

Matematyka dyskretna

Skrypt dydaktyczny dla studentów Informatyki

Grzegorz Dymek

KUL 2025

Wstęp

Skrypt ten powstał na podstawie wykładów z Matematyki dyskretnej, które autor prowadzi na Katolickim Uniwersytecie Lubelskim Jana Pawła II dla studentów informatyki. Kurs Matematyka dyskretna obejmuje dużo różnych ciekawych tematów. Omawiane są elementy logiki i teorii zbiorów, relacje, funkcje, moce zbiorów, indukcja, rekurencja, kombinatoryka, podstawy teorii grafów, algebry Boole’a, funkcje booleowskie i wreszcie na koniec absolutne podstawy teorii automatów. Przedstawione pojęcia są podane w przystępny i zrozumiały sposób oraz zilustrowane są wieloma przykładami. Autor ma nadzieję, że skrypt ten będzie dla studentów pozycją wartościową i pomocną w zrozumieniu pojęć matematyki dyskretnej.

1. Elementy logiki

Definicja. *Zdaniem* w sensie logicznym nazywa się zdanie oznajmujące, któremu można przypisać jedną z dwóch wartości – **prawda** lub **fałsz**. Przyjmuje się, że symbolem prawdy jest 1, a symbolem fałszu 0.

Definicja. *Zmienną zdaniową (logiczną)* nazywa się zmienną, w miejsce której wstawia się zdanie (prawdziwe lub fałszywe) otrzymując zdanie w sensie logicznym.

Uwaga. Najczęściej zmienne zdaniowe oznaczają się małymi literami: $p, q, r, \dots$

Wartość logiczną zdania p oznacza się symbolem $w(p)$. Zatem

$w(p) = 0$ oznacza, że p jest zdaniem fałszywym,

$w(p) = 1$ oznacza, że p jest zdaniem prawdziwym.

Ze zdań logicznych można tworzyć zdania złożone przy pomocy *spójników logicznych* (*funktorów zdaniotwórczych*). W matematyce używane są następujące spójniki logiczne:

- *negacja*: $\sim p$ lub $\neg p$, czyta się „nie p ” lub „nieprawda, że p ”,
- *koniunkcja*: $p \wedge q$, czyta się „ p i q ”,
- *alternatywa*: $p \vee q$, czyta się „ p lub q ”,
- *implikacja*: $p \Rightarrow q$, czyta się „ p implikuje q ” lub „jeżeli p , to q ”,
- *równoważność*: $p \Leftrightarrow q$, czyta się „ p jest równoważne q ” lub „ p wtedy i tylko wtedy, gdy q ”.

Uwaga. W implikacji $p \Rightarrow q$ zdanie p nazywa się *poprzednikiem*, a zdanie q nazywa się *następnikiem* tej implikacji.

Uwaga. Negacja to funktor jednoargumentowy, a koniunkcja, alternatywa, implikacja i równoważność to funktory dwuargumentowe.

Definicje wartościowań spójników logicznych:

$w(p)$	$w(\sim p)$
0	1
1	0

$w(p)$	$w(q)$	$w(p \wedge q)$	$w(p \vee q)$	$w(p \Rightarrow q)$	$w(p \Leftrightarrow q)$
0	0	0	0	1	1
0	1	0	1	1	0
1	0	0	1	0	0
1	1	1	1	1	1

Definicja. *Formułą logiczną* nazywa się wyrażenie zbudowane ze zmiennych zdaniowych, funktorów zdaniotwórczych oraz nawiasów w następujący sposób:

- zdania proste i zmienne zdaniowe są formułami logicznymi,
- jeżeli P i Q są formułami logicznymi, to $\sim P$, $(P \wedge Q)$, $(P \vee Q)$, $(P \Rightarrow Q)$ oraz $(P \Leftrightarrow Q)$ są również formułami logicznymi.

Definicja. *Tautologią* lub *prawem rachunku zdań* nazywa się formułę logiczną, która jest prawdziwa bez względu na wartość logiczną występujących w niej zmiennych zdaniowych.

Metody sprawdzania, czy formuła jest tautologią:

1. *Metoda zero-jedynkowa:* buduje się tablicę dla wszystkich układów wartościowań zmiennych zdaniowych, a następnie kolejno ocenia prawdziwość zdań złożonych. Formuła jest tautologią, jeśli kolumna odpowiadająca formule jest złożona z samych jedynek.
2. *Skrócona metoda sprawdzania wyrażen rachunku zdań:* zakłada się, że formuła jest fałszywa i na podstawie tabelki wartościowań poszczególnych funktorów zdaniotwórczych sprawdza się dla jakich wartościowań zmiennych logicznych jest to możliwe. Jeśli takich wartościowań nie ma, to formuła jest tautologią.

Przykład. Sprawdźmy metodą zero-jedynkową, czy formuła

$$p \Rightarrow (\sim p \vee q)$$

jest tautologią. Mamy

$w(p)$	$w(q)$	$w(\sim p)$	$w(\sim p \vee q)$	$w(p \Rightarrow (\sim p \vee q))$
0	0	1	1	1
0	1	1	1	1
1	0	0	0	0
1	1	0	1	1

Zatem formuła ta nie jest tautologią.

Przykład. Sprawdzimy metodą skróconą, czy formuła

$$[(p \vee q) \wedge \sim p] \Rightarrow q$$

jest tautologią. Mamy

1. $w([(p \vee q) \wedge \sim p] \Rightarrow q) = 0$	{zał.}
2. $w((p \vee q) \wedge \sim p) = 1$	{1.}
3. $w(q) = 0$	{1.}
4. $w(p \vee q) = 1$	{2.}
5. $w(\sim p) = 1$	{2.}
6. $w(p) = 0$	{5.}
7. $w(q) = 1$	{4., 6.}
8. sprzeczność	{3., 7.}

Zatem formuła ta jest tautologią.

Wybrane prawa rachunku zdań:

- prawa przemienności

$$(p \wedge q) \Leftrightarrow (q \wedge p),$$

$$(p \vee q) \Leftrightarrow (q \vee p),$$

- prawa łączności

$$[(p \wedge q) \wedge r] \Leftrightarrow [p \wedge (q \wedge r)],$$

$$[(p \vee q) \vee r] \Leftrightarrow [p \vee (q \vee r)],$$

- prawa rozdzielności

$$[(p \wedge q) \vee r] \Leftrightarrow [(p \vee r) \wedge (q \vee r)],$$

$$[(p \vee q) \wedge r] \Leftrightarrow [(p \wedge r) \vee (q \wedge r)],$$

- prawa idempotentności

$$(p \wedge p) \Leftrightarrow p,$$

$$(p \vee p) \Leftrightarrow p,$$

- prawo podwójnego przeczenia

$$\sim (\sim p) \Leftrightarrow p,$$

- prawo wyłączonego środka

$$p \vee \sim p,$$

- prawo sprzeczności

$$\sim (p \wedge \sim p),$$

- prawa De Morgana

$$\sim (p \wedge q) \Leftrightarrow (\sim p \vee \sim q),$$

$$\sim (p \vee q) \Leftrightarrow (\sim p \wedge \sim q),$$

- określenie równoważności

$$(p \Leftrightarrow q) \Leftrightarrow [(p \Rightarrow q) \wedge (q \Rightarrow p)],$$

- określenie implikacji

$$(p \Rightarrow q) \Leftrightarrow (\sim p \vee q),$$

- zaprzeczenie implikacji

$$\sim (p \Rightarrow q) \Leftrightarrow (p \wedge \sim q),$$

- prawo transpozycji

$$(p \Rightarrow q) \Leftrightarrow (\sim q \Rightarrow \sim p),$$

- prawo sylogizmu

$$[(p \Rightarrow q) \wedge (q \Rightarrow r)] \Rightarrow (p \Rightarrow r),$$

- reguła odrywania

$$[p \wedge (p \Rightarrow q)] \Rightarrow q,$$

- prawo sygnifikacji

$$p \Rightarrow (q \Rightarrow p),$$

- prawo Duns Scotusa

$$\sim p \Rightarrow (p \Rightarrow q),$$

- prawo Claviusa

$$(\sim p \Rightarrow p) \Rightarrow p,$$

- prawo Fregego

$$[p \Rightarrow (q \Rightarrow r)] \Rightarrow [(p \Rightarrow q) \Rightarrow (p \Rightarrow r)].$$

Uwaga. Prawdziwość wszystkich powyższych praw łatwo sprawdza się metodą zero-jedynkową lub metodą skróconą.

Definicja. *Funkcją zdaniową* nazywa się wyrażenie zawierające zmienną zdaniową, które staje się zdaniem w sensie logicznym, jeżeli w miejsce zmiennej podstawimy nazwę przedmiotu.

Definicja. Z każdą funkcją zdaniową związany jest zbiór, który zawiera nazwy wszystkich przedmiotów, które wolno podstawiać w miejsce zmiennej zdaniowej. Zbiór ten nazywa się *zakres zmiennej zdaniowej*.

Uwaga. Prawdziwość funkcji zdaniowej może zależeć od zakresu zmiennej zdaniowej w niej występującej.

Definicja. Przedmiot spełnia funkcję zdaniową, jeżeli po wstawieniu jego nazwy w miejsce zmiennej zdaniowej w funkcji zdaniowej otrzymuje się zdanie prawdziwe.

Przykład. Rozważmy następującą funkcję zdaniową w zbiorze $\mathbb{R}$ wszystkich liczb rzeczywistych:

$$x^2 - 4x + 3 = 0.$$

Zakresem zmiennej x występującej w tej funkcji jest zatem zbiór $\mathbb{R}$. Element 3 spełnia tę funkcję zdaniową, a na przykład element 4 jej nie spełnia. Ponadto funkcja ta nie jest prawdziwa w zbiorze $\mathbb{R} \setminus \{1, 3\}$.

Uwaga. Rozpatruje się również funkcje zdaniowe więcej niż jednej zmiennej zdaniowej. Analogicznie mówimy wówczas o zakresach zmiennych i spełnianiu funkcji zdaniowej.

Uwaga. Funkcje zdaniowe oznacza się przez: $\varphi(x), \psi(x, y), \dots$. Funkcje $\varphi, \psi, \dots$ nazywa się *predykaty*.

Definicja. *Kwantyfikatory* to funktory zdaniotwórcze, które przekształcają funkcje zdaniowe w zdanie w sensie logicznym.

Niech zbiór X będzie zakresem zmiennej zdaniowej x . Rozróżniamy:

- *kwantyfikator ogólny (uniwersalny)*:

$$\bigwedge_{x \in X} \varphi(x) \quad \text{lub} \quad \forall x \in X \varphi(x),$$

czyta się „dla każdego x spełniona jest funkcja $\varphi(x)$ ”.

- *kwantyfikator szczegółowy (egzystencjalny)*:

$$\bigvee_{x \in X} \varphi(x) \quad \text{lub} \quad \exists x \in X \varphi(x),$$

czyta się „istnieje x taki, że spełniona jest funkcja $\varphi(x)$ ”.

Uwaga. Kwantyfikator ogólny jest uogólnieniem koniunkcji, zaś szczegółowy – alternatywy. Niech $X = \{x_1, x_2, \dots, x_n\}$ będzie zakresem zmiennej zdaniowej x w funkcji zdaniowej $\varphi(x)$. Wówczas

- $\left(\bigwedge_{x \in X} \varphi(x) \right) \Leftrightarrow (\varphi(x_1) \wedge \varphi(x_2) \wedge \dots \wedge \varphi(x_n)),$
- $\left(\bigvee_{x \in X} \varphi(x) \right) \Leftrightarrow (\varphi(x_1) \vee \varphi(x_2) \vee \dots \vee \varphi(x_n)).$

Definicja. *Zasięgiem* kwantyfikatora występującego w danym wyrażeniu jest ta część tego wyrażenia, będąca również wyrażeniem zdaniowym, ujęta w parę jednakowych nawiasów, z których otwierający występuje bezpośrednio po danym kwantyfikatorze.

Definicja. Zmienna zdaniowa x nazywa się *zmienna wolna* wtedy i tylko wtedy, gdy nie występuje w zasięgu żadnego kwantyfikatora, który odnosi się do zmiennej x .

Definicja. Zmienna zdaniowa x nazywa się *zmienna związana* przez dany kwantyfikator wtedy i tylko wtedy, gdy kwantyfikator odnosi się do zmiennej x , zmienna x występuje w zasięgu tego kwantyfikatora i jest w tym zasięgu zmienną wolną.

Umowa. Opuszczamy nawiasy w ciągu kilku występujących po sobie kwantyfikatorów oraz w prostych funkcjach zdaniowych występujących bezpośrednio po kwantyfikatorze.

Reguła podstawiania:

Jeśli P jest prawem rachunku zdań, to podstawiając w P za zmienne zdaniowe dowolne funkcje zdaniowe lub zdania zbudowane z funkcji zdaniowych za pomocą funktorów zdaniotwórczych i kwantyfikatorów otrzymuje się funkcję zdaniową lub zdanie prawdziwe (w odpowiednim zbiorze).

Przykład. Weźmy następujące prawo rachunku zdań:

$$P : (p \Rightarrow q) \Rightarrow (\sim p \vee q).$$

Podstawmy:

$$p : \bigwedge_x \varphi(x, y) \quad \text{oraz} \quad q : \bigvee_x \psi(x, y).$$

Otrzymujemy:

$$P^* : \left(\bigwedge_x \varphi(x, y) \Rightarrow \bigvee_x \psi(x, y) \right) \Rightarrow \left(\sim \bigwedge_x \varphi(x, y) \vee \bigvee_x \psi(x, y) \right).$$

Wybrane prawa rachunku kwantyfikatorów:

Niech φ i ψ będą funkcjami zdaniowymi jednej zmiennej. Niech x będzie zmienną zdaniową, której zakresem jest zbiór X . Mamy

•

$$\bigwedge_{x \in X} \varphi(x) \Rightarrow \varphi(y), \quad y \in X,$$

•

$$\bigwedge_{x \in X} \varphi(x) \Rightarrow \bigvee_{x \in X} \varphi(x),$$

•

$$\varphi(y) \Rightarrow \bigvee_{x \in X} \varphi(x), \quad y \in X,$$

- prawa De Morgana dla kwantyfikatorów

$$\begin{aligned}\sim \left(\bigwedge_{x \in X} \varphi(x) \right) &\Leftrightarrow \left(\bigvee_{x \in X} \sim \varphi(x) \right), \\ \sim \left(\bigvee_{x \in X} \varphi(x) \right) &\Leftrightarrow \left(\bigwedge_{x \in X} \sim \varphi(x) \right),\end{aligned}$$

- prawa włączania-wyłączania dla kwantyfikatorów

$$\begin{aligned}\bigwedge_{x \in X} (\varphi(x) \vee \psi) &\Leftrightarrow \left(\bigwedge_{x \in X} \varphi(x) \vee \psi \right), \\ \bigvee_{x \in X} (\varphi(x) \vee \psi) &\Leftrightarrow \left(\bigvee_{x \in X} \varphi(x) \vee \psi \right), \\ \bigwedge_{x \in X} (\varphi(x) \wedge \psi) &\Leftrightarrow \left(\bigwedge_{x \in X} \varphi(x) \wedge \psi \right), \\ \bigvee_{x \in X} (\varphi(x) \wedge \psi) &\Leftrightarrow \left(\bigvee_{x \in X} \varphi(x) \wedge \psi \right), \\ \bigwedge_{x \in X} (\varphi(x) \Rightarrow \psi) &\Leftrightarrow \left(\bigvee_{x \in X} \varphi(x) \Rightarrow \psi \right), \\ \bigvee_{x \in X} (\varphi(x) \Rightarrow \psi) &\Leftrightarrow \left(\bigwedge_{x \in X} \varphi(x) \Rightarrow \psi \right), \\ \bigwedge_{x \in X} (\psi \Rightarrow \varphi(x)) &\Leftrightarrow \left(\psi \Rightarrow \bigwedge_{x \in X} \varphi(x) \right), \\ \bigvee_{x \in X} (\psi \Rightarrow \varphi(x)) &\Leftrightarrow \left(\psi \Rightarrow \bigvee_{x \in X} \varphi(x) \right),\end{aligned}$$

gdzie ψ jest funkcją zdaniową, w której zmienna x nie występuje jako zmienna wolna,

- prawa rozdzielności kwantyfikatorów

$$\begin{aligned}\left(\bigwedge_{x \in X} (\varphi(x) \wedge \psi(x)) \right) &\Leftrightarrow \left(\bigwedge_{x \in X} \varphi(x) \wedge \bigwedge_{x \in X} \psi(x) \right), \\ \left(\bigvee_{x \in X} (\varphi(x) \vee \psi(x)) \right) &\Leftrightarrow \left(\bigvee_{x \in X} \varphi(x) \vee \bigvee_{x \in X} \psi(x) \right), \\ \left(\bigvee_{x \in X} (\varphi(x) \wedge \psi(x)) \right) &\Rightarrow \left(\bigvee_{x \in X} \varphi(x) \wedge \bigvee_{x \in X} \psi(x) \right), \\ \left(\bigwedge_{x \in X} \varphi(x) \vee \bigwedge_{x \in X} \psi(x) \right) &\Rightarrow \left(\bigwedge_{x \in X} (\varphi(x) \vee \psi(x)) \right),\end{aligned}$$

- prawa przestawiania dla kwantyfikatorów

$$\begin{aligned} \left(\bigwedge_{x \in X} \bigwedge_{y \in Y} \varphi(x, y) \right) &\Leftrightarrow \left(\bigwedge_{y \in Y} \bigwedge_{x \in X} \varphi(x, y) \right), \\ \left(\bigvee_{x \in X} \bigvee_{y \in Y} \varphi(x, y) \right) &\Leftrightarrow \left(\bigvee_{y \in Y} \bigvee_{x \in X} \varphi(x, y) \right), \\ \left(\bigvee_{x \in X} \bigwedge_{y \in Y} \varphi(x, y) \right) &\Rightarrow \left(\bigwedge_{y \in Y} \bigvee_{x \in X} \varphi(x, y) \right), \end{aligned}$$

gdzie φ jest teraz funkcją zdaniową dwu zmiennych, a y jest zmienną zdaniową, której zakresem jest zbiór Y .

Dowody powyższych praw nie są trudne. Udowodnijmy dla przykładu prawo:

$$\bigwedge_{x \in X} \varphi(x) \Rightarrow \varphi(y), \quad y \in X.$$

Niech a będzie dowolnym elementem zbioru X . Gdyby zdanie $\bigwedge_{x \in X} \varphi(x) \Rightarrow \varphi(a)$ było fałszywe, to wtedy $\varphi(a)$ byłoby zdaniem fałszywym i $\bigwedge_{x \in X} \varphi(x)$ byłoby zdaniem prawdziwym. Ale $\bigwedge_{x \in X} \varphi(x)$ jest zdaniem prawdziwym tylko wtedy, gdy dla każdego $b \in X$ zdanie $\varphi(b)$ jest prawdziwe. Otrzymalibyśmy sprzeczność z założeniem, że zdanie $\varphi(a)$ jest fałszywe. A więc każdy element $a \in X$ spełnia powyższą funkcję zdaniową, czyli jest ona prawem rachunku kwantyfikatorów.

2. Elementy teorii zbiorów

Zbiór oraz relacja należenia do zbioru są w teorii zbiorów pojęciami pierwotnymi, tzn., takimi, których się nie definiuje.

Przedmioty, które należą do danego zbioru nazywa się jego *elementami*. Zbiory oznaczają się dużymi literami, zaś elementy zbioru małymi. Fakt, że a jest elementem zbioru A zapisuje się symbolicznie

$$a \in A.$$

i czyta się „ a należy do A ”. W przypadku występowania kilku elementów należących do tego samego zbioru często pisze się $a, b, c \in A$, zamiast $a \in A, b \in A, c \in A$. Symbol $a \notin A$ oznacza, że element a nie należy do zbioru A . Mamy

$$a \notin A \Leftrightarrow \sim (a \in A).$$

Sposoby określania zbiorów:

- zbiór można określić wprost poprzez podanie wszystkich jego elementów, np. $\{1, 2, 3\}$,
- zbiór można określić poprzez podanie kilku jego elementów, o ile na tej podstawie znana jest reguła klasyfikująca elementy tego zbioru, np. $\{1, 4, 7, 10, \dots\}$,
- zbiór można określić poprzez podanie własności, którą musi posiadać element, żeby należeć do zbioru, np. $\{x : x \in \mathbb{N} \wedge x < 7\}$. Jest to tzw. opis zbioru.

Uwaga. Kolejność elementów w definicji zbioru oraz powtarzanie się elementów nie mają znaczenia. Istotne jest tylko to, czy element należy do zbioru czy też nie. Np. $\{1, 2, 3\} = \{3, 1, 2\} = \{1, 3, 2, 1, 3\}$.

Zbiory liczbowe:

- zbiór liczb naturalnych

$$\mathbb{N} = \{1, 2, \dots\},$$

•

$$\mathbb{N}_0 = \{0, 1, 2, \dots\},$$

- zbiór liczb całkowitych

$$\mathbb{Z} = \{0, 1, -1, 2, -2, \dots\},$$

- zbiór liczb wymiernych

$$\mathbb{Q} = \left\{ \frac{p}{q} : p, q \in \mathbb{Z} \wedge q \neq 0 \right\},$$

- zbiór liczb rzeczywistych

$$\mathbb{R},$$

- zbiór liczb zespolonych

$$\mathbb{C}.$$

Definicja. *Zbiorem pustym* nazywa się zbiór, do którego nie należy żaden element. Zbiór ten oznacza się symbolem $\emptyset$.

Definicja. Dwa zbiory A i B są *równe* wtedy i tylko wtedy, gdy mają te same elementy, tzn.

$$A = B \Leftrightarrow \bigwedge_x (x \in A \Leftrightarrow x \in B).$$

Definicja. Zbiór A jest *podzbiorem* zbioru B wtedy i tylko wtedy, gdy każdy element zbioru A jest również elementem zbioru B . Mówi się też wtedy inaczej, że zbiór A *zawiera się* w zbiorze B . Oznaczenie: $A \subseteq B$. Zatem

$$A \subseteq B \Leftrightarrow \bigwedge_x (x \in A \Rightarrow x \in B).$$

Symbol $\subseteq$ nazywa się znakiem *inkluzji*. Jeżeli A nie jest podzbiorem zbioru B , to pisze się $A \not\subseteq B$. Jeżeli $A \subseteq B$ oraz $A \neq B$, to zbiór A nazywa się *podzbiór właściwy* zbioru B .

Definicja. Zbiór wszystkich podzbiorów zbioru A nazywa się *zbiór potęgowy* zbioru A i oznacza się przez $\mathcal{P}(A)$ lub 2^A .

Twierdzenie. (Własności inkluzji) Niech A, B, C będą dowolnymi zbiorami. Wtedy

$$1) \emptyset \subseteq A,$$

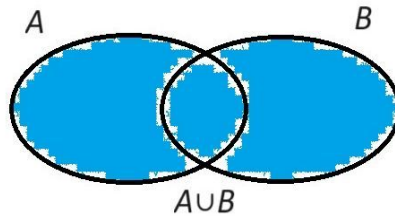
- 2) $A \subseteq A$,
- 3) $(A \subseteq B \wedge B \subseteq C) \Rightarrow A \subseteq C$,
- 4) $(A \subseteq B \wedge B \subseteq A) \Rightarrow A = B$,
- 5) $A \neq B \Rightarrow (A \not\subseteq B \vee B \not\subseteq A)$.

Dowód. Łatwy. $\square$

Definicja. Sumą zbiorów A i B nazywa się zbiór tych wszystkich elementów, które należą do zbioru A lub należą do zbioru B i żadnych innych, tzn.

$$x \in (A \cup B) \Leftrightarrow (x \in A \vee x \in B).$$

Sumę $A \cup B$ zbiorów A i B przedstawia następujący rysunek:



Twierdzenie. (Własności sumy zbiorów) Niech A, B, C, D będą dowolnymi zbiorami. Wtedy

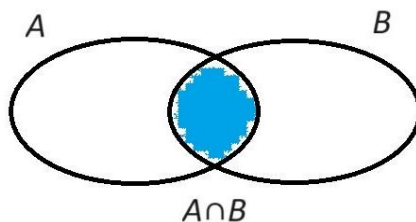
- 1) $A \cup B = B \cup A$,
- 2) $A \cup \emptyset = A$,
- 3) $A \cup A = A$,
- 4) $A \cup (B \cup C) = (A \cup B) \cup C$,
- 5) $A \subseteq A \cup B$,
- 6) $(A \subseteq C \wedge B \subseteq C) \Rightarrow A \cup B \subseteq C$,
- 7) $(A \subseteq B \wedge C \subseteq D) \Rightarrow A \cup C \subseteq B \cup D$,
- 8) $A \subseteq B \Leftrightarrow A \cup B = B$.

Dowód. Dowody punktów 1)–7) są łatwe. Udowodnimy punkt 8). Załóżmy, że $A \subseteq B$. Ponieważ jednocześnie $B \subseteq B$, więc z punktów 6) i 3) otrzymujemy $A \cup B \subseteq B$. Ale $B \subseteq A \cup B$ na mocy punktów 1) i 5). Na mocy punktu 4) poprzedniego twierdzenia dostajemy żadaną równość. Załóżmy teraz, że $A \cup B = B$. Stąd i z punktu 5) otrzymujemy $A \subseteq B$. $\square$

Definicja. Iloczynem zbiorów A i B nazywa się zbiór tych wszystkich elementów, które należą do zbioru A i należą do zbioru B i żadnych innych, tzn.

$$x \in (A \cap B) \Leftrightarrow (x \in A \wedge x \in B).$$

Iloczyn $A \cap B$ zbiorów A i B przedstawia następujący rysunek:



Definicja. Zbiory A i B są *rozłączne*, jeżeli nie mają elementów wspólnych, tzn. $A \cap B = \emptyset$.

Twierdzenie. (Własności iloczynu zbiorów) Niech A, B, C, D będą dowolnymi zbiorami. Wtedy

- 1) $A \cap B = B \cap A$,
- 2) $A \cap \emptyset = \emptyset$,
- 3) $A \cap A = A$,
- 4) $A \cap (B \cap C) = (A \cap B) \cap C$,
- 5) $A \cap B \subseteq A$,
- 6) $(A \subseteq B \wedge A \subseteq C) \Rightarrow A \subseteq B \cap C$,
- 7) $(A \subseteq B \wedge C \subseteq D) \Rightarrow A \cap C \subseteq B \cap D$,
- 8) $A \subseteq B \Leftrightarrow A \cap B = A$,
- 9) $A \cap (A \cup B) = A$,
- 10) $(A \cap B) \cup B = B$,
- 11) $A \cap (B \cup C) = (A \cap B) \cup (A \cap C)$,
- 12) $A \cup (B \cap C) = (A \cup B) \cap (A \cup C)$.

Dowód. Dowody punktów 1)–7) są łatwe.

8) Załóżmy, że $A \subseteq B$. Ponieważ jednocześnie $A \subseteq A$, więc $A \subseteq A \cap B$ na mocy punktów 3) i 6). Stąd i z punktu 5) dostajemy $A \cap B = A$. Załóżmy teraz, że $A \cap B = A$. Na mocy punktów 1) i 5) otrzymujemy $A \subseteq B$.

9) Na mocy punktu 8) równość 9) jest równoważna inkluzji $A \subseteq A \cup B$, która jest zawsze spełniona (zobacz własności sumy zbiorów).

10) Podobny.

11) Załóżmy, że $x \in A \cap (B \cup C)$. Zatem $x \in A$ i $x \in B \cup C$. Stąd $x \in A$ i x jest elementem co najmniej jednego ze zbiorów B, C . Jeżeli $x \in B$, to $x \in A \cap B$, czyli $x \in (A \cap B) \cup (A \cap C)$. Podobnie, jeżeli $x \in C$, to $x \in A \cap C$, czyli $x \in (A \cap B) \cup (A \cap C)$. W konsekwencji,

$$A \cap (B \cup C) \subseteq (A \cap B) \cup (A \cap C).$$

Założmy teraz, że $x \in (A \cap B) \cup (A \cap C)$. Wynika stąd, że x jest elementem co najmniej jednego ze składników tej sumy. Jeżeli $x \in A \cap B$, to $x \in A$ i $x \in B$. Stąd $x \in A$ i $x \in B \cup C$, a więc $x \in A \cap (B \cup C)$. Podobnie, jeżeli $x \in A \cap C$, to $x \in A$ i $x \in C$, skąd $x \in A$ i $x \in B \cup C$, czyli $x \in A \cap (B \cup C)$. Zatem

$$(A \cap B) \cup (A \cap C) \subseteq A \cap (B \cup C).$$

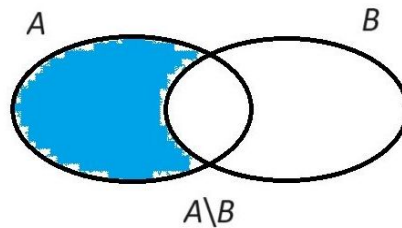
Z dwóch powyższych inkluzji otrzymujemy żadaną równość.

12) Podobny. $\square$

Definicja. *Różnicę* zbiorów A i B nazywa się zbiór tych wszystkich elementów, które należą do zbioru A i nie należą do zbioru B i żadnych innych, tzn.

$$x \in (A \setminus B) \Leftrightarrow (x \in A \wedge x \notin B).$$

Różnicę $A \setminus B$ zbiorów A i B przedstawia następujący rysunek:



Twierdzenie. (Własności różnicy zbiorów) Niech A, B, C, D będą dowolnymi zbiorami. Wtedy

- 1) $A \setminus B \subseteq A$,
- 2) $(A \subseteq B \wedge C \subseteq D) \Rightarrow A \setminus D \subseteq B \setminus C$,
- 3) $C \subseteq D \Rightarrow A \setminus D \subseteq A \setminus C$,
- 4) $A \subseteq B \Leftrightarrow A \setminus B = \emptyset$,
- 5) $A \setminus (B \cup C) = (A \setminus B) \cap (A \setminus C)$,
- 6) $A \setminus (B \cap C) = (A \setminus B) \cup (A \setminus C)$,
- 7) $A \cup (B \setminus A) = A \cup B$,
- 8) $A \subseteq B \Rightarrow A \cup (B \setminus A) = B$,
- 9) $A \setminus (A \setminus B) = A \cap B$,
- 10) $A \setminus (B \cup C) = (A \setminus B) \setminus C$.

Dowód. Dowody wszystkich punktów tego twierdzenia przeprowadza się podobnie. Udowodnijmy dla przykładu punkty 2) i 5).

2) Założmy, że $A \subseteq B$ i $C \subseteq D$. Jeśli $x \in A \setminus D$, to $x \in A$ i $x \notin D$. Stąd $x \in B$ i $x \notin C$. Zatem $x \in B \setminus C$.

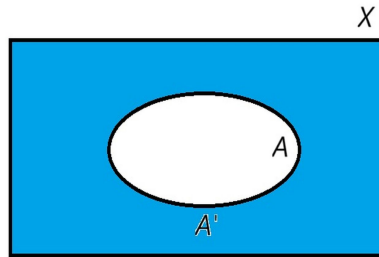
5) Zauważmy, że $x \in A \setminus (B \cup C)$ wtedy i tylko wtedy, gdy $x \in A$ i $x \notin B \cup C$. Ale $x \notin B \cup C$ wtedy i tylko wtedy, gdy $x \notin B$ i $x \notin C$. Warunek $x \in A$ i $x \notin B$ i $x \notin C$ jest równoważny warunkowi $x \in A \setminus B$ i $x \in A \setminus C$, który z kolei jest równoważny warunkowi $x \in (A \setminus B) \cap (A \setminus C)$ co dowodzi punktu 5). $\square$

Uwaga. W zastosowaniach teorii zbiorów rozważa się przeważnie tylko takie zbiory, które są podzbiorem pewnego ustalonego zbioru zwanego przestrzenią. Przestrzeń oznacza się na ogół symbolem X .

Definicja. Niech X będzie ustaloną przestrzenią. *Dopełnieniem* zbioru $A \subseteq X$ w przestrzeni X nazywa się zbiór $X \setminus A$ i oznacza przez A' , czyli $A' = X \setminus A$. Mamy

$$x \in A' \Leftrightarrow x \in (X \setminus A) \Leftrightarrow x \notin A.$$

Dopełnienie A' zbioru A w przestrzeni X przedstawia następujący rysunek:



Twierdzenie. (Własności przestrzeni i dopełnienia zbioru) Niech A, B będą dowolnymi podzbiorem przestrzeni X . Wtedy

- 1) $X \cap A = A$,
- 2) $X \cup A = X$,
- 3) $X' = \emptyset$,
- 4) $\emptyset' = X$,
- 5) $A \cup A' = X$,
- 6) $A \cap A' = \emptyset$,
- 7) $(A')' = A$,
- 8) $A \subseteq B \Leftrightarrow B' \subseteq A'$,
- 9) $(A \cup B)' = A' \cap B'$,
- 10) $(A \cap B)' = A' \cup B'$,
- 11) $A \setminus B = A \cap B'$,
- 12) $A \setminus B = (A' \cup B)'$,
- 13) $A \subseteq B \Leftrightarrow A \cap B' = \emptyset$,
- 14) $A \subseteq B \Leftrightarrow A' \cup B = X$.

