemy względem punktu $\boldsymbol{\mu}$**)

$$\vec{x} - \vec{\mu} = \mathbf{R}(\vec{x}' - \vec{\mu}')$$

$$f(\vec{x}') dx'_1 dx'_2 = \frac{1}{\|\mathbf{J}\|} f(\vec{x}) dx_1 dx_2 \quad \|\mathbf{J}\| = \det(\mathbf{R}) = 1$$

ponieważ obrót nie zmienia układu współrzędnych więc

$$dx_1 = dx'_1, \quad dx_2 = dx'_2 \quad \rightarrow f(\vec{x}) = f(\vec{x}')$$

$$f(\vec{x}') = \frac{1}{2\pi\sqrt{|\mathbf{D}|}} e^{-(\mathbf{A}(\vec{x}'-\vec{\mu}'))^T \mathbf{D}^{-1} \mathbf{A}(\vec{x}'-\vec{\mu}')/2} = \frac{1}{2\pi\sqrt{|\mathbf{D}|}} e^{-(\vec{x}'-\vec{\mu}')^T \mathbf{A}^T \mathbf{D}^{-1} \mathbf{A}(\vec{x}'-\vec{\mu}')/2}$$

Obrót w 2D

$$\mathbf{R}_\theta = \begin{bmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{bmatrix}$$

$$\mathbf{R}_\theta \mathbf{R}_\theta^{-1} = \mathbf{I} \quad \rightarrow \quad \mathbf{R}_\theta^{-1} = \mathbf{R}_\theta^T$$

$$\vec{V} = \mathbf{R}_\theta \vec{V}'$$

$$\mathbf{R}_\theta^{-1} \vec{V} = \mathbf{R}_\theta^T \vec{V} = \vec{V}'$$

Iloczyn macierzy można zapisać w zwartej postaci jeśli znajdziemy jego macierz odwrotną

wykorzystujemy własność macierzy obrotu

$$\mathbf{R}\mathbf{R}^{-1} = \mathbf{1} \quad \rightarrow \quad \mathbf{R}^{-1} = \mathbf{R}^T$$

$$\mathbf{R}^T \mathbf{D}^{-1} \mathbf{R} (\mathbf{R}^T \mathbf{D} \mathbf{R}) = \mathbf{1}$$

$$\mathbf{R}^T \mathbf{D}^{-1} \mathbf{R} = (\mathbf{R}^T \mathbf{D} \mathbf{R})^{-1} = \mathbf{C}^{-1}$$

C - to **macierz kowariancji**, zawiera informacje o zależnościach pomiędzy zmiennymi

Ponieważ transformacja $\mathbf{R}^T \mathbf{D}^{-1} \mathbf{R} = (\mathbf{R}^T \mathbf{D} \mathbf{R})^{-1} = \mathbf{C}^{-1}$

jest unitarna więc wyznacznik macierzy nie ulega zmianie

$$|\mathbf{D}| = |\mathbf{C}|$$

Po zmianie oznaczeń otrzymujemy ogólny wielowymiarowy (**n-wymiarowy**) rozkład normalny

$$f(\vec{x}) = \frac{1}{\sqrt{(2\pi)^n |\mathbf{C}|}} e^{-\frac{1}{2} (\vec{x} - \vec{\mu})^T \mathbf{C}^{-1} (\vec{x} - \vec{\mu})}$$

Na czym polega różnica?

- Macierz **D** parametryzowała rozkład zmiennych niezależnych (brak elementów pozadiagonalnych)
- Pojawienie się elementów pozadiagonalnych w macierzy **C** oznacza bezpośrednią korelację pomiędzy zmiennymi – o indeksach odpowiadających numerom wierszy/kolumn z niezerowymi elementami.

Jawna postać macierzy kowariancji

$$\mathbf{C} = \begin{bmatrix} \sigma_{x_1}^2 & \sigma_{x_1, x_2}^2 \\ \sigma_{x_1, x_2}^2 & \sigma_{x_2}^2 \end{bmatrix}$$

$$\sigma_{x_1}^2 = \langle (x_1 - \mu_1)^2 \rangle$$

$$\sigma_{x_1, x_2}^2 = \langle (x_1 - \mu_1)(x_2 - \mu_2) \rangle$$

μ_1, μ_2 - to wartości oczekiwane
skorelowanych zmiennych x_1 i x_2

$$\det(\mathbf{C}) = \sigma_{x_1}^2 \sigma_{x_2}^2 - \sigma_{x_1, x_2}^2 \geq 0 \quad \text{-macierz nieujemnieokreślona}$$

$$\sigma_{x_1}^2 \sigma_{x_2}^2 \geq \sigma_{x_1, x_2}^2 \rightarrow \frac{\sigma_{x_1, x_2}^2}{\sigma_{x_1}^2 \sigma_{x_2}^2} \leq 1 \rightarrow -1 \leq \frac{\sigma_{x_1, x_2}}{\sigma_{x_1} \sigma_{x_2}} \leq 1$$

