



H. STEINHAUS (Wrocław)

McKinsey o grach

Amerykański logik J. C. C. McKinsey, uczeń Tarskiego i profesor filozofii w kalifornijskim uniwersytecie Stanford, wydał w roku 1952 wstęp do teorii gier — książka dotarła do nas po dwóch latach a niemal równocześnie polscy matematycy dowiedzieli się o samobójstwie jej autora.

Nie jest to pierwsza książka o grach. Jakkolwiek nowoczesna teoria gier nie ma więcej niż 30 lat, doczekała się obszernej monografii w postaci dzieła J. von Neumanna i C. Morgensterna *Theory of games and economic behavior*, które wyszło już w drugim wydaniu (1944, 1947). Obciążenie zbytnią ogólnością, nieprzejrzystością układu i duże rozmiary książki Neumanna i Morgensterna sprawiają, że czytelnik nieobeznany z tą niezmiernie interesującą dziedziną wiedzy niełatwo dotrze do niej poprzez wspólne dzieło pryncetonskiego matematyka i wiedeńskiego ekonomisty, chociaż własny wkład v. Neumanna do teorii gier jest kapitalny, o czym powiemy w toku tego streszczenia. Dlatego należy powitać książkę McKinsey'a, której kryształowa jasność, umiar, prostota i przystępność mogą służyć za wzór każdemu, kto chce pisać podręczniki dla publiczności matematycznej składającej się nie tylko ze współspecjalistów autora. Przegląd jej treści nie jest recenzją w zwykłym sensie — jest raczej sposobnością do poinformowania, chociaż dość pobieżnego, o samej materii, w Polsce dotychczas mało znanej poza nieliczną grupą osób.

Tytuł książki brzmi w oryginale: *Introduction to the theory of games*, The RAND Corporation, first edition, New York 1952. Książka składa się z 18 rozdziałów, przedmowy, bibliografii i skorowidza (X+371 str.). Do jej studiowania nie jest konieczna obszerna wiedza matematyczna — na obwołucie czytamy zapewnienie, że prócz klasycznego rachunku różniczkowego i całkowego oraz klasycznej algebry czytelnik nie musi niczego poznać przed lekturą, i rzeczywiście to „studium badawcze”, jak je nazywa firma wydawnicza RAND, spełnia obietnicę. Tak np., gdy autorowi trzeba całek Stieltjesa, poświęca im osobny rozdział (IX), w którym nie tylko podaje definicje, twierdzenia i dowody, ale także tłumaczy, do czego i gdzie to uogólnienie zwykłej całki będzie niezbędne.

Pierwszy rozdział wprowadza w teorię gier prostokątnych — schemat oznaczony tą nazwą będzie służył za model gry przez cztery rozdziały.

Nazwijmy pierwszego gracza X , a drugiego Y . Gracz X nadaje zmiennej x dowolną wartość z zapasu obejmującego liczby naturalne od 1 do m włącznie, a jego przeciwnik Y nadaje zmiennej y dowolną wartość z zapasu obejmującego liczby naturalne od 1 do n włącznie.

W zbiorze par (i, j) ($1 \leq i \leq m, 1 \leq j \leq n$) określona jest funkcja rzeczywista $f(x, y)$, którą będziemy nazywali *funkcją wypłaty*.

Regulamin gry każe wybrać x i y równocześnie, po czym następuje wypłata kwoty $f(x, y)$ na rzecz X -a z kieszeni Y -a. Funkcja f jest znana partnerom przed rozpoczęciem partii; może ona przybierać wartości ujemne — wypłatę kwoty ujemnej — K ($K > 0$) X -owi należy rozumieć jako wypłatę przez X -a Y -owi kwoty K .

Jeżeli zamiast $f(i, j)$ napiszemy a_{ij} , to symbolem gry prostokątnej będzie macierz $\|a_{ij}\|$ o m wierszach i n kolumnach, a partia będzie polegała na tym, że gracz X wybiera wiersz, a gracz Y kolumnę tej macierzy, po czym X otrzymuje od Y -a kwotę wyrażoną w jednostkach pieniężnych przez element macierzy należący do wybranego wiersza i wybranej kolumny.

Ważnym problemem, który powstaje przy rozpatrywaniu gier prostokątnych, jest znalezienie optymalnej metody gry. Wobec tego powstaje kwestia ustalenia racjonalnego kryterium pozwalającego na wyróżnienie pewnej taktyki na niekorzyść innych. Zakłada się przy tym nieznaną aktualnej decyzji przeciwnika, jak również jego przyzwyczajęń.

Jednym z takich kryteriów jest *kryterium minimaks*, które zaleca X -owi wybór i -tego wiersza, przy czym i jest tą liczbą z zapasu $1, 2, \dots, m$, która nadaje największą wartość wyrażeniu $\min_j a_{ij}$ (symbol $\min_j a_{ij}$ oznacza najmniejszą z wartości $a_{i1}, a_{i2}, \dots, a_{in}$). Liczbę $w = \max_i \min_j a_{ij}$

definiuje się jako wartość gry dla X -a; łatwo wykazać, że X może zapewnić sobie przez wyżej zalecony wybór wygraną równą liczbie w i że żaden inny wybór nie zagwarantuje mu większej wygranej. W analogicznym sensie wartością gry dla Y -a jest wyrażenie $W = \min_j \max_i a_{ij}$: jeżeli Y wybierze j tak, by nadać najmniejszą wartość wyrażeniu $\max_i a_{ij}$, to zapewni

sobie, że nie przegra więcej niż W , a żaden inny wybór nie zagwarantuje mu mniejszej przegranej niż W . Już z tego wynika nierówność $W \geq w$, którą spełnia każda macierz. McKinsey nie uważa jednak, że w przypadku $W > w$ rozwiązanie wyżej znalezione jest ostateczne. Mówiąc popularnie,