Dowód. Dowody punktów 1)–8) są łatwe. Dla dowodu punktów 9) i 10) wystarczy we własnościach 5) i 6) różnicy zbiorów podstawić za A przestrzeń X , za B zbiór A i za C zbiór B i zastosować definicję dopełnienia zbioru.

11) Zauważmy, że $x \in A \setminus B$ wtedy i tylko wtedy, gdy $x \in A$ i $x \notin B$, czyli gdy $x \in A$ i $x \in B'$, co jest równoważne $x \in A \cap B'$.

12) Wynika z punktów 11), 7) i 9).

13) Wynika z własności 4) różnicy zbiorów oraz z punktu 11).

14) Wynika punktów 13), 10), 7) i 4). $\square$

Niech $X \neq \emptyset$ będzie dowolną przestrzenią i niech $\mathcal{P}(X)$ będzie rodziną wszystkich podzbiorów przestrzeni X . Niech $T \neq \emptyset$ będzie dowolnym zbiorem.

Definicja. Funkcja $f : T \rightarrow \mathcal{P}(X)$ nazywa się *indeksowana rodzina zbiorów* (lub inaczej *indeksowana rodzina podzbiorów przestrzeni X*).

Uwaga. Zamiast $f(t)$ pisze się A_t , gdzie $A_t \subseteq X$ dla każdego $t \in T$. Indeksowaną rodzinę zbiorów oznacza się symbolem $(A_t)_{t \in T}$.

Definicja. Niech $(A_t)_{t \in T}$ będzie dowolną indeksowaną rodziną podzbiorów przestrzeni X . Sumą zbiorów A_t , gdzie $t \in T$, nazywa się zbiór $\bigcup_{t \in T} A_t$, dla którego spełniony jest warunek

$$x \in \bigcup_{t \in T} A_t \Leftrightarrow \bigvee_{t \in T} x \in A_t.$$

Definicja. Niech $(A_t)_{t \in T}$ będzie dowolną indeksowaną rodziną podzbiorów przestrzeni X . Iloczynem zbiorów A_t , gdzie $t \in T$, nazywa się zbiór $\bigcap_{t \in T} A_t$, dla którego spełniony jest warunek

$$x \in \bigcap_{t \in T} A_t \Leftrightarrow \bigwedge_{t \in T} x \in A_t.$$

Uwaga.

1. Jeśli $T = \{1, 2, \dots, n\}$, to

$$\bigcup_{t \in T} A_t = A_1 \cup A_2 \cup \dots \cup A_n$$

oraz

$$\bigcap_{t \in T} A_t = A_1 \cap A_2 \cap \dots \cap A_n.$$

2. Jeśli $T = \mathbb{N}$, to zamiast $\bigcup_{t \in T} A_t$ pisze się $\bigcup_{n=1}^{\infty} A_n$ oraz zamiast $\bigcap_{t \in T} A_t$ pisze się $\bigcap_{n=1}^{\infty} A_n$.

Uwaga. Sumy i iloczyny indeksowanej rodziny zbiorów nazywa się też *sumy i iloczyny uogólnione zbiorów*.

Twierdzenie. (Własności sum i iloczynów uogólnionych) Niech A będzie dowolnym zbiorem oraz niech $(A_t)_{t \in T}$ i $(B_t)_{t \in T}$ będą dowolnymi indeksowanymi rodzinami zbiorów. Wtedy

- 1) $\bigwedge_{t \in T} \left(A_t \subseteq \bigcup_{t \in T} A_t \right),$
- 2) $\bigwedge_{t \in T} \left(\bigcap_{t \in T} A_t \subseteq A_t \right),$
- 3) $\left(\bigwedge_{t \in T} A_t \subseteq A \right) \Rightarrow \left(\bigcup_{t \in T} A_t \subseteq A \right),$
- 4) $\left(\bigwedge_{t \in T} A \subseteq A_t \right) \Rightarrow \left(A \subseteq \bigcap_{t \in T} A_t \right),$
- 5) $A \cup \bigcup_{t \in T} A_t = \bigcup_{t \in T} (A \cup A_t),$
- 6) $A \cap \bigcup_{t \in T} A_t = \bigcup_{t \in T} (A \cap A_t),$
- 7) $A \cup \bigcap_{t \in T} A_t = \bigcap_{t \in T} (A \cup A_t),$
- 8) $A \cap \bigcap_{t \in T} A_t = \bigcap_{t \in T} (A \cap A_t),$
- 9) $\left(\bigwedge_{t \in T} A_t \subseteq B_t \right) \Rightarrow \left(\bigcup_{t \in T} A_t \subseteq \bigcup_{t \in T} B_t \right),$
- 10) $\left(\bigwedge_{t \in T} A_t \subseteq B_t \right) \Rightarrow \left(\bigcap_{t \in T} A_t \subseteq \bigcap_{t \in T} B_t \right),$
- 11) $\bigcup_{t \in T} A_t \cup \bigcup_{t \in T} B_t = \bigcup_{t \in T} (A_t \cup B_t),$
- 12) $\bigcap_{t \in T} A_t \cap \bigcap_{t \in T} B_t = \bigcap_{t \in T} (A_t \cap B_t),$
- 13) $\bigcup_{t \in T} (A_t \cap B_t) \subseteq \bigcup_{t \in T} A_t \cap \bigcup_{t \in T} B_t,$
- 14) $\bigcap_{t \in T} A_t \cup \bigcap_{t \in T} B_t \subseteq \bigcap_{t \in T} (A_t \cup B_t),$
- 15) $A \setminus \bigcup_{t \in T} A_t = \bigcap_{t \in T} (A \setminus A_t),$
- 16) $A \setminus \bigcap_{t \in T} A_t = \bigcup_{t \in T} (A \setminus A_t).$

Dowód. Dowody wszystkich punktów tego twierdzenia przeprowadza się podobnie. Udowodnijmy dla przykładu punkt 11). Zauważmy, że $x \in \bigcup_{t \in T} A_t \cup \bigcup_{t \in T} B_t$ wtedy i tylko wtedy, gdy $x \in \bigcup_{t \in T} A_t$ lub $x \in \bigcup_{t \in T} B_t$, czyli gdy $x \in A_{t_1}$ dla pewnego $t_1 \in T$ lub $x \in B_{t_2}$ dla pewnego $t_2 \in T$, co jest równoważne $x \in \bigcup_{t \in T} (A_t \cup B_t)$. $\square$

Definicja. Parą uporządkowaną (x, y) nazywa się zbiór $\{x, \{x, y\}\}$. Element x nazywa się *poprzednik* pary (x, y) , zaś element y nazywa się *następnik* tej pary. Równość par

uporządkowanych określa się następująco:

$$(x, y) = (w, z) \Leftrightarrow (x = w \wedge y = z).$$

Definicja. *Iloczynem kartezjańskim* zbiorów X i Y nazywa się zbiór

$$X \times Y = \{(x, y) : x \in X \wedge y \in Y\}.$$

Uwaga. Iloczyn kartezjański nie jest przemienny, tzn. na ogół $X \times Y \neq Y \times X$.

Definicja. *Iloczynem kartezjańskim* zbiorów $X_1, X_2, \dots, X_n$ nazywa się zbiór

$$X_1 \times X_2 \times \dots \times X_n = \{(x_1, x_2, \dots, x_n) : x_1 \in X_1 \wedge x_2 \in X_2 \wedge \dots \wedge x_n \in X_n\}.$$

Jeśli $X_1 = X_2 = \dots = X_n$, to iloczyn kartezjański $X \times X \times \dots \times X$ nazywa się *n -tą potęgą kartezjańską* zbioru X i oznacza się przez X^n .

Przykład. $\mathbb{R} \times \mathbb{R} \times \dots \times \mathbb{R} = \mathbb{R}^n = \{(x_1, x_2, \dots, x_n) : x_1 \in \mathbb{R} \wedge x_2 \in \mathbb{R} \wedge \dots \wedge x_n \in \mathbb{R}\}.$

3. Relacje

Definicja. Relacją dwuczłonową $\mathcal{R}$ określoną w iloczynie kartezjańskim $X \times Y$ nazywa się dowolny podzbiór tego iloczynu, tzn.

$$\mathcal{R} \subseteq X \times Y.$$

Jeżeli $(x, y) \in \mathcal{R}$, to mówi się, że x jest w relacji $\mathcal{R}$ z y i pisze się $x\mathcal{R}y$.

Definicja. Niech $\mathcal{R} \subseteq X \times Y$. Zbiór

$$D(\mathcal{R}) = \left\{ x \in X : \bigvee_{y \in Y} x\mathcal{R}y \right\}$$

nazywa się *dziedzina* relacji $\mathcal{R}$, a zbiór

$$P(\mathcal{R}) = \left\{ y \in Y : \bigvee_{x \in X} x\mathcal{R}y \right\}$$

nazywa się *przeciwdziedzina* relacji $\mathcal{R}$. Zbiór $D(\mathcal{R}) \cup P(\mathcal{R})$ nazywa się *pole* relacji $\mathcal{R}$.

Definicja. Relacją odwrotną do relacji $\mathcal{R} \subseteq X \times Y$ nazywa się zbiór

$$\mathcal{R}^{-1} = \{(y, x) : y\mathcal{R}x\}.$$

Definicja. Niech $\mathcal{R} \subseteq D(\mathcal{R}) \times P(\mathcal{R})$ i $\mathcal{S} \subseteq D(\mathcal{S}) \times P(\mathcal{S})$, gdzie $D(\mathcal{S}) \subseteq P(\mathcal{R})$. Złożeniem relacji $\mathcal{R}$ i $\mathcal{S}$ nazywa się relację $\mathcal{S} \circ \mathcal{R} \subseteq D(\mathcal{R}) \times P(\mathcal{S})$ spełniającą warunek

$$x(\mathcal{S} \circ \mathcal{R})z \Leftrightarrow \bigvee_{y \in D(\mathcal{S})} x\mathcal{R}y \wedge y\mathcal{S}z.$$

Uwaga. Złożenie relacji nie jest przemienne, tzn. na ogół $\mathcal{S} \circ \mathcal{R} \neq \mathcal{R} \circ \mathcal{S}$.

Twierdzenie. Niech $\mathcal{R}_1 \subseteq D(\mathcal{R}_1) \times P(\mathcal{R}_1)$, $\mathcal{R}_2 \subseteq D(\mathcal{R}_2) \times P(\mathcal{R}_2)$ i $\mathcal{R}_3 \subseteq D(\mathcal{R}_3) \times P(\mathcal{R}_3)$, gdzie $D(\mathcal{R}_2) \subseteq P(\mathcal{R}_1)$ oraz $D(\mathcal{R}_3) \subseteq P(\mathcal{R}_2)$. Wtedy

$$\mathcal{R}_3 \circ (\mathcal{R}_2 \circ \mathcal{R}_1) = (\mathcal{R}_3 \circ \mathcal{R}_2) \circ \mathcal{R}_1.$$

Dowód. Załóżmy, że $x(\mathcal{R}_3 \circ (\mathcal{R}_2 \circ \mathcal{R}_1))z$. Istnieje y takie, że $x(\mathcal{R}_2 \circ \mathcal{R}_1)y$ i $y\mathcal{R}_3z$. Stąd istnieje t takie, że $x\mathcal{R}_1t$ i $t\mathcal{R}_2y$. Ponieważ $t\mathcal{R}_2y$ i $y\mathcal{R}_3z$, więc $t(\mathcal{R}_3 \circ \mathcal{R}_2)z$. Stąd i z $x\mathcal{R}_1t$

wynika, że $x((\mathcal{R}_3 \circ \mathcal{R}_2) \circ \mathcal{R}_1)z$. Zatem

$$\mathcal{R}_3 \circ (\mathcal{R}_2 \circ \mathcal{R}_1) \subseteq (\mathcal{R}_3 \circ \mathcal{R}_2) \circ \mathcal{R}_1.$$

Dowód inkluzji w drugą stronę jest analogiczny. Stąd otrzymujemy równość. $\square$

Twierdzenie. Niech $\mathcal{R} \subseteq D(\mathcal{R}) \times P(\mathcal{R})$ i $\mathcal{S} \subseteq D(\mathcal{S}) \times P(\mathcal{S})$, gdzie $D(\mathcal{S}) \subseteq P(\mathcal{R})$. Wtedy

$$(\mathcal{S} \circ \mathcal{R})^{-1} = \mathcal{R}^{-1} \circ \mathcal{S}^{-1}.$$

Dowód. Niech $x(\mathcal{S} \circ \mathcal{R})^{-1}z$. Stąd $z(\mathcal{S} \circ \mathcal{R})x$ oraz istnieje y takie, że $z\mathcal{R}y$ i $y\mathcal{S}x$. Zatem $x\mathcal{S}^{-1}y$ i $y\mathcal{R}^{-1}z$, czyli $x(\mathcal{R}^{-1} \circ \mathcal{S}^{-1})z$. Stąd

$$(\mathcal{S} \circ \mathcal{R})^{-1} \subseteq \mathcal{R}^{-1} \circ \mathcal{S}^{-1}.$$

Dowód inkluzji w drugą stronę jest analogiczny. Zatem dostajemy równość. $\square$

Najczęściej rozpatruje się relacje $\mathcal{R}$, w których $X = Y$, czyli $\mathcal{R} \subseteq X^2$. Mówi się wówczas, że relacja $\mathcal{R}$ jest określona w zbiorze X . Wśród wielu takich relacji szczególną rolę odgrywają relacje o następujących własnościach:

- *zwrotność*

$$\bigwedge_{x \in X} x\mathcal{R}x,$$

- *symetryczność*

$$\bigwedge_{x \in X} \bigwedge_{y \in X} (x\mathcal{R}y \Rightarrow y\mathcal{R}x),$$

- *antysymetryczność*

$$\bigwedge_{x \in X} \bigwedge_{y \in X} [(x\mathcal{R}y \wedge y\mathcal{R}x) \Rightarrow x = y],$$

- *przechodniość*

$$\bigwedge_{x \in X} \bigwedge_{y \in X} \bigwedge_{z \in X} [(x\mathcal{R}y \wedge y\mathcal{R}z) \Rightarrow x\mathcal{R}z],$$

- *spójność*

$$\bigwedge_{x \in X} \bigwedge_{y \in X} (x\mathcal{R}y \vee y\mathcal{R}x).$$

Twierdzenie. Relacja $\mathcal{R} \subseteq X^2$ jest symetryczna wtedy i tylko wtedy, gdy

$$\mathcal{R} \subseteq \mathcal{R}^{-1}.$$

Dowód. Załóżmy, że $\mathcal{R}$ jest symetryczna i niech $x\mathcal{R}y$. Z symetryczności mamy $y\mathcal{R}x$, czyli $x\mathcal{R}^{-1}y$. Stąd $\mathcal{R} \subseteq \mathcal{R}^{-1}$.

Założmy teraz, że $\mathcal{R} \subseteq \mathcal{R}^{-1}$ i niech $x\mathcal{R}y$. Stąd $x\mathcal{R}^{-1}y$, czyli $y\mathcal{R}x$. Zatem relacja $\mathcal{R}$ jest symetryczna. $\square$

Twierdzenie. Relacja $\mathcal{R} \subseteq X^2$ jest przechodnia wtedy i tylko wtedy, gdy

$$\mathcal{R} \circ \mathcal{R} \subseteq \mathcal{R}.$$

Dowód. Załóżmy, że $\mathcal{R}$ jest przechodnia i niech $x(\mathcal{R} \circ \mathcal{R})z$. Istnieje y takie, że $x\mathcal{R}y$ i $y\mathcal{R}z$. Z przechodniości mamy $x\mathcal{R}z$, czyli $\mathcal{R} \circ \mathcal{R} \subseteq \mathcal{R}$.

Założmy teraz, że $\mathcal{R} \circ \mathcal{R} \subseteq \mathcal{R}$ i niech $x\mathcal{R}y$ i $y\mathcal{R}z$. Wtedy $x(\mathcal{R} \circ \mathcal{R})z$. Stąd $x\mathcal{R}z$, czyli $\mathcal{R}$ jest przechodnia. $\square$

Definicja. Relacja $\mathcal{R} \subseteq X^2$ nazywa się *relacją równoważności*, jeśli jest zwrotna, symetryczna i przechodnia.

Definicja. Klasę równoważności elementu $a \in X$ względem relacji równoważności $\mathcal{R}$, oznaczaną $[a]_{\mathcal{R}}$, nazywa się zbiór wszystkich elementów zbioru X , które są w relacji $\mathcal{R}$ z elementem a . Zatem

$$[a]_{\mathcal{R}} = \{x \in X : x\mathcal{R}a\}.$$

Element a nazywa się *reprezentant* klasy równoważności $[a]_{\mathcal{R}}$.

Twierdzenie. Niech $\mathcal{R} \subseteq X^2$ będzie relacją równoważności. Wtedy

- 1) $\bigcup_{a \in X} [a]_{\mathcal{R}} = X$,
- 2) jeżeli $a\mathcal{R}b$, to $[a]_{\mathcal{R}} = [b]_{\mathcal{R}}$,
- 3) jeżeli $\neg(a\mathcal{R}b)$, to $[a]_{\mathcal{R}} \cap [b]_{\mathcal{R}} = \emptyset$.

Dowód. Łatwy. $\square$

Definicja. Zbiorem ilorazowym $X/\mathcal{R}$ relacji równoważności $\mathcal{R}$ nazywa się zbiór wszystkich klas równoważności relacji $\mathcal{R}$, tzn.

$$X/\mathcal{R} = \{[a]_{\mathcal{R}} : a \in X\}.$$

Definicja. Relacja $\mathcal{R} \subseteq X^2$ nazywa się *relacją częściowego porządku* w zbiorze X , jeśli jest zwrotna, antysymetryczna i przechodnia.

Definicja. Jeśli relacja częściowego porządku w zbiorze X jest dodatkowo spójna, to nazywa się *relacją liniowego porządku* w zbiorze X .

Definicja. Zbiorem częściowo uporządkowanym nazywa się parę $(X, \mathcal{R})$, gdzie $\mathcal{R}$ jest relacją częściowego porządku w zbiorze X .

Definicja. Zbiorem liniowo uporządkowanym nazywa się parę $(X, \mathcal{R})$, gdzie $\mathcal{R}$ jest relacją liniowego porządku w zbiorze X .

Definicja. Niech $(X, \mathcal{R})$ będzie zbiorem częściowo uporządkowanym. Element $a \in X$ nazywa się:

- *maksymalny* w zbiorze $(X, \mathcal{R})$, jeśli

$$\bigwedge_{x \in X} (a \mathcal{R} x \Rightarrow x = a),$$

- *minimalny* w zbiorze $(X, \mathcal{R})$, jeśli

$$\bigwedge_{x \in X} (x \mathcal{R} a \Rightarrow x = a),$$

- *największy* w zbiorze $(X, \mathcal{R})$, jeśli

$$\bigwedge_{x \in X} x \mathcal{R} a,$$

- *najmniejszy* w zbiorze $(X, \mathcal{R})$, jeśli

$$\bigwedge_{x \in X} a \mathcal{R} x.$$

Twierdzenie.

- 1) W zbiorze częściowo uporządkowanym istnieje co najwyżej jeden element największy i co najwyżej jeden element najmniejszy.
- 2) Element największy w zbiorze częściowo uporządkowanym jest jedynym elementem maksymalnym w tym zbiorze.
- 3) Element najmniejszy w zbiorze częściowo uporządkowanym jest jedynym elementem minimalnym w tym zbiorze.

Dowód. Niech $(X, \mathcal{R})$ będzie zbiorem częściowo uporządkowanym.

- 1) Niech $a \in X$ będzie dowolnym elementem. Jasne jest, że $x \mathcal{R} a$ nie musi zachodzić dla każdego $x \in X$, czyli element największy może nie istnieć w zbiorze X . Załóżmy, że w X

istnieje element największy a . Wtedy

$$\bigwedge_{x \in X} x \mathcal{R} a,$$

Niech $b \in X$ będzie również elementem największym w zbiorze X . Wtedy

$$\bigwedge_{x \in X} x \mathcal{R} b.$$

Stąd mamy $a \mathcal{R} b$ i $b \mathcal{R} a$. Z antysymetryczności relacji $\mathcal{R}$ wynika $a = b$. Zatem, jeżeli w X istnieje element największy, to istnieje taki element dokładnie jeden. Podobnie pokazuje się w przypadku elementu najmniejszego.

2) Załóżmy, że $a \in X$ jest elementem największym w zbiorze X . Wtedy

$$\bigwedge_{x \in X} x \mathcal{R} a. \quad (*)$$

Wiemy, że element $b \in X$ jest elementem maksymalnym w X , jeśli

$$\bigwedge_{x \in X} (b \mathcal{R} x \Rightarrow x = b). \quad (**)$$

Z $(*)$ mamy $b \mathcal{R} a$. Na mocy $(**)$ wynika stąd $a = b$.

3) Pokazuje się podobnie jak punkt 2). $\square$

Definicja. Łańcuchem w zbiorze częściowo uporządkowanym $(X, \mathcal{R})$ nazywa się każdy niepusty podzbiór $Y \subseteq X$ taki, że $(Y, \mathcal{R} \cap Y^2)$ jest zbiorem liniowo uporządkowanym.

Uwaga. Jeżeli $\mathcal{L}(X)$ jest zbiorem wszystkich łańcuchów w zbiorze częściowo uporządkowanym $(X, \mathcal{R})$, to $(\mathcal{L}(X), \subseteq)$ jest również zbiorem częściowo uporządkowanym.

Definicja. Łańcuchem maksymalnym w zbiorze częściowo uporządkowanym $(X, \mathcal{R})$ nazywa się element maksymalny w zbiorze $(\mathcal{L}(X), \subseteq)$, tzn. taki łańcuch, który nie jest właściwym podzbiorem żadnego innego łańcucha.

Definicja. Niech $(X, \mathcal{R})$ będzie zbiorem częściowo uporządkowanym i niech $A \subseteq X$. Element $g \in X$ nazywa się *ograniczenie górne* zbioru A , jeśli spełniony jest warunek

$$\bigwedge_{a \in A} a \mathcal{R} g.$$

Ponadto, jeśli zbiór ograniczeń górnych zbioru A ma element najmniejszy, to element ten nazywa się *kres górny* zbioru A i oznacza się przez $\sup(A)$.

Definicja. Niech $(X, \mathcal{R})$ będzie zbiorem częściowo uporządkowanym i niech $A \subseteq X$. Element $d \in X$ nazywa się *ograniczenie dolne* zbioru A , jeśli spełniony jest warunek

$$\bigwedge_{a \in A} d \mathcal{R} a.$$

Ponadto, jeśli zbiór ograniczeń dolnych zbioru A ma element największy, to element ten nazywa się *kres dolny* zbioru A i oznacza się przez $\inf(A)$.

4. Funkcje

Definicja. Relacja $f \subseteq X \times Y$ nazywa się *funkcją* odwzorowującą zbiór X w zbiór Y , jeśli

- 1) $\bigwedge_{x \in X} \bigvee_{y \in Y} xfy$,
- 2) $\bigwedge_{x \in X} \bigwedge_{y_1 \in Y} \bigwedge_{y_2 \in Y} (xfy_1 \wedge xfy_2 \Rightarrow y_1 = y_2)$.

Tradycyjnie, zamiast pisać xfy pisze się $y = f(x)$. Element x nazywa się *argument* funkcji f , zaś element y nazywa się *wartość funkcji* dla argumentu x .

Definicja. Zbiory $D(f) = X$ oraz $P(f) \subseteq Y$ nazywają się odpowiednio *dziedzina* i *przeciwdziedzina* funkcji f .

Uwaga. Funkcję odwzorowującą zbiór X w zbiór Y oznacza się symbolem $f : X \rightarrow Y$.

Definicja. Niech dana będzie funkcja $f : X \rightarrow Y$ oraz zbiór $A \subseteq X$. Zbiór

$$f(A) = \left\{ y \in Y : \bigvee_{x \in A} y = f(x) \right\} = \{ f(x) : x \in A \}$$

nazywa się *obraz* zbioru A poprzez funkcję f .

Definicja. Niech dana będzie funkcja $f : X \rightarrow Y$ oraz zbiór $B \subseteq Y$. Zbiór

$$f^{-1}(B) = \left\{ x \in X : \bigvee_{y \in B} f(x) = y \right\}$$

nazywa się *przeciwwobraz* zbioru B poprzez funkcję f .

Twierdzenie. Niech dana będzie funkcja $f : X \rightarrow Y$ oraz zbiory $A_1, A_2 \subseteq X$. Wtedy

- 1) $f(A_1 \cup A_2) = f(A_1) \cup f(A_2)$,
- 2) $f(A_1 \cap A_2) \subseteq f(A_1) \cap f(A_2)$,
- 3) $f(A_1) \setminus f(A_2) \subseteq f(A_1 \setminus A_2)$.

Dowód. 1) Jeśli $y \in f(A_1 \cup A_2)$, to istnieje $x \in A_1 \cup A_2$ takie, że $y = f(x)$. Ponieważ $x \in A_1 \cup A_2$, więc $x \in A_1$ lub $x \in A_2$. Stąd $y \in f(A_1)$ lub $y \in f(A_2)$, czyli $y \in f(A_1) \cup f(A_2)$. Zatem $f(A_1 \cup A_2) \subseteq f(A_1) \cup f(A_2)$. Dowód inkluzji w drugą stronę jest podobny.

2) Jeśli $y \in f(A_1 \cap A_2)$, to istnieje $x \in A_1 \cap A_2$ takie, że $y = f(x)$. Ponieważ $x \in A_1 \cap A_2$, więc $x \in A_1$ i $x \in A_2$. Stąd $y \in f(A_1)$ i $y \in f(A_2)$, czyli $y \in f(A_1) \cap f(A_2)$.

3) Załóżmy, że $y \in f(A_1) \setminus f(A_2)$. Wtedy $y \in f(A_1)$ i $y \notin f(A_2)$. Stąd wynika, że istnieje $x \in A_1$ takie, że $y = f(x)$. Zatem $x \in A_1 \setminus A_2$, czyli $y \in f(A_1 \setminus A_2)$. $\square$

Uwaga. W punktach 2) i 3) powyższego twierdzenia inkluzja nie może być zastąpiona równością, co pokazuje przykład funkcji $f : \mathbb{R} \rightarrow \mathbb{R}$ takiej, że $f(x) = x^2$ oraz $A_1 = [-1, 0]$, $A_2 = [0, 1]$. Wtedy

$$f(A_1 \cap A_2) = f(\{0\}) = \{0\} \neq f(A_1) \cap f(A_2) = [0, 1]$$

oraz

$$f(A_1) \setminus f(A_2) = \emptyset \neq f(A_1 \setminus A_2) = f([-1, 0)) = (0, 1].$$

Twierdzenie. Niech dana będzie funkcja $f : X \rightarrow Y$ oraz zbiory $B_1, B_2 \subseteq Y$. Wtedy

- 1) $f^{-1}(B_1 \cup B_2) = f^{-1}(B_1) \cup f^{-1}(B_2)$,
- 2) $f^{-1}(B_1 \cap B_2) \subseteq f^{-1}(B_1) \cap f^{-1}(B_2)$,
- 3) $f^{-1}(B_1 \setminus B_2) = f^{-1}(B_1) \setminus f^{-1}(B_2)$.

Dowód. Łatwy. $\square$

Definicja. Niech dane będą dwie funkcje $f : X \rightarrow Y$ i $g : Y \rightarrow Z$. *Złożeniem* funkcji f i g nazywa się funkcję $g \circ f : X \rightarrow Z$ taką, że

$$\bigwedge_{x \in X} (g \circ f)(x) \stackrel{df}{=} g(f(x)).$$

Definicja. Funkcja $f : X \rightarrow Y$ nazywa się *różnowartościowa*, jeśli

$$\bigwedge_{x_1 \in X} \bigwedge_{x_2 \in X} (f(x_1) = f(x_2) \Rightarrow x_1 = x_2).$$

Uwaga. Powyższy warunek jest równoważny warunkowi

$$\bigwedge_{x_1 \in X} \bigwedge_{x_2 \in X} (x_1 \neq x_2 \Rightarrow f(x_1) \neq f(x_2)).$$

Twierdzenie. Jeżeli relacja f jest funkcją różnowartościową, to relacja f^{-1} jest również funkcją różnowartościową.

Dowód. Łatwy. $\square$

Uwaga. Funkcja f^{-1} nazywa się *funkcją odwrotną* do funkcji f . Jeśli $f : X \rightarrow Y$ i $P(f) \subseteq Y$ jest przeciwdziedzina funkcji f , to $f^{-1} : P(f) \rightarrow X$.

Definicja. Funkcja $f : X \rightarrow Y$ nazywa się *funkcją zbioru X na zbiór Y* (krótko *funkcją „na”*), jeśli

$$\bigwedge_{y \in Y} \bigvee_{x \in X} y = f(x).$$

Uwaga. Jeśli funkcja $f : X \rightarrow Y$ jest „na”, to $P(f) = Y$. Dla oznaczenia funkcji „na” pisze się też czasem $f : X \xrightarrow{na} Y$.

Twierdzenie. Niech $f : X \rightarrow Y$ i $g : Y \rightarrow Z$ będą dowolnymi funkcjami. Wtedy

- 1) jeżeli f i g są funkcjami „na”, to $g \circ f : X \rightarrow Z$ jest „na”,
- 2) jeżeli f i g są różnowartościowe, to $g \circ f : X \rightarrow Z$ jest różnowartościowa.

Dowód. 1) Załóżmy, że funkcje f i g są „na”. Niech $z \in Z$. Istnieje wówczas $y \in Y$ takie, że $z = g(y)$. Jednocześnie istnieje $x \in X$ takie, że $y = f(x)$. Stąd $z = g(y) = g(f(x)) = (g \circ f)(x)$. Zatem funkcja $g \circ f$ jest „na”.

2) Załóżmy, że funkcje f i g są różnowartościowe. Jeżeli dla $x_1, x_2 \in X$ spełniony jest warunek $x_1 \neq x_2$, to $f(x_1) \neq f(x_2)$. Stąd $g(f(x_1)) \neq g(f(x_2))$, czyli $(g \circ f)(x_1) \neq (g \circ f)(x_2)$. Zatem $g \circ f$ jest funkcją różnowartościową. $\square$

Definicja. Funkcja nazywa się *bijekcją* (funkcja *wzajemnie jednoznaczna*), jeśli jest różnowartościowa i „na”.

Definicja. Niech $f : X \rightarrow Y$ będzie funkcją i niech $A \subseteq X$. *Obcięciem* funkcji f do zbioru A nazywa się funkcję $f|A$ taką, że

$$(f|A)(x) = f(x) \quad \text{dla } x \in A.$$

Definicja. Funkcję *identycznościową* (*identycznością*) na zbiorze X nazywa się funkcję $id_X : X \rightarrow X$ taką, że

$$id_X(x) = x \quad \text{dla każdego } x \in X.$$

Definicja. Funkcję $f : X \rightarrow Y$ nazywa się *funkcją stałą*, jeśli istnieje element $y_0 \in Y$ taki, że

$$f(x) = y_0 \quad \text{dla każdego } x \in X.$$

Uwaga. Funkcję nazywa się też inaczej *odwzorowanie* lub *przekształcenie*.

5. Teoria mocy zbiorów

Definicja. Zbiory X i Y nazywa się *równoliczne* wtedy i tylko wtedy, gdy istnieje bijekcja zbioru X na zbiór Y .

Przykłady.

1. Zbiór liczb naturalnych jest równoliczny ze zbiorem liczb naturalnych parzystych. Wystarczy zauważyć, że funkcja $f(n) = 2n$ jest różnowartościowa i odwzorowuje zbiór liczb naturalnych na zbiór liczb naturalnych parzystych.
2. Przedział $(0, 1)$ jest równoliczny z przedziałem $(1, +\infty)$. Funkcja $f(x) = \frac{1}{x}$ jest bijekcją przedziału $(0, 1)$ na przedział $(1, +\infty)$.
3. Przedział $(-\frac{\pi}{2}, \frac{\pi}{2})$ jest równoliczny ze zbiorem $\mathbb{R}$. Wystarczy zauważyć, że funkcja $f(x) = \operatorname{tg}(x)$ jest różnowartościowa i odwzorowuje przedział $(-\frac{\pi}{2}, \frac{\pi}{2})$ na zbiór $\mathbb{R}$.