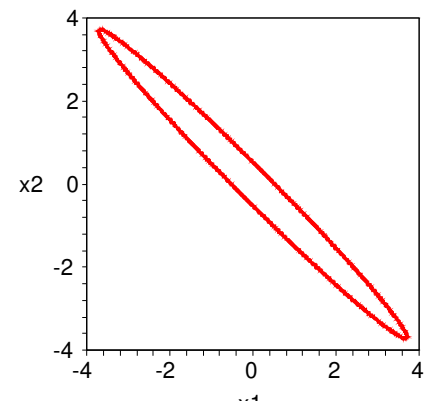
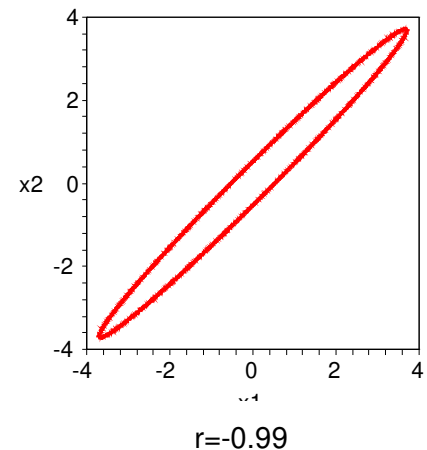
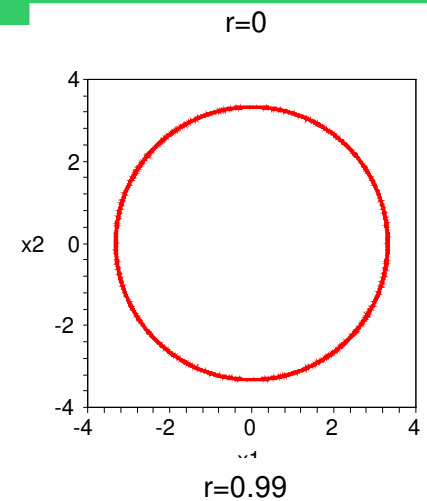
$$r = \frac{\sigma_{x_1, x_2}}{\sigma_{x_1} \sigma_{x_2}} \in [-1, 1] \quad \text{- współczynnik korelacji}$$

$r = 1$ - całkowita korelacja

$r = -1$ - całkowita antykorelacja

$-1 < r < 1$ - zmienne częściowo skorelowane

Korelacja/antykorelacja oznacza że wylosowanie zmiennej x_1 określa jednoznacznie wartość drugiej zmiennej x_2 .



Profile dwuwymiarowego rozkładu normalnego o współczynniku korelacji: $r=0, 0.99, -0.99$.

Przykład losowania pary zmiennych skorelowanych o rozkładzie normalnym

dana jest macierz kowariancji oraz pierwsze momenty rozkładu

$$C = \begin{bmatrix} 4 & 2 \\ 2 & 3 \end{bmatrix} \quad C^{-1} = \begin{bmatrix} 3/8 & -1/4 \\ -1/4 & 1/2 \end{bmatrix} \quad \det(C) = 8 \quad \mu_1 = \mu_2 = 0$$

$$f(\vec{x}) = \frac{1}{2\pi\sqrt{8}} e^{-\vec{x}^T C^{-1} \vec{x}/2} = \frac{1}{2\pi\sqrt{8}} e^{-Q/2} \quad \text{Q - forma kwadratowa}$$

postać formy kwadratowej Q (wykładnik)

$$Q = x_1 (3/8 x_1 - 1/4 x_2) + x_2 (-1/4 x_1 + 1/2 x_2)$$

dokonujemy redukcji/marginalizacji zmiennej x_2 - całkujemy

$$f(x_1) = \int_{-\infty}^{\infty} dx_2 f(x_1, x_2) = \frac{1}{2\sqrt{2\pi}} e^{-\frac{x_1^2}{2 \cdot 2^2}}$$

jest to rozkład normalny o parametrach

$$\mu = 0, \quad \sigma = 2$$

losujemy zmienną z rozkładu normalnego

$$x_1 \in N(0, 2)$$

Mając określoną wartość x_1 , wyznaczamy x_2 posługując się **prawdopodobieństwem warunkowym**