w przypadku $W > w$ obydwaj partnerzy mogą mieć zał do doradcy, który im wskazał metody wyżej określone. W przykładzie macierzy

$$(I) \quad X \left\{ \begin{array}{c} 1 \\ 2 \\ 3 \end{array} \right\} \left\| \begin{array}{ccc} -1 & 2 & 3 \\ & 3 & 2 & -1 \\ & 2 & 1 & 0 \end{array} \right\|$$

$$\quad \quad \quad \underbrace{\begin{array}{ccc} 1 & 2 & 3 \end{array}}_Y$$

wartością gry dla X jest 0, wartością gry dla Y jest 2; idąc za doradcą X wybrał trzeci wiersz, a Y drugą kolumnę i X wygrał 1 zł. Wiedząc jednak, że najlepszą (według słów doradcy) taktyką dla Y -a jest wybór drugiej kolumny, X żałuje, że nie wybrał pierwszego wiersza, bo byłby wygrał 2 zł zamiast 1 zł. Autor nie wchodzi bliżej w rozdziale I w te wywody i zajmuje się przede wszystkim grami prostokątnymi spełniającymi $w = W$, to jest

$$(1) \quad \max_i \min_j a_{ij} = \min_j \max_i a_{ij}.$$

Odejźmy od terminologii autora i nazwijmy *zamkniętą* taką grę, która spełnia warunek (1). Przykład takiej gry daje macierz

$$(II) \quad X \left\{ \begin{array}{c} 1 \\ 2 \\ 3 \end{array} \right\} \left\| \begin{array}{ccc} 6 & -1 & -3 \\ & 2 & 0 & 1 \\ & 0 & -1 & -1 \end{array} \right\|$$

$$\quad \quad \quad \underbrace{\begin{array}{ccc} 1 & 2 & 3 \end{array}}_Y$$

Tutaj wartością gry dla obu stron jest zero i w porównaniu z przykładem (I) sytuacja zmienia się radykalnie. Metodami zalecanymi przez kryterium minimaks są tu: dla gracza X wybór drugiego wiersza, a dla gracza Y wybór drugiej kolumny. Jeżeli więc Y wybrał drugą kolumnę, to jakkolwiek wiersz wybierze X , nie wygra więcej niż zero. Wybór drugiego wiersza gwarantuje mu wartość zero, jest więc racjonalna przyczyna nazwania tej taktyki najlepszą. Analogiczna sytuacja powstaje dla gracza Y . Dzieje się to dlatego, że element a_{22} leżący na przecięciu drugiego wiersza i drugiej kolumny jest najmniejszy w swoim wierszu i największy w swojej kolumnie. Ogólnie: Punkt (x_0, y_0) nazywa się *punktem siodłowym* na mapie powierzchni $z = f(x, y)$, jeżeli

$$(2) \quad \max_x f(x, y_0) = f(x_0, y_0), \quad \min_y f(x_0, y) = f(x_0, y_0),$$

a sama powierzchnia nazywa się *siodłem* (zwróćmy uwagę na asymetrię tej definicji!). Warunkiem koniecznym i dostatecznym na to, by gra prostokątna była zamknięta, jest, by powierzchnia określona przez wypłatę f była siodłem. Jeżeli więc napiszemy warunek (2) w transkrypcji macierzowej, to okaże się równoważność tego warunku z warunkiem (1). Wtedy najlepszą taktyką dla X -a jest wybór wiersza zawierającego punkt siodłowy, a dla Y -a wybór kolumny zawierającej taki punkt. Gdy obydwaj partnerzy trzymają się tych taktyk, rezultat partii realizuje wartość gry dla jednego i drugiego. Ta wartość jest równa obu stronom wzoru (1), a zarazem — w pisowni funkcyjnej — równa $f(x_0, y_0)$, czyli wysokości punktu siodłowego (wyobraźmy sobie oś z skierowaną w górę). Byłoby jednak błędne mniemanie, jakoby każdy wyraz a_{rs} równy wyrażeniom (1) określał najlepsze taktyki jako wybór r -tego wiersza i s -tej kolumny — już przykład macierzy (II) przeczy temu, bo tu $a_{31} = a_{22} = 0$, jednak wybór trzeciego wiersza ($r = 3$) przez X -a naraża go na przegranie 1 zł, co nie może się stać, gdy zgodnie z teorią wybierze drugi wiersz.

Gdy wiadomo, jakie jest posunięcie (czyli wybór) Y -a, łatwo określić najlepsze posunięcie X -a, i na odwrót. Przypuśćmy, że wybór X -a jest najlepszy przy założeniu jakiegoś określonego wyboru Y -a i że ten wybór Y -a jest najlepszy przy założeniu rozważanego wyboru X -a. Istnienie takich wzajemnie najlepszych taktyk jest równoważne z istnieniem punktu siodłowego, a same taktyki są również najlepsze w sensie naszej teorii. Nazwijmy *wariantem głównym* partię graną w taki sposób. Przykładem mogą być końcówki szachowe, gdyż mieszczą się one w schemacie gier prostokątnych: Jeśli weźmiemy np. pięciochodówkę i oznaczmy przez Y gracza, który ma dać mata drugiemu (X -owi), a przez x i y ich taktyki, to definiując funkcję wypłaty $f(x, y)$ jako ilość posunięć prowadzących do mata, otrzymamy grę prostokątną i, jak się później dowiemy, zamkniętą. Ma więc ona wariant główny — czasopisma szachowe często ograniczają się do podania wariantu głównego jako rozwiązania⁽¹⁾. Warto wspomnieć mimochodem, że przyczynek szachistów do teorii gier jest znikomy — nawet E. Lasker, który miał tytuł mistrza świata z początkiem bieżącego stulecia i był matematykiem, nie stanowi wyjątku, o czym świadczy jego książka o grach.