Uwaga. Jeżeli zbiory X i Y są równoliczne, to pisze się $X \sim Y$.

Twierdzenie. Niech X, Y, Z będą dowolnymi zbiorami. Wtedy

- 1) $X \sim X$,
- 2) $X \sim Y \Rightarrow Y \sim X$,
- 3) $X \sim Y \wedge Y \sim Z \Rightarrow X \sim Z$.

Dowód. 1) Ponieważ odwzorowanie identycznościowe $id_X : X \rightarrow X$ jest różnowartościowe i „na”, więc $X \sim X$.

2) Załóżmy, że $X \sim Y$. Zatem istnieje funkcja $f : X \rightarrow Y$ różnowartościowa i „na”. Wtedy funkcja odwrotna $f^{-1} : Y \rightarrow X$ jest również różnowartościowa i „na”. Stąd $Y \sim X$.

3) Załóżmy, że $X \sim Y$ i $Y \sim Z$. Istnieją zatem funkcje $f : X \rightarrow Y$ i $g : Y \rightarrow Z$ różnowartościowe i „na”. Stąd funkcja $g \circ f : X \rightarrow Z$ jest różnowartościowa i „na”. Zatem $X \sim Z$. $\square$

Wniosek. Relacja równoliczności zbiorów jest relacją równoważności.

Definicja. Zbiorem *nieskończonym* nazywa się zbiór równoliczny ze swoim podzbiorem właściwym.

Wniosek. Zbiory $\mathbb{N}$ i $\mathbb{R}$ są nieskończone.

Definicja. Zbiorem *przeliczalnym* nazywa się zbiór skończony lub równoliczny ze zbiorem liczb naturalnych $\mathbb{N}$.

Wniosek. Zbiór liczb naturalnych parzystych jest przeliczalny.

Definicja. Zbiorem *nieprzeliczalnym* nazywa się zbiór, który nie jest przeliczalny.

Twierdzenie. (O zbiorze przeliczalnym) Zbiór $A \neq \emptyset$ jest przeliczalny wtedy i tylko wtedy, gdy jest on zbiorem wyrazów pewnego ciągu nieskończonego, czyli wtedy i tylko wtedy, gdy istnieje funkcja odwzorowująca zbiór $\mathbb{N}$ na zbiór A .

Dowód. Załóżmy, że $A \neq \emptyset$ jest zbiorem przeliczalnym. Jeżeli A jest zbiorem skończonym $\{a_1, \dots, a_m\}$, to jest zbiorem wyrazów ciągu nieskończonego $(b_n)_{n \in \mathbb{N}}$ określonego następująco: $b_n = a_n$ dla $n = 1, \dots, m$ i $b_n = a_m$ dla $n > m$. Jeżeli A jest zbiorem przeliczalnym nieskończonym, to jest równoliczny z $\mathbb{N}$. Z definicji równoliczności istnieje funkcja $f : \mathbb{N} \rightarrow A$ przekształcająca zbiór $\mathbb{N}$ na A , a zatem A jest zbiorem wyrazów ciągu $(b_n)_{n \in \mathbb{N}}$ takiego, że $b_n = f(n)$ dla $n \in \mathbb{N}$.

Założmy teraz, że A jest zbiorem wyrazów pewnego ciągu $(b_n)_{n \in \mathbb{N}}$. Jeżeli A jest zbiorem skończonym, to z definicji jest zbiorem przeliczalnym. Jeżeli A jest zbiorem nieskończonym, to przyjmując $f(1) = b_1$ i $f(k)$ równe pierwszemu wyrazowi ciągu $(b_n)_{n \in \mathbb{N}}$ różnemu od $f(k - i)$ dla $1 \leq i < k$ otrzymujemy funkcję $f : \mathbb{N} \rightarrow A$ różnowartościową przekształcającą $\mathbb{N}$ na A . Zatem zbiór A jest równoliczny z $\mathbb{N}$, więc A jest zbiorem przeliczalnym. $\square$

Twierdzenie. Podzbiór zbioru przeliczalnego jest zbiorem przeliczalnym.

Dowód. Załóżmy, że A jest zbiorem przeliczalnym i że $B \subseteq A$. Jeżeli B jest zbiorem skończonym, to z definicji jest zbiorem przeliczalnym. Jeżeli $B = A$, to oczywiście B jest zbiorem przeliczalnym. Załóżmy więc, że B jest zbiorem nieskończonym i $B \neq A$. Ponieważ A jest zbiorem przeliczalnym, więc istnieje funkcja $f : \mathbb{N} \rightarrow A$ różnowartościowa przekształcająca $\mathbb{N}$ na A . Niech $g : \mathbb{N} \rightarrow B$ będzie funkcją określoną następująco: $g(1) = f(k_1)$, gdzie k_1 jest najmniejszą liczbą naturalną taką, że $f(k_1) \in B$; $g(n) = f(k_n)$ dla $n > 1$, gdzie k_n jest najmniejszą liczbą naturalną taką, że $k_n > k_{n-1}$ i $f(k_n) \in B$. Funkcja g jest różnowartościowa i przekształca $\mathbb{N}$ na B . Zatem B jest zbiorem równolicznym z $\mathbb{N}$, a więc przeliczalnym nieskończonym. $\square$

Twierdzenie. Suma dwóch zbiorów przeliczalnych jest zbiorem przeliczalnym.

Dowód. Przypadek, gdy jeden z danych zbiorów przeliczalnych jest pusty jest oczywisty. Załóżmy więc, że oba zbiory są niepuste. Niech $f : \mathbb{N} \rightarrow A$ przekształca zbiór $\mathbb{N}$ na zbiór A i niech $g : \mathbb{N} \rightarrow B$ przekształca zbiór $\mathbb{N}$ na zbiór B . Zdefiniujmy funkcję $h : \mathbb{N} \rightarrow A \cup B$ następująco: $h(2n - 1) = f(n)$ i $h(2n) = g(n)$ dla $n \in \mathbb{N}$. Funkcja h przekształca zbiór $\mathbb{N}$ na zbiór $A \cup B$. Zatem, na mocy Twierdzenia o zbiorze przeliczalnym, zbiór $A \cup B$ jest przeliczalny. $\square$

Wniosek. Suma przeliczalnej ilości zbiorów przeliczalnych jest zbiorem przeliczalnym.

Uwaga. Dowód powyższego faktu dostaje się przez indukcję.

Twierdzenie. Iloczyn kartezjański dwóch zbiorów przeliczalnych jest zbiorem przeliczalnym.

Dowód. Jeżeli jeden ze zbiorów jest pusty, to iloczyn kartezjański jest zbiorem pustym. Możemy więc założyć, że oba zbiory są niepuste. Niech zbiory A i B będą przeliczalne. Na mocy Twierdzenia o zbiorze przeliczalnym istnieją funkcje „na” $f : \mathbb{N} \rightarrow A$ i $g : \mathbb{N} \rightarrow B$. Zauważmy, że wśród par uporządkowanych $(f(m), g(n))$, gdzie $m, n \in \mathbb{N}$, występują wszystkie elementy iloczynu $A \times B$. Ustawmy zbiór par uporządkowanych $(f(m), g(n))$, gdzie $m, n \in \mathbb{N}$ w następującą tablicę nieskończoną:

$$\begin{array}{ccccccc} (f(1), g(1)) & (f(1), g(2)) & (f(1), g(3)) & \dots & (f(1), g(n)) & \dots \\ (f(2), g(1)) & (f(2), g(2)) & (f(2), g(3)) & \dots & (f(2), g(n)) & \dots \\ (f(3), g(1)) & (f(3), g(2)) & (f(3), g(3)) & \dots & (f(3), g(n)) & \dots \\ \vdots & \vdots & \vdots & & \vdots & \end{array}$$

Zdefiniujmy teraz przekształcenie h zbioru $\mathbb{N}$ na zbiór wszystkich par $(f(m), g(n))$, gdzie $m, n \in \mathbb{N}$, następująco:

$$\begin{aligned} h(1) &= (f(1), g(1)) \\ h(2) &= (f(1), g(2)) \\ h(3) &= (f(2), g(1)) \\ h(4) &= (f(1), g(3)) \\ h(5) &= (f(2), g(2)) \\ h(6) &= (f(3), g(1)) \\ h(7) &= (f(1), g(4)) \\ &\vdots \end{aligned}$$

Widzimy, że każda para $(f(m), g(n))$ jest obrazem pewnej liczby naturalnej przy przekształceniu h . Stąd h odwzorowuje zbiór $\mathbb{N}$ na zbiór $A \times B$. Zatem, na mocy Twierdzenia o zbiorze przeliczalnym, zbiór $A \times B$ jest przeliczalny. $\square$

Twierdzenie. Zbiór liczb całkowitych $\mathbb{Z}$ jest przeliczalny.

Dowód. Funkcja $f(n) = -n$ przekształca wzajemnie jednoznacznie zbiór liczb naturalnych $\mathbb{N}$ na zbiór liczb całkowitych ujemnych $\mathbb{Z}_-$. Stąd zbiór $\mathbb{Z}_-$ jest przeliczalny oraz $\mathbb{Z} = \mathbb{N}_0 \cup \mathbb{Z}_-$ jest sumą dwu zbiorów przeliczalnych, więc jest zbiorem przeliczalnym. $\square$

Twierdzenie. Zbiór liczb wymiernych $\mathbb{Q}$ jest przeliczalny.

Dowód. Ponieważ $\mathbb{Q} = \mathbb{Z} \times \mathbb{N}$ jest iloczynem kartezjańskim dwu zbiorów przeliczalnych, więc jest zbiorem przeliczalnym. $\square$

Twierdzenie. Jeżeli zbiór A jest nieprzeliczalny i $A \subseteq B$, to B jest nieprzeliczalny.

Dowód. Wiemy, że podzbiór zbioru przeliczalnego jest przeliczalny. Gdyby więc B był przeliczalny, to i A musiałby być przeliczalny. $\square$

Twierdzenie. Przedział $[0, 1]$ jest zbiorem nieprzeliczalnym.

Dowód. Niech $(c_n)_{n \in \mathbb{N}}$ będzie dowolnym ciągiem spełniającym warunek

$$0 \leq c_n < 1 \text{ dla każdego } n \in \mathbb{N}.$$

Oznaczmy przez $[a_0, b_0]$ przedział $[0, 1]$. Wybierzmy spośród przedziałów $[0, \frac{1}{3}]$, $[\frac{1}{3}, \frac{2}{3}]$, $[\frac{2}{3}, 1]$ taki, do którego nie należy c_1 i oznaczmy go przez $[a_1, b_1]$. Spełnione są wówczas warunki

$$c_1 \notin [a_1, b_1], \quad b_1 - a_1 = \frac{1}{3} \quad \text{ i } \quad [a_1, b_1] \subseteq [0, 1] = [a_0, b_0].$$

Podobnie postępując z przedziałem $[a_1, b_1]$ wyznaczamy przedział $[a_2, b_2]$ spełniający warunki

$$c_2 \notin [a_2, b_2], \quad b_2 - a_2 = \frac{1}{3^2} \quad \text{ i } \quad [a_2, b_2] \subseteq [a_1, b_1].$$

Ogólnie, mając już określony przedział $[a_{n-1}, b_{n-1}]$, $n \geq 2$, taki, że

$$c_{n-1} \notin [a_{n-1}, b_{n-1}], \quad b_{n-1} - a_{n-1} = \frac{1}{3^{n-1}} \quad \text{ i } \quad [a_{n-1}, b_{n-1}] \subseteq [a_{n-2}, b_{n-2}]$$

wyznaczamy, kontynuując to postępowanie, przedział $[a_n, b_n]$ spełniający warunki

$$c_n \notin [a_n, b_n], \quad b_n - a_n = \frac{1}{3^n} \quad \text{ i } \quad [a_n, b_n] \subseteq [a_{n-1}, b_{n-1}]. \quad (*)$$

W ten sposób określamy ciąg przedziałów $([a_n, b_n])_{n \in \mathbb{N}}$ spełniający warunki $(*)$ dla każdego $n \geq 1$. Z $(*)$ wynika, że dla każdego $n \in \mathbb{N}$ zachodzą nierówności

$$0 \leq a_n \leq a_{n+1} < b_{n+1} \leq b_n \leq 1.$$

Ciągi $(a_n)_{n \in \mathbb{N}}$ i $(b_n)_{n \in \mathbb{N}}$ są więc monotoniczne i ograniczone, a zatem są zbieżne. Jednocześnie, na mocy $(*)$, $b_n - a_n = \frac{1}{3^n}$ dla każdego $n \in \mathbb{N}$. Stąd $\lim_{n \rightarrow \infty} (b_n - a_n) = 0$, czyli $\lim_{n \rightarrow \infty} a_n = \lim_{n \rightarrow \infty} b_n = c$, gdzie c jest liczbą rzeczywistą z przedziału $[0, 1]$. Liczba c należy do każdego z przedziałów $[a_n, b_n]$; jest więc różna od każdego z c_n , ponieważ, na mocy $(*)$, $c_n \notin [a_n, b_n]$ dla każdego $n \in \mathbb{N}$. Ponieważ ciąg $(c_n)_{n \in \mathbb{N}}$ jest dowolny, więc przedział $[0, 1]$ nie jest zbiorem wyrazów żadnego ciągu nieskończonego. Zatem, na mocy Twierdzenia o zbiorze przeliczalnym, nie jest on przeliczalny, czyli jest nieprzeliczalny. $\square$

Na mocy dwóch ostatnich twierdzeń otrzymujemy następujący wniosek.

Wniosek. Zbiór liczb rzeczywistych $\mathbb{R}$ jest nieprzeliczalny.

Przypomnijmy, że relacja $\sim$ równoliczności zbiorów jest relacją równoważności. Stąd mamy następującą definicję.

Definicja. Każdej klasie równoważności względem relacji $\sim$ równoliczności zbiorów przyporząkowuje się pewien obiekt nazywany *liczba kardynalna* i oznaczany przez $\mathfrak{m}, \mathfrak{n}, \dots$. Zatem zbiorom równolicznym przyporządkowuje się tę samą liczbę kardynalną. Nazywa się ją *moc zbioru* i oznacza przez $|X|$ lub $\overline{X}$.

Uwaga. Moc zbioru skończonego jest równa liczbie jego elementów.

Dla zbiorów nieskończonych wprowadza się nowe liczby:

- $\aleph_0$ (alef zero) – oznacza moc zbioru liczb naturalnych,
- $\mathfrak{c}$ (continuum) – oznacza moc zbioru liczb rzeczywistych.

Uwaga. Wiemy już, że przedział $(-\frac{\pi}{2}, \frac{\pi}{2})$ jest równoliczny ze zbiorem $\mathbb{R}$. Stąd mamy następujący wniosek.

Wniosek. Przedział otwarty $(-\frac{\pi}{2}, \frac{\pi}{2})$ jest mocy continuum.

Twierdzenie. Każdy przedział otwarty (a, b) , gdzie $a < b$, jest mocy continuum.

Dowód. Funkcja $f : (-\frac{\pi}{2}, \frac{\pi}{2}) \rightarrow (a, b)$ dana wzorem

$$f(x) = \frac{b-a}{\pi} \left(x + \frac{\pi}{2} \right) + a$$

jest funkcją liniową, więc różnowartościową i „na”. Ustala więc równoliczność tych przedziałów. Zatem, na mocy powyższego wniosku, przedział (a, b) , gdzie $a < b$, jest mocy continuum. $\square$

Wśród liczb kardynalnych wprowadza się następującą relację:

$$\mathfrak{m} \preceq \mathfrak{n} \Leftrightarrow \text{istnieją zbiory } X \text{ i } Y \text{ takie, że } |X| = \mathfrak{m}, |Y| = \mathfrak{n} \text{ i } X \subseteq Y.$$

Relacja $\preceq$ jest relacją częściowego porządku w zbiorze liczb kardynalnych. Warunek antysymetrii relacji $\preceq$ nosi nazwę Twierdzenia Cantora–Bernsteina. Ponadto mamy

$$\mathfrak{m} \prec \mathfrak{n} \Leftrightarrow \mathfrak{m} \preceq \mathfrak{n} \wedge \mathfrak{m} \neq \mathfrak{n}.$$

Twierdzenie. (Cantor)

$$\aleph_0 \prec \mathfrak{c}.$$

(bez dowodu)

Twierdzenie. (Cantor)

$$|X| \prec |\mathcal{P}(X)|.$$

(bez dowodu)

Wniosek. Istnieje więc nieskończenie wiele liczb kardynalnych.

Konstrukcja zbioru liczb naturalnych $\mathbb{N}_0$

Przyjmując teorię mnogości za teorię pierwotną można liczby naturalne zdefiniować w następujący sposób:

$$0 = |\emptyset|$$

$$1 = |\{\emptyset\}|$$

$$2 = |\{\emptyset, \{\emptyset\}\}|$$

$$3 = \left| \left\{ \emptyset, \{\emptyset\}, \{\emptyset, \{\emptyset\}\} \right\} \right|$$

itd.

Wówczas suma liczb naturalnych odpowiada mocy sumy odpowiednich zbiorów, zaś iloczyn – mocy iloczynu kartezjańskiego tych zbiorów.

6. Indukcja i rekurencja

Zbiór liczb naturalnych można też zdefiniować na drodze aksjomatycznej. Elementarna arytmetyka liczb naturalnych używa języka, w którym oprócz stałej 0 występuje jednoargumentowa funkcja s , nazywana *następnik*, i relacja równości.

Aksjomaty Peano liczb naturalnych $\mathbb{N}_0$:

A1. $0 \in \mathbb{N}_0$.

A2. Dla każdego $n \in \mathbb{N}_0$ istnieje dokładnie jedno $s(n) \in \mathbb{N}_0$ takie, że $s(n)$ jest następnikiem n .

A3. Zero nie jest następnikiem żadnej liczby należącej do $\mathbb{N}_0$.

A4. Jeżeli $k = s(n)$ i $k = s(m)$, to $n = m$.

A5. Zasada indukcji matematycznej: Jeżeli $A \subseteq \mathbb{N}_0$ i

1) $0 \in A$,

2) dla każdego $n \in \mathbb{N}_0$, jeśli $n \in A$, to $s(n) \in A$,

to $A = \mathbb{N}_0$.

Uwaga. Zamiast symbolu $s(n)$ pisze się często $n + 1$.

Uwaga. Aksjomat **A2** definiuje funkcję jednoargumentową s określoną na $\mathbb{N}_0$ i nazywaną *funkcją następnika*. Aksjomat **A3** wyróżnia 0 jako taki element zbioru $\mathbb{N}_0$, który nie jest wartością funkcji s . Aksjomat **A4** zapewnia, że funkcja s jest różnowartościowa, tzn. różnym liczbom naturalnym przyporządkowuje różne następniki. Wreszcie aksjomat **A5** mówi jak można zbudować zbiór liczb naturalnych z zera przez sukcesywne zastosowanie funkcji następnika: każda liczba naturalna n jest otrzymana z zera przez n -krotne wykonanie operacji następnika,

$$1 \stackrel{\text{df}}{=} s(0), \quad 2 \stackrel{\text{df}}{=} s(s(0)), \quad 3 \stackrel{\text{df}}{=} s(s(s(0))), \dots$$

Przykład. Pokażemy, że każda liczba postaci $n^5 - n$ dla $n \in \mathbb{N}_0$ jest podzielna przez 5. Niech

$$A = \{n \in \mathbb{N}_0 : n^5 - n \text{ jest podzielna przez } 5\}$$

Udowodnimy indukcyjnie, że $A = \mathbb{N}_0$. Mamy

1) $0 \in A$, bo $0^5 - 0 = 0$ i jest podzielna przez 5.

2) Niech $n \in A$. Stąd istnieje $k \in \mathbb{N}_0$ takie, że $n^5 - n = 5k$. Wtedy $(n+1)^5 - (n+1)$ też jest podzielna przez 5, bo

$$\begin{aligned}(n+1)^5 - (n+1) &= n^5 + 5n^4 + 10n^3 + 10n^2 + 5n + 1 - n - 1 \\&= (n^5 - n) + 5(n^4 + 2n^3 + 2n^2 + n) \\&= 5(k + n^4 + 2n^3 + 2n^2 + n).\end{aligned}$$

Stąd $n+1 \in A$. Ponieważ oba założenia zasady indukcji matematycznej zostały spełnione, więc możemy wywnioskować, że $A = \mathbb{N}_0$. Oznacza to, że dla dowolnej liczby naturalnej n liczba $n^5 - n$ jest podzielna przez 5.

Twierdzenie. (Zasada minimum) W każdym niepustym zbiorze liczb naturalnych istnieje liczba najmniejsza.

Dowód. Niech $A \neq \emptyset$ i $A \subseteq \mathbb{N}_0$. Załóżmy, że A nie ma liczby najmniejszej. Stąd w szczególności $0 \notin A$. Rozważmy zbiór

$$B = \left\{ n \in \mathbb{N}_0 : \bigwedge_{m \leq n} m \notin A \right\}.$$

Wykorzystując zasadę indukcji matematycznej pokażemy, że wtedy $B = \mathbb{N}_0$. Ponieważ $0 \notin A$, więc z definicji zbioru B wynika, że $0 \in B$. Załóżmy, że $n \in B$. Wtedy $m \notin A$ dla każdego $m \leq n$. Gdyby więc $n+1 \in A$, to byłby to element najmniejszy w A , co nie jest możliwe. Zatem $n+1 \notin A$, a stąd $n+1 \in B$. Ponieważ wykazaliśmy, że oba założenia zasady indukcji są spełnione, więc $B = \mathbb{N}_0$. Stąd $A = \emptyset$ i dostajemy sprzeczność. Zatem w A istnieje liczba najmniejsza. $\square$

Podamy teraz inne sformułowanie zasady indukcji matematycznej.

Definicja. (Zasada indukcji matematycznej) Niech W będzie pewną własnością liczb naturalnych, np. tych które należą do $A \subseteq \mathbb{N}_0$. Zapis $W(n)$ oznacza wtedy, że liczba $n \in \mathbb{N}_0$ posiada własność W . Jeżeli

- 1) $W(0)$, tzn. 0 ma własność W ,
- 2) dla każdego $n \in \mathbb{N}_0$, jeśli $W(n)$, to $W(n+1)$,

to dla każdego n , $W(n)$, tzn. każda liczba naturalna ma własność W .

Twierdzenie. Liczba wszystkich podzbiorów zbioru n -elementowego równa się 2^n , czyli dla dowolnego zbioru X , jeśli $|X| = n$, to $|\mathcal{P}(X)| = 2^n$.

Dowód. Oznaczmy przez W zdanie wyrażające własność liczb naturalnych taką, że

$$W(n) \Leftrightarrow \text{liczba podzbiorów zbioru } n\text{-elementowego równa się } 2^n.$$

Ponieważ zbiór pusty ma dokładnie jeden podzbiór, więc jeśli $|X| = 0$, to $|\mathcal{P}(X)| = 1 = 2^0$. Wynika stąd, że liczba 0 ma własność W . Załóżmy, że wszystkie zbiory k -elementowe mają własność W , tzn. liczba wszystkich podzbiorów zbioru k -elementowego równa się 2^k . Pokażemy, że zbiór $(k+1)$ -elementowy ma też własność W . Weźmy zbiór $X = \{x_1, \dots, x_k, x_{k+1}\}$. Rozpatrzmy najpierw podzbiory zbioru X , w których nie występuje element x_{k+1} . Widzimy, że są to wszystkie podzbiory zbioru k -elementowego, więc jest ich 2^k na mocy założenia indukcyjnego. Dalej zauważmy, że podzbiorów zbioru X , w których występuje element x_{k+1} jest tyle ile podzbiorów zbioru k -elementowego $X \setminus \{x_{k+1}\}$ (do każdego z tych podzbiorów dołączamy następnie element x_{k+1}), czyli również 2^k . Ponieważ podzbiory zbioru X albo zawierają element x_{k+1} albo go nie zawierają, więc wszystkich podzbiorów zbioru X jest $2^k + 2^k = 2^{k+1}$. Stąd własność W jest prawdziwa dla $k+1$. Zatem, na mocy zasady indukcji, zdanie $W(n)$ jest prawdziwe dla każdego $n \in \mathbb{N}_0$. $\square$

Czasami zdarza się, że zdanie określające własność liczb naturalnych jest prawdziwe dla wszystkich liczb naturalnych poczynając od pewnej ustalonej liczby. Zdarza się też, że chcemy udowodnić twierdzenie dotyczące liczb całkowitych większych od pewnej ustalonej liczby m . Mamy wtedy następującą postać zasady indukcji matematycznej.

Twierdzenie. Niech $m \in \mathbb{Z}$ i niech $\alpha(n)$ będzie zdaniem określonym na zbiorze $\{n \in \mathbb{Z} : n \geq m\}$. Jeśli

- 1) zdanie $\alpha(m)$ jest prawdziwe,
- 2) jeśli wszystkie zdania $\alpha(m), \alpha(m+1), \dots, \alpha(k-1)$ dla pewnego $k > m$ są prawdziwe, to $\alpha(k)$ też jest zdaniem prawdziwym,

to $\alpha(n)$ jest zdaniem prawdziwym dla dowolnych liczb całkowitych $n \geq m$.

Dowód. Niech własność $W(n)$ oznacza zdanie $\alpha(m+n)$. Z warunku 1) spełniona jest własność $W(0)$. Załóżmy, że spełnione są własności $W(0), \dots, W(k-1)$, czyli zdania $\alpha(m), \alpha(m+1), \dots, \alpha(m+k-1)$ są prawdziwe. Zatem, z warunku 2), zdanie $\alpha(m+k)$ jest prawdziwe. Stąd własność $W(k)$ jest spełniona. Pokazaliśmy, że jeśli własności $W(0), \dots, W(k-1)$ są spełnione dla pewnego k , to własność $W(k)$ jest spełniona. Zatem, na mocy zasady indukcji, wszystkie własności $W(n)$ są spełnione, czyli zdania $\alpha(n)$ są prawdziwe dla dowolnego $n \geq m$. $\square$

Twierdzenie. (O definiowaniu przez indukcję) Załóżmy, że A jest dowolnym niepustym zbiorem. Jeśli $\varphi : \mathbb{N}_0^r \rightarrow A$ oraz $\psi : \mathbb{N}_0 \times \mathbb{N}_0^r \times A \rightarrow A$ są dowolnymi funkcjami, to istnieje dokładnie jedna funkcja $f : \mathbb{N}_0 \times \mathbb{N}_0^r \rightarrow A$ o tej własności, że dla dowolnych liczb naturalnych $m_1, m_2, \dots, m_r$ zachodzą warunki

$$f(0, m_1, m_2, \dots, m_r) = \varphi(m_1, m_2, \dots, m_r)$$

i

$$f(s(n), m_1, m_2, \dots, m_r) = \psi(n, m_1, m_2, \dots, m_r, f(n, m_1, m_2, \dots, m_r)).$$

(bez dowodu)

Uwaga. Pierwszy z powyższych warunków nazywa się *początkowy*, a drugi *rekurencyjny*. Przedstawiony w twierdzeniu schemat definiowania funkcji nazywa się też *definiowanie rekurencyjne*, a taką definicję funkcji – *definicja rekurencyjna funkcji*.

Wniosek. Jeżeli A jest dowolnym niepustym zbiorem, a $\varphi : \mathbb{N}_0 \rightarrow A$ i $\psi : \mathbb{N}_0 \times \mathbb{N}_0 \times A \rightarrow A$ są dowolnymi funkcjami, to istnieje dokładnie jedna funkcja $f : \mathbb{N}_0 \times \mathbb{N}_0 \rightarrow A$ taka, że dla dowolnych $m, n \in \mathbb{N}_0$ zachodzą

$$f(0, m) = \varphi(m)$$

oraz

$$f(s(n), m) = \psi(n, m, f(n, m)).$$

Niekiedy korzysta się ze schematu definiowania przez indukcję w jeszcze prostszej postaci.

Wniosek. Załóżmy, że A jest dowolnym niepustym zbiorem. Niech $a \in A$ i niech dana będzie funkcja $\gamma : A \rightarrow A$. Wtedy istnieje dokładnie jedna funkcja $f : \mathbb{N}_0 \rightarrow A$ taka, że

$$f(0) = a$$

i

$$f(s(n)) = \gamma(f(n)).$$

Korzystając z pierwszego z powyższych wniosków można w zbiorze liczb naturalnych $\mathbb{N}_0$ zdefiniować działania dodawania i mnożenia.

Twierdzenie. Istnieje funkcja $+$: $\mathbb{N}_0 \times \mathbb{N}_0 \rightarrow \mathbb{N}_0$ o tej własności, że dla dowolnych $m, n \in \mathbb{N}_0$ zachodzą

$$+(0, m) = m$$

i

$$+(s(n), m) = s(+(n, m)).$$

Funkcja ta jest wyznaczona jednoznacznie.

Dowód. Należy zastosować pierwszy z powyższych wniosków do funkcji $\varphi : \mathbb{N}_0 \rightarrow \mathbb{N}_0$ takiej, że $\varphi(m) = m$ oraz do funkcji $\psi : \mathbb{N}_0 \times \mathbb{N}_0 \times \mathbb{N}_0 \rightarrow \mathbb{N}_0$ takiej, że $\psi(j, k, l) = s(l)$. Istnieje dokładnie jedna funkcja $+: \mathbb{N}_0 \times \mathbb{N}_0 \rightarrow \mathbb{N}_0$ taka, że

$$+(0, m) = \varphi(m) = m$$

i

$$+(s(n), m) = \psi(n, m, +(n, m)) = s(+(n, m))$$

dla dowolnych $m, n \in \mathbb{N}_0$. $\square$

Uwaga. Zamiast symbolu $+(n, m)$ pisze się $n + m$. Liczba $n + m$ nazywa się *suma* liczb n i m .

Twierdzenie. Istnieje dokładnie jedna funkcja $\cdot : \mathbb{N}_0 \times \mathbb{N}_0 \rightarrow \mathbb{N}_0$ taka, że dla dowolnych $m, n \in \mathbb{N}_0$ zachodzą

$$\cdot(0, m) = 0$$

i

$$\cdot(s(n), m) = \cdot(n, m) + m.$$

Dowód. Łatwy. $\square$

Uwaga. Zamiast symbolu $\cdot(n, m)$ pisze się $n \cdot m$ (lub nm). Liczba nm nazywa się *iloczyn* liczb n i m .

Między innymi rekurencyjnie (przez indukcję) definiuje się funkcje $f : \mathbb{N}_0 \rightarrow \mathbb{C}$, czyli ciągi. Mamy mianowicie.

Definicja. Ciąg jest zdefiniowany rekurencyjnie, jeśli

- 1) określony jest pewien skończony zbiór wyrazów ciągu,
- 2) pozostałe wyrazy ciągu są zdefiniowane za pomocą poprzednich wyrazów ciągu.

Przykłady.

1. Rekurencyjna definicja silni:

$$\begin{aligned} 0! &= 1, \\ (n+1)! &= n! \cdot (n+1) \quad \text{dla } n \in \mathbb{N}_0. \end{aligned}$$

2. Rekurencyjna definicja potęgi ($a \in \mathbb{R} \setminus \{0\}$):

$$\begin{aligned}a^0 &= 1, \\ a^{n+1} &= a^n \cdot a \quad \text{dla } n \in \mathbb{N}_0.\end{aligned}$$

3. Rekurencyjna definicja ciągu Fibonacciego:

$$\begin{aligned}f_0 &= 1, \\ f_1 &= 1, \\ f_n &= f_{n-1} + f_{n-2} \quad \text{dla } n \in \mathbb{N}_0, \quad n \geq 2.\end{aligned}$$

Dla ciągów zdefiniowanych rekurencyjnie wyznacza się wzór jawny w następujący sposób. Dany jest ciąg rekurencyjny:

$$\begin{cases} s_0, s_1 - \text{dane}, \\ s_n = as_{n-1} + bs_{n-2}, \quad \text{gdzie } n \in \mathbb{N}_0, \quad n \geq 2, \end{cases}$$

natomiast $a, b \in \mathbb{R} \setminus \{0\}$ są stałymi.

Definicja. Równanie $x^2 - ax - b = 0$ nazywa się *równanie charakterystyczne* powyższej zależności rekurencyjnej.

Twierdzenie. 1) Jeżeli równanie charakterystyczne ma dwa pierwiastki (rzeczywiste lub zespolone) z_1 i z_2 , to

$$s_n = c_1(z_1)^n + c_2(z_2)^n$$

dla pewnych stałych c_1 i c_2 .

2) Jeżeli równanie charakterystyczne ma jeden pierwiastek r , to

$$s_n = c_1 r^n + c_2 n r^n$$

dla pewnych stałych rzeczywistych c_1 i c_2 .