$$f(x_2 | \underbrace{x_1}_{=const}) = \frac{f(x_1, x_2)}{f(x_1)} = \frac{1}{\sqrt{2}\sqrt{2\pi}} e^{-\frac{(x_2 - x_1/2)^2}{2 \cdot (\sqrt{2})^2}}$$

to także jest rozkład normalny ale o innych parametrach

$$\mu = \frac{x_1}{2}, \quad \sigma = \sqrt{2}$$

Drugą zmienną losujemy z tego rozkładu

$$x_2 \in N\left(\frac{x_1}{2}, \sqrt{2}\right)$$

Uwagi:

- dwukrotnie posłużyliśmy się rozkładem normalnym, ale o parametrach określonych przez macierz kowariancji
- metoda łatwa w implementacji dla 2 zmiennych, dla większej liczby staje się mało praktyczna (potrzebny jest inny sposób losowania – łatwy do implementacji niezależnie od liczby zmiennych np. 100 ?)

Metoda transformacji osi głównych

$$f(\vec{x}) = \frac{1}{\sqrt{(2\pi)^n |\mathbf{C}|}} e^{-(\vec{x}-\vec{\mu})^T \mathbf{C}^{-1} (\vec{x}-\vec{\mu}) / 2}$$

Macierz kowariancji $\mathbf{C}$ uzyskaliśmy dokonując przekształcenia macierzy $\mathbf{D}$

$$\mathbf{D} = \begin{bmatrix} \sigma_{x_1}^2 & 0 \\ 0 & \sigma_{x_2}^2 \end{bmatrix} \quad \mathbf{R}^T \mathbf{D}^{-1} \mathbf{R} = (\mathbf{R}^T \mathbf{D} \mathbf{R})^{-1} = \mathbf{C}^{-1}$$

Obie macierze są podobne i mają identyczne widmo wartości własnych: to nieujemne elementy diagonalne $\mathbf{D}$.
Zatem macierz $\mathbf{C}$ jest symetryczna i dodatnio-określona.

Można więc łatwo (numerycznie) wyznaczyć jej rozkład **Banachiewicza-Choleskyego** ($\mathbf{L}\mathbf{L}^T$)

$$\mathbf{C} = \mathbf{L}\mathbf{L}^T \quad \mathbf{C}\mathbf{C}^{-1} = \mathbf{L}\mathbf{L}^T (\mathbf{L}^T)^{-1} \mathbf{L}^{-1} = \mathbf{I} \quad \mathbf{C}^{-1} = (\mathbf{L}^T)^{-1} \mathbf{L}^{-1}$$

podstawiamy $\vec{x} - \vec{\mu} = \vec{y}$

i przekształcamy wykładnik $\vec{y}^T \mathbf{C}^{-1} \vec{y} = \vec{y}^T (\mathbf{L}\mathbf{L}^T)^{-1} \vec{y} = \underbrace{\vec{y}^T (\mathbf{L}^T)^{-1}}_{\vec{z}^T} \underbrace{\mathbf{L}^{-1} \vec{y}}_{\vec{z}} = \vec{z}^T \mathbf{I} \vec{z}$

Po transformacji, macierz kowariancji dla nowych wektorów $\mathbf{z}$ jest macierzą jednostkową $\mathbf{I}$:

- brak korelacji między nowymi zmiennymi
- wariancja każdej zmiennej jest równa 1
- **nowe zmienne mają rozkład normalny $N(0,1)$ i możemy je losować niezależnie**

$$\mathbf{L}^{-1} \vec{y} = \vec{z} \quad \rightarrow \quad \vec{y} = \mathbf{L} \vec{z} \quad \rightarrow \quad \vec{x} = \mathbf{L} \vec{z} - \vec{\mu} \quad z_i \in N(0, 1)$$

Wektory losowe $\mathbf{x}$ mają rozkład zadany macierzą kowariancji $\mathbf{C}$.

Zalety metody: łatwa, skuteczna, wydajna, niezależnie od liczby zmiennych

Testowanie generatorów liczb pseudolosowych

Jakość wyników generowanych w metodzie Monte Carlo w dużej mierze zależy od jakości użytego generatora. Należy używać generatorów „dobrej jakości” tj. mających duży okres oraz dobre własności statystyczne (o znikomej korelacji elementów ciągu liczb).