Nic łatwiejszego jak zbudowanie macierzy bez punktu siodłowego — podaliśmy taki przykład (I). Taka macierz określa *grę otwartą* (nasza terminologia) — tak będziemy nazywać każdą grę, dla której $W > w$ (w grze zamkniętej $W = w$).

⁽¹⁾ Uważny czytelnik spostrzeże, że powszechnie przyjęte przez szachistów definicje dwuchodówki, trójdchodówki etc. są dowodem naturalności zasady minims.

Nowoczesna teoria gier szczyści się teorematem, który autor nazywa fundamentalnym, a który odkrył J. v. Neumann. Dzięki temu twierdzeniu teoretycy gier uważają za opanowane zagadnienie gier dwuosobowych o sumie zerowej — tak nazywają grę, w której suma wypłat po każdej partii jest zerem; zauważmy, że ten warunek nie ma nic wspólnego z równością szans ani też ze sprawiedliwością gry w jakimkolwiek popularnym sensie. Jak widzieliśmy, już w grach prostokątnych pojawiają się gry otwarte i na nich wyjaśnimy teoremat fundamentalny.

Wyobraźmy sobie, że X reprezentuje drużynę złożoną ze 100 graczy $X_1, X_2, \dots, X_{100}$, i podobnie Y reprezentuje $Y_1, Y_2, \dots, Y_{100}$. X wybiera w imieniu swoich graczy wiersze: i_1 dla X_1 , i_2 dla X_2 , ..., i_{100} dla X_{100} , i analogicznie Y wybiera kolumny dla swoich graczy. Przeciwnicy nie znają nawzajem swoich decyzji. Partia X -a kontra Y polega na rozegraniu 100 partii X_n kontra Y_n ($n = 1, 2, \dots, 100$), a kwota, którą Y wypłaci X -owi, będzie równa średniej kwot należnych wszystkim X_n -om od ich partnerów, czyli $\sum_{n=1}^{100} f(i_n, j_n)/100$. Tym razem taktyka X -a polega na tym, ilu swoim graczom przydzieli wiersz pierwszy, ilu drugi, ... itd., ilu ostatni wiersz macierzy. Analogicznie określona jest taktyka Y -a. Ale o wyniku partii zadecydują nie tylko obie taktyki, ale także ponumerowanie graczy w drużynach. Z tą chwilą wygrana X -a staje się zmienną losową, a jej wartość oczekiwana jest zupełnie określoną wielkością zależną wyłącznie od obu taktyk. Przez to określiliśmy nową grę i o tej grze orzeka teoremat fundamentalny, że jest zamknięta, a raczej, że byłoby tak, gdyby liczbę 100 zamienić na nieskończoność. (Za tę nieścisłość nie jest odpowiedzialny McKinsey, który nie posługuje się naszym obrazem drużyn). Ścisłe ujęcie uzyskuje się przez wprowadzenie aparatu do losowania tak zbudowanego, że wyrzuca on liczby $1, 2, \dots, m$ z prawdopodobieństwami $p_1, p_2, \dots, p_m$ (taktykę X -a oznaczmy przez P). Analogicznie definiuje się taktykę Q dla Y -a. Zarówno P jak Q są wektorami; każdy wektor P (spełniający warunki $\sum_{i=1}^m p_i = 1$, $p_i \geq 0$) symbolizuje „kostkę” (aparat do losowania) dla X -a, każdy wektor Q kostkę dla Y -a. Wynik partii dla X -a jest a_{ij} , jeżeli jego kostka padła na i -tą ścianę, a kostka przeciwnika na j -tą. Wynik jest zatem zmienną losową, która jest zdefiniowana przez P i Q . Tę zmienną możemy oznaczyć przez $F(P, Q)$, a jej wartość oczekiwaną $EF(P, Q)$ krótko przez $E(P, Q)$. Wielkość E już nie jest zmienną losową, lecz zwykłą funkcją wektorów P, Q , którą można wyrachować ze schematu a_{ij} . Jest mianowicie

$$(3) \quad E(P, Q) = \sum_{i=1}^m \sum_{j=1}^n a_{ij} p_i q_j.$$

Teoremat fundamentalny orzeka, że gra określona przez funkcję (3) w tym sensie, że przyznaje X -owi kwotę $E(P, Q)$ od Y -a, jeżeli X wybrał określony wektor P , Y zaś Q , spełnia warunek

$$(4) \quad \max_P \min_Q E(P, Q) = \min_Q \max_P E(P, Q),$$

a więc jest zamknięta. Stąd wynika, że w tej nowej grze istnieje wariant główny określony przez najlepsze taktyki P_0, Q_0 . Wartością nowej gry dla obu stron jest $E(P_0, Q_0)$. Jest jednak różnica między zwykłą grą prostokątną a nową grą; teraz partner X nie jest pewny, że taktyka polegająca na konstrukcji kostki P_0 zagwarantuje mu w każdej partii z osobna wartość gry $E(P_0, Q_0)$. Byłoby nawet błędne twierdzenie (i McKinsey tego błędu nie popełnia!), że X trzymając się stale taktyki P_0 uzyska co najmniej wartość gry jako przeciętny zysk przy nieskończeniu wielu partiach. To twierdzenie staje się jednak prawdziwe, gdy przeciwnik też trzyma się stale jednej i tej samej taktyki — gdy jest nią taktyka optymalna, to X (i Y) uzyska przeciętnie wartość gry. Dodajmy od siebie, że gdy Y zmienia taktykę, a X trzyma się taktyki P_0 , to limes inferior średnich wygranych X -a w nieskończonym ciągu partii jest co najmniej taki, jak wartość gry.