(bez dowodu)

7. Elementy kombinatoryki

Twierdzenie. (Prawo dodawania) Jeżeli $A_1, A_2, \dots, A_n$ są dowolnymi zbiorami takimi, że $A_i \cap A_j = \emptyset$ dla $i \neq j$ oraz $i, j = 1, 2, \dots, n$, to

$$\left| \bigcup_{i=1}^n A_i \right| = \sum_{i=1}^n |A_i|.$$

Dowód. Przeprowadzimy indukcyjnie względem n . Dla $n = 2$ mamy dwa zbiory A_1, A_2 . Niech $|A_1| = n$ i $|A_2| = m$. Istnieją bijekcje $f_1 : A_1 \rightarrow \{1, 2, \dots, n\}$ oraz $f_2 : A_2 \rightarrow \{1, 2, \dots, m\}$. Wtedy bijekcją jest również odwzorowanie $f : A_1 \cup A_2 \rightarrow \{1, 2, \dots, n + m\}$ takie, że

$$f(x) = \begin{cases} f_1(x) & \text{dla } x \in A_1, \\ f_2(x) + n & \text{dla } x \in A_2. \end{cases}$$

Stąd

$$|A_1 \cup A_2| = n + m = |A_1| + |A_2|.$$

Założmy teraz, że równość jest prawdziwa dla pewnego n , tzn.

$$\left| \bigcup_{i=1}^n A_i \right| = \sum_{i=1}^n |A_i|.$$

Wówczas dla $n + 1$ zbiorów $A_1, A_2, \dots, A_n, A_{n+1}$ mamy

$$\begin{aligned} \left| \bigcup_{i=1}^{n+1} A_i \right| &= \left| \bigcup_{i=1}^n A_i \cup A_{n+1} \right| \\ &= \left| \bigcup_{i=1}^n A_i \right| + |A_{n+1}| \\ &= \sum_{i=1}^n |A_i| + |A_{n+1}| \\ &= \sum_{i=1}^{n+1} |A_i|. \end{aligned}$$

Zatem żądana równość zachodzi dla dowolnego n . $\square$

Twierdzenie. Niech A i B będą zbiorami skończonymi. Wówczas

$$|A \times B| = |A| \cdot |B|.$$

Dowód. Niech $A = \{x_1, x_2, \dots, x_n\}$ i $B = \{y_1, y_2, \dots, y_m\}$. Wtedy $|A| = n$ i $|B| = m$. Zastosujemy indukcję względem m . Dla $m = 1$ mamy $B = \{y_1\}$ oraz

$$A \times B = \{(x_1, y_1), (x_2, y_1), \dots, (x_n, y_1)\},$$

czyli

$$|A \times B| = n = n \cdot 1 = |A| \cdot |B|.$$

Założmy teraz, że równość zachodzi dla pewnego m . Wtedy dla zbioru $(m+1)$ -elementowego $B = \{y_1, y_2, \dots, y_{m+1}\}$ zachodzi równość

$$A \times B = A \times B_1 \cup A \times B_2,$$

gdzie $B_1 = \{y_1, y_2, \dots, y_m\}$ i $B_2 = \{y_{m+1}\}$. Ponieważ zbiory $A \times B_1$ i $A \times B_2$ są rozłączne, więc na mocy prawa dodawania mamy

$$|A \times B| = |A \times B_1| + |A \times B_2| = |A| \cdot |B_1| + |A \times B_2| = n \cdot m + n = n \cdot (m + 1) = |A| \cdot |B|.$$

Zatem, na mocy zasady indukcji, równość zachodzi dla dowolnego m . $\square$

Twierdzenie. (Zasada mnożenia) Niech $A_1, A_2, \dots, A_n$ będą zbiorami skończonymi. Wówczas liczba ciągów $(a_1, a_2, \dots, a_n)$, gdzie $a_i \in A_i$ dla $i = 1, 2, \dots, n$, wynosi

$$|A_1| \cdot |A_2| \cdots |A_n|.$$

Dowód. Łatwy. Korzystamy z zasady indukcji i poprzedniego twierdzenia. $\square$

Uwaga. Powyższa zasada pozwala na „przeliczanie” wielu fundamentalnych struktur pojawiających się w kombinatoryce.

Przykład. Ciągami binarnymi nazywa się ciągi złożone z zer i jedynek. Długością ciągu binarnego nazywa się liczbę jego elementów. Stosując zasadę mnożenia łatwo widać, że

istnieje 2^k ciągów binarnych długości k . Ogólniej, istnieje n^k ciągów długości k na zbiorze n symboli (wystarczy w zasadzie mnożenia przyjąć $A_1 = A_2 = \dots = A_n = \{0, 1, 2, \dots, n-1\}$).

Twierdzenie. (Ogólna zasada mnożenia) Jeżeli pewna procedura może być rozbita na k kolejnych kroków, z r_i wynikami w i -tym kroku dla $i = 1, 2, \dots, k$, to w całej procedurze jest

$$r_1 \cdot r_2 \cdots r_k$$

łączych wyników (rozumianych jako uporządkowane ciągi wyników cząstkowych).

(bez dowodu)

Twierdzenie. (Zasada bijekcji) Niech A i B będą zbiorami skończonymi. Jeśli istnieje bijekcja $f : A \rightarrow B$, to

$$|A| = |B|.$$

Dowód. Oczywisty. $\square$

Uwaga. Sztuka w posługiwaniu się zasadą bijekcji w konkretnych sytuacjach polega na tym, żeby dostrzec bijekcję pomiędzy interesującym nas zbiorem a zbiorem, którego liczbę elementów już znamy.

Twierdzenie. Zbiór n -elementowy ma 2^n podzbiorów.

Dowód. Niech $X = \{x_1, x_2, \dots, x_n\}$ i niech $\mathcal{B}_n$ będzie zbiorem ciągów binarnych długości n . Określmy funkcję $f : \mathcal{P}(X) \rightarrow \mathcal{B}_n$ następująco:

$$f(S) = b_1 b_2 \dots b_n,$$

gdzie $S \subseteq X$ oraz

$$b_i = 1 \Leftrightarrow x_i \in S.$$

Jasne jest, że f jest bijekcją. Zatem

$$|\mathcal{P}(X)| = |\mathcal{B}_n| = 2^n. \quad \square$$

Definicja. *Permutacją n -wyrazową zbioru n -elementowego X nazywa się każdy n -wyrazowy ciąg różnych elementów wybranych ze zbioru X .*

Na mocy zasady mnożenia łatwo otrzymuje się następujące twierdzenie.

Twierdzenie. Permutacji n -wyrazowych zbioru n -elementowego jest

$$P_n = n!.$$

Definicja. Niech $X = \{x_1, x_2, \dots, x_k\}$ będzie zbiorem k różnych elementów. *Permutacją n -wyrazową z powtórzeniami*, w której x_i powtarza się n_i razy dla $i = 1, 2, \dots, k$, przy czym $n_1 + n_2 + \dots + n_k = n$, nazywa się każdy n -wyrazowy ciąg, w którym poszczególne elementy zbioru X powtarzają się wskazaną liczbę razy.

Twierdzenie. Permutacji n -wyrazowych z powtórzeniami jest

$$P_{n_1, n_2, \dots, n_k}^n = \frac{n!}{n_1! \cdot n_2! \cdot \dots \cdot n_k!}.$$

Dowód. Weźmy zbiór $X = \{x_1, x_2, \dots, x_k\}$. Rozpatrzmy n -wyrazowe ciągi elementów ze zbioru X , w których element x_i powtarza się n_i razy dla $i = 1, 2, \dots, k$, przy czym $n_1 + n_2 + \dots + n_k = n$. Gdyby wszystkie wyrazy takich ciągów były rozróżnialne (gdybyśmy potrafili jakoś rozróżniać między sobą powtarzające się wyrazy), to liczba takich ciągów wynosiłaby $n!$. Z drugiej strony można obliczyć tę liczbę w następujący sposób. Wszystkich takich ciągów mamy $P_{n_1, n_2, \dots, n_k}^n$ (tak oznaczyliśmy liczbę permutacji n -wyrazowych z powtórzeniami). Potem powodujemy, że powtarzające się wyrazy stają się rozróżnialne (np. numerujemy je). Wtedy takich ciągów mamy $n_1! \cdot n_2! \cdot \dots \cdot n_k!$ (bo elementy i -tego rodzaju w momencie gdy stają się rozróżnialne mogą zostać ustawione w ciąg na $n_i!$ sposobów). Zatem

$$n! = P_{n_1, n_2, \dots, n_k}^n \cdot n_1! \cdot n_2! \cdot \dots \cdot n_k!,$$

czyli

$$P_{n_1, n_2, \dots, n_k}^n = \frac{n!}{n_1! \cdot n_2! \cdot \dots \cdot n_k!}. \quad \square$$

Przykład. Policzmy liczbę różnych wyrazów z sensem lub bez sensu, które można utworzyć ze wszystkich liter występujących w słowie MATEMATYKA. Widzimy, że $n = 10$ oraz litera A występuje 3 razy, litery M, T – po 2 razy i litery E, Y, K – po razie. Stąd

$$P_{3,2,2,1,1,1}^{10} = \frac{10!}{3! \cdot 2! \cdot 2! \cdot 1! \cdot 1! \cdot 1!} = 151\,200.$$

Definicja. *Wariacją k -wyrazową bez powtórzeń zbioru n -elementowego X , gdzie $0 \leq k \leq n$, nazywa się każdy k -wyrazowy ciąg różnych elementów ze zbioru X .*

Na mocy ogólnej zasady mnożenia otrzymuje się następujące twierdzenie.

Twierdzenie. Wariacji k -wyrazowych bez powtórzeń zbioru n -elementowego, gdzie $0 \leq k \leq n$, jest

$$V_n^k = \frac{n!}{(n-k)!}.$$

Definicja. *Wariacją k -wyrazową z powtórzeniami zbioru n -elementowego X nazywa się każdy k -wyrazowy ciąg elementów ze zbioru X .*

Na mocy zasady mnożenia otrzymuje się następujące twierdzenie.

Twierdzenie. Wariacji k -wyrazowych z powtórzeniami zbioru n -elementowego jest

$$\overline{V}_n^k = n^k.$$

Definicja. *Kombinacją k -elementową bez powtórzeń zbioru n -elementowego X , gdzie $0 \leq k \leq n$, nazywa się każdy k -elementowy podzbiór zbioru X .*

Twierdzenie. Kombinacji k -elementowych bez powtórzeń zbioru n -elementowego, gdzie $0 \leq k \leq n$, jest

$$C_n^k = \frac{n!}{k!(n-k)!} \stackrel{\text{ozn.}}{=} \binom{n}{k}.$$

Dowód. Ponieważ $V_n^k = C_n^k P_k$ (bo każdy ciąg k -wyrazowy można otrzymać wybierając podzbiór i ustawiając jego elementy w ciąg), więc

$$C_n^k = \frac{V_n^k}{P_k} = \frac{n!}{(n-k)! \cdot k!} = \frac{n!}{k!(n-k)!} = \binom{n}{k}. \quad \square$$

Uwaga. Symbol $\binom{n}{k}$ nosi nazwę *współczynnik dwumianowy* lub *symbol Newtona*.

Wybrane tożsamości kombinatoryczne:

1. $\binom{n}{k} = \binom{n}{n-k}$,
2. $\binom{n}{k} = \binom{n-1}{k} + \binom{n-1}{k-1}$,
3. $k\binom{n}{k} = n\binom{n-1}{k-1}$,
4. $\binom{m+n}{k} = \sum_{l=0}^k \binom{m}{l} \binom{n}{k-l}$,
5. $(x+y)^n = \sum_{k=0}^n \binom{n}{k} x^k y^{n-k}$ dla $x, y \in \mathbb{R}$.

Definicja. Niech X będzie dowolnym zbiorem i $k \in \mathbb{N}$. *Multizbiorem k -elementowym* o elementach ze zbioru X nazywa się dowolny zbiór postaci

$$\{(x_1, r_1), (x_2, r_2), \dots, (x_m, r_m)\},$$

gdzie $m \in \mathbb{N}$, $x_1, x_2, \dots, x_m$ są parami różnymi elementami zbioru X , $r_1, r_2, \dots, r_m \in \mathbb{N}$ i $r_1 + r_2 + \dots + r_m = k$. Wtedy x_j jest r_j -krotnym elementem tego multizbioru dla $j = 1, 2, \dots, m$. Krotność elementu $x \in X \setminus \{x_1, x_2, \dots, x_m\}$ w tym multizbiorze wynosi 0.

Uwaga. Multizbiór k -elementowy $\{(x_1, r_1), (x_2, r_2), \dots, (x_m, r_m)\}$ oznacza się jako $\{c_1, c_2, \dots, c_k\}$, gdzie $(c_i)_{i=1}^k$ jest dowolnym ciągiem, w którym wyraz x_j występuje dokładnie r_j razy dla $j = 1, 2, \dots, m$. Na przykład każdy z zapisów $\{1, 1, 1, 2, 2\}$, $\{2, 2, 1, 1, 1\}$, $\{1, 2, 1, 2, 1\}$ oznacza ten sam multizbiór $\{(1, 3), (2, 2)\}$.

Uwaga. Multizbiorem 0-elementowym o elementach ze zbioru X jest (z definicji) tylko zbiór pusty.

Definicja. *Kombinacją k -elementową z powtórzeniami* zbioru n -elementowego X nazywa się każdy k -elementowy multizbiór o elementach z X .

Twierdzenie. Kombinacji k -elementowych z powtórzeniami zbioru n -elementowego jest

$$\overline{C}_n^k = \binom{n+k-1}{k}.$$

Dowód. Rozważmy zbiór $\{1, 2, \dots, n\}$. Każdej jego k -elementowej kombinacji z powtórzeniami możemy w sposób wzajemnie jednoznaczny przypisać następujące ustawienie w rzędzie $n-1$ kresek i k kropek:

r_1 kropek, kreska, r_2 kropek, kreska, $\dots$, r_{n-1} kropek, kreska, r_n kropek,

gdzie r_i to krotność elementu i -tego w tej kombinacji dla $i = 1, 2, \dots, n$. Z drugiej strony każdemu takiemu ustawieniu $n - 1$ kresek i k kropek możemy bijektywnie przypisać k -elementową kombinację bez powtórzeń zbioru $\{1, 2, \dots, n + k - 1\}$ w ten sposób, że elementami kombinacji będą „numery pozycji”, na których są kropki w danym ustawieniu. Oznacza to, że k -elementowych kombinacji z powtórzeniami zbioru $\{1, 2, \dots, n\}$ jest tyle samo, co k -elementowych kombinacji bez powtórzeń zbioru $\{1, 2, \dots, n + k - 1\}$, a zatem

$$\overline{C}_n^k = \binom{n+k-1}{k}. \quad \square$$

Uwaga. Dla przykładu dla $n = 4$ i $k = 2$ mamy

2-elementowe kombinacje z powtórzeniami zbioru $\{1, 2, 3, 4\}$	ustawienia 3 kresek i 2 kropek	2-elementowe kombinacje bez powtórzeń zbioru $\{1, 2, 3, 4, 5\}$
$\{1, 1\}$	• •	$\{1, 2\}$
$\{1, 2\}$	• •	$\{1, 3\}$
$\{1, 3\}$	• •	$\{1, 4\}$
$\{1, 4\}$	• •	$\{1, 5\}$
$\{2, 2\}$	• •	$\{2, 3\}$
$\{2, 3\}$	• •	$\{2, 4\}$
$\{2, 4\}$	• •	$\{2, 5\}$
$\{3, 3\}$	• •	$\{3, 4\}$
$\{3, 4\}$	• •	$\{3, 5\}$
$\{4, 4\}$	• •	$\{4, 5\}$

Przykład. Obliczymy liczbę rozwiązań w zbiorze $\mathbb{N}_0$ równania $x_1 + x_2 + \dots + x_n = k$, gdzie $k \in \mathbb{N}_0$ i $n \in \mathbb{N}$. Równanie to ma $\overline{C}_n^k = \binom{n+k-1}{k}$ rozwiązań w zbiorze $\mathbb{N}_0$. Istotnie, każdemu takiemu rozwiązaniu $(x_1, x_2, \dots, x_n)$ można wzajemnie jednoznacznie przyporządkować k -elementową kombinację z powtórzeniami o elementach ze zbioru $\{1, 2, \dots, n\}$ taką, że i jest jej x_i -krotnym elementem dla $i = 1, 2, \dots, n$.

Przykład. Mamy 4 puszki z farbami różnych kolorów ponumerowane cyframi 1, 2, 3, 4. Malujemy 6 kul po prostu wrzucając je do tych puszek. Obliczymy liczbę sposobów pomalowania tych kul. Mamy tutaj 6-elementowe multizbiory (czyli kombinacje z powtórzeniami) o elementach ze zbioru 4-elementowego $\{1, 2, 3, 4\}$. Liczba tych multizbiorów (czyli sposobów pomalowania tych kul) jest równa

$$\overline{C}_4^6 = \binom{4+6-1}{6} = 84.$$

Twierdzenie. (Zasada włączania-wyłączania) Niech $A_1, A_2, \dots, A_n$ będą zbiorami skończonymi. Wtedy

$$\begin{aligned} \left| \bigcup_{i=1}^n A_i \right| &= \sum_{i=1}^n |A_i| - \sum_{1 \leq i < j \leq n} |A_i \cap A_j| + \\ &+ \sum_{1 \leq i < j < k \leq n} |A_i \cap A_j \cap A_k| - \dots + \\ &+ (-1)^{n+1} \cdot |A_1 \cap A_2 \cap \dots \cap A_n|. \end{aligned}$$

Dowód. Zastosujemy indukcję względem n . Niech $n = 2$. Ponieważ

$$A_1 \cup A_2 = [A_1 \setminus (A_1 \cap A_2)] \cup [A_2 \setminus (A_1 \cap A_2)] \cup [A_1 \cap A_2],$$

więc

$$\begin{aligned} |A_1 \cup A_2| &= |A_1 \setminus (A_1 \cap A_2)| + |A_2 \setminus (A_1 \cap A_2)| + |A_1 \cap A_2| \\ &= |A_1| - |A_1 \cap A_2| + |A_2| - |A_1 \cap A_2| + |A_1 \cap A_2| \\ &= |A_1| + |A_2| - |A_1 \cap A_2|. \end{aligned}$$

Wykorzystaliśmy tutaj prawo dodawania, bo zbiory $A_1 \setminus (A_1 \cap A_2)$, $A_2 \setminus (A_1 \cap A_2)$ oraz $A_1 \cap A_2$ są parami rozłączne. Ponadto $A_1 \cap A_2 \subseteq A_1$ i $A_1 \cap A_2 \subseteq A_2$, więc

$$|A_i \setminus (A_1 \cap A_2)| = |A_i| - |A_1 \cap A_2| \quad \text{dla } i = 1, 2.$$

Założmy teraz, że wzór zachodzi dla $n \in \mathbb{N}$. Wtedy

$$\begin{aligned}
\left| \bigcup_{i=1}^{n+1} A_i \right| &= \left| \bigcup_{i=1}^n A_i \right| + |A_{n+1}| - \left| A_{n+1} \cap \bigcup_{i=1}^n A_i \right| \\
&= \sum_{i=1}^n |A_i| - \sum_{1 \leq i < j \leq n} |A_i \cap A_j| + \dots + \\
&\quad + (-1)^{n+1} \cdot |A_1 \cap A_2 \cap \dots \cap A_n| + |A_{n+1}| - \\
&\quad - \left[\sum_{i=1}^n |A_i \cap A_{n+1}| - \sum_{1 \leq i < j \leq n} |A_i \cap A_j \cap A_{n+1}| + \dots + \right. \\
&\quad \left. + (-1)^{n+1} \cdot |A_1 \cap A_2 \cap \dots \cap A_n \cap A_{n+1}| \right] \\
&= \sum_{i=1}^{n+1} |A_i| - \sum_{1 \leq i < j \leq n+1} |A_i \cap A_j| + \dots + \\
&\quad + (-1)^{n+1} \cdot |A_1 \cap A_2 \cap \dots \cap A_{n+1}|. \quad \square
\end{aligned}$$

Uwaga. Szczególne przypadki:

1. Dla $n = 2$ mamy

$$|A_1 \cup A_2| = |A_1| + |A_2| - |A_1 \cap A_2|,$$

czyli dodając do siebie liczby elementów dwóch zbiorów dwukrotnie liczymy część wspólną, więc trzeba ją odjąć.

2. Dla $n = 3$ mamy

$$|A_1 \cup A_2 \cup A_3| = |A_1| + |A_2| + |A_3| - |A_1 \cap A_2| - |A_1 \cap A_3| - |A_2 \cap A_3| + |A_1 \cap A_2 \cap A_3|,$$

czyli teraz odejmując trzykrotnie części wspólne par zbiorów odejmujemy o jeden raz za dużo część wspólną wszystkich trzech zbiorów, więc trzeba ją dodać.

Przykład. W pewnym klubie jest 10 osób grających w szachy i 15 osób grających w brydża, przy czym 6 spośród nich gra w obie te gry. Policzmy liczbę członków tego klubu. Oznaczmy przez A zbiór szachistów i przez B zbiór brydżystów. Szukamy $|A \cup B|$. Mamy

$$|A \cup B| = |A| + |B| - |A \cap B| = 10 + 15 - 6 = 19.$$

Zadanie. W pewnym klubie jest 10 osób grających w szachy, 15 osób grających w brydża i 12 osób grających w pokera. Spośród nich 5 gra w szachy i brydża, 4 w brydża i pokera, 3 w szachy i pokera, a tylko 2 grają we wszystkie te gry. Ile osób jest w tym klubie?

Permutację n -wyrazową zbioru n -elementowego X można też zdefiniować jako dowolną bijekcję $\sigma : X \rightarrow X$. Niech $X = \{x_1, x_2, \dots, x_n\}$. Zamiast pisać $\sigma : X \rightarrow X$ można też pisać

$$\begin{pmatrix} x_1 & x_2 & \dots & x_n \\ \sigma(x_1) & \sigma(x_2) & \dots & \sigma(x_n) \end{pmatrix}.$$

Definicja. Niech X będzie zbiorem n -elementowym. Element $x \in X$ nazywa się *punkt stały* permutacji $\sigma : X \rightarrow X$, jeśli

$$\sigma(x) = x.$$

Definicja. Niech X będzie zbiorem n -elementowym. Permutacja $\sigma : X \rightarrow X$ nazywa się *nieporządek*, jeśli nie posiada punktów stałych.

Przykład. Niech $X = \{1, 2, 3, 4, 5\}$. Nieporządkiem jest na przykład permutacja

$$\begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 5 & 4 & 2 & 3 & 1 \end{pmatrix}.$$

Twierdzenie. Liczba nieporządków D_n za zbiorze n -elementowym X dana jest wzorem

$$D_n = n! \sum_{k=2}^n \frac{(-1)^k}{k!}.$$

Dowód. Zauważmy najpierw, że nie istnieją nieporządki na zbiorze jednoelementowym X . Załóżmy więc, że $n > 1$. Niech $X = \{x_1, x_2, \dots, x_n\}$. Dla $k = 1, 2, \dots, n$ oznaczmy przez A_k zbiór permutacji, dla których x_k jest punktem stałym. Wtedy zbiór $\bigcup_{k=1}^n A_k$ zawiera permutacje, które mają co najmniej jeden punkt stały. Ponieważ

$$\begin{aligned}
|A_k| &= (n-1)!, \quad k = 1, 2, \dots, n, \\
|A_k \cap A_l| &= (n-2)!, \quad 1 \leq k < l \leq n, \\
|A_k \cap A_l \cap A_m| &= (n-3)!, \quad 1 \leq k < l < m \leq n, \\
&\vdots
\end{aligned}$$

więc, na mocy zasady włączania-wyłączania, mamy

$$\left| \bigcup_{k=1}^n A_k \right| = \binom{n}{1}(n-1)! - \binom{n}{2}(n-2)! + \binom{n}{3}(n-3)! - \dots + (-1)^{n-1} \binom{n}{n}.$$

Zatem

$$\begin{aligned}
D_n &= n! - \binom{n}{1}(n-1)! + \binom{n}{2}(n-2)! - \binom{n}{3}(n-3)! + \dots - (-1)^{n-1} \binom{n}{n} \\
&= n! - \frac{n!}{1!(\cancel{n-1})!} \cancel{(n-1)!} + \frac{n!}{2!(\cancel{n-2})!} \cancel{(n-2)!} - \frac{n!}{3!(\cancel{n-3})!} \cancel{(n-3)!} + \\
&\quad \dots - (-1)^{n-1} \frac{\cancel{n!}}{\cancel{n!} \cdot 0!} \\
&= \cancel{n!} - \cancel{n!} + \frac{n!}{2!} - \frac{n!}{3!} + \dots + (-1)^n \\
&= n! \left(\frac{1}{2!} - \frac{1}{3!} + \dots + \frac{(-1)^n}{n!} \right) \\
&= n! \sum_{k=2}^n \frac{(-1)^k}{k!}. \quad \square
\end{aligned}$$

Uwaga. Z powyższego wzoru wynika, że przy $n \rightarrow \infty$ nieporządki stanowią $e^{-1} = 0,36788\dots$ wszystkich permutacji (bo $e^x = \sum_{k=0}^{\infty} \frac{x^k}{k!}$).

Twierdzenie. (Zasada szufladkowa) Jeżeli skończony zbiór X jest podzielony na n podzbiorów, to co najmniej jeden z podzbiorów zawiera nie mniej niż $\frac{|X|}{n}$ elementów.

Dowód. Gdyby każdy z podzbiorów, na które podzielony jest zbiór X miał mniej niż $\frac{|X|}{n}$ elementów, to suma tych podzbiorów, czyli X zawierałaby mniej niż $|X|$ elementów. $\square$

Uwaga. Z powyższego twierdzenia wynika na przykład, że jeżeli rozmieści się $n+1$ przedmiotów w n szufladkach, to co najmniej jedna szufladka będzie zawierać więcej niż jeden element.

Przykład. Pewne 9 osób waży razem 810 kg. Sprawdźmy, czy dowolna trójka tych osób

może wsiąść do windy o udźwigu 250 kg. Oznaczmy te 9 osób przez $O_1, O_2, \dots, O_9$. Rozważmy tablicę:

$$\begin{array}{ccccccc} O_1 & O_2 & O_3 & \dots & O_7 & O_8 & O_9 \\ O_2 & O_3 & O_4 & \dots & O_8 & O_9 & O_1 \\ O_3 & O_4 & O_5 & \dots & O_9 & O_1 & O_2 \end{array}$$

Suma wag w każdym wierszu tablicy wynosi 810 kg, czyli razem w trzech wierszach mamy 2430 kg. Kolumny tej tablicy tworzą 9 różnych trójek. Na mocy zasady szufladkowej, przynajmniej w jednej kolumnie suma wag wynosi co najmniej $\frac{2430}{9}$ kg = 270 kg. Czyli istnieje (co najmniej jedna) taka trójka osób, które nie mogą razem wsiąść do windy.

Następne twierdzenie jest nieco inną postacią zasady szufladkowej.

Twierdzenie. Niech dana będzie funkcja $f : X \rightarrow Y$ oraz niech $|X| > k \cdot |Y|$. Wówczas co najmniej jeden ze zbiorów $f^{-1}(y)$ ma więcej niż k elementów.

Dowód. Zbiór X jest więc podzielony na n podzbiorów $f^{-1}(y)$, gdzie $n \leq |Y|$. Na mocy zasady szufladkowej, pewien zbiór $f^{-1}(y)$ ma co najmniej $\frac{|X|}{n}$ elementów. Ponieważ

$$\frac{|X|}{n} \geq \frac{|X|}{|Y|} > k,$$

więc zbiór taki ma więcej niż k elementów. $\square$

Przykład. Pewna uczelnia zatrudnia 5 profesorów reprezentujących 4 różne specjalności. Przewiduje się, że na koniec roku akademickiego 40 studentów tej uczelni będzie chciało bronić swoich prac magisterskich. W każdej komisji egzaminacyjnej, zgodnie z przyjętymi ustaleniami, powinno brać udział 3 profesorów reprezentujących różne specjalności. Pokażemy, że przynajmniej jedna z komisji będzie musiała uczestniczyć w co najmniej 6 egzaminach. Dwóch spośród 5 profesorów reprezentuje tę samą specjalność. Oznaczmy ich przez A i B . Wszystkie komisje trzy osobowe możemy podzielić na 3 kategorie:

- 1) takie, w których nie uczestniczą ani A , ani B ,
- 2) takie, w których uczestniczy A , a nie uczestniczy B ,

oraz

- 3) takie, w których uczestniczy B , a nie uczestniczy A .

Łatwo widać, że jest tylko jedna komisja pierwszego rodzaju oraz 3 komisje drugiego i 3 komisje trzeciego rodzaju (wybieramy dwóch pozostałych członków komisji spośród 3 profesorów, czyli $\binom{3}{2}$ komisji). Możliwych komisji jest więc 7. Określmy funkcję f , która każdemu magistrantowi przyporządkowuje komisję egzaminacyjną. Ponieważ $40 > 5 \cdot 7$,

więc, na mocy zasady szufladkowej, co najmniej jeden ze zbiorów $f^{-1}(y)$ ma więcej niż 5 elementów, co oznacza, że co najmniej jedna komisja będzie musiała uczestniczyć w co najmniej 6 egzaminach.

Udowodnimy teraz następujące twierdzenie.

Twierdzenie. Jeżeli $|X| = n$ i $|Y| = m$, to wszystkich funkcji $f : X \rightarrow Y$ jest m^n .

Dowód. Przy tworzeniu funkcji $f : X \rightarrow Y$ dowolny element $x \in X$ trafia do dowolnego jednego ze zbiorów $f^{-1}(y)$. Ponieważ zbiorów tych jest m , a elementów $x \in X$ jest n , więc wszystkich sposobów zdefiniowania funkcji f jest

$$\underbrace{m \cdot m \cdot \dots \cdot m}_{n \text{ razy}} = m^n. \quad \square$$

Definicja. Niech $\pi = \{A_1, A_2, \dots, A_k\}$ będzie rodziną podzbiorów zbioru skończonego X taką, że $A_1 \cup A_2 \cup \dots \cup A_k = \bigcup_{i=1}^k A_i = X$, $A_i \cap A_j = \emptyset$ dla $i \neq j$ oraz $A_i \neq \emptyset$ dla $i = 1, 2, \dots, k$. Rodzina π nazywa się *podział* zbioru X na k podzbiorów. Podzbiory $A_1, A_2, \dots, A_k$ nazywa się *bloki* podziału π . Zbiór wszystkich podziałów zbioru X na k bloków oznacza się przez $\Pi_k(X)$, a zbiór wszystkich podziałów zbioru X przez $\Pi(X)$.

Definicja. Liczbą *Stirlinga drugiego rodzaju* $S(n, k)$ nazywa się liczbę podziałów zbioru n -elementowego na k bloków, czyli

$$S(n, k) = |\Pi_k(X)|,$$

gdzie $|X| = n$.

Uwaga. Definiuje się również liczby Stirlinga pierwszego rodzaju, jednak liczby Stirlinga drugiego rodzaju pojawiają się częściej oraz są bardziej interesujące.

Przykład. Dla zbioru 4-elementowego $X = \{1, 2, 3, 4\}$ mamy $S(4, 2) = 7$, ponieważ istnieją tylko następujące podziały zbioru X na 2 bloki:

- 1) $\{\{1, 2, 3\}, \{4\}\},$
- 2) $\{\{1, 2, 4\}, \{3\}\},$
- 3) $\{\{1, 3, 4\}, \{2\}\},$
- 4) $\{\{2, 3, 4\}, \{1\}\},$
- 5) $\{\{1, 2\}, \{3, 4\}\},$
- 6) $\{\{1, 3\}, \{2, 4\}\},$
- 7) $\{\{1, 4\}, \{2, 3\}\}.$

Twierdzenie. Prawdziwe są równości:

- 1) $S(0, 0) = 1$,
- 2) $S(n, 0) = 0$ dla $n > 0$,
- 3) $S(n, k) = 0$ dla $k < 0$,
- 4) $S(n, k) = 0$ dla $k > n$,
- 5) $S(n, 1) = 1$ dla $n > 0$,
- 6) $S(n, n) = 1$,
- 7) $S(n, 2) = 2^{n-1} - 1$ dla $n > 0$,
- 8) $S(n, n-1) = \binom{n}{2}$.

Dowód. 1) Pusta rodzina zbiorów jest jedynym podziałem zbioru pustego.

2) Niepusty zbiór nie może zostać podzielony na zero bloków.

3) Oczywiście.

4) Nie jest możliwe podzielenie zbioru n -elementowego na więcej niż n niepustych parami rozłącznych zbiorów.