Generalnie testowanie generatorów może odbywać się na dwa sposoby:

1) **testy analityczne** – polegają na określeniu własności generatora tylko na podstawie analizy jego konstrukcji, taka analiza możliwa jest jedynie dla generatorów o prostej budowie (generatory liniowe i złożone)

2) **testy aplikacyjne (empiryczne)** – testowane są ciągi generowanych na komputerze liczb, takie testy mają różną postać, ponieważ pojedynczy test jest ukierunkowany na badanie jednej lub kilku cech generatora mających wpływ na korelację elementów ciągu czy poprawności generowanego rozkładu (lub inaczej – wielkości odchylenia od zakładanego rozkładu)

Testowane są głównie generatory o rozkładzie jednorodnym, gdyż są one podstawą działania generatorów o dowolnym rozkładzie prawdopodobieństwa.

Testowanie generatorów o innym rozkładzie wykonuje się w celu przetestowania poprawności implementowanej metody (odwracania dystrybucji czy eliminacji).

Dobry generator nie musi spełniać wszystkich testów, idealny generator nie istnieje. Testy których nie spełnia dają informację czy i w jakich zastosowaniach należy spodziewać się problemów, których po takiej analizie możemy uniknąć.

Testy zgodności z zadaniem rozkładem - test chi-kwadrat

Jest to prosty i zarazem najpopularniejszy test. Badamy w nim hipotezę, że generowana zmienna losowa X ma rozkład prawdopodobieństwa o dystrybuancie F , czyli jest miarą odchylenia generowanego numerycznie rozkładu od rozkładu zadanego.

Jeżeli $F(a)=0$ $F(b)=1$

to możemy dokonać następującego podziału zbioru wartości zmiennej X

$$a < a_1 < a_2 < \dots < a_k = b \quad \Longrightarrow \quad p_i = P(a_{i-1} \leq x \leq a_i) = \int_{a_{i-1}}^{a_i} f(x) dx$$

Generujemy ciąg n liczb

$$X_1, X_2, \dots, X_n$$

Sprawdzamy ile z nich spełnia warunek

$$a_{i-1} < X \leq a_i$$

ich liczbę oznaczamy n_i .

Statystyką testu jest

$$\chi^2_{k-1} = \sum_{i=1}^k \frac{(n_i - np_i)^2}{np_i}$$

Dla dużego n statystyka ta ma rozkład χ^2 o $(k-1)$ stopniach swobody.

Dlaczego $k-1$ stopni swobody?

Znamy wartość sumy oraz zliczenia w k podprzedziałów, zatem jedna z nich jest liniowo zależna od pozostałych – 1 stopień swobody mniej.

Test χ^2 jest miarą jakości całego generowanego rozkładu – to odróżnia go od testów cząstkowych polegających na sprawdzeniu poprawności generowanych kolejnych momentów rozkładu.

W tablicach statystycznych podawane są zazwyczaj wartości prawdopodobieństwa dla warunku

$$P(\chi^2 \geq \chi_0^2)$$

zatem aby uzyskać prawdopodobieństwo dla warunku przeciwnego musimy użyć

$$P(\chi^2 \leq \chi_0^2) = 1 - P(\chi^2 \geq \chi_0^2)$$

Wynik porównujemy z wartościami granicznymi: dolną i górną dla zadanego przedziału ufności $(1-\alpha)$

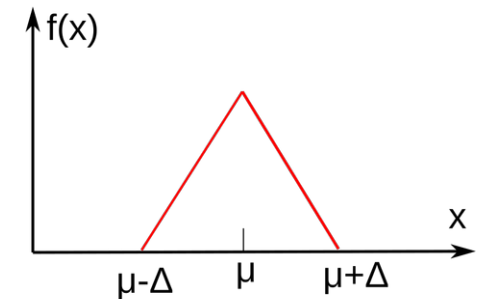
$$\chi^2(\alpha) < \chi^2 < \chi^2(1 - \alpha)$$

jeśli	$\chi^2 < \chi^2(\alpha)$	- hipotezę odrzucamy, bo jest „zbyt dobry”, generowany rozkład jest zbyt bliski rozkładowi zadanemu
	$\chi^2 > \chi^2(1 - \alpha)$	- hipotezę odrzucamy, bo rozkład bardzo odbiega od rozkładu zadanego

Przykład. Test χ^2 generatora o rozkładzie trójkątnym

$$f(x; \mu, \Delta) = -\frac{|x - \mu|}{\Delta^2} + \frac{1}{\Delta}$$

mamy 2 parametry
więc 2 stopnie swobody
mniej:
(k-1-2)