Bibliografia na końcu rozdziału II informuje, że v. Neumann oparł swój pierwszy dowód teorematu fundamentalnego na twierdzeniu Brouwera o punkcie stałym.

Autor po każdym rozdziale podaje zadania. Tak np. każe czytelnikowi zbudować przykład, w którym P_0, Q_0 nie określają wariantu głównego, chociaż

$$E(P_0, Q_0) = \max_P \min_Q E(P, Q) = \min_Q \max_P E(P, Q)$$

— jak wiadomo ta relacja jest konieczna do tego, by (P_0, Q_0) określało wariant główny, lecz nie jest dostateczna, jak to pokazaliśmy na przykładzie gry (II). Intuicyjnie sprawa jest jasna już przy zwykłych grach zamkniętych: Wartością końcówki szachowej jest liczba 5, jeżeli jest to pięciochódka; może się zdarzyć, że Białe (Y) grając źle dadzą mata Czarnym (X -owi) w pięciu posunięciach dlatego, że Czarne też grają źle i nie umieją wyzyskać błędów przeciwnika. W geometrycznej interpretacji zadanie wymaga znalezienia powierzchni, która ma punkt siodłowy i punkt na równej wysokości z siodłowym, który sam nie jest siodłowy — powierzchnia $z = xy$ jest łatwym przykładem.

Trzeci rozdział zajmuje się techniką rozwiązywania gier; teorematy o istnieniu nie zawsze wskazują drogę rachunkową, a już przy łatwych grach otwartych nie można liczyć na jej znalezienie bez wskazówek.

Aparaturą jest tutaj teoria macierzy prostokątnych, a więc rozdział algebry klasycznej.

Jak widać od razu, trudność polega na tym, że trzeba szukać ekstremów (4) formy kwadratowej (3), w której zmienne p i q są związane dwoma równaniami i $m+n$ nierównościami. Znajdujemy tu uwagę, że w grach otwartych (po domknięciu, które wyżej opisano) może zajść ewentualność wyboru wśród nieskończenie wielu najlepszych taktyk, a to dlatego, że w ogóle taktyk jest nieskończenie wiele (tyle ile wektorów P, Q , a więc kontinuum).

Wtedy można jeszcze do tych najlepszych taktyk wprowadzić dyskryminację za pomocą pojęcia dominacji: *Taktyka P_1 dominuje taktykę P_2* , jeżeli przy każdej taktyce „ j ” (to jest wyborze kolumny — j -tej przez Y -a) P_1 daje co najmniej taką wygraną X -owi jak P_2 .

Liczne przykłady objaśniają tekst. Jedno z zadań daje schemat graficzny złożony z 7 punktów; niektóre są połączone strzałkami; każdy gracz wybiera punkt. Jeżeli wybrana para jest połączona, to grót wskazuje, kto ma otrzymać 1 zł od partnera, jeżeli nie są połączone, to wypłata jest 0. Nietrudno stwierdzić, że grę tę można zredukować do schematu prostokątnego.

Trzeba stwierdzić, że dotąd nie znaleziono dostatecznie wygodnej metody wyznaczania wszystkich najlepszych taktyk mieszanych (tj. wektorów P_0, Q_0). Metoda podana w tym rozdziale dla gier prostokątnych, chociaż ogólna, prowadzi w przypadku macierzy wypłaty o dużej ilości wierszy i kolumn do niesłychanie żmudnych (choć prostych) rachunków.

Używając nawet elektronowych maszyn do rachowania niemożliwe jest na ogół znalezienie wszystkich rozwiązań już np. dla macierzy kwadratowych rzędu 100. Trzeba mianowicie rozpatrywać wszystkie podwyznaczniki kwadratowe i sprawdzać, czy spełniają one pewne wzory.

Żanim pójdziemy dalej, przypomnijmy raz jeszcze, że autor uważa zamkniętą grę prostokątną za banalną i w dalszych rozdziałach książki uważa redukcję do takiej gry za zupełne rozwiązanie. Jego terminologia nazywa *mieszaną strategią* postępowanie polegające na rzucaniu odpowiednich kostek, którym powierza się decyzję w grach otwartych. Nawiasem mówiąc, tej ilustracji autor nie używa. Mogłoby się wydawać, że w praktyce te sytuacje, którym współczesna teoria gier przypisuje takie znaczenie, nigdy się nie zdarzają. Otóż zacytujmy z innego źródła przykład realny:

Torpedy sterowane radioelektrycznie przestają działać — w pewnej frakcji — gdy przeciwnik ma urządzenia głośzące; istnieją aparaty, które osłabiają urządzenia głośzące przeciwnika. Jedne i drugie obniżają sprawność bojową okrętów i nie przynoszą żadnej korzyści, gdy przeciw-

nik nie używa aparatów, które mają być przez nie zwalczane. Liczbowy obraz korzyści i strat jest schematem prostokątnym otwartym:

$$X \begin{cases} 1. \text{ Okręty mające aparaty do sterowania} \\ 2. \text{ Okręty nie mające tych aparatów} \end{cases} \left\| \begin{array}{cc} \alpha & \gamma + \delta \\ \alpha + \beta & \gamma \end{array} \right\|$$

$\underbrace{\begin{array}{cc} \text{Okręty mające} & \text{Okręty nie mające} \\ \text{urządzenia głośzące} & \text{tych urządzeń} \\ 1. & 2. \end{array}}_Y$