5) Istnieje tylko jeden podział niepustego zbioru na jeden blok – blok ten musi być całym dzielonym zbiorem.

6) Jedyne podział spełniający tę własność to podział na bloki 1-elementowe.

7) Niech $|X| = n$ i niech $x \in X$. Zauważmy, że podział na dwa bloki jest zdeterminowany jednym z tych bloków – drugi to po prostu dopełnienie pierwszego. Niech więc blokiem determinującym podział będzie blok zawierający x . Bloków takich jest tyle ile podzbiorów $(n-1)$ -elementowego zbioru $X \setminus \{x\}$, poza podzbiorem pełnym, gdyż wtedy drugi blok byłby pusty, do których dokładamy element x . Zatem jest dokładnie $2^{n-1} - 1$ możliwości wyboru bloku dla x , i tym samym tyle jest podziałów zbioru X .

8) Dzielimy tutaj zbiór n -elementowy na $n-1$ bloków. Zatem jeden z tych bloków musi być 2-elementowy. Po wybraniu takiego bloku pozostałe bloki będą 1-elementowe. Podziałów jest zatem tyle na ile sposobów można wybrać 2-elementowy podzbiór zbioru n -elementowego, czyli $\binom{n}{2}$. $\square$

Twierdzenie. Liczby Stirlinga drugiego rodzaju spełniają następującą zależność rekurencyjną:

$$\begin{aligned} S(n, k) &= S(n-1, k-1) + k \cdot S(n-1, k) \quad \text{dla } 0 < k < n, \\ S(n, n) &= 1 \quad \text{dla } n \geq 0, \\ S(n, 0) &= 0 \quad \text{dla } n > 0. \end{aligned}$$

Dowód. Rozważmy zbiór wszystkich podziałów zbioru n -elementowego $\{1, 2, \dots, n\}$ na k bloków. Zbiór ten możemy podzielić na dwie rozłączne klasy:

1. klasa podziałów, które zawierają blok $\{n\}$,

2. klasa podziałów, w których n jest elementem bloku co najmniej 2-elementowego.

Liczność pierwszej klasy jest równa $S(n-1, k-1)$, bo każdy podział zbioru $(n-1)$ -elementowego na $k-1$ bloków można jednoznacznie przekształcić do podziału z pierwszej klasy dodając blok $\{n\}$. Liczność drugiej klasy wynosi $k \cdot S(n-1, k)$, ponieważ w każdym podziale zbioru $(n-1)$ -elementowego na k bloków dokładamy element n do jednego z bloków, a możliwości takiego dokładania dla danego podziału jest k . Zatem

$$S(n, k) = S(n-1, k-1) + k \cdot S(n-1, k) \quad \text{dla } 0 < k < n.$$

Ponadto, z poprzedniego twierdzenia, wiemy, że $S(n, n) = 1$ dla $n \geq 0$ oraz $S(n, 0) = 0$ dla $n > 0$. $\square$

Twierdzenie. Jeśli $|X| = n$ i $|Y| = m$, to liczba wszystkich funkcji $f : X \xrightarrow{na} Y$ jest równa

$$s_{n,m} = \sum_{j=0}^{m-1} (-1)^j \binom{m}{j} (m-j)^n.$$

Dowód. Niech $Y = \{y_1, y_2, \dots, y_m\}$ oraz niech $A_i = \{f : y_i \notin f(X)\}$. Wtedy

$$f(X) = Y \Leftrightarrow f \notin \bigcup_{i=1}^m A_i.$$

Wszystkich funkcji $f : X \rightarrow Y$ jest m^n . Szukamy więc $|A_1 \cup A_2 \cup \dots \cup A_m|$. Skorzystamy z zasady włączania-wyłączania. Rozpatrzmy iloczyn $A_{k_1} \cap A_{k_2} \cap \dots \cap A_{k_j}$ dla dowolnego ciągu $1 \leq k_1 < k_2 < \dots < k_j \leq m$. Iloczyn ten jest zbiorem wszystkich funkcji $f : X \rightarrow Y \setminus \{y_{k_1}, y_{k_2}, \dots, y_{k_j}\}$. Zatem jego moc jest równa $(m-j)^n$. Ciąg $1 \leq k_1 < k_2 < \dots < k_j \leq m$ można wybrać na $\binom{m}{j}$ sposobów. Stąd

$$\begin{aligned} \left| \bigcup_{i=1}^m A_i \right| &= \sum_{i=1}^m |A_i| - \sum_{1 \leq k_i < k_j \leq m} |A_{k_i} \cap A_{k_j}| + \sum_{1 \leq k_i < k_j < k_l \leq m} |A_{k_i} \cap A_{k_j} \cap A_{k_l}| \\ &\quad - \dots + (-1)^{m+1} \cdot |A_1 \cap A_2 \cap \dots \cap A_m| = m \cdot (m-1)^n - \binom{m}{2} (m-2)^n \\ &\quad + \dots + (-1)^m \binom{m}{m-1} (m - (m-1))^n + (-1)^{m+1} \binom{m}{m} (m-m)^n \\ &= \binom{m}{1} (m-1)^n - \binom{m}{2} (m-2)^n + \dots + (-1)^m \binom{m}{m-1} (m - (m-1))^n \\ &= \sum_{j=1}^{m-1} (-1)^{j+1} \binom{m}{j} (m-j)^n. \end{aligned}$$

Zatem

$$\begin{aligned} s_{n,m} &= m^n - \left| \bigcup_{i=1}^m A_i \right| = m^n - \sum_{j=1}^{m-1} (-1)^{j+1} \binom{m}{j} (m-j)^n \\ &= \sum_{j=0}^{m-1} (-1)^j \binom{m}{j} (m-j)^n. \quad \square \end{aligned}$$

Z drugiej strony zachodzi następujące twierdzenie.

Twierdzenie. Jeśli $|X| = n$ i $|Y| = m$, to liczba wszystkich funkcji $f : X \xrightarrow{na} Y$ jest równa

$$s_{n,m} = m! \cdot S(n, m).$$

Dowód. Niech $Y = \{y_1, y_2, \dots, y_m\}$. Dowolna funkcja $f : X \xrightarrow{na} Y$ wyznacza podział postaci $\{f^{-1}(y_1), f^{-1}(y_2), \dots, f^{-1}(y_m)\}$ zbioru X na m bloków. Funkcji f jest więc tyle ile podziałów takiej postaci. Wszystkich podziałów zbioru n -elementowego X na m bloków jest $S(n, m)$. Bloki w podziale $\{f^{-1}(y_1), f^{-1}(y_2), \dots, f^{-1}(y_m)\}$ są wyznaczone przez elementy $y_1, y_2, \dots, y_m$. Elementy te do poszczególnych bloków można wybrać na $m!$ sposobów. Zatem liczba podziałów żądanej postaci, a więc również funkcji $f : X \xrightarrow{na} Y$ jest równa

$$s_{n,m} = m! \cdot S(n, m). \quad \square$$

Z dwóch ostatnich twierdzeń wynika następujące twierdzenie.

Twierdzenie. Liczba podziałów zbioru n -elementowego na k bloków jest równa

$$S(n, k) = \frac{1}{k!} \sum_{j=0}^{k-1} (-1)^j \binom{k}{j} (k-j)^n.$$

Uwaga. Liczby Stirlinga drugiego rodzaju $S(n, k)$ dla $n, k = 2, 3, \dots, 10$ przedstawia poniższa tabelka:

$n \backslash k$	2	3	4	5	6	7	8	9	10
2	1	0	0	0	0	0	0	0	0
3	3	1	0	0	0	0	0	0	0
4	7	6	1	0	0	0	0	0	0
5	15	25	10	1	0	0	0	0	0
6	31	90	65	15	1	0	0	0	0
7	63	301	350	140	21	1	0	0	0
8	127	966	1701	1050	266	28	1	0	0
9	255	3025	7770	6951	2646	462	36	1	0
10	511	9330	34105	42525	22827	5880	750	45	1

Interpretacje kombinatoryczne liczb Stirlinga drugiego rodzaju:

Ilekoć poniżej będzie mowa o przypisywaniu ludzi do stolików lub do pokoi, to przyjmuje się, że

- osoby są rozróżnialne,
- pokoje są rozróżnialne, np. ponumerowane,
- stoliki są nierozróżnialne (identyczne).

1. Liczba rozmieszczeń n różnych przedmiotów w k identycznych pojemnikach, przy założeniu, że żaden pojemnik nie pozostanie pusty, jest równa $S(n, k)$, gdzie $n \geq k$. Podobnie, n osób można rozsadzić przy dokładnie k stolikach na $S(n, k)$, jeśli przy stoliku może siedzieć nieograniczona liczba osób i sposób ich usadzenia przy danym stoliku nie ma znaczenia. Istotnie, wystarczy przedmioty lub osoby potraktować jak elementy zbioru, a pojemniki lub stoliki jak bloki podziału i skorzystać z zasady bijekcji oraz definicji liczb Stirlinga drugiego rodzaju.

2. Liczba sposobów ulokowania n osób w k pokojach, przy założeniu, że w każdym z pokoi jest co najmniej jedna osoba jest równa $k! \cdot S(n, k)$. Istotnie, wystarczy podzielić osoby na k grup, a następnie grupy przyporządkować w sposób wzajemnie jednoznaczny do pokoi.

3. W kryptografii i kryptoanalizie klasyfikuje się słowa według ich tzw. ciągów modelowych. Polega to na tym, że litery słowa czytane od lewej do prawej są kodowane liczbami $1, 2, 3, \dots$, np. słowo KOMBINATORYKA będzie kodowane ciągiem

1, 2, 3, 4, 5, 6, 7, 8, 2, 9, 10, 1, 7,

a słowo MATEMATYKA ciągiem

1, 2, 3, 4, 1, 2, 3, 5, 6, 2.

Liczba ciągów modelowych odpowiadających słowom n -literowym (czyli długości n) składającym się z k różnych liter jest równa $S(n, k)$. Istotnie, powtarzające się litery wkłada się do pojemnika z liczbą.

Definicja. Niech $|X| = n$. Liczbę *Bella* definiuje się wzorem

$$B_n = \sum_{k=0}^n S(n, k),$$

czyli $B_n = |\Pi(X)|$.

Twierdzenie. Liczby Bella spełniają następującą zależność rekurencyjną:

$$B_{n+1} = \sum_{i=0}^n \binom{n}{i} B_i,$$
$$B_0 = 1.$$

Dowód. Niech X będzie zbiorem $(n+1)$ -elementowym i niech $x \in X$. Wyznamy liczbę podziałów zbioru X , w których blok zawierający element x składa się z $i+1$ elementów dla pewnego $i = 0, 1, \dots, n$. W takim bloku pozostałe i elementów można wybrać ze zbioru $X \setminus \{x\}$ na $\binom{n}{i}$ sposobów. Pozostałe $n-i$ elementów dzielimy na bloki otrzymując rozważany podział. Podział taki jest możliwy na B_{n-i} sposobów. Zatem, wykorzystując równość $\binom{n}{k} = \binom{n}{n-k}$, mamy

$$\begin{aligned} B_{n+1} &= \sum_{i=0}^n \binom{n}{i} B_{n-i} \\ &= \binom{n}{0} B_n + \binom{n}{1} B_{n-1} + \dots + \binom{n}{n-1} B_1 + \binom{n}{n} B_0 \\ &= \binom{n}{0} B_0 + \binom{n}{1} B_1 + \dots + \binom{n}{n-1} B_{n-1} + \binom{n}{n} B_n \\ &= \sum_{i=0}^n \binom{n}{i} B_i \quad \square \end{aligned}$$

Uwaga. Liczby Bella B_n dla $n = 0, 1, \dots, 10$ przedstawia poniższa tabelka:

n	0	1	2	3	4	5	6	7	8	9	10	$\dots$
B_n	1	1	2	5	15	52	203	877	4140	21147	115975	$\dots$

Zauważmy jeszcze, że od podziałów $\pi = \{A_1, A_2, \dots, A_k\}$ zbioru n -elementowego X nie wymagano, żeby były uporządkowane, tzn. takie, że $|A_i| = n_i$ dla $i = 1, 2, \dots, k$, gdzie $n_1 + n_2 + \dots + n_k = n$. Dla podziałów uporządkowanych prawdziwe jest następujące twierdzenie.

Twierdzenie. Niech X będzie zbiorem n -elementowym i niech $n_1 + n_2 + \dots + n_k = n$. Wtedy istnieje

$$P_{n_1, n_2, \dots, n_k}^n = \frac{n!}{n_1! \cdot n_2! \cdot \dots \cdot n_k!}$$

podziałów uporządkowanych $\{A_1, A_2, \dots, A_k\}$ zbioru X takich, że $|A_i| = n_i$ dla $i = 1, 2, \dots, k$.

Dowód. Łatwy. $\square$

8. Podstawy teorii grafów

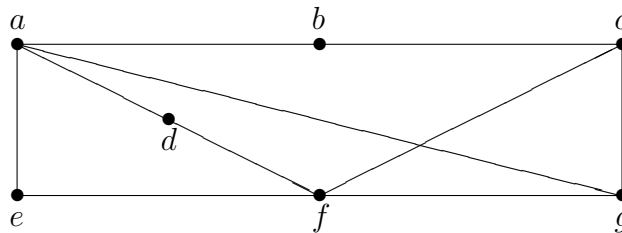
Definicja. Grafem $G = (V, E)$ nazywa się parę składającą się ze skończonego zbioru wierzchołków V oraz ze zbioru krawędzi E , gdzie krawędzie to pary wierzchołków:

$$E \subseteq \{\{u, v\} : u, v \in V, u \neq v\}.$$

O krawędzi $e = \{u, v\}$ mówi się, że *łączy wierzchołki u i v* , a o wierzchołkach u i v , że są *końcami krawędzi e* . Krawędź $\{u, v\}$ jest również oznaczana przez uv .

Uwaga. Graf często przedstawiany jest na rysunku jako zbiór punktów połączonych odcinkami (lub łukami).

Przykład. Graf $G = (V, E)$ ze zbiorem wierzchołków $V = \{a, b, c, d, e, f, g\}$ i zbiorem krawędzi $E = \{ab, ad, ae, ag, bc, cg, cf, df, ef, fg\}$ przedstawiony jest na poniższym rysunku:



Definicja. Grafem *pełnym* nazywa się graf, w którym każde dwa wierzchołki są połączone krawędzią.

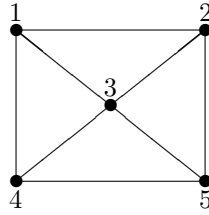
Definicja. Grafem *pustym* nazywa się graf, którego zbiór krawędzi jest zbiorem pustym.

Definicja. Niech $e = uv$ będzie krawędzią grafu G . Mówi się wtedy, że *wierzchołki u i v sąsiadują ze sobą*, oraz że *krawędź e jest incydentna z wierzchołkami u i v* .

Uwaga. Do przechowywania dużego grafu w pamięci komputera wygodniejszym sposobem, niż w postaci rysunku, jest reprezentacja macierzowa.

Definicja. Niech G będzie grafem, którego wierzchołki są oznakowane liczbami ze zbioru $\{1, 2, \dots, n\}$. *Macierzą sąsiedztwa* nazywa się macierz $A = [a_{ij}]$ wymiaru $n \times n$, której wyraz a_{ij} jest równy liczbie krawędzi łączących wierzchołki i -ty z wierzchołkiem j -tym.

Przykład. Macierzą sąsiedztwa grafu

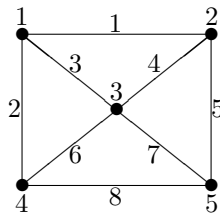


jest macierz

$$A = \begin{bmatrix} 0 & 1 & 1 & 1 & 0 \\ 1 & 0 & 1 & 0 & 1 \\ 1 & 1 & 0 & 1 & 1 \\ 1 & 0 & 1 & 0 & 1 \\ 0 & 1 & 1 & 1 & 0 \end{bmatrix}.$$

Definicja. Niech G będzie grafem, którego wierzchołki są oznakowane liczbami ze zbioru $\{1, 2, \dots, n\}$ oraz krawędzie są oznakowane liczbami ze zbioru $\{1, 2, \dots, m\}$. *Macierzą incydencji* nazywa się macierz $B = [b_{ij}]$ wymiaru $n \times m$, której wyraz b_{ij} jest równy 1, jeśli wierzchołek i -ty jest incydentny z krawędzią j -tą, oraz równy 0 w przeciwnym przypadku.

Przykład. Macierzą incydencji grafu



jest macierz

$$A = \begin{bmatrix} 1 & 1 & 1 & 0 & 0 & 0 & 0 & 0 \\ 1 & 0 & 0 & 1 & 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 1 & 0 & 1 & 1 & 0 \\ 0 & 1 & 0 & 0 & 0 & 1 & 0 & 1 \\ 0 & 0 & 0 & 0 & 1 & 0 & 1 & 1 \end{bmatrix}.$$

Definicja. Różne krawędzie grafu G łączące te same dwa wierzchołki nazywa się *krawędzie wielokrotne*.

Definicja. *Pętlą* nazywa się krawędź z tylko jednym końcem.

Definicja. *Grafem prostym* nazywa się graf bez krawędzi wielokrotnych i pętli.

Definicja. Liczba krawędzi wychodzących z wierzchołka v grafu G nazywa się *stopień wierzchołka v* i oznacza jest przez $d(v)$. Maksymalny spośród stopni wierzchołków grafu G oznacza się przez $\Delta(G)$.

Definicja. Wierzchołek stopnia 0 nazywa się *wierzchołek izolowany*.

Definicja. Graf, w którym wszystkie wierzchołki mają ten sam stopień nazywa się *graf regularny*. Jeżeli stopień ten jest równy k , to graf ten nazywa się *graf regularny stopnia k* .

Twierdzenie. W dowolnym grafie $G = (V, E)$ zachodzi wzór

$$\sum_{v \in V} d(v) = 2 \cdot |E|.$$

Dowód. Rzeczywiście, licząc krawędzie wychodzące z wierzchołków v grafu G , każdą krawędź liczymy dwa razy. $\square$

Wniosek. W dowolnym grafie liczba wierzchołków o nieparzystych stopniach jest parzysta.

Definicja. *Ciąg stopni grafu* składa się ze stopni wierzchołków w kolejności wzrastającej, przy czym uwzględnia się w nim powtórzenia.

Definicja. Graf $G_1 = (V_{G_1}, E_{G_1})$ nazywa się *podgraf* grafu $G_2 = (V_{G_2}, E_{G_2})$, jeżeli

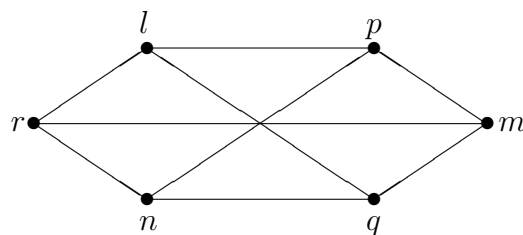
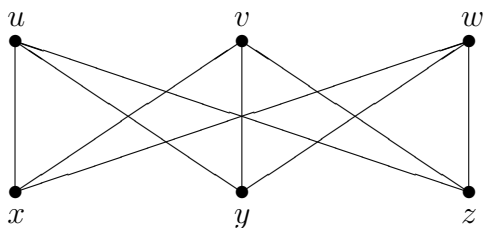
$$V_{G_1} \subseteq V_{G_2} \text{ i } E_{G_1} \subseteq E_{G_2}.$$

Definicja. Dwa grafy $G_1 = (V_{G_1}, E_{G_1})$ i $G_2 = (V_{G_2}, E_{G_2})$ nazywa się *izomorficzne*, jeśli istnieje bijekcja $f : V_{G_1} \rightarrow V_{G_2}$ taka, że

$$uv \in E_{G_1} \Leftrightarrow f(u)f(v) \in E_{G_2}.$$

Funkcja f nazywa się wtedy *izomorfizm grafów*.

Przykład. Dwa grafy pokazane na poniższym rysunku są izomorficzne przy bijekcji f określonej następująco: $f(u) = l$, $f(v) = m$, $f(w) = n$, $f(x) = p$, $f(y) = q$ i $f(z) = r$:



Uwaga. W wielu zadaniach nazwy wierzchołków są niepotrzebne, więc się je pomija. Mówi się wtedy, że dwa „grafy nieoznakowane” są izomorficzne, jeśli można wierzchołkom tak przyporządkować nazwy, żeby otrzymane „grafy oznakowane” były izomorficzne.

Uwaga. Różnica między grafami oznakowanymi i nieoznakowanymi staje się wyraźniejsza, gdy próbuje się je zliczać. Zagadnieniem tym zajmiemy się dokładniej przy omawianiu drzew.

Definicja. *Trasą* w grafie G nazywa się skończony ciąg krawędzi grafu G postaci

$$v_0v_1, v_1v_2, \dots, v_{k-1}v_k,$$

w którym każde dwie kolejne krawędzie są albo sąsiednie albo identyczne. Taka trasa wyznacza ciąg wierzchołków $v_0, v_1, \dots, v_k$. O trasie $v_0v_1, v_1v_2, \dots, v_{k-1}v_k$ mówi się, że *łączy wierzchołki* v_0 i v_k . Mówi się także, że *wierzchołek* v_k *jest osiągalny z wierzchołka* v_0 . Liczba krawędzi na trasie nazywa się *długość trasy*.

Uwaga. Trasę $v_0v_1, v_1v_2, \dots, v_{k-1}v_k$ zapisujemy też następująco:

$$v_0 \rightarrow v_1 \rightarrow \dots \rightarrow v_k.$$

Definicja. Trasa, w której wszystkie krawędzie są różne nazywa się *ścieżką*. Jeżeli ponadto wierzchołki $v_0, v_1, \dots, v_k$ są różne (z wyjątkiem, być może, równości $v_0 = v_k$), to ścieżka nazywa się *drogą*. Droga lub ścieżka nazywa się *zamknięta*, jeśli $v_0 = v_k$.

Definicja. *Cyklem* nazywa się drogę zamkniętą zawierającą co najmniej jedną krawędź.

Uwaga. Każda pętla lub para krawędzi wielokrotnych jest cyklem.

Twierdzenie. Jeżeli w grafie istnieją dwa różne wierzchołki, które można połączyć dwiema różnymi drogami, to w grafie tym istnieje cykl.

Dowód. Oczywisty. $\square$

Wniosek. Graf nie posiada cykli wtedy i tylko wtedy, gdy każde dwa wierzchołki grafu można połączyć co najwyżej jedną drogą.

Definicja. *Obwodem* grafu G nazywa się długość najkrótszego cyklu w G .

Konstrukcje nowych grafów z danego grafu G :

Niech v będzie wierzchołkiem, a e krawędzią grafu G .

1. Usunięcie z grafu G krawędzi e powoduje powstanie nowego grafu oznaczanego przez $G - e$. Graf $G - e$ jest podgrafem grafu G .
2. Usunięcie z grafu G wierzchołka v wraz ze wszystkimi krawędziami incydentnymi z v powoduje powstanie nowego grafu oznaczanego przez $G - v$. Graf $G - v$ jest podgrafem grafu G .
3. Ściągnięcie krawędzi e w grafie G polega na usunięciu krawędzi e w taki sposób, że dwa wierzchołki, które łączyła staną się teraz jednym wierzchołkiem. Ściągnięcie w grafie G krawędzi e powoduje powstanie nowego grafu oznaczanego przez $G \setminus e$.
4. Niech $e = ab$. Zastąpienie krawędzi e dwiema krawędziami ax i xb , gdzie x jest nowym wierzchołkiem, nazywa się *podział krawędzi* i powoduje powstanie nowego grafu. Dwa grafy, które mogą być otrzymane z tego samego grafu przez podziały jego krawędzi nazywa się *homeomorficzne*.

Definicja. Graf G nazywa się *spójny*, jeżeli dla każdych dwóch wierzchołków $u, v \in V_G$ istnieje droga łącząca u i v . W przeciwnym przypadku graf G nazywa się *niespójny*. Każdy niespójny graf G można rozbić na *składowe*, czyli maksymalne spójne podgrafy.

Twierdzenie. Niech G będzie grafem prostym mającym n wierzchołków. Jeśli graf G ma k składowych, to liczba m jego krawędzi spełnia nierówność

$$n - k \leq m \leq \frac{1}{2}(n - k)(n - k + 1).$$

Dowód. Udowodnimy najpierw nierówność $m \geq n - k$ przez indukcję względem liczby krawędzi m grafu G . Dla grafu pustego nierówność ta jest trywialna. Załóżmy, że nierówność ta jest prawdziwa dla grafu mającego $m - 1$ krawędzi i k składowych. Jeśli graf G ma najmniejszą możliwą liczbę krawędzi m , to usunięcia którejkolwiek krawędzi powoduje zwiększenie liczby składowych o 1 i powstały w ten sposób graf ma n wierzchołków, $k + 1$ składowych oraz $m - 1$ krawędzi. Z założenia indukcyjnego wynika, że

$$m - 1 \geq n - (k + 1),$$

skąd

$$m \geq n - k.$$

Zatem nierówność ta jest prawdziwa dla każdego m .

Pokażemy teraz nierówność $m \leq \frac{1}{2}(n-k)(n-k+1)$. Możemy założyć, że każda składowa grafu G jest grafem pełnym. Przypuśćmy, że istnieją dwie składowe C_i oraz C_j mające odpowiednio n_i i n_j wierzchołków, przy czym $n_i \geq n_j \geq 1$. Jeśli zastąpimy składowe C_i oraz C_j grafami pełnymi mającymi odpowiednio $n_i + 1$ i $n_j - 1$ wierzchołków, to całkowita liczba wierzchołków nie zmieni się, a liczba krawędzi zmieni się o wielkość

$$\begin{aligned} & \left(\binom{n_i+1}{2} - \binom{n_i}{2} \right) + \left(\binom{n_j-1}{2} - \binom{n_j}{2} \right) \\ &= \frac{1}{2} \left((n_i+1)n_i - n_i(n_i-1) \right) + \frac{1}{2} \left((n_j-1)(n_j-2) - n_j(n_j-1) \right) \\ &= \frac{1}{2} n_i (\cancel{n_i} + 1 - \cancel{n_i} + 1) + \frac{1}{2} (n_j - 1) (\cancel{n_j} - 2 - \cancel{n_j}) \\ &= n_i - n_j + 1, \end{aligned}$$

która jest liczbą dodatnią. Zatem liczba krawędzi zwiększy się. Wynika stąd, że aby graf G miał maksymalną liczbę krawędzi, musi składać się z grafu pełnego mającego $n - k + 1$ wierzchołków i z $k - 1$ wierzchołków izolowanych. Stąd wynika teza twierdzenia. $\square$

Wniosek. Jeśli G jest grafem prostym i spójnym mającym n wierzchołków, to liczba m jego krawędzi spełnia nierówności

$$n - 1 \leq m \leq \frac{1}{2}(n-1)n.$$

Twierdzenie. Każdy graf prosty, który ma n wierzchołków i więcej niż $\frac{1}{2}(n-1)(n-2)$ krawędzi jest spójny.

Dowód. Zauważmy, że

$$\binom{n-1}{2} = \frac{(n-1)!}{2!(n-3)!} = \frac{1}{2}(n-1)(n-2).$$

Wynika stąd, że graf prosty i pełny mający $n-1$ wierzchołków ma dokładnie $\frac{1}{2}(n-1)(n-2)$ krawędzi. Jeżeli teraz dołożymy do tego grafu 1 wierzchołek i założymy, że ma on więcej niż $\frac{1}{2}(n-1)(n-2)$ krawędzi, to przynajmniej jedna krawędź musi łączyć ten nowy wierzchołek

z którymś z pozostałych $n - 1$ wierzchołków. Stąd ten nowy wierzchołek nie może być wierzchołkiem izolowanym, czyli taki graf jest spójny. $\square$

Definicja. *Zbiorem rozspajającym* grafu spójnego G nazywa się zbiór krawędzi, których usunięcie spowoduje, że graf G przestanie być spójny.

Definicja. *Rozcięciem* nazywa się zbiór rozspajający, którego żaden podzbiór właściwy nie jest już zbiorem rozspajającym.

Definicja. *Mostem* nazywa się rozcięcie składające się z jednej krawędzi.

Definicja. Niech G będzie grafem spójnym. *Spójnością krawędziową* $\lambda(G)$ nazywa się liczbę krawędzi należących do najmniej licznego rozcięcia grafu G .

Definicja. Graf spójny G nazywa się *k -spójny krawędziowo*, jeśli $\lambda(G) \geq k$.

Definicja. *Zbiorem rozdziłającym* grafu spójnego G nazywa się zbiór wierzchołków, których usunięcie powoduje, że graf przestaje być spójny.

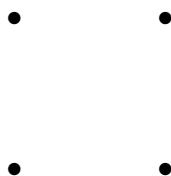
Uwaga. Usuując wierzchołek z grafu, usuwamy również wszystkie krawędzie z nim incydentne.

Definicja. Niech graf G będzie spójny i nie będzie grafem pełnym. *Spójnością wierzchołkową* $\kappa(G)$ nazywa się liczbę elementów najmniej licznego zbioru rozdziłającego grafu G .

Definicja. Niech graf G będzie spójny i nie będzie grafem pełnym. Graf G nazywa się *k -spójny*, jeśli $\kappa(G) \geq k$.

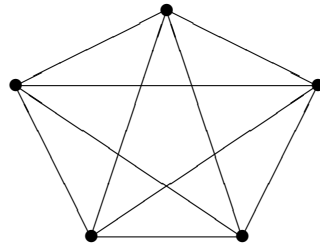
Definicja. *Grafem pustym* nazywa się graf, którego zbiór krawędzi jest zbiorem pustym. Graf pusty mający n wierzchołków oznacza się przez N_n .

Przykład. Poniższy graf przedstawia graf N_4 :



Definicja. *Grafem pełnym* nazywa się graf prosty, w którym każde dwa różne wierzchołki są połączone krawędzią. Graf pełny mający n wierzchołków oznacza się przez K_n .

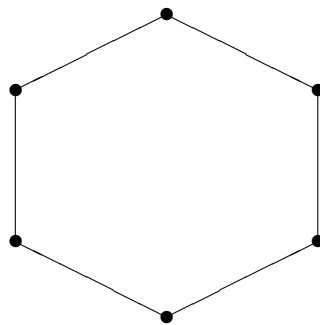
Przykład. Poniższy rysunek przedstawia graf K_5 :



Uwaga. Graf pełny K_n jest grafem spójnym i regularnym stopnia $n - 1$.

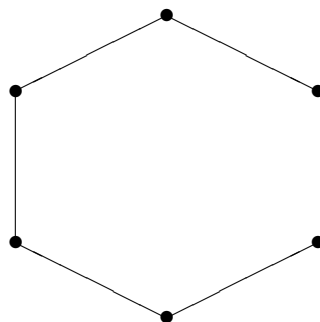
Definicja. *Grafem cyklicznym* nazywa się graf, który jest spójny i regularny stopnia 2. Graf cykliczny mający n wierzchołków oznacza się przez C_n .

Przykład. Poniższy rysunek przedstawia graf C_6 :



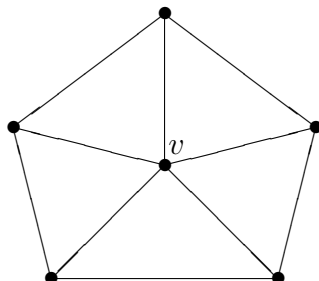
Definicja. *Grafem liniowym* nazywa się graf otrzymany z grafu C_n przez usunięcie jednej krawędzi. Graf liniowy mający n wierzchołków oznacza się przez P_n .