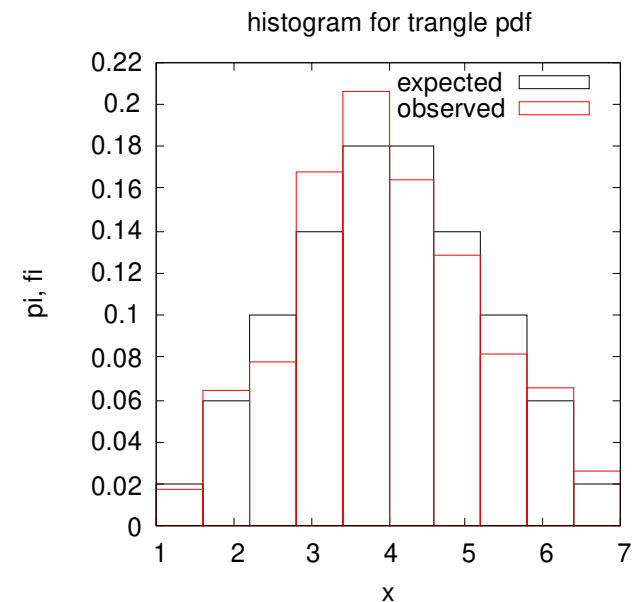
$$F(a) = P(x < a) = \int_{\mu-\Delta}^a f(x; \mu, \Delta) dx = \begin{cases} -\frac{1}{\Delta^2} \left(-\frac{x^2}{2} + \mu x \right) + \frac{x}{\Delta}, & x \leq \mu \\ -\frac{1}{\Delta^2} \left(\frac{x^2}{2} - \mu x + \mu^2 \right) + \frac{x}{\Delta}, & x > \mu \end{cases}$$

$$x = \mu + (\xi_1 + \xi_2 - 1) \cdot \Delta, \quad \xi_1, \xi_2 \in U(0, 1)$$

Procedura testowa:

1. Generujemy serię N liczb z rozkładu
2. Grupujemy liczby w k przedziałach
3. Liczymy wartość statystyki testowej χ^2
4. Dla zadanego poziomu α sprawdzamy czy

$$\chi^2 < \text{wartości granicznej}$$



$$2.167 (\alpha = 0.05) < \chi_{k-2-1}^2 = 4.52 < 14.06 (\alpha = 0.95) \quad - \text{nie odrzucamy } H_0$$

Test sum – generator U(0,1)

Definiujemy funkcję

$$y = X_1 + X_2 + \dots + X_m$$

wygenerowana zmienna losowa Y ma rozkład o wielomianowej gęstości prawdopodobieństwa

dla $m=2$

$$g_2(y) = \begin{cases} y & 0 \leq y \leq 1 \\ 2 - y & 1 < y \leq 2 \end{cases}$$

dla $m=3$

$$g_3(y) = \begin{cases} \frac{y^2}{2} & 0 \leq y \leq 1 \\ \frac{1}{2}(y^2 - 3(y-1)^2) & 1 < y \leq 2 \\ \frac{1}{2}(y^2 - 3((y-1)^2 + 3(y-2)^2)) & 2 < y \leq 3 \end{cases}$$

Testowanie generatora polega na weryfikacji hipotezy, że zmienna losowa Y ma rozkład zgodny z $g_m(y)$.

Zazwyczaj m nie przekracza 5.

Test niezależności elementów ciągu – generator U(0,1)

Chcemy zbadać czy w generowanym ciągu istnieją korelacje.

Ustalamy zasięg korelacji r i generujemy $(n+r)$ liczb. Następnie wyznaczamy współczynnik korelacji dla próbki

$$R_r(n) = \frac{\frac{1}{n} \sum_{i=1}^n x_i x_{i+r} - \left(\frac{1}{n} \sum_{i=1}^n x_i \right)^2}{\frac{1}{n} \sum_{i=1}^n x_i^2 - \left(\frac{1}{n} \sum_{i=1}^n x_i \right)^2} \quad r = 1, 2, 3, \dots$$

Jeśli w ciągu nie występuje korelacja to $R_r(n)$ ma rozkład normalny

$$R_r(n) \in N \left(-\frac{1}{n}, \frac{1}{\sqrt{n}} \right)$$

Test polega na sprawdzeniu tej hipotezy.

Test graficzny - jakościowy

Inną formą testu jest wizualizacja, robimy wykres punktowy w 2D dla odpowiednio dobranych par

$$(x_i, x_{i+r}), \quad i = 1, 2, 3, \dots, n$$

Wystąpienie korelacji powoduje układanie się punktów wzdłuż kilku wybranych kierunków
- tworzy się regularna struktura/siatka (przykład: generator **RANDU**).