(α , β , γ , δ są to np. pewne stałe dodatnie). Każde spotkanie statku ze statkiem wrogiem jest partią. X realizuje mieszaną strategię $P = (p_1, p_2)$ w ten sposób, że tylko część swoich okrętów zaopatruje w aparaty do sterowania — mianowicie tyle, by frakcja takich okrętów wśród okrętów X -a była równa p_1 . Analogicznie Y może stosować mieszaną strategię $Q = (q_1, q_2)$. Domknięcie tej gry w myśl twierdzenia fundamentalnego prowadzi do optymalnych taktyk mieszanych dla obu stron. Łatwo sprawdzić, że jeśli tylko $\alpha + \beta > \gamma$ (w innym przypadku macierz wypłaty ma punkt siodłowy), to taktyki te są dane przez

$$p_1 = (\alpha + \beta - \gamma)/(\beta + \delta), \quad q_1 = \delta/(\beta + \delta).$$

Rozdział czwarty mówi o aproksymatywnym znajdowaniu wartości gier na przykładzie otwartej gry prostokątnej

$$\left\| \begin{array}{ccc} 1 & 2 & 3 \\ 4 & 0 & 1 \\ 2 & 3 & 0 \end{array} \right\|.$$

Aproksymacja polega na tym, że w pierwszej partii partnerzy wybierają dowolne posunięcia, w drugiej X gra optymalnie przy hipotezie, że Y zagra tak, jak zagrał w pierwszej partii, Y zaś postępuje analogicznie do X -a. Rezultaty się zapisuje i zawsze po n partiach X gra tak, żeby jego posunięcie powtórzone n razy dało mu maksymalny zysk łączny przy założeniu, że Y będzie grał w partii $(n+k)$ -ej tak jak grał w k -tej partii ($k = 1, 2, \dots, n$); ale naprawdę Y będzie grał inaczej, a mianowicie on także zagra w $(n+1)$ -ej partii według zasady analogicznej do tej, według której wybrał właśnie X swoje n plus pierwsze posunięcie. Tak powstaje cały ciąg partii zupełnie zdeterminowany, który oczywiście można zbudować teoretycznie, tj. bez realnych żywych graczy. Autor podaje bez dowodu twierdzenie, że maksymalny zysk łączny przewidziany przez X -a (przy założeniu, że obaj grają w sposób wyżej zdefiniowany) podzielony przez n daje ciąg o wyrazach niemniejszych od prawdziwej wartości gry i że ta wartość jest dolnym kresem wyrazów tego ciągu. Analogiczne twier-

dzenie zachodzi dla Y -a. Dlatego można po dostatecznie długiej konstrukcji ująć wartość gry w ciasne granice.

Znajdowanie najlepszych taktyk mieszanych wymaga obliczenia, jaka jest relatywna frekwencja posunięć w owym ciągu partii. Jeżeli X ma dokładnie jedną optymalną taktykę mieszaną, to ciąg relatywnych frekwencji jest zbieżny do tej taktyki. Gdy X ma więcej optymalnych taktyk, ciąg frekwencji może nie być zbieżny. W tym przypadku każdy podciąg zbieżny dąży do pewnej optymalnej taktyki.

Należy dodać, że podana wyżej metoda jest dość ekonomiczna w przypadku, gdy ilość wierszy i kolumn macierzy wypłaty jest bardzo duża. Praca potrzebna do wyznaczenia wartości gry (przy danej dokładności) jest bowiem w przybliżeniu liniową funkcją ilości wierszy i kolumn.

Rozdział piąty mówi o grach w ogólnej postaci, to znaczy podanych inaczej niż przez macierz prostokątną. Okazuje się, że każda gra, której zakresami wyborów są zbiory skończone i z góry określone, da się sprowadzić do formy normalnej, to jest prostokątnej. Na przykład, niech takie będą reguły gry: X i Y wybierają kolor biały (B) lub czarny (C) i to naprzód X jawnie, potem Y jawnie, potem X jawnie (autor nie omieszkiał zaznaczyć, że X przy drugim wyborze nie tylko wie, co wybrał Y , ale także nie zapomina, co sam wybrał wcześniej); do reguł należy funkcja $M(, ,)$, znana obu graczom, w którą wstawia się symbole kolorów wybranych, a która podaje wypłatę należną X -owi. Autor podaje konkretny przykład:

$$\begin{aligned}
 & M(B, B, B) = -2, & M(B, B, C) = -2, \\
 (III) \quad & M(C, B, B) = 5, & M(C, B, C) = 2, \\
 & M(B, C, B) = 3, & M(B, C, C) = -4, \\
 & M(C, C, B) = 2, & M(C, C, C) = 6.
 \end{aligned}$$

Bezpośrednio nie jest to gra prostokątna. Gracz X ma 32 możliwe czyste taktyki, a Y ma ich 4. Schemat prostokątny o 32 wierszach i 4 kolumnach jest grą zamkniętą o 4 punktach siodłowych; X ma 4 taktyki optymalne, a Y dwie. Tak np. optymalną taktyką dla X -a jest wybór czarnego koloru w pierwszym posunięciu, a w drugim naśladowanie Y -a. Dla Y -a jest optymalną taktyką wybór biały, jak również wybór przeciwny do X -a. Wartość gry jest 5.

Inny przykład: Przeciw X -owi gra spółka Y złożona z dwóch osób, Y_1 i Y_2 . Gracze X , Y_1 i Y_2 mają osobne pokoje i nie mogą się porozumiewać podczas partii, gra polega znowu na wyborze kolorów B lub C; X wybiera pierwszy. Jeżeli X wybrał B, arbiter idzie do Y_1 , który wybiera jako drugi, a potem do Y_2 , który wybiera jako ostatni. Jeżeli jednak X wybrał C, to arbiter odwiedza naprzód Y_2 , a potem Y_1 . Znowu jest dana funkcja

$M(, ,)$, w którą wstawia się B lub C na każde miejsce według kolejności wyborów — ona określa zapłatę dla X -a. Tutaj dwuosobowy Y nie wie, co wybrał X , a żadna z jego części nie wie, co wybrała druga, a nawet nie wie, czy jej wybór dotyczy drugiej, czy trzeciej zmiennej w M .