Przykład. Poniższy rysunek przedstawia graf P_6 :



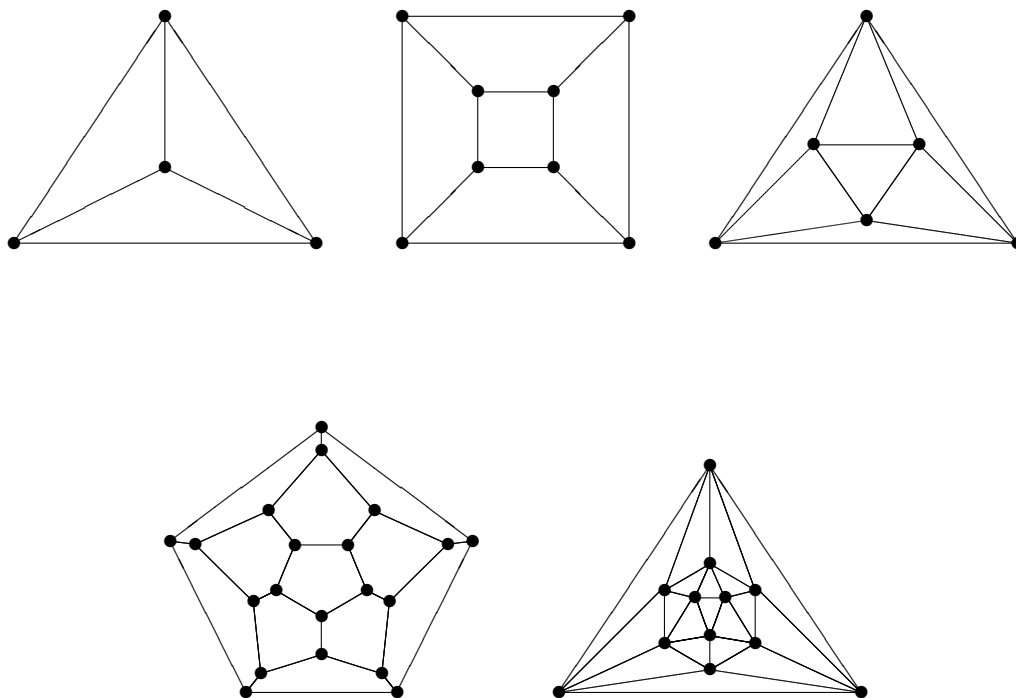
Definicja. Graf powstający z grafu C_{n-1} przez połączenie każdego wierzchołka z nowym wierzchołkiem v nazywa się *koło o n wierzchołkach*. Koło o n wierzchołkach oznacza się przez W_n .

Przykład. Poniższy rysunek przedstawia graf W_6 :



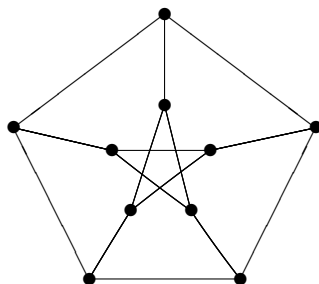
Definicja. *Grafy platońskie* to grafy utworzone z wierzchołków i krawędzi pięciu wielościanów foremnych (platońskich) – czworościanu, sześcianu, ośmiościanu, dwunastościanu i dwudziestościanu.

Przykład. Poniższy rysunek przedstawia grafy platońskie:



Definicja. *Grafem kubicznym* nazywa się graf regularny stopnia 3.

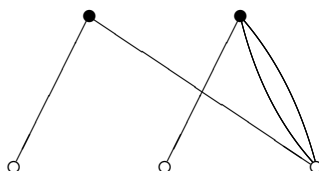
Przykład. Przykładem grafu kubicznego jest graf Petersena przedstawiony na poniższym rysunku:



Definicja. *Grafem dwudzielnym* nazywa się graf, którego zbiór wierzchołków można podzielić na dwa podzbiory w taki sposób, że każda krawędź łączy wierzchołek z pierwszego podzbioru z wierzchołkiem z drugiego podzbioru.

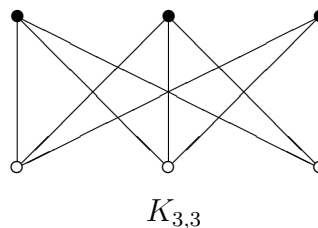
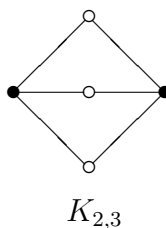
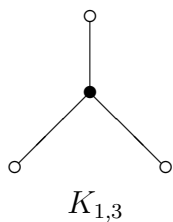
Uwaga. Równoważnie, graf dwudzielnym to taki graf, którego wierzchołki można pokolorować dwoma kolorami, czarnym i białym, w taki sposób, aby każda krawędź łączyła wierzchołek czarny z wierzchołkiem białym.

Przykład. Poniższy rysunek przedstawia przykład grafu dwudzielnego:



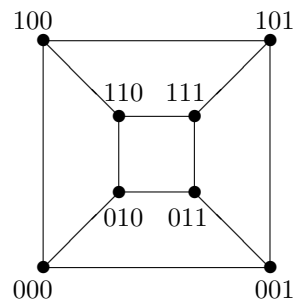
Definicja. *Grafem pełnym dwudzielnym* nazywa się graf dwudzielnym, w którym każdy wierzchołek z pierwszego podzbioru jest połączony dokładnie jedną krawędzią z każdym wierzchołkiem z drugiego podzbioru. Graf pełny dwudzielnym mający m czarnych i n białych wierzchołków oznacza się przez $K_{m,n}$.

Przykład. Poniższy rysunek przedstawia grafy $K_{1,3}$, $K_{2,3}$ i $K_{3,3}$:



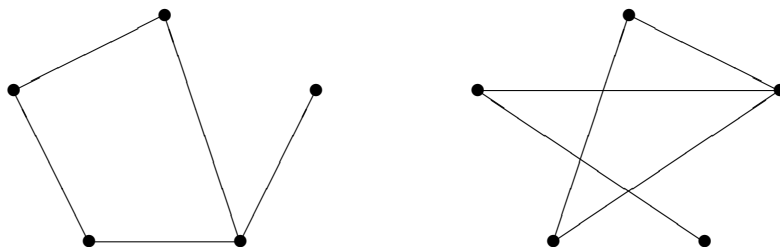
Definicja. k -kostką (lub po prostu *kostką*) Q_k nazywa się graf, którego wierzchołki odpowiadają ciągom $(a_1, a_2, \dots, a_k)$ takim, że każdy wyraz a_i jest równy 0 lub 1, i którego krawędzie łączą ciągi różniące się dokładnie jednym wyrazem.

Przykład. Poniższy rysunek przedstawia kostkę Q_3 :



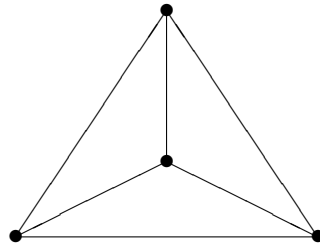
Definicja. Niech G będzie grafem prostym, którego zbiorem wierzchołków jest V . *Dopełnieniem* $\overline{G}$ grafu G nazywa się graf prosty, którego zbiorem wierzchołków jest V i w którym dwa wierzchołki są sąsiednie wtedy i tylko wtedy, gdy nie są sąsiednie w grafie G .

Przykład. Poniższy rysunek przedstawia pewien graf i jego dopełnienie:



Definicja. *Grafem planarnym* nazywa się graf, który może być narysowany na płaszczyźnie tak, że dowolne dwie jego krawędzie nie przecinają się w żadnym punkcie płaszczyzny (poza ewentualnym wspólnym wierzchołkiem). Taki rysunek grafu nazywa się jego *planarną reprezentacją*.

Przykład. Graf pełny K_4 jest planarny o czym świadczy chociażby rysunek pierwszego grafu platońskiego:

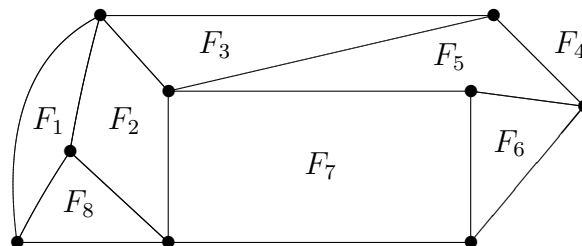


Twierdzenie. (Fáry-Wagner) Dowolny planarny graf prosty posiada planarną reprezentację, w której każda krawędź jest odcinkiem linii prostej.

(bez dowodu)

Definicja. Niech G będzie grafem planarnym. Każda planarna reprezentacja grafu G dzieli zbiór punktów płaszczyzny, które nie leżą na G , na obszary nazywane *ścianami*. Zbiór ścian grafu G oznacza się przez F .

Przykład. Graf z poniższą planarną reprezentacją ma 8 ścian:



Lemat. Niech G będzie grafem spójnym niezawierającym cykli i mającym n wierzchołków. Wtedy G ma $n - 1$ krawędzi.

Dowód. Dowód przeprowadzimy indukcyjnie względem liczby wierzchołków n . Dla $n = 1$ i $n = 2$ teza jest oczywista. Załóżmy, że teza lematu jest prawdziwa dla każdego $k \leq n - 1$. Weźmy graf spójny G niezawierający cykli i mający n wierzchołków. Ponieważ G nie zawiera cykli, więc usunięcie którejkolwiek krawędzi musi spowodować rozpad G na dwa grafy, z których każdy jest spójny i nie zawiera cykli. Z założenia indukcyjnego wynika, że liczba krawędzi w każdym z tych grafów jest o jeden mniejsza od liczby wierzchołków. Stąd wynika, że liczba wszystkich krawędzi grafu G jest równa $n - 1$. $\square$

Twierdzenie. (Euler) Niech $G = (V, E)$ będzie spójnym grafem planarnym oraz niech F będzie zbiorem ścian grafu G . Wtedy

$$|V| - |E| + |F| = 2.$$

Dowód. Dowód przeprowadzimy przez indukcję względem liczby krawędzi grafu G . Niech $|E| = m$. Jeśli $m = 0$, to $|V| = 1$ (bo graf G jest spójny) oraz $|F| = 1$. Twierdzenie zatem jest prawdziwe. Załóżmy teraz, że twierdzenie jest prawdziwe dla wszystkich grafów mających co najwyżej $m - 1$ krawędzi i niech G będzie grafem mającym m krawędzi. Jeśli graf G nie zawiera cykli, to z Lematu, $m = |V| - 1$ i $|F| = 1$, a więc rzeczywiście $|V| - |E| + |F| = 2$. Jeśli graf G zawiera cykle, to niech e będzie dowolną krawędzią jakiegoś cyklu w G . Wtedy graf $G - e$ jest spójnym grafem planarnym mającym $|V|$ wierzchołków, $m - 1$ krawędzi i $|F| - 1$ ścian, a więc z założenia indukcyjnego wynika, że zachodzi równość

$$|V| - (m - 1) + |F| - 1 = 2.$$

Stąd

$$|V| - |E| + |F| = 2. \quad \square$$

Kolejne trzy twierdzenia wynikają z twierdzenia Eulera.

Twierdzenie. Niech $G = (V, E)$ będzie spójnym planarnym grafem prostym. Wtedy

$$|E| \leq 3|V| - 6.$$

Dowód. Załóżmy, dla ułatwienia argumentacji, że graf G nie zawiera wierzchołków stopnia 1. Niech $F = \{F_1, F_2, \dots, F_k\}$ będzie zbiorem wszystkich ścian grafu G . Niech $|E_i|$ oznacza liczbę krawędzi ściany F_i . Ponieważ każda krawędź rozgranicza dwie ściany, więc zliczając krawędzie kolejnych ścian otrzymujemy

$$|E_1| + |E_2| + \dots + |E_k| = 2|E|.$$

Z drugiej strony każda ściana jest ograniczona przez co najmniej 3 krawędzie, a więc $|E_i| \geq 3$ i stąd

$$2|E| = |E_1| + |E_2| + \dots + |E_k| \geq 3k = 3|F|,$$

skąd

$$3|F| \leq 2|E|.$$

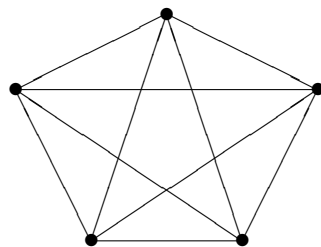
Po zastosowaniu wzoru z twierdzenia Eulera otrzymujemy

$$3(|E| - |V| + 2) \leq 2|E|,$$

czyli

$$|E| \leq 3|V| - 6. \quad \square$$

Przykład. Pokażemy, że graf pełny K_5 nie jest planarny. Przypomnijmy graf K_5 :



Graf ten jest więc spójnym grafem prostym, w którym

$$|E| = 10 > 3|V| - 6 = 3 \cdot 5 - 6 = 9.$$

Zatem, na mocy powyższego twierdzenia, graf K_5 nie jest planarny.

Twierdzenie. Niech $G = (V, E)$ będzie planarnym grafem dwudzielnym. Wtedy

$$|E| \leq 2|V| - 4.$$

Dowód. Przy oznaczeniach z poprzedniego dowodu zauważmy, że tym razem $|E_i| \geq 4$, bo w grafie dwudzielnym nie ma cykli długości nieparzystej. Stąd teraz

$$4|F| \leq 2|E|$$

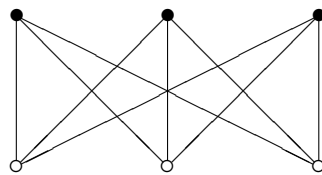
i na mocy wzoru z twierdzenia Eulera dostajemy

$$4(|E| - |V| + 2) \leq 2|E|,$$

skąd

$$|E| \leq 2|V| - 4. \quad \square$$

Przykład. Pokażemy, że graf pełny dwudzielnym $K_{3,3}$ nie jest planarny. Przypomnijmy graf $K_{3,3}$:



Dla grafu $K_{3,3}$ mamy więc

$$|E| = 9 > 2|V| - 4 = 2 \cdot 6 - 4 = 8.$$

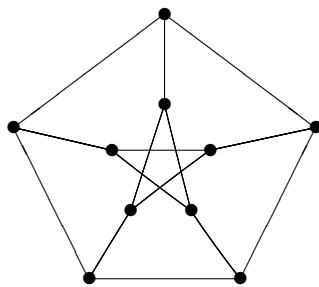
Zatem, na mocy powyższego twierdzenia, nie jest on planarny.

Twierdzenie. Niech $G = (V, E)$ będzie spójnym grafem planarnym o obwodzie r . Wtedy

$$(r - 2)|E| \leq r(|V| - 2).$$

Dowód. Dowód przebiega tak samo jak dwa poprzednie, przy czym teraz mamy $|E_i| \geq r$. $\square$

Przykład. Pokażemy, że graf Petersena nie jest planarny. Przypomnijmy graf Petersena:



Graf ten jest więc grafem spójnym, w którym

$$(r - 2)|E| = (5 - 2) \cdot 15 = 45 > r(|V| - 2) = 5 \cdot (10 - 2) = 40.$$

Zatem, na mocy powyższego twierdzenia, graf Petersena nie jest planarny.

Uwaga. Z powyższych rozważań wiadomo już, że grafy K_5 i $K_{3,3}$ nie są planarne. Jak się okazuje są to w pewnym sensie jedyne takie grafy, bo prawdziwe są następujące twierdzenia.

Twierdzenie. (Kuratowski) Graf jest planarny wtedy i tylko wtedy, gdy nie zawiera podgrafu homeomorficznego z grafem K_5 lub z grafem $K_{3,3}$.

(bez dowodu)

Twierdzenie. Graf jest planarny wtedy i tylko wtedy, gdy nie zawiera podgrafu ściągającego do grafu K_5 lub do grafu $K_{3,3}$.

(bez dowodu)

Twierdzenie. Każdy planarny graf prosty zawiera wierzchołek stopnia co najwyżej 5.

Dowód. Bez straty ogólności możemy założyć, że graf $G = (V, E)$ jest spójny i ma co najmniej trzy wierzchołki. Wiemy, że wtedy $|E| \leq 3|V| - 6$. Gdyby każdy wierzchołek miał stopień co najmniej 6, to mielibyśmy nierówność

$$6|V| \leq 2|E|,$$

czyli

$$3|V| \leq |E|.$$

Stąd

$$3|V| \leq 3|V| - 6,$$

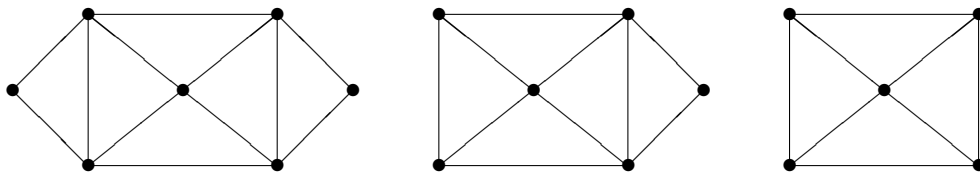
co jest niemożliwe. $\square$

Definicja. *Grafem eulerowskim* nazywa się taki graf spójny G , w którym istnieje zamknięta ścieżka zawierająca każdą krawędź grafu G . Taka ścieżka nazywa się *cykl Euler*.

Uwaga. Zauważmy, że cykl Euler przechodzi przez każdą krawędź dokładnie raz.

Definicja. Graf G , który nie jest grafem eulerowskim, nazywa się *grafem półeulerowskim*, jeśli istnieje ścieżka zawierająca każdą krawędź grafu G .

Przykład. Poniższy rysunek przedstawia przykłady grafów: eulerowskiego, półeulerowskiego oraz grafu, który nie jest grafem półeulerowskim:



Lemat. Jeśli w grafie G każdy wierzchołek ma stopień równy co najmniej 2, to graf G zawiera cykl.

Dowód. Jeśli graf G ma pętle lub krawędzie wielokrotne, to lemat jest oczywisty. Założmy, że graf G jest grafem prostym. Niech v będzie dowolnym wierzchołkiem grafu G . Tworzymy trasę

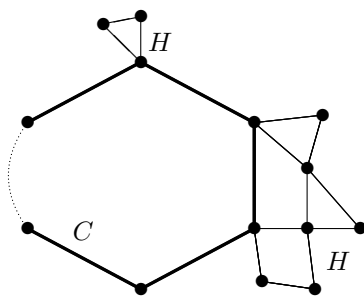
$$v \rightarrow v_1 \rightarrow v_2 \rightarrow \dots$$

przez indukcję, wybierając jako v_1 dowolny wierzchołek sąsiadujący z wierzchołkiem v , a następnie, dla każdego $i > 1$, wybierając jako wierzchołek v_{i+1} dowolny wierzchołek sąsiadujący z v_i różny od v_{i-1} – istnienie takiego wierzchołka wynika z założenia o stopniach wierzchołków grafu G . Ponieważ graf G ma tylko skończenie wiele wierzchołków, więc w końcu musimy wybrać wierzchołek wybrany już wcześniej. Jeśli v_k jest pierwszym takim wierzchołkiem, to część trasy znajdująca się między pierwszym i drugim wystąpieniem wierzchołka v_k jest szukanim cyklem. $\square$

Twierdzenie. (Euler) Graf spójny G jest grafem eulerowskim wtedy i tylko wtedy, gdy stopień każdego wierzchołka grafu G jest liczbą parzystą.

Dowód. Niech spójny graf G będzie grafem eulerowskim. Niech P będzie cyklem Eulera w grafie G . Za każdym razem, gdy cykl P przechodzi przez wierzchołek, udział krawędzi należących do tego cyklu w stopniu tego wierzchołka zwiększa się o 2. Ponieważ każda krawędź występuje w P dokładnie raz, więc stopień każdego wierzchołka musi być liczbą parzystą.

Niech teraz stopień każdego wierzchołka grafu spójnego G będzie liczbą parzystą. Pokażemy, że G jest grafem eulerowskim. Dowód przeprowadzimy przez indukcję względem liczby krawędzi grafu G . Ponieważ graf G jest spójny, więc każdy wierzchołek ma stopień równy co najmniej 2, a zatem z powyższego lematu wynika, że graf G ma cykl C . Jeśli cykl C zawiera każdą krawędź grafu G , to twierdzenie jest udowodnione. W przeciwnym przypadku usuwamy z grafu G krawędzie należące do C , tworząc w ten sposób nowy graf H , być może niespójny, który ma mniej krawędzi niż graf G i w którym stopień każdego wierzchołka jest nadal liczbą parzystą. Z założenia indukcyjnego wynika, że każda składowa grafu H ma cykl Eulera. Ponieważ, na mocy spójności, każda składowa grafu H ma co najmniej jeden wspólny wierzchołek z C , więc szukany cykl Eulera w G otrzymujemy przechodząc krawędzie cyklu C dotąd, aż napotkamy nieizolowany wierzchołek grafu H , przechodząc następnie przez cykl Eulera w tej składowej grafu H , która zawiera napotkany wierzchołek, potem podążając znów wzdłuż cyklu C aż do napotkania następnej składowej grafu H i tak dalej. Ilustruje to poniższy rysunek:



Postępowanie zakończy się, gdy powrócimy do wierzchołka wyjściowego. $\square$

Łatwe modyfikacje powyższego dowodu pokazują prawdziwość następujących dwóch wniosków.

Wniosek. Graf spójny jest grafem eulerowskim wtedy i tylko wtedy, gdy jego zbiór krawędzi można podzielić na rozłączne cykle.

Wniosek. Graf spójny jest grafem półeulerowskim wtedy i tylko wtedy, gdy ma dokładnie dwa wierzchołki nieparzystych stopni.

Uwaga. W grafie półeulerowskim każda ścieżka Eulera musi zaczynać się w jednym wierzchołku nieparzystego stopnia i kończyć w drugim takim wierzchołku.

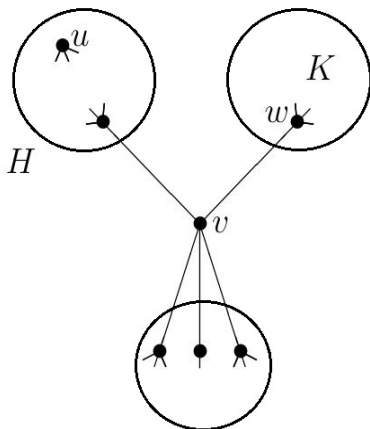
Kolejne twierdzenie podaje algorytm służący do konstrukcji cyklu Eulera w danym grafie eulerowskim.

Twierdzenie. (Algorytm Fleury’ego) Niech G będzie grafem eulerowskim. Wtedy następująca konstrukcja jest wykonalna i daje w wyniku cykl Eulera w grafie G :

Zaczynij cykl w dowolnym wierzchołku u i przechodź krawędzie w dowolnej kolejności, dbając jedynie o zachowanie następujących zasad:

- 1) *usuwaj z grafu przechodzone krawędzie i wierzchołki izolowane powstające w wyniku usuwania tych krawędzi;*
- 2) *w każdym momencie przechodź przez most tylko wtedy, gdy nie masz innej możliwości.*

Dowód. Pokażemy najpierw, że ta konstrukcja może być wykonana w każdym momencie. Przypuśćmy, że dotarliśmy właśnie do wierzchołka v . Jeśli $v \neq u$, to podgraf H , który pozostał po usunięciu przebytych krawędzi, jest spójny i zawiera tylko dwa wierzchołki nieparzystych stopni: u i v . Aby pokazać, że następny krok konstrukcji może być wykonany musimy dowieść, że usunięcie następnej krawędzi nie rozspójni grafu H , lub równoważnie, że wierzchołek v jest incydentny z co najwyżej jednym mostem. Ale gdyby tak nie było, to istniałby most vw taki, że składowa K grafu $H - vw$ zawierająca wierzchołek w nie zawierałaby wierzchołka u :



Ponieważ stopień wierzchołka w jest liczbą nieparzystą w składowej K , więc stopień pewnego innego wierzchołka w tej składowej też musi być liczbą nieparzystą, co prowadzi do sprzeczności. Jeśli natomiast $v = u$, to dowód jest w zasadzie taki sam, chyba, że nie ma już więcej krawędzi incydentnych z wierzchołkiem u .

Należy jeszcze pokazać, że ta konstrukcja zawsze daje w wyniku cykl Eulera. Ale jest to oczywiste, ponieważ w grafie G nie mogą pozostać krawędzie nieprzebyte, gdy została usunięta ostatnia krawędź incydentna z wierzchołkiem u . Gdyby tak się stało, to usunięcie

pewnej wcześniejszej krawędzi incydentnej z którąś z pozostałych krawędzi rozpójniłoby graf, co jednak przeczy zasadzie 2). $\square$

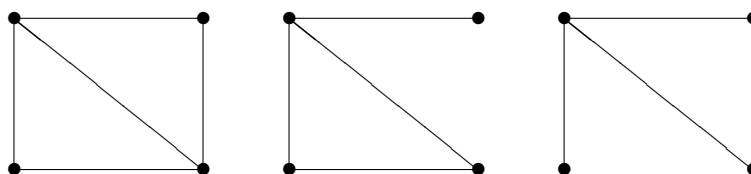
Definicja. *Grafem hamiltonowskim* nazywa się taki graf spójny G , w którym istnieje zamknięta droga przechodząca przez każdy wierzchołek grafu G .

Uwaga. Oczywiście zamknięta droga istniejąca w grafie hamiltonowskim przechodzi przez każdy wierzchołek dokładnie jeden raz.

Definicja. Graf G , który nie jest grafem hamiltonowskim, nazywa się *grafem półhamiltonowskim*, jeśli istnieje droga przechodząca przez każdy wierzchołek grafu G .

Uwaga. Oczywiście znowu droga ta przechodzi przez każdy wierzchołek dokładnie jeden raz.

Przykład. Poniższy rysunek przedstawia przykłady grafów: hamiltonowskiego, półhamiltonowskiego oraz grafu, który nie jest grafem półhamiltonowskim:



Twierdzenie. (Ore) Jeśli graf prosty G ma n wierzchołków, gdzie $n \geq 3$, oraz

$$d(u) + d(v) \geq n$$

dla każdej pary wierzchołków niesąsiednich u i v , to graf G jest hamiltonowski.

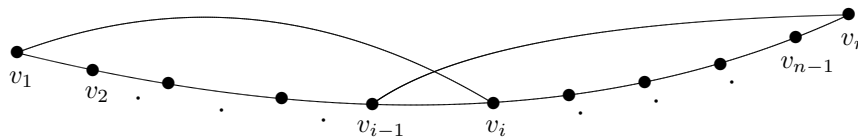
Dowód. Załóżmy, że graf G nie jest hamiltonowski. Wtedy dodając do G odpowiednią ilość krawędzi możemy otrzymać graf G_1 , który jest już grafem hamiltonowskim. Usuwa-
jąc z G_1 jedną krawędź uzyskujemy graf G_2 , który jest półhamiltonowski. Niech drogą istniejącą w takim grafie będzie droga

$$v_1 \rightarrow v_2 \rightarrow \dots \rightarrow v_n.$$

Oczywiście wierzchołki v_1 i v_n nie są sąsiednie. Na mocy założenia

$$d(v_1) + d(v_n) \geq n,$$

ponieważ dodanie krawędzi do grafu G zachowuje nierówności zawarte w założeniu twierdzenia. Z powyższej nierówności wynika, że istnieją wierzchołki v_i oraz v_{i-1} takie, że v_i sąsiaduje z wierzchołkiem v_1 , natomiast v_{i-1} sąsiaduje z wierzchołkiem v_n :



Ale wtedy droga

$$v_1 \rightarrow v_2 \rightarrow \dots \rightarrow v_{i-1} \rightarrow v_n \rightarrow v_{n-1} \rightarrow \dots \rightarrow v_i \rightarrow v_1$$

jest zamknięta, co jest sprzeczne z założeniem, że graf nie jest hamiltonowski. $\square$

Z twierdzenia Orego wynika natychmiast twierdzenie.

Twierdzenie. (Dirac) Jeśli graf prosty G ma n wierzchołków, gdzie $n \geq 3$, oraz $d(v) \geq \frac{n}{2}$ dla każdego wierzchołka v grafu G , to graf G jest hamiltonowski.

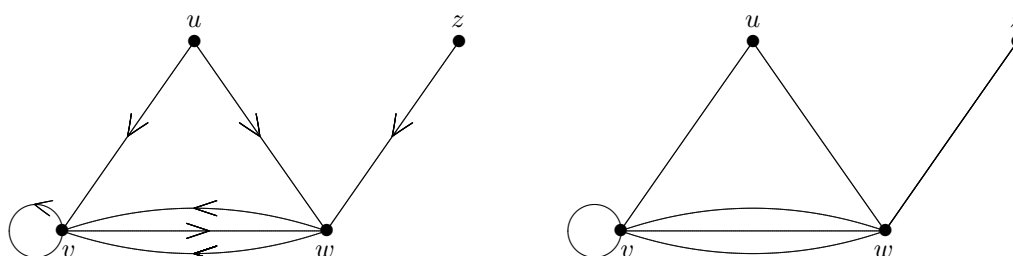
Definicja. *Digraf* $D = (V, A)$ jest to para składająca się ze skończonego zbioru wierzchołków V oraz ze skończonego zbioru łuków A , gdzie łuki to uporządkowane pary wierzchołków:

$$A \subseteq \{(u, v) : u, v \in V, u \neq v\}.$$

Uwaga. Łuk (u, v) zwykle zapisuje się jako uv . Na rysunku kolejność wierzchołków w danym łuku jest wskazywana za pomocą strzałki.

Definicja. *Szkieletem* digrafu D nazywa się graf otrzymany z D przez zastąpienie każdego łuku uv odpowiadającą mu krawędzią uv .

Przykład. Poniższy rysunek przedstawia przykładowy digraf oraz jego szkielet:



Definicja. *Digrafem prostym* nazywa się digraf D , w którym wszystkie łuki są różne oraz w D nie ma pętli.

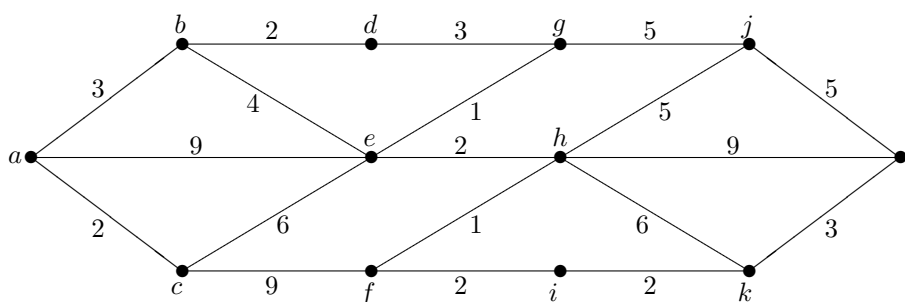
Uwaga. Szkielet digrafu prostego nie musi być grafem prostym.

Uwaga. Definicje wielu pojęć dotyczących grafów w naturalny sposób uogólniają się na przypadek digrafów.

Definicja. *Grafem z wagami* nazywa się graf spójny, w którym każdej krawędzi przypisano pewną liczbę nieujemną. Liczba ta nazywa się *waga* danej krawędzi e i jest oznaczona przez $w(e)$.

Uwaga. Wagi mogą oznaczać długości dróg, czas podróży, koszt przejazdu, itp.

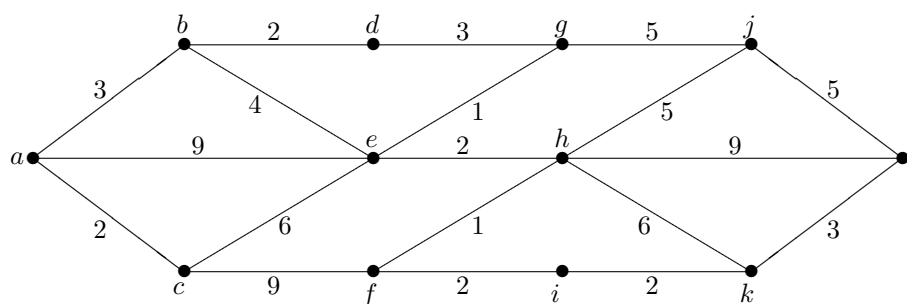
Przykład. Poniższy rysunek przedstawia przykładowy graf z wagami:



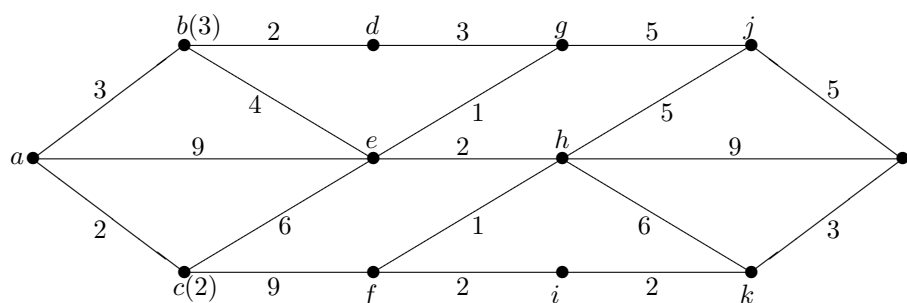
Definicja. *Zagadnienie najkrótszej drogi* polega na znalezieniu drogi łączącej dwa punkty grafu z wagami mającej najmniejszą całkowitą wagę. Na przykład dla grafu z wagami z powyższego przykładu zagadnienie polega na znalezieniu drogi z A do L mającej najmniejszą całkowitą wagę.

Uwaga. Algorytm rozwiązania zagadnienia najkrótszej drogi polega na przesuwaniu się wzdłuż grafu od wierzchołka wyjściowego do kolejnego wierzchołka przypisując każdemu wierzchołkowi v liczbę wskazującą najmniejszą wagę od wierzchołka wyjściowego do wierzchołka v .