Tę grę i wiele innych można przedstawić graficznie przez dendryt, którego rozgałęzienia są miejscami wyboru dalszej drogi. W przykładzie autora wartości M (w porządku alfabetycznym) są 0, 2, 6, 8, 4, 0, 5, 6. Gra po znormalizowaniu, to znaczy po sprowadzeniu do schematu prostokątnego, daje macierz o dwóch wierszach i czterech kolumnach. Gra jest otwarta. X ma tylko dwa posunięcia: B lub C. Y ma cztery taktyki:

$$\begin{array}{ll} 1) \left\{ \begin{array}{l} Y_1 \text{ wybiera B,} \\ Y_2 \text{ wybiera B,} \end{array} \right. & 2) \left\{ \begin{array}{l} Y_1 \text{ wybiera B,} \\ Y_2 \text{ wybiera C,} \end{array} \right. \\ 3) \left\{ \begin{array}{l} Y_1 \text{ wybiera C,} \\ Y_2 \text{ wybiera B,} \end{array} \right. & 4) \left\{ \begin{array}{l} Y_1 \text{ wybiera C,} \\ Y_2 \text{ wybiera C.} \end{array} \right. \end{array}$$

Wartość gry jest $\frac{12}{5}$, optymalna taktyka mieszana dla X -a jest 40% B, 60% C, a dla Y -a 60% 1), 0% 2), 40% 3), 0% 4).

Metodę dendrytów stosuje autor także do gier o posunięciach losowych, a więc takich jak karty. Nie zapominajmy, że losowość zwykłych gier nie jest tego rodzaju, co losowość kostek, które definiują taktyki mieszane. Kostki poprawiają taktykę, czego nie można powiedzieć o losie w grach hazardowych — tam kostkę konstruuje partner, w grach hazardowych kostka jest dana i należy do regulaminu gry. Losowość może interweniować przy określaniu kwoty wygranej, przy konstrukcji zbioru możliwych posunięć i przy kolejności posunięć — autor podaje przykłady na te trzy rodzaje losowości. Za każdym razem jest możliwa normalizacja gry. Autor podaje przykłady dendrytów nie dających się znormalizować.

Rozdział szósty jest poświęcony ogólnej teorii skończonych, n -osobowych gier w postaci ogólnej. Za podstawę definicji takiej gry bierze autor dendryt; jego struktura topologiczna wraz z kwotami na szczytach, z regułą kolejności posunięć i z podziałem rozgałęzień pomiędzy obszary informacyjne (*information sets*) określa kompletnie grę. Autor uważa, że dendryt jest pojęciem dostatecznie ogólnym do objęcia wszystkich typów gier mających sens — przykłady podane na końcu poprzedniego rozdziału pokazują nawet, że pojęcie dendrytu jest zbyt ogólne. To też autor poddaje obszary informacyjne czterem postulatam, które nie wszystkie są oczywiste.

Paragraf 2 tego rozdziału dotyczy gier o kompletnej informacji. Gry te charakteryzują się tym, że w każdej partii gracz mający ruch wie, jakie posunięcia zrobiono przedtem — zna aktualną sytuację. Ścisła definicja opiera się na pojęciu obszarów informacyjnych. Przykładami

takich gier są szachy, warcaby, halma — ale nie bridge lub poker. Autor podaje następujący przykład: W pierwszym posunięciu gracz X wybiera kulę białą (B) lub czarną (C), w drugim posunięciu kulę białą lub czarną wyrzuca mechanizm losowy (i to w ten sposób, że prawdopodobieństwo wyrzucenia kuli białej jest równe $\frac{1}{5}$ — obaj gracze wiedzą o tym). W przypadku (C) ostatnie posunięcie ma gracz X , który znów wybiera (B) lub (C). Jeśli natomiast mechanizm wyrzucił kulę białą, to samo zamiast X -a czyni Y . Wypłata M zdefiniowana jest przez (III). Z jakim spośród trzech przypadków losowości, wymienionych na poprzedniej stronie, mamy do czynienia?

Głównym wynikiem paragrafu 2 jest twierdzenie, że matryca znormalizowanej gry o sumie zerowej i kompletnej informacji ma punkt siodłowy, a więc jest grą zamkniętą i mającą optymalne taktyki czyste (tj. niemieszane) dla obu stron. Dowód idzie przez lemat o istnieniu punktu równowagi, czyli takiego punktu $(x_1, x_2, \dots, x_n)$ w produkcie kartezjańskim taktyk, że jeżeli każdy partner X_k ($k = 1, 2, \dots, i-1, i+1, \dots, n$) trzyma się taktyki x_k , to taktyka x_i daje partnerowi X_i maksymalną wygraną. Innymi słowy: zespół taktyk $(x_1, x_2, \dots, x_n)$ jest wzajemnie optymalny w tym sensie, że każda z nich jest optymalna względem zespołu pozostałych. Dla $n = 2$ to twierdzenie redukuje się do istnienia punktu siodłowego określonego w pierwszym rozdziale. Wynikające stąd twierdzenie o grach dwuosobowych, które je trywializuje, wcale nie jest znane powszechnie, nawet w kołach szachistów. Udowodnił je Kalmar w roku 1929 i można postawić książce Neumanna i Morgensterna zarzut, że trudno się go tam doczytać. Ale w samej rozprawie Kalmara jest ono tak sformułowane, że niektórzy matematycy sądzą, iż *explicite* tam nie figuruje; charakterystyczne jest też pominięcie nazwiska Kalmara w książce dwóch autorów. Natomiast McKinsey (str. 126 i dalsze) podaje zarówno sformułowanie słowne, popularnie i bez pretensji do ścisłości, jak i ścisłe oraz zaopatrzone w kompletny dowód. Wyżej zacytowany lemat o punkcie równowagi potrzebny jest do tego tylko dla $n = 2$. Indukcja przenosi go na wszelkie n większe od 2. Z tego uogólnienia mógłby ktoś wysnuć wniosek o tym, że także dla $n > 2$ gry n -osobowe o sumie zerowej i kompletnej informacji są zamknięte, a więc mają czyste taktyki optymalne. W jednym z dalszych rozdziałów okaże się, dlaczego autor tego wniosku nie wysnuwa.

W tymże rozdziale jest też mowa o grach z zupełną pamięcią; ścisła definicja jest trudna, ale na przykładzie gier w karty łatwa. Bridge nie jest grą z zupełną pamięcią, bo tam gracz składa się z dwóch osób i nie ma żadnej osoby, która by wiedziała to wszystko, co te dwie osoby wiedzą. Natomiast znane gry dwuosobowe są zwykle grami z pamięcią zupełną. *Taktyką behawiorystyczną* nazywa się taktykę polegającą na loso-

waniu posunięć (mieszane taktyki polegały na losowaniu czystych taktyk, a nie poszczególnych posunięć!). Taktyki behawiorystyczne nie domykają na ogół gry, a więc pod tym względem są gorsze od taktyk mieszanych. Jeżeli jednak pamięć jest zupełna, to dla tych taktyk istnieje twierdzenie analogiczne do teorematu fundamentalnego. Na korzyść ich przemawia fakt, że ilość parametrów potrzebnych do określenia takiej taktyki jest na ogół mniejsza niż dla taktyk mieszanych. Wyjaśnimy to na przykładzie: X i Y wybierają trzy liczby ze zbioru $\{1, 2\}$. Robią to w następujący sposób: Najpierw X wybiera u , potem Y wybiera v , w końcu X nie znając posunięcia Y -a, lecz pamiętając swoje, wybiera w . Następuje wypłata sumy $M(u, v, w)$ przez Y -a na rzecz X -a. W tej grze X ma cztery czyste taktyki, tj. każdą taktykę mieszaną definiuje wektor P o czterech składowych. Jest więc ona funkcją trzech parametrów (czwarty można wyrachować z wzoru $\sum_{i=1}^4 p_i = 1$). Taktyka behawiorystyczna natomiast będzie funkcją tylko dwóch parametrów, gdyż do jej zdefiniowania potrzeba znać tylko prawdopodobieństwa, z jakimi X wybiera liczbę 1 (lub 2) w swoim pierwszym i drugim posunięciu. Dla gier, w których ilość posunięć jest duża, liczba parametrów może różnić się znacznie.

Rozdział siódmy mówi o grach, w których wybiera się posunięcia ze zbiorów nieskończonych. Tu należą gry ciągłe, w których zbiory są odcinkami, a wybiera się punkty x, y z tych odcinków. Jeżeli $M(x, y)$ jest wypłatą dla X -a, to może się zdarzyć, że obie lub jedna z wielkości $\max \min M(x, y)$, $\min \max M(x, y)$ nie istnieje. To nie wyłącza istnienia punktu siodłowego, jak widać na przykładzie „ $M(x, y) = 1/x - 1/y$ dla $x \neq 0, y \neq 0, M(x, 0) = M(0, x) = 0$ ”, gdzie $(0, 0)$ jest takim punktem. W przypadku nieistnienia wielkości $\min \max M(x, y)$ lub $\max \min M(x, y)$ może się zdarzyć, że istnieją i są równe wielkości $\inf \sup M(x, y)$ i $\sup \inf M(x, y)$. Również i w tej sytuacji wspólną wartość tych wyrażeń autor nazywa wartością gry v i chociaż teraz jeden lub obaj gracze mogą nie mieć najlepszych, tj. gwarantujących wartość v czystych taktyk, to do każdej liczby $\varepsilon > 0$ istnieje dla każdego gracza taktyka (czysta) gwarantująca liczbę, w przypadku X -a co najwyżej o ε mniejszą, a w przypadku Y -a co najwyżej o ε większą niż v .

Pomijamy tu rozdział ósmy o dystrybuantach i dziewiąty o całce Stieltjesa, gdyż stanowią one tylko przygotowanie matematyczne do rozdziału następnego i zawierają wiadomości, które można znaleźć w wielu podręcznikach rachunku prawdopodobieństwa. Nie jest to zarzut, gdyż czytelnik może polegać zarówno na ścisłości wywodów, jak i na idealnej oszczędności materiału.

Rozdział dziesiąty traktuje teoremat fundamentalny dla gier ciągłych. Orzeka on, że zawsze istnieją optymalne taktyki mieszane. Taktyka

jest tu dystrybuantą, którą gracz wybiera dowolnie; przez to zmienne x, y zamieniają się na zmienne losowe o rozkładach wybranych przez partnerów. Matematycznie teoremat pisze się tak (str. 186, § 3, tw. 19.4):

Dla funkcji $M(x, y)$ ciągłej w $\langle 0, 1 \rangle^2$ istnieją i są równe wyrażenia

$$\max_{F \in D} \min_{G \in D} \int_0^1 \int_0^1 M(x, y) dF(x) dG(y), \quad \min_{G \in D} \max_{F \in D} \int_0^1 \int_0^1 M(x, y) dF(x) dG(y).$$

Tutaj D oznacza zbiór wszystkich dystrybuant. Dowód redukuje całki do sum aproksymacyjnych, po czym wyzyskuje lematy z algebry form liniowych. Przypominamy, że nieciągłe funkcje M wymagałyby całek Stieltjesa i Lebesgue'a — dlatego autor ogranicza się do M ciągłych, nie chcąc ograniczyć koła czytelników do specjalistów.