Przykład. Dla grafu z wagami z poprzedniego przykładu znajdziemy najkrótszą drogę z punktu a do punktu l . Mamy:



1) do punktu b można dojść z punktu a najkrótszą drogą $a \rightarrow b$, a do c najkrótszą drogą $a \rightarrow c$. Zatem punktowi b przyporządkowujemy liczbę 3, a punktowi c liczbę 2:



2) do punktu d najkrótsza droga wiedzie przez punkt b :

$$b \rightarrow d \quad 5,$$

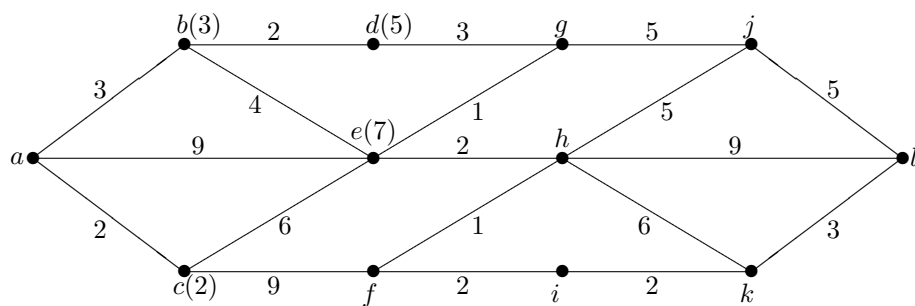
a do punktu e bezpośrednia droga nie jest najkrótsza, bo mamy:

$$a \rightarrow e \quad 9,$$

$$b \rightarrow e \quad 7,$$

$$c \rightarrow e \quad 8.$$

Stąd punktowi e przyporządkowujemy liczbę 7 i mamy:

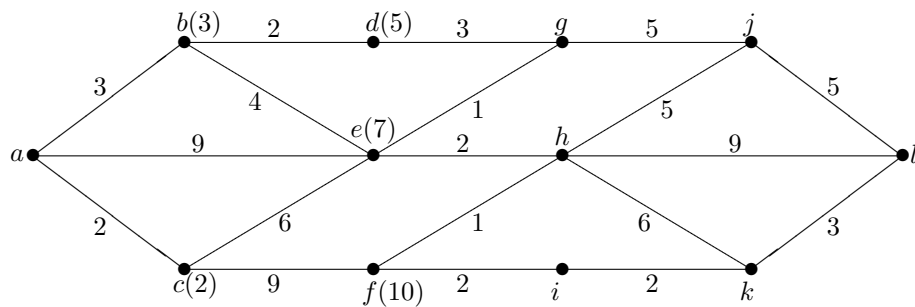


3) do punktu f można dojść z punktów c lub e drogami:

$$c \rightarrow f \quad 11,$$

$$e \rightarrow h \rightarrow f \quad 10.$$

Punktowi f przyporządkowujemy więc liczbę 10:

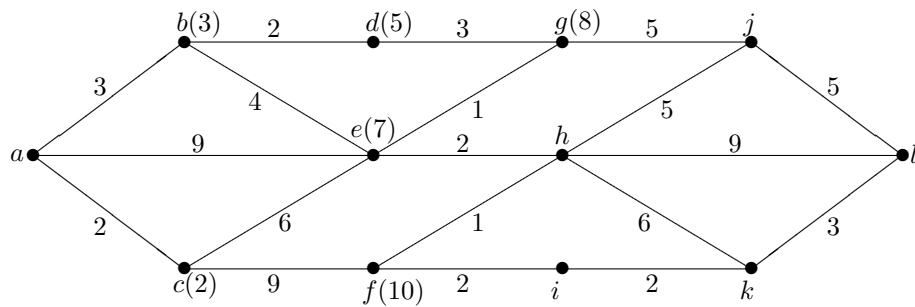


4) do punktu g można dojść z punktów d lub e drogami:

$$d \rightarrow g \quad 8,$$

$$e \rightarrow g \quad 8.$$

Zatem punktowi g przyporządkowujemy liczbę 8:

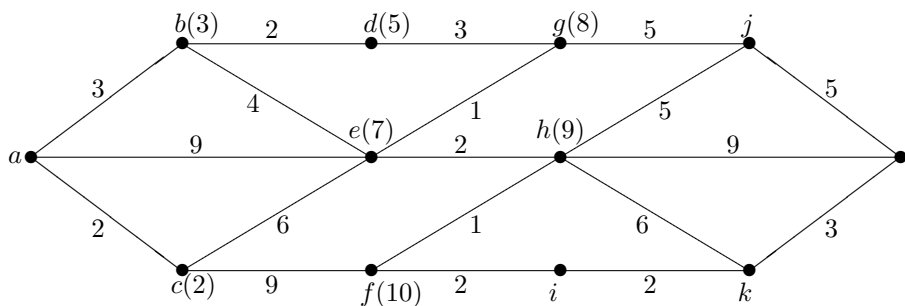


5) do punktu h można dojść z punktów e lub f drogami:

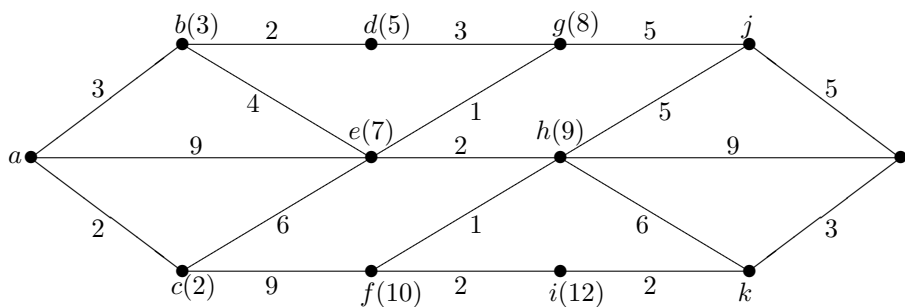
$$e \rightarrow h \quad 9,$$

$$f \rightarrow h \quad 11.$$

Przyporządkowujemy punktowi h liczbę 9:



6) do punktu i najkrótsza droga wiedzie przez f i ma wagę 12. Stąd:



7) do punktu j można dojść z punktów g lub h drogami:

$$g \rightarrow j \quad 13,$$

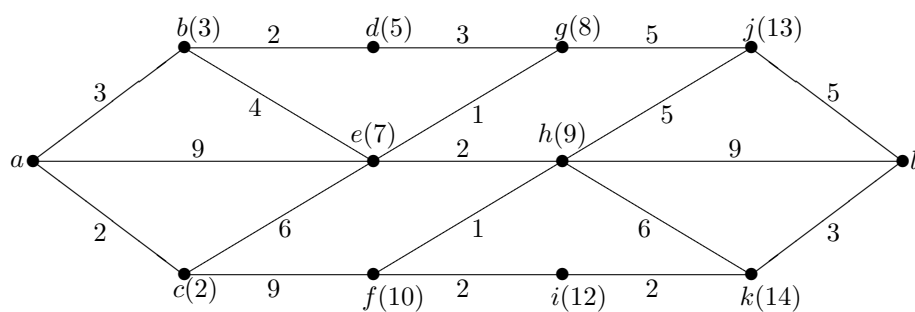
$$h \rightarrow j \quad 14,$$

a do punktu j można dojść z punktów h lub i drogami:

$$h \rightarrow k \quad 15,$$

$$i \rightarrow k \quad 14.$$

Stąd



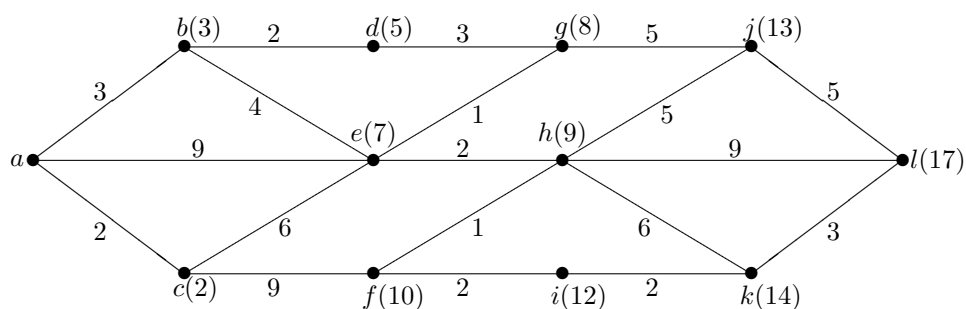
8) do punktu l można dojść z punktów h , j lub k drogami:

$$h \rightarrow l \quad 18,$$

$$j \rightarrow l \quad 18,$$

$$k \rightarrow l \quad 17.$$

Przyporządkowujemy punktowi l liczbę 17:

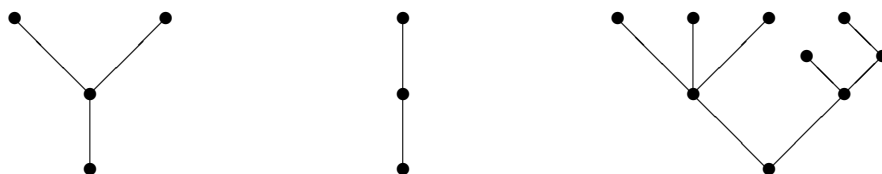


Realizując opisany algorytm przeszliśmy jednocześnie najkrótszą drogę z a do l . Możemy ją zapisać:

$$a \rightarrow b \rightarrow e \rightarrow h \rightarrow f \rightarrow i \rightarrow k \rightarrow l.$$

Definicja. *Drzewem* nazywa się każdy graf spójny nie zawierający cykli.

Przykład. Poniższy rysunek przedstawia przykładowe drzewa:



Twierdzenie. Niech T będzie grafem mającym n wierzchołków. Wtedy następujące warunki są równoważne:

- 1) T jest drzewem,
- 2) T nie zawiera cykli i ma $n - 1$ krawędzi,
- 3) T jest grafem spójnym i ma $n - 1$ krawędzi,
- 4) T jest grafem spójnym i każda krawędź jest mostem,
- 5) każde dwa wierzchołki grafu T są połączone dokładnie jedną drogą,
- 6) T nie zawiera cykli, ale po dodaniu dowolnej nowej krawędzi otrzymany graf ma dokładnie jeden cykl.

Dowód. Jeśli $n = 1$, to twierdzenie jest trywialne. Załóżmy, że $n \geq 2$.

(1) $\Rightarrow$ (2). Patrz lemat przed twierdzeniem Eulera.

(2) $\Rightarrow$ (3). Załóżmy, że graf T jest niespójny. Wtedy każda jego składowa jest grafem spójnym bez cykli, czyli, na mocy (2), w każdej składowej liczba wierzchołków jest o jeden większa od liczby krawędzi. Stąd całkowita liczba wierzchołków grafu T jest większa od liczby krawędzi o co najmniej dwa, co przeczy założeniu, że graf T ma $n - 1$ krawędzi.

(3) $\Rightarrow$ (4). Usunięcie którejkolwiek krawędzi powoduje powstanie grafu mającego n wierzchołków i $n - 2$ krawędzi. Wiemy, że liczba krawędzi m grafu prostego mającego n wierzchołków i k składowych spełnia nierówność

$$m \geq n - k.$$

Stąd nasz graf ma dwie składowe, czyli jest niespójny. Zatem każda krawędź grafu T jest mostem.

(4) $\Rightarrow$ (5). Ponieważ graf T jest spójny, więc każde dwa wierzchołki są połączone co najmniej jedną drogą. Gdyby istniała para wierzchołków, które łączą dwie drogi, to drogi te tworzyłyby cykl, czyli nie każda krawędź byłaby mostem.

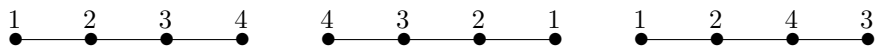
(5) $\Rightarrow$ (6). Gdyby graf T zawierał cykl, to każde dwa wierzchołki w tym grafie byłyby połączone co najmniej dwiema drogami, co przeczy założeniu. Jeśli dodamy krawędź e do grafu T , to zostanie utworzony cykl, bo wierzchołki incydentne z e były już połączone w grafie T . Gdyby powstały dwa takie cykle zawierające krawędź e , to w grafie byłby też cykl niezawierający krawędzi e , co przeczy pierwszej części punktu (6).

(6) $\Rightarrow$ (1). Na mocy punktu (6) graf T nie zawiera cykli. Wystarczy więc pokazać, że T jest spójny. Załóżmy, że graf T jest niespójny, czyli ma co najmniej dwie składowe. Dodajmy więc krawędź łączącą wierzchołek należący do jednej składowej z wierzchołkiem należącym do innej składowej. Oczywiście nie powstanie wtedy cykl, co przeczy założeniu. Stąd graf T jest spójny. Zatem T jest drzewem. $\square$

Przypomnijmy, że *grafem oznakowanym* nazywa się graf, który ma oznakowane (np. ponumerowane) wierzchołki. Dwa grafy oznakowane (ponumerowane) będą *izomorficzne*,

gdy dwa wierzchołki jednakowo ponumerowane w obydwu grafach (i mające ten sam stopień) będą na siebie przechodziły przy tym izomorfizmie. Zajmiemy się teraz zagadnieniem zliczania różnych (nieizomorficznych) drzew oznakowanych.

Przykład. Na poniższym rysunku pokazano trzy sposoby oznakowania drzewa mającego cztery wierzchołki:



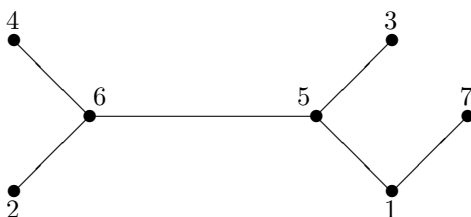
Zauważmy, że dwa pierwsze drzewa są izomorficzne oraz żadne z nich nie jest izomorficzne z trzecim drzewem (porównaj stopnie wierzchołka 3).

Twierdzenie. (Cayley) Istnieje n^{n-2} różnych drzew oznakowanych mających n wierzchołków.

Uwaga. Jeden z dowodów twierdzenia Cayleya, pochodzący od Prüfera, polega na wskazaniu bijekcji między zbiorem drzew oznakowanych mających n wierzchołków i zbiorem ciągów skończonych postaci $(a_1, a_2, \dots, a_{n-2})$ takich, że $1 \leq a_i \leq n$ dla $i = 1, 2, \dots, n-2$. Ciąg taki nazywa się *kod Prüfera* danego drzewa oznakowanego. Weźmy drzewo oznakowane T mające n wierzchołków. Pokażemy jak określić odpowiadający mu kod Prüfera. Niech b_1 będzie najmniejszą liczbą przypisaną wierzchołkowi końcowemu i niech a_1 będzie nazwą wierzchołka, z którym wierzchołek b_1 sąsiaduje. Usuwamy teraz wierzchołek b_1 i krawędź incydentną z nim, otrzymując drzewo mające $n-1$ wierzchołków. Niech następnie b_2 będzie najmniejszą liczbą będącą nazwą wierzchołka końcowego tego nowego drzewa i niech a_2 będzie nazwą wierzchołka sąsiadującego z wierzchołkiem b_2 . Usuwamy, tak jak poprzednio, wierzchołek b_2 i krawędź incydentną z nim. Postępujemy w ten sposób dotąd aż pozostaną tylko dwa wierzchołki i ostatecznie otrzymamy ciąg $(a_1, a_2, \dots, a_{n-2})$.

Aby teraz otrzymać przyporządkowanie odwrotne, weźmy ciąg $(a_1, a_2, \dots, a_{n-2})$. Niech b_1 będzie najmniejszą liczbą naturalną, która nie występuje w tym ciągu. Otrzymujemy pierwszą krawędź $a_1 b_1$. Usuwamy teraz liczbę a_1 z naszego ciągu, pomijamy liczbę b_1 w dalszych rozważaniach i kontynuujemy to postępowanie. W ten sposób budujemy drzewo krok po kroku, dodając po jednej krawędzi.

Przykład. Wyznamy kod Prüfera drzewa przedstawionego na poniższym rysunku:



Mamy:

$$b_1 = 2, a_1 = 6,$$

$$b_2 = 3, a_2 = 5,$$

$$b_3 = 4, a_3 = 6,$$

$$b_4 = 6, a_4 = 5,$$

$$b_5 = 5, a_5 = 1.$$

Zostały już tylko dwa wierzchołki, więc to kończy procedurę. Otrzymujemy ciąg $(6, 5, 6, 5, 1)$.

Jeśli teraz weźmiemy ciąg $(6, 5, 6, 5, 1)$, to dla liczb naturalnych $1, 2, 3, 4, 5, 6, 7$ otrzymujemy kolejno

$$b_1 = 2, b_2 = 3, b_3 = 4, b_4 = 6, b_5 = 5$$

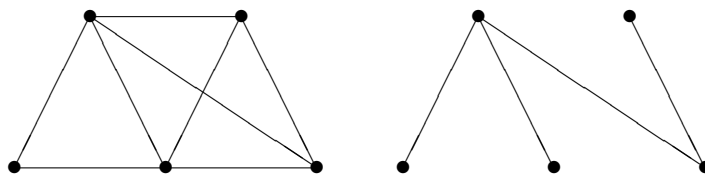
i odpowiednimi krawędziami będą:

$$62, 53, 64, 56, 15 \text{ oraz } 17.$$

Definicja. Drzewo $T = (V_T, E_T)$ nazywa się *drzewo spinające* grafu spójnego $G = (V_G, E_G)$, jeśli T jest podgrafem grafu G oraz $V_T = V_G$.

Uwaga. Zauważmy, że aby otrzymać drzewo spinające danego grafu spójnego G należy usunąć dowolną jedną krawędź z każdego cyklu grafu G tak, żeby w G nie było już cykli.

Przykład. Poniższy rysunek przedstawia graf i jedno z jego drzew spinających:



Twierdzenie. Jeśli T jest drzewem spinającym grafu spójnego G , to

- 1) każde rozcięcie grafu G ma wspólną krawędź z T ,
- 2) każdy cykl w grafie G ma wspólną krawędź z dopełnieniem T .

Dowód. 1) Niech S będzie rozcięciem grafu G , czyli zbiorem krawędzi, których usunięcie powoduje rozpad grafu G na dwa podgrafy G_1 i G_2 . Ponieważ T jest drzewem spinającym grafu G , więc T musi zawierać krawędź łączącą pewien wierzchołek należący do G_1 z wierzchołkiem należącym do G_2 ; jest to szukana krawędź.

2) Niech C będzie cyklem w grafie G niemającym wspólnych krawędzi z dopełnieniem drzewa T . Wtedy cykl C musi być zawarty w drzewie T , co jest niemożliwe. $\square$

Na mocy twierdzenia Cayleya otrzymuje się natychmiast następujące twierdzenie.

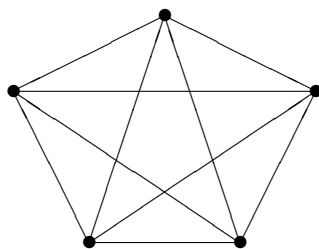
Twierdzenie. Graf pełny K_n ma n^{n-2} drzew spinających.

Następne twierdzenie można wykorzystać do obliczania liczby drzew spinających dowolnego spójnego grafu prostego.

Twierdzenie. (Macierzowe twierdzenie o drzewach) Niech G będzie spójnym grafem prostym, którego zbiorem wierzchołków jest $\{v_1, v_2, \dots, v_n\}$ i niech $M = [m_{ij}]_{n \times n}$ będzie macierzą taką, że $m_{ii} = d(v_i)$, $m_{ij} = -1$, gdy wierzchołki v_i i v_j są sąsiednie, oraz $m_{ij} = 0$ w przeciwnym przypadku. Wtedy liczba drzew spinających grafu G jest równa dopełnieniu algebraicznemu dowolnego elementu macierzy M .

(bez dowodu)

Przykład. Weźmy graf pełny K_5 :



Z twierdzenia wiemy, że ma on $5^{5-2} = 125$ drzew spinających. Zauważmy, że na mocy macierzowego twierdzenia o drzewach macierz M ma następującą postać:

$$M = \begin{bmatrix} 4 & -1 & -1 & -1 & -1 \\ -1 & 4 & -1 & -1 & -1 \\ -1 & -1 & 4 & -1 & -1 \\ -1 & -1 & -1 & 4 & -1 \\ -1 & -1 & -1 & -1 & 4 \end{bmatrix}$$

Wyznamy dla przykładu dopełnienie algebraiczne drugiego elementu z pierwszego wiersza:

$$\begin{aligned}
M_{12} &= (-1)^3 \cdot \begin{vmatrix} -1 & -1 & -1 & -1 \\ -1 & 4 & -1 & -1 \\ -1 & -1 & 4 & -1 \\ -1 & -1 & -1 & 4 \end{vmatrix} \stackrel{w_2=w_1}{=} \stackrel{w_3=w_1}{=} \stackrel{w_4=w_1}{=} -1 \cdot \begin{vmatrix} -1 & -1 & -1 & -1 \\ 0 & 5 & 0 & 0 \\ 0 & 0 & 5 & 0 \\ 0 & 0 & 0 & 5 \end{vmatrix} \\
&= -1 \cdot (-1) \cdot 5 \cdot 5 \cdot 5 = 125.
\end{aligned}$$

Nietrudno jest sprawdzić, że dopełnienie algebraiczne dowolnego elementu macierzy M jest równe 125. Zatem, na mocy macierzowego twierdzenia o drzewach, graf K_5 faktycznie ma 125 drzew spinających.

Następne twierdzenie podaje procedurę znajdowania drzewa spinającego T danego grafu z wagami tak, aby całkowita waga $w(T)$ była jak najmniejsza. Procedura ta nazywa się *algorytm zachłanny*. Algorytm zachłanny może być użyty do rozwiązania na przykład zagadnienia najkrótszej drogi.

Twierdzenie. (Algorytm zachłanny) Niech G będzie grafem z wagami mającym n wierzchołków. Wtedy następująca konstrukcja daje w wyniku drzewo spinające grafu G o najmniejszej całkowitej wadze:

- 1) *wybierz krawędź e_1 o najmniejszej wadze,*
- 2) *definiuj krawędzie $e_2, e_3, \dots, e_{n-1}$ wybierając za każdym razem nową krawędź o najmniejszej możliwej wadze, która nie tworzy cyklu z dotychczas wybranymi krawędziami e_i .*

Podgraf T grafu G , którego krawędziami są $e_1, e_2, \dots, e_{n-1}$, jest szukany drzewem spinającym.

Dowód. Graf T ma n wierzchołków, $n-1$ krawędzi i nie zawiera cykli. Stąd T jest drzewem spinającym grafu G . Wystarczy jeszcze pokazać, że T ma minimalną wagę całkowitą. Załóżmy więc, że T' jest drzewem spinającym grafu G takim, że

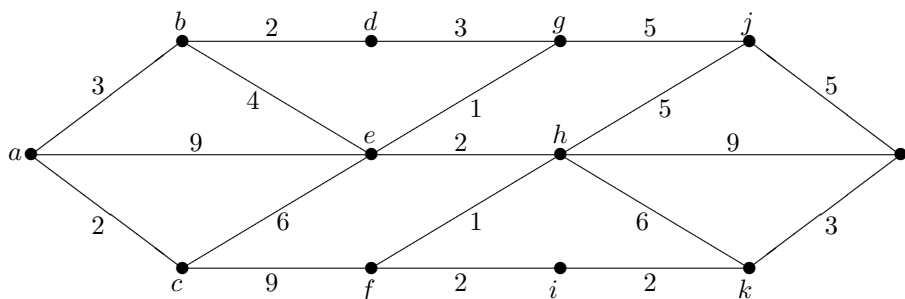
$$w(T') < w(T).$$

Niech e_k będzie pierwszą krawędzią drzewa T , która nie należy do T' . Wtedy podgraf grafu G utworzony z drzewa T' przez dodanie krawędzi e_k zawiera dokładnie jeden cykl C przechodzący przez e_k . Ponieważ cykl C zawiera krawędź e należącą do T' i nienależącą do T , więc podgraf T'' otrzymany z T' przez zastąpienie krawędzi e krawędzią e_k jest drzewem spinającym. Ale z konstrukcji ciągu krawędzi wynika, że $w(e_k) \leq w(e)$, a więc $w(T'') \leq w(T')$. Drzewo T'' ma jedną więcej wspólną krawędź z drzewem T . Wynika stąd, że powtarzając to postępowanie możemy przekształcić drzewo T' w drzewo T wymieniając po jednej krawędzi i niepowiększając łącznej wagi drzewa w żadnym kroku. Zatem

$$w(T) \leq w(T'),$$

co przeczy założeniu. $\square$

Przykład. Wyznamy drzewo spinające grafu z wagami



o najmniejszej możliwej wadze. Drzewo spinające naszego grafu ma 12 wierzchołków i 11 krawędzi. Jako krawędź e_1 możemy wziąć krawędzie eg lub fh , bo mają najniższą wagę równą 1. Nie ma znaczenia którą weźmiemy jako e_1 , bo tę drugą weźmiemy jako e_2 . Niech więc

$$e_1 = eg, \quad e_2 = fh.$$

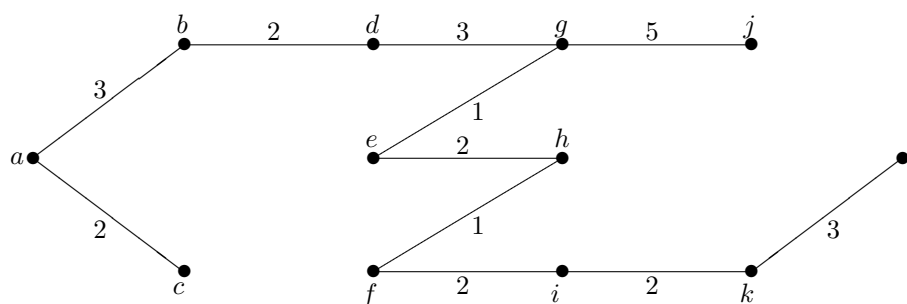
Następnie bierzemy pod uwagę krawędzie ac , bd , eh , fi oraz ik o kolejnej najniższej wadze równej 2. Ponieważ wybierając te krawędzie w żadnym momencie nie utworzymy cyklu, więc przyjmujemy

$$e_3 = ac, \quad e_4 = bd, \quad e_5 = eh, \quad e_6 = fi, \quad e_7 = ik.$$

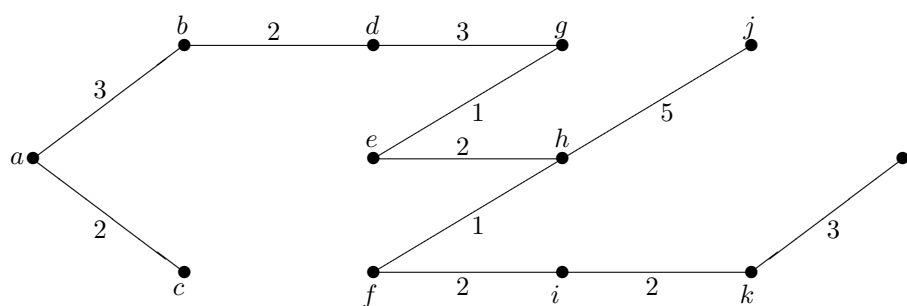
Wagę równą 3 mają krawędzie ab , dg i kl . Wybranie tych krawędzi nie utworzy cyklu, więc mamy

$$e_8 = ab, \quad e_9 = dg, \quad e_{10} = kl.$$

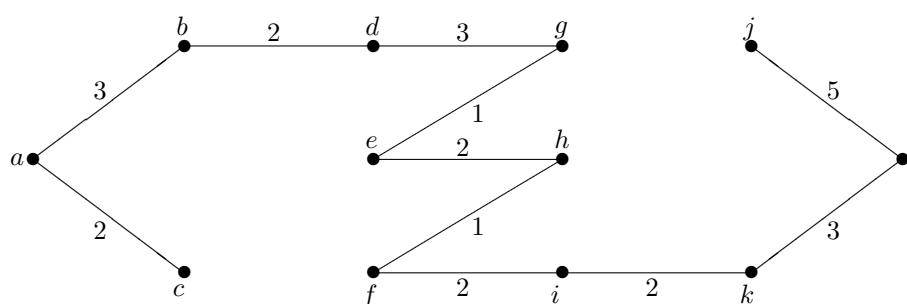
Pozostaje jeszcze wybrać jedną krawędź. Nie może to być krawędź be o wadze równej 4, bo utworzy ona cykl z już wybranymi krawędziami bd , dg i eg . Stąd jako e_{11} możemy wybrać krawędzie gj , hj lub jl o wadze równej 5. Zatem nasz graf z wagami ma trzy drzewa spinające o najmniejszej wadze równej 26:



lub

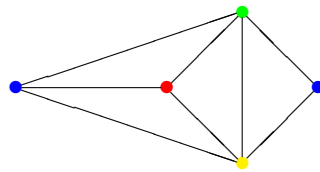


lub



Definicja. Graf bez pętli G jest k -kolorowalny, jeśli każdemu wierzchołkowi można przypisać jeden z k kolorów w taki sposób, żeby sąsiednie wierzchołki miały różne kolory. Jeżeli graf G jest k -kolorowalny, ale nie jest $(k - 1)$ -kolorowalny, to mówi się, że jest on k -chromatyczny lub że jego liczba chromatyczna jest równa k ; pisze się wtedy $\chi(G) = k$.

Przykład. Poniższy rysunek przedstawia graf G , dla którego $\chi(G) = 4$:



Graf ten jest zatem k -kolorowalny dla $k \geq 4$.

Uwagi.

1. Zakładamy, że grafy, którymi się zajmujemy są grafami prostymi oraz spójnymi.
2. Oczywiście jest, że $\chi(K_n) = n$ oraz że istnieją grafy mające dowolnie duże liczby chromatyczne.
3. $\chi(G) = 1$ wtedy i tylko wtedy, gdy G jest grafem pustym.
4. $\chi(G) = 2$ wtedy i tylko wtedy, gdy G jest niepustym grafem dwudzielnym.
5. Każde drzewo jest 2-kolorowalne, podobnie graf cykliczny C_n , gdy n jest parzyste.
6. Graf Petersena oraz dla n nieparzystego grafy C_n i W_n są 3-chromatyczne.

Twierdzenie. Jeśli G jest grafem prostym, w którym największym stopniem wierzchołka jest Δ , to graf G jest $(\Delta + 1)$ -kolorowalny.

Dowód. Dowód prowadzimy przez indukcję względem liczby wierzchołków grafu G . Jeśli graf G ma 1 lub 2 wierzchołki, to twierdzenie jest oczywiste. Załóżmy, że twierdzenie jest prawdziwe dla każdej liczby wierzchołków $k \leq n - 1$. Niech graf prosty G ma n wierzchołków. Jeśli usuniemy z niego dowolny wierzchołek v wraz z krawędziami z nim incydującymi, to otrzymany graf będzie grafem prostym, który będzie miał $n - 1$ wierzchołków oraz największy stopień wierzchołka będzie nie większy niż Δ . Z założenia indukcyjnego wynika, że ten graf jest $(\Delta + 1)$ -kolorowalny. Wtedy $(\Delta + 1)$ -kolorowanie grafu G otrzymujemy nadając wierzchołkowi v kolor inny niż kolory (co najwyżej Δ) wierzchołków sąsiadujących z v . $\square$

Zachodzi również nieco silniejsze twierdzenie.

Twierdzenie. (Brooks) Jeśli G jest spójnym grafem prostym niebędącym grafem pełnym i jeśli największy stopień wierzchołka grafu G wynosi Δ , gdzie $\Delta \geq 3$, to graf G jest Δ -kolorowalny.

(bez dowodu)

Ograniczmy teraz nasze rozważania do grafów planarnych. Łatwo pokazuje się, że każdy planarny graf prosty jest 6-kolorowalny. Bardziej wnikliwe rozumowanie pokazuje,

że taki graf jest nawet 5-kolorowalny. Mamy natomiast następujące twierdzenie o czterech barwach.

Twierdzenie. (O czterech barwach dla grafów planarnych, Appel i Haken)

Każdy planarny graf prosty jest 4-kolorowalny.

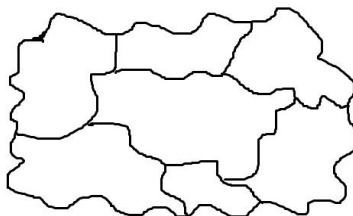
(bez dowodu)

Zajmiemy się teraz zagadnieniem kolorowania ścian grafu planarnego.

Definicja. *Mapą* nazywa się graf planarny 3-spójny.

Uwaga. Bezpośrednio z definicji wynika, że mapa nie zawiera rozcięć mających 1 lub 2 krawędzie, a więc w szczególności nie ma wierzchołków stopnia 1 lub 2.

Przykład. Poniższy rysunek przedstawia przykład mapy:

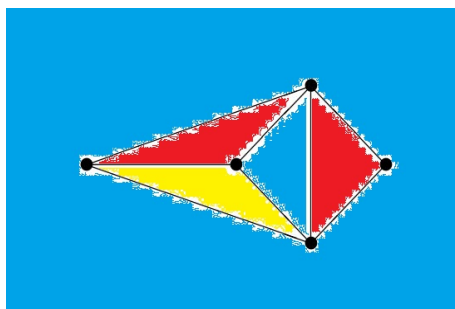


Uwaga. Pojęcie mapy zgadza się z potocznym rozumieniem mapy przedstawiającej państwa.

Definicja. Mapa jest k -kolorowalna(f), jeśli jej ściany można pokolorować k kolorami w taki sposób, żeby żadne dwie ściany ograniczone wspólną krawędzią nie były pokolorowane tym samym kolorem.

Uwaga. Zagadnienie polega więc na takim pokolorowaniu mapy przedstawiającej wiele państw, żeby żadne dwa państwa mające wspólną granicę nie były pokolorowane tym samym kolorem.

Przykład. Poniższy rysunek przedstawia mapę, która jest 3-kolorowalna(f):



Twierdzenie. (O czterech barwach dla map)

Każda mapa jest 4-kolorowalna(f).

(bez dowodu)

Pokazuje się, że obie postacie twierdzenia o czterech barwach są równoważne.

Twierdzenie. Twierdzenie o czterech barwach dla map jest równoważne z twierdzeniem o czterech barwach dla grafów planarnych.

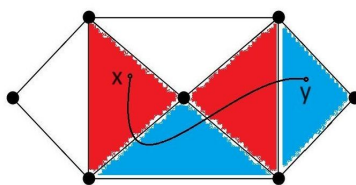
(bez dowodu)

Podamy teraz warunki, w których mapę można pokolorować 2 lub 3 kolorami.

Twierdzenie. Mapa G jest 2-kolorowalna(f) wtedy i tylko wtedy, gdy graf G jest grafem eulerowskim.

Dowód. Załóżmy najpierw, że mapa G jest 2-kolorowalna(f). Dla każdego wierzchołka grafu G liczba ścian otaczających ten wierzchołek musi być parzysta, ponieważ mogą one być pokolorowane dwoma kolorami. Wynika stąd, że każdy wierzchołek ma stopień parzysty. Zatem, na mocy twierdzenia Eulera, graf G jest grafem eulerowskim.

Założmy teraz, że graf G jest grafem eulerowskim. Ściany grafu G kolorujemy dwoma kolorami w następujący sposób. Wybieramy dowolną ścianę F i kolorujemy ją na czerwono. Rysujemy krzywe z pewnego punktu x na ścianie F do pewnych punktów na wszystkich pozostałych ścianach, nieprzechodzące przez żaden wierzchołek grafu G . Jeśli taka krzywa przecina parzystą liczbę krawędzi, to ścianę, do której dochodzi kolorujemy na czerwono; w przeciwnym razie na niebiesko. Ilustruje to poniższy rysunek:



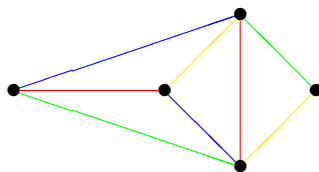
Ten sposób kolorowania jest dobrze określony, o czym można się przekonać wybierając „cykl” złożony z dwóch takich krzywych i dowodząc, że przecina on parzystą liczbę krawędzi grafu G – korzystamy z tego, że każdy wierzchołek jest incydentny z parzystą liczbą krawędzi. $\square$

Twierdzenie. Niech G będzie mapą kubiczną. Wówczas mapa G jest 3-kolorowalna(f) wtedy i tylko wtedy, gdy każda jej ściana jest ograniczona parzystą liczbą krawędzi.

(bez dowodu)

Definicja. Graf G jest k -kolorowalny(e) (lub k -kolorowalny krawędziowo), jeśli krawędzie można pokolorować k kolorami w taki sposób, żeby żadne dwie sąsiednie krawędzie nie miały tego samego koloru. Jeżeli graf G jest k -kolorowalny(e), ale nie jest $(k - 1)$ -kolorowalny(e), to mówi się, że jego *indeks chromatyczny* wynosi k i pisze się $\chi'(G) = k$.

Przykład. Poniższy rysunek przedstawia graf G , dla którego $\chi'(G) = 4$:



Uwagi.

1. $\chi'(C_n) = 2$ dla n parzystych oraz $\chi'(C_n) = 3$ dla n nieparzystych.
2. $\chi'(W_n) = n - 1$ dla $n \geq 4$.
3. $\chi'(K_n) = n$ dla n nieparzystych ($n \neq 1$) oraz $\chi'(K_n) = n - 1$ dla n parzystych.

Następne twierdzenie podaje bardzo dobre oszacowanie indeksu chromatycznego grafu prostego.

Twierdzenie. (Vizing) Jeśli G jest grafem prostym, w którym największy stopień wierzchołka wynosi Δ , to

$$\Delta \leq \chi'(G) \leq \Delta + 1.$$

(bez dowodu)

Poniższe twierdzenie podaje związek między twierdzeniem o czterech barwach i kolorowaniu krawędzi grafu.

Twierdzenie. Twierdzenie o czterech barwach jest równoważne stwierdzeniu, że $\chi'(G) = 3$ dla każdej mapy kubicznej G .

(bez dowodu)

Na koniec sformułujemy twierdzenie o indeksie chromatycznym grafu dwudzielnego.

Twierdzenie. (König) Jeśli w grafie dwudzielnym G największy stopień wierzchołka wynosi Δ , to $\chi'(G) = \Delta$.

(bez dowodu)

Wniosek. $\chi'(K_{m,n}) = \max(m, n)$.

9. Algebry Boole'a i funkcje booleowskie

Definicja. Niech A będzie dowolnym niepustym zbiorem. *Działaniem n -argumentowym* w zbiorze A nazywa się dowolną funkcję $f : A^n \rightarrow A$.

Uwaga. Interesują nas działania n -argumentowe tylko dla $n = 0, 1, 2$, przy czym działanie 0-argumentowe określa pewien ustalony element zbioru A nazywany elementem wyróżnionym.

Przykład. W zbiorze liczb rzeczywistych $\mathbb{R}$

- 1) działaniem 0-argumentowym może być ustalona liczba rzeczywista, na przykład 0 lub 1,
- 2) działaniem jednoargumentowym może być przyporządkowanie liczbie rzeczywistej liczby do niej przeciwnej,
- 3) działaniem dwuargumentowym może być przyporządkowanie dwu liczbom rzeczywistym ich sumy.

Definicja. *Algebrą Boole'a* nazywa się zbiór A z dwoma działaniami dwuargumentowymi $\vee$ i $\wedge$, działaniem jednoargumentowym $'$ oraz różnymi elementami 0 i 1, spełniającymi następujące prawa:

- 1) $x \vee y = y \vee x$ i $x \wedge y = y \wedge x$ (prawa przemienności),
- 2) $(x \vee y) \vee z = x \vee (y \vee z)$ i $(x \wedge y) \wedge z = x \wedge (y \wedge z)$ (prawa łączności),
- 3) $x \vee (y \wedge z) = (x \vee y) \wedge (x \vee z)$ i $x \wedge (y \vee z) = (x \wedge y) \vee (x \wedge z)$ (prawa rozdzielności),
- 4) $x \vee 0 = x$ i $x \wedge 1 = x$ (prawa identyczności),
- 5) $x \vee x' = 1$ i $x \wedge x' = 0$ (prawa dopełnienia).

Działanie $\vee$ nazywa się *suma*, $\wedge$ – *iloczyn*, a działanie jednoargumentowe $'$ – *dopełnienie*.

Przykłady.

1. Zbiór $\mathcal{P}(X)$ wszystkich podzbiorów niepustego zbioru X wraz z działaniami sumy $\cup$, iloczynu $\cap$ i dopełnienia $'$ jest algebrą Boole'a, w której zerem jest $\emptyset$, a jednością zbiór X .
2. Zbiór $\mathbb{B} = \{0, 1\}$ z działaniami logicznymi $\vee$, $\wedge$ i $'$ traktowanymi jako alternatywa, koniunkcja i negacja zdań o wartościach logicznych 0,1, jest algebrą Boole'a. Poniższe tablice podają opis działań $\vee$, $\wedge$ i $'$:

$\vee$	0	1	$\wedge$	0	1	$'$
0	0	0	0	0	0	1
1	1	1	1	0	1	0

3. Zbiór $\mathbb{B}^n = \mathbb{B} \times \mathbb{B} \times \dots \times \mathbb{B}$ (z n czynnikami $\mathbb{B}$) ciągów zerojedynekowych długości n wraz z działaniami booleowskimi $\vee$, $\wedge$ i $'$ pochodzącymi od odpowiednich działań w zbiorze $\mathbb{B}$, jest algebrą Boole'a. Działania $\vee$, $\wedge$ i $'$ są wtedy zdefiniowane po współrzędnych. Na przykład:

$$(a_1, a_2, \dots, a_n) \vee (b_1, b_2, \dots, b_n) = (a_1 \vee b_1, a_2 \vee b_2, \dots, a_n \vee b_n).$$

Twierdzenie. Następujące własności zachodzą w każdej algebrze Boole'a:

- 1) $x \vee x = x$ i $x \wedge x = x$ (prawa idempotentności),
- 2) $x \vee 1 = 1$ i $x \wedge 0 = 0$ (następne prawa identyczności),
- 3) $(x \wedge y) \vee x = x$ i $(x \vee y) \wedge x = x$ (prawa pochłaniania),
- 4) $(x \vee y)' = x' \wedge y'$ i $(x \wedge y)' = x' \vee y'$ (prawa De Morgana).

Dowód. 1) Mamy

$$x = x \vee 0 = x \vee (x \wedge x') = (x \vee x) \wedge (x \vee x') = (x \vee x) \wedge 1 = x \vee x.$$

Podobnie pokazuje się drugie prawo idempotentności.

2) Łatwo widzieć, że

$$x \vee 1 = x \vee (x \vee x') = (x \vee x) \vee x' = x \vee x' = 1.$$

Podobnie pokazuje się drugie prawo identyczności.

3) Teraz mamy

$$(x \wedge y) \vee x = (x \wedge y) \vee (x \wedge 1) = x \wedge (y \vee 1) = x \wedge 1 = x.$$

Podobnie pokazuje się drugie prawo pochłaniania.

4) Zauważmy, że

$$\begin{aligned} (x \vee y) \vee (x' \wedge y') &= ((x \vee y) \vee x') \wedge ((x \vee y) \vee y') \\ &= ((x \vee x') \vee y) \wedge (x \vee (y \vee y')) \\ &= (1 \vee y) \wedge (x \vee 1) \\ &= 1 \wedge 1 = 1 \end{aligned}$$

oraz

$$\begin{aligned}
(x \vee y) \wedge (x' \wedge y') &= (x \wedge (x' \wedge y')) \vee (y \wedge (x' \wedge y')) \\
&= ((x \wedge x') \wedge y') \vee (x' \wedge (y \wedge y')) \\
&= (0 \wedge y') \vee (x' \wedge 0) \\
&= 0 \vee 0 = 0.
\end{aligned}$$

Na mocy powyższych równości mamy

$$\begin{aligned}
(x \vee y)' &= (x \vee y)' \vee 0 = (x \vee y)' \vee ((x \vee y) \wedge (x' \wedge y')) \\
&= ((x \vee y)' \vee (x \vee y)) \wedge ((x \vee y)' \vee (x' \wedge y')) \\
&= 1 \wedge ((x \vee y)' \vee (x' \wedge y')) \\
&= ((x \vee y) \vee (x' \wedge y')) \wedge ((x \vee y)' \vee (x' \wedge y')) \\
&= ((x \vee y) \wedge (x \vee y)') \vee (x' \wedge y') \\
&= 0 \vee (x' \wedge y') \\
&= x' \wedge y'.
\end{aligned}$$

Podobnie pokazuje się drugie prawo De Morgana. $\square$

Definicja. Dla dowolnych elementów x, y algebry Boole'a definiuje się relację $\leq$ następująco:

$$x \leq y \Leftrightarrow x \vee y = y.$$

Ponadto

$$x < y \Leftrightarrow x \leq y \text{ i } x \neq y.$$

Twierdzenie. Relacja $\leq$ określona powyżej jest relacją częściowego porządku w algebrze Boole'a.

Dowód. Relacja $\leq$ jest zwrotna, bo $x \vee x = x$. Załóżmy, że $x \leq y$ i $y \leq x$, czyli $x \vee y = y$ i $y \vee x = x$. Wobec przemienności działania $\vee$ otrzymujemy $x = y$. Zatem relacja $\leq$ jest antysymetryczna. Załóżmy teraz, że $x \leq y$ i $y \leq z$, czyli $x \vee y = y$ i $y \vee z = z$. Wtedy

$$x \vee z = x \vee (y \vee z) = (x \vee y) \vee z = y \vee z = z.$$

Stąd $x \leq z$, czyli relacja $\leq$ jest również przechodnia. $\square$

Uwaga. Łatwo widać na mocy praw pochłaniania, że

$$x \leq y \Leftrightarrow x \wedge y = y.$$

Definicja. *Atomem* w algebrze Boole'a nazywa się niezerowy element a , który nie może być przedstawiony w postaci sumy dwóch elementów różnych od niego.

Przykłady.

1. Atomami w algebrze $\mathcal{P}(X)$ są zbiory jednoelementowe $\{x\}$.
2. Jedynym atomem algebry $\mathbb{B}$ jest 1.
3. Atomami algebry $\mathbb{B}^n$ są te ciągi n -elementowe, w których dokładnie jeden wyraz jest równy 1, a pozostałe są równe 0.

Twierdzenie. Niezerowy element x algebry Boole'a A jest atomem wtedy i tylko wtedy, gdy nie istnieje element $y \in A$ taki, że $0 < y < x$.

Dowód. Niech $x \neq 0$ będzie atomem algebry Boole'a A . Załóżmy, że istnieje element $y \in A$ taki, że $0 < y < x$. Wtedy $x \vee y = x$ oraz $x \wedge y = y$. Rozpatrzmy element $x \wedge y'$. Gdyby $x \wedge y' = 0$, to

$$y = y \vee 0 = y \vee (x \wedge y') = (y \vee x) \wedge (y \vee y') = x \wedge 1 = x,$$

co jest niemożliwe. Ponadto, gdyby $x \wedge y' = x$, to

$$y = x \wedge y = (x \wedge y') \wedge y = x \wedge (y' \wedge y) = x \wedge 0 = 0,$$

co również jest niemożliwe. Zatem mamy $x \wedge y' \neq 0$, $x \wedge y' \neq x$ oraz $x = y \vee (x \wedge y')$. Stąd x jest sumą dwóch elementów różnych od x ; nie może więc być atomem i otrzymujemy sprzeczność.

Założmy teraz, że nie istnieje $y \in A$ taki, że $0 < y < x$. Przypuśćmy, że $x \neq 0$ nie jest atomem algebry A . Wtedy istnieją $a, b \in A$ takie, że $x = a \vee b$ i $a \neq x$ i $b \neq x$. Stąd $a \neq 0$. Z prawa pochłaniania, $a \wedge (a \vee b) = a$, czyli $a \leq a \vee b = x$. Stąd

$$0 < a < x$$

i mamy sprzeczność. Zatem x jest atomem algebry Boole'a A . $\square$

Twierdzenie. Niech B będzie skończoną algebrą Boole'a, której zbiorem atomów jest zbiór $A = \{a_1, a_2, \dots, a_n\}$. Każdy niezerowy element x algebry B może być przedstawiony w postaci sumy różnych atomów:

$$x = a_{i_1} \vee a_{i_2} \vee \dots \vee a_{i_k}.$$

Ponadto ten sposób przedstawienia jest jednoznaczny z dokładnością do kolejności atomów.

Dowód. Niech $x \in B$ i $x \neq 0$. Jeżeli $x \in A$, to twierdzenie jest oczywiste. Załóżmy więc, że x nie jest atomem algebry B . Przypuśćmy, że x nie jest sumą atomów. Wtedy

$$x = y \vee z,$$

gdzie $0 < y < x$, $0 < z < x$ oraz co najmniej jeden z elementów y i z nie jest sumą atomów. Powyższe rozumowanie możemy powtórzyć dla tego elementu y lub z , który nie jest sumą atomów. W konsekwencji otrzymujemy ciąg malejący

$$\cdots < x_2 < x_1 < x$$

elementów algebry B niebędących sumami atomów. Ciąg ten jest skończony, bo algebra B jest skończona, więc istnieje w tym ciągu element najmniejszy. Element ten jest atomem; z drugiej strony, wobec założenia dotyczącego x , nie jest on atomem. Otrzymaliśmy sprzeczność.

Pokażemy teraz jednoznaczność przedstawienia x w postaci sumy atomów. Załóżmy, że

$$x = y_1 \vee y_2 \vee \cdots \vee y_p$$

oraz

$$x = z_1 \vee z_2 \vee \cdots \vee z_q,$$

gdzie $y_1, y_2, \dots, y_p, z_1, z_2, \dots, z_q$ są atomami. Wtedy

$$x = x \wedge x = (y_1 \wedge z_1) \vee (y_1 \wedge z_2) \vee \cdots \vee (y_1 \wedge z_q) \vee \cdots \vee (y_p \wedge z_q).$$

Zauważmy, że

$$y_i \wedge z_j = 0 \text{ lub } y_i = z_j \text{ dla } i = 1, 2, \dots, p, j = 1, 2, \dots, q,$$

bo y_i oraz z_j są atomami. W ten sposób twierdzenie zostało udowodnione. $\square$

Definicja. Funkcję booleowską n -argumentową nazywa się funkcję

$$f : \mathbb{B}^n \rightarrow \mathbb{B}.$$

Zbiór wszystkich n -argumentowych funkcji booleowskich oznacza się przez $BOOL(n)$.

Uwaga. Funkcję booleowską można traktować jako kolumnę w macierzy logicznej. Poniżej

przykładowa matryca dla funkcji booleowskiej trójargumentowej:

x	y	z	f
0	0	0	1
0	0	1	1
0	1	0	0
0	1	1	1
1	0	0	0
1	0	1	0
1	1	0	0
1	1	1	1

Zauważmy przy tym, że możliwych funkcji trójargumentowych jest $2^8 = 256$, czyli $|BOOL(3)| = 256$.

Uwaga. Zauważmy, że zbiór $BOOL(n)$ z działaniami booleowskimi $\vee$, $\wedge$ i $'$ zdefiniowanymi dla $f, g \in BOOL(n)$ następująco:

$$\begin{aligned}(f \vee g)(x) &= f(x) \vee g(x), \\ (f \wedge g)(x) &= f(x) \wedge g(x), \\ f'(x) &= (f(x))'\end{aligned}$$

jest algebrą Boole'a. Zauważmy ponadto, że atomami w tej algebrze są te i tylko te funkcje booleowskie $f : \mathbb{B}^n \rightarrow \mathbb{B}$, które przyjmują wartość 1 w dokładnie jednym punkcie algebry $\mathbb{B}^n$.

Definicja. Wyrażenia booleowskie n zmiennych $x_1, x_2, \dots, x_n$ definiuje się rekurencyjnie w następujący sposób:

- symbole 0,1 oraz $x_1, x_2, \dots, x_n$ są wyrażeniami booleowskimi zmiennych $x_1, x_2, \dots, x_n$,
- jeżeli E_1 i E_2 są wyrażeniami booleowskimi zmiennych $x_1, x_2, \dots, x_n$, to $(E_1 \vee E_2)$, $(E_1 \wedge E_2)$ oraz E_1' są również wyrażeniami booleowskimi.

Uwaga. W praktyce opuszcza się zewnętrzne nawiasy oraz korzysta się z praw łączności.

Przykłady.

1. Przykładowe wyrażenia booleowskie trzech zmiennych:

$$(x \vee y) \wedge (x' \vee z) \wedge 1; \quad (x' \wedge z) \vee (x' \wedge y) \vee z'.$$

2. Przykładowe wyrażenie booleowskie n zmiennych:

$$(x_1 \wedge x_2 \wedge \cdots \wedge x_n) \vee (x'_1 \wedge x_2 \wedge \cdots \wedge x_n) \vee (x_1 \wedge x'_2 \wedge \cdots \wedge x_n).$$

Uwaga. Używanie obu symboli $\vee$ i $\wedge$ prowadzi do powstawania długich i nieczytelnych wyrażeń booleowskich, więc zazwyczaj zastępuje się spójnik $\wedge$ kropką lub pomija całkowicie.

Przykład. Pierwsze z wyrażeń booleowskich z poprzedniego przykładu przyjmuje postać:

$$(x \vee y)(x' \vee z)1.$$

Definicja. Jeśli E jest wyrażeniem booleowskim n zmiennych $x_1, x_2, \dots, x_n$, to E definiuje funkcję booleowską $f : \mathbb{B}^n \rightarrow \mathbb{B}$, której wartością dla argumentu $(a_1, a_2, \dots, a_n)$ jest element algebry $\mathbb{B}$ otrzymany przez zastąpienie w wyrażeniu E zmiennej x_i wartością a_i dla $i = 1, 2, \dots, n$.

Przykład. Wartości funkcji booleowskiej $f : \mathbb{B}^3 \rightarrow \mathbb{B}$, odpowiadające wyrażeniu $x'z \vee x'y \vee z'$ przedstawia poniższa tabela:

x	y	z	$x'z$	$x'y$	z'	$x'z \vee x'y \vee z'$
0	0	0	0	0	1	1
0	0	1	1	0	0	1
0	1	0	0	1	1	1
0	1	1	1	1	0	1
1	0	0	0	0	1	1
1	0	1	0	0	0	0
1	1	0	0	0	1	1
1	1	1	0	0	0	0

Uwaga. Często używa się wyrażeń booleowskich jako nazw funkcji booleowskich, które one definiują.

Przykład. Funkcja należąca do zbioru $BOOL(3)$ mająca nazwę $xy'z$ jest zdefiniowana wzorem $xy'z(a, b, c) = ab'c$ dla wszystkich $(a, b, c) \in \mathbb{B}^3$, a więc

$$xy'z(a, b, c) = \begin{cases} 1, & \text{jeśli } a = 1, b = 0, c = 1, \\ 0 & \text{w przeciwnym przypadku.} \end{cases}$$

Zauważmy przy tym, że ponieważ funkcja ta przyjmuje wartość 1 tylko w jednym punkcie algebry $\mathbb{B}^3$, więc jest atomem algebry Boole'a $BOOL(3)$. Pozostałymi siedmioma atomami

algebry Boole'a $BOOL(3)$ są

$$xyz, \quad xyz', \quad xy'z', \quad x'yz, \quad x'yz', \quad x'y'z \quad \text{oraz} \quad x'y'z'.$$

Definicja. Dwa wyrażenia booleowskie nazywa się *równoważne*, jeśli odpowiadające im funkcje booleowskie są takie same.

Przykład. Wyrażenia $x(y \vee z)$ i $(xy) \vee (xz)$ są równoważne (co łatwo można sprawdzić).

Uwaga. Pisze się $E = F$, jeśli dwa wyrażenia booleowskie E i F są równoważne.

Definicja. *Symbolem atomowym* nazywa się wyrażenie booleowskie składające się tylko z jednej zmiennej lub jej dopełnienia (np. x czy y').

Definicja. *Iloczynem minimalnym n zmiennych* nazywa się iloczyn dokładnie n symboli atomowych, z których każdy zawiera inną zmienną.

Uwaga. Każdy atom algebry $BOOL(n)$ odpowiada pewnemu iloczynowi minimalnemu.

Przykłady.

1. Wyrażenie $xy'z$ jest iloczynem minimalnym 3 zmiennych, bo jest atomem algebry Boole'a $BOOL(3)$.
2. Wyrażenia

$$x', \quad xz', \quad xyx'z$$

nie są iloczynami minimalnymi 3 zmiennych.

Wiemy, że każdy niezerowy element skończonej algebry Boole'a może być przedstawiony w postaci sumy różnych atomów, przy czym przedstawienie to jest jednoznaczne, z dokładnością do kolejności atomów. Ponieważ $BOOL(n)$ jest skończoną algebrą Boole'a, więc każde wyrażenie booleowskie n zmiennych jest równoważne z sumą różnych iloczynów minimalnych, przy czym takie przedstawienie w postaci sumy jest jednoznaczne, z dokładnością do porządku, w jakim są wypisane te iloczyny minimalne.

Definicja. *Postacią kanoniczną* wyrażenia booleowskiego E nazywa się sumę iloczynów minimalnych równoważną z E .

Przykład. Znajdziemy postać kanoniczną wyrażenia (3 zmiennych)

$$(x \vee yz')(yz)'.$$

Metoda I

Wypisujemy w tabelce wszystkie wartości funkcji booleowskiej, której odpowiada to wyrażenie:

x	y	z	$f(x, y, z) = (x \vee yz')(yz)'$
0	0	0	0
0	0	1	0
0	1	0	1
0	1	1	0
1	0	0	1
1	0	1	1
1	1	0	1
1	1	1	0

Bierzemy pod uwagę te wiersze, w których w ostatniej kolumnie jest 1. Patrzymy dla jakich wartości x, y, z mamy 1. Otrzymujemy postać kanoniczną:

$$x'yz' \vee xy'z' \vee xy'z \vee xyz'.$$

Metoda II

Zajmiemy się samym wyrażeniem i korzystając z praw algebry Boole'a spróbujemy doprowadzić je do postaci kanonicznej:

$$\begin{aligned}
 (x \vee yz')(yz)' &= (x \vee yz')(y' \vee z') && \text{z prawa De Morgana} \\
 &= (x(y' \vee z')) \vee ((yz')(y' \vee z')) && \text{z prawa rozdzielności} \\
 &= (xy' \vee xz') \vee (yz'y' \vee yz'z') && \text{z prawa rozdzielności} \\
 &= (xy' \vee xz') \vee (0 \vee yz') && \text{z równości } yy' = 0, z'z' = z' \\
 &= xy' \vee xz' \vee yz' && \text{z prawa łączności i własności zera} \\
 &= xy'1 \vee x1z' \vee 1yz' && \text{z prawa identyczności} \\
 &= xy'(z \vee z') \vee x(y \vee y')z' \vee (x \vee x')yz' && \text{z prawa dopełnienia} \\
 &= xy'z \vee xy'z' \vee xyz' \vee xy'z' \vee xyz' \vee x'yz' && \text{z prawa rozdzielności} \\
 &= xy'z \vee xy'z' \vee xyz' \vee x'yz' && \text{z prawa idempotentności}
 \end{aligned}$$

10. Elementy teorii automatów

Rozważać będziemy matematyczne modele skończonych maszyn sekwencyjnych. Skończone maszyny sekwencyjne to maszyny akceptujące skończone ciągi symboli wejściowych. W określonej jednostce czasu maszyna taka znajdować się może w jednym ze skończonej ilości stanów. Istnieje również skończony zbiór symboli wyjściowych. Stan maszyny, jak i symbol wyjściowy zależą od wprowadzonego symbolu wejściowego i stanu, w jakim była maszyna poprzednio. Maszyny cyfrowe i ich części są takimi skończonymi maszynami sekwencyjnymi.

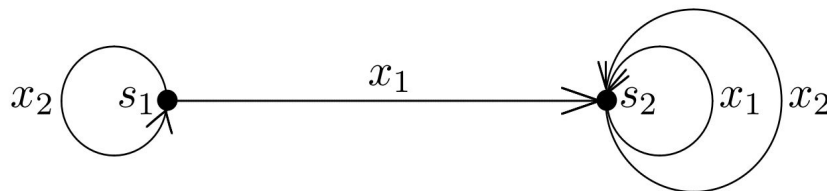
Definicja. *Półautomatem* nazywa się trójkę (S, X, δ) , gdzie $S = \{s_1, s_2, \dots, s_n\}$ jest skończonym zbiorem stanów, $X = \{x_1, x_2, \dots, x_k\}$ jest skończonym zbiorem niepustym zwanym *zbiorem symboli wejściowych*, a $\delta : S \times X \rightarrow S$ jest funkcją zwaną *funkcją zmiany stanów*.

Przykład. (Pułapka na myszy)

Niech $X = \{x_1, x_2\}$ i $S = \{s_1, s_2\}$. Symbol x_1 oznacza mysz w zasięgu pułapki, natomiast x_2 oznacza brak myszy w zasięgu pułapki. Pułapka jest w stanie s_1 , jeśli jest nastawiona, a w stanie s_2 , jeśli nie jest nastawiona. Funkcję zmiany stanów przedstawia poniższa tabelka:

δ	x_1	x_2
s_1	s_2	s_1
s_2	s_2	s_2

Działanie pułapki można też przedstawić za pomocą poniższego schematu:



Definicja. *Algebrą półautomatu* (S, X, δ) nazywa się algebrę $(S, (f_x)_{x \in X})$, gdzie $f_x : S \rightarrow S$ i $f_x(s) = \delta(s, x)$ dla każdego $x \in X$.

Przykład. Algebrą Pułapki na myszy jest algebra (S, f_{x_1}, f_{x_2}) , gdzie $f_{x_1}, f_{x_2} : S \rightarrow S$ oraz $f_{x_1}(s_1) = f_{x_1}(s_2) = s_2$ i $f_{x_2}(s_1) = s_1, f_{x_2}(s_2) = s_2$.

Definicja. Półautomat $A' = (S', X, \delta')$ nazywa się *podpółautomatem* półautomatu $A = (S, X, \delta)$ i pisze się $A' \subseteq A$, jeśli $S' \subseteq S$ i $\delta' = \delta|_{S' \times X}$.

Definicja. Półautomat (S, X, δ) nazywa się *iloczynem prostym* półautomatów (S', X', δ') i (S'', X'', δ'') , jeśli $S = S' \times S''$, $X = X' \times X''$ i $\delta : S' \times S'' \times X' \times X'' \rightarrow S' \times S''$ oraz

$$\delta(s', s'', x', x'') = (\delta'(s', x'), \delta''(s'', x'')) = (f'_{x'}(s'), f''_{x''}(s'')) =: f_{(x', x'')}(s', s'').$$

Uwaga. Iloczyn prosty dwóch półautomatów nazywa się również ich *układem równoległym*. W teorii automatów rozważa się również półautomaty zwane *układami szeregowymi* półautomatów. Mają one jednak znacznie naturalniejszą interpretację w przypadku automatów (zob. definicję poniżej).

Uwaga. W algebraicznej teorii automatów istnieją metody konstruowania półautomatów z układów równoległych i szeregowych półautomatów o szczególnie prostej postaci. Z drugiej strony istnieje twierdzenie (Krohn-Rodes) mówiące o tym, że każdy półautomat można „rozłożyć” używając tych dwóch konstrukcji na półautomaty takiej prostej postaci.

Uzupełniając półautomat o zbiór symboli wyjściowych otrzymuje się automat.

Definicja. *Automatem* nazywa się piątkę $(S, X, \delta, Y, \lambda)$ taką, że (S, X, δ) jest półautomatem, Y jest niepustym skończonym zbiorem zwanym *zbiorem symboli wyjściowych*, a $\lambda : S \times X \rightarrow Y$ jest funkcją zwaną *funkcją wyjścia*.

Uwaga. Układ równoległy automatów definiuje się podobnie jak układ równoległy półautomatów.

Definicja. *Układ szeregowy* $A_1 \parallel A_2$ automatów A_1 i A_2 jest to automat $(S_1 \times S_2, X_1, \delta, Y_2, \lambda)$, gdzie $X_2 = Y_1$ oraz

$$\begin{aligned}\delta((s_1, s_2), x_1) &= (\delta_1(s_1, x_1), \delta_2(s_2, \lambda_1(s_1, x_1))), \\ \lambda((s_1, s_2), x_1) &= \lambda_2(s_2, \lambda_1(s_1, x_1)).\end{aligned}$$

Uwaga. Wiele pojęć i własności dotyczących półautomatów można przenieść na automaty. Definiuje się również pojęcie podautomatu. Dużą rolę odgrywa pojęcie równoważności automatów (te same wejścia określają w obu automatach te same wyjścia). Jedno z ważnych twierdzeń teorii automatów mówi, że dla każdego automatu istnieje dokładnie jeden równoważny mu automat minimalny (tzn. z minimalną liczbą stanów).

Literatura

- [1] J. Grygiel, *Wprowadzenie do matematyki dyskretnej*, Akademicka Oficyna Wydawnicza EXIT, Warszawa 2007.
- [2] H. Rasiowa, *Wstęp do matematyki współczesnej*, PWN, Warszawa 2013.
- [3] K.A. Ross, C.R.B. Wright, *Matematyka dyskretna*, PWN, Warszawa 2008.
- [4] A. Szepietowski, *Matematyka dyskretna*, Wydawnictwo GU, Gdańsk, 2004.
- [5] R.J. Wilson, *Wprowadzenie do teorii grafów*, PWN, Warszawa 2012.