Efektywne wyznaczanie optymalnych taktyk w grach ciągłych jest trudne w całej rozciągłości. Autor podaje jednak kilka teorematów pozwalających np. na stwierdzenie, czy dana taktyka jest optymalna, lub jaki jest związek między optymalnymi taktykami obu graczy. Pozwolę sobie zacytować jedno twierdzenie (str. 204): Jeżeli $M(x, y)$ jest funkcją wypłaty i jeżeli dla wszystkich x i y z przedziału $\langle 0, 1 \rangle$ $M(x, y) = -M(y, x)$, to gra ma wartość 0 i każda optymalna taktyka dla jednego gracza jest również optymalna dla drugiego.

W następnych rozdziałach autor rozważa pewne klasy gier ciągłych, dla których uzyskano ogólne wyniki. W rozdziale XI omawia klasę gier rozdzielnych, a w rozdziale XII klasę gier wypukłych. Te przymiotniki odnoszą się właściwie do wypłaty $M(x, y)$. Mówi się o niej, że jest rozdzielna, jeżeli wyraża się przez skończoną sumę iloczynów funkcji ciągłych jednej zmiennej:

$$(5) \quad M(x, y) = \sum_{i=1}^m \sum_{j=1}^n r_i(x) s_j(y).$$

Wtedy optymalne dystrybuanty dla X i Y są funkcjami schodkowymi i to odpowiednio o co najwyżej m i n stopniach. Wynika to z tego, że jeżeli liczba n jest zdefiniowana wzorem (5), to każda mieszana taktyka X -a jest równoważna (w sensie funkcji wypłaty) z pewną funkcją schodkową o co najwyżej n stopniach. Autor podaje ogólną metodę znajdowania optymalnych taktyk dla gier rozdzielnych, jest ona jednak skomplikowana.

Założenie wypukłości lub wklęsłości funkcji ze względu na każdą zmienną daje oczywiście punkt siodłowy, gdy M jest wypukła względem X , a wklęsła względem Y . W rozdziale XII autor robi jednak tylko jedno z tych założeń i otrzymuje niebanalny wynik, a mianowicie dystrybuantę jednoschodkową dla jednego gracza, a dwuschodkową dla drugiego.

PRZYKŁAD: Niech $M(x, y) = \sin \frac{1}{2}\pi(x+y)$ oraz $0 \leq x \leq 1$ i $0 \leq y \leq 1$. Łatwo sprawdzić, że funkcja ta jest wklęsłą funkcją zmiennej x dla każdego y i wklęsłą funkcją zmiennej y dla każdego x . Ponieważ procedura prowadząca do znalezienia optymalnych taktyk jest dość długa, podamy je bez rachunków. Niech $I_u(x)$ oznacza funkcję równą 0 dla $x \leq u$ i równą 1 dla $x > u$. Pokażemy, że dla X -a najlepszą taktyką jest $F_0(x) = I_{1/2}(x)$, a dla Y -a $G_0(y) = 1/2I_0(y) + 1/2I_1(y)$. Oznaczmy przez

$$E(F, G) = \int_0^1 \int_0^1 \sin \frac{1}{2}\pi(x+y) dF(x) dG(y)$$

wartość oczekiwaną wypłaty, zatem

$$E(F_0, G_0) = \frac{1}{2}\sqrt{2}, \quad E(F_0, G) = \int_0^1 \sin \frac{1}{2}\pi(\frac{1}{2}+y) dG(y) \geq \frac{1}{2}\sqrt{2},$$

$$\begin{aligned} E(F, G_0) &= \frac{1}{2} \int_0^1 (\sin \frac{1}{2}\pi x + \sin \frac{1}{2}\pi(1+x)) dF(x) = \\ &= \frac{1}{2}\sqrt{2} \int_0^1 \sin \frac{1}{2}\pi(\frac{1}{2}+x) dF(x) \leq \frac{1}{2}\sqrt{2}. \end{aligned}$$

Taktyki F_0 i G_0 gwarantują tę samą wartość gry $W = w = \frac{1}{2}\sqrt{2}$ dla obu graczy, są więc optymalne. Łatwo widzieć, że optymalna taktyka dla X -a jest jednoschodkowa, a dla Y -a dwuschodkowa.

W rozdziale trzynastym są zastosowania do wnioskowania statystycznego. Zagadnienie wyjaśnimy na przykładzie wnioskowania o rozkładzie pewnej cechy C w populacji na podstawie zbadania jej wielkości u sztuk próbnych wylosowanych z tejże populacji. Tutaj statystyk nie poprzestaje na wniosku, lecz decyduje się postąpić w pewien sposób x z całą populacją. Jeżeli prawdziwy rozkład cechy C w populacji jest y , to decyzja x da X -owi pewną korzyść (która może być także ujemna) i jest znaną funkcją $f(x, y)$ dwóch zmiennych. Tutaj X wybiera x , a natura Y wybiera y . Tym sposobem wchodzimy w teorię gier. Pomysł wyszedł z tej szkoły, która potępiła regułę Bayes'a, przede wszystkim od przedwcześnie zmarłego A. Walda. Przeniesienie teorii gier na problematykę decyzji nie jest jednak formalną transkrypcją, ponieważ drugim partnerem jest przyroda, której nie można posądzać o świadome działanie na szkodę człowieka. Stąd różne ograniczenia, które wprowadzają są korzystne dla człowieka, ale wyłączają ogólne stosowanie teorematu fundamentalnego (np. wtedy, gdy wiadomo, że natura nie może użyć dowolnej spośród taktyk mieszanych). Dla przykładu podam sformułowanie konkretnego i prostego zagadnienia: Urna zawiera białe i czarne kule; frakcja y białych kul zależy od przyrody Y . Statystyk X ma prawo wyjąć n kul z urny i orzec potem x ; taktyką czystą

X -a jest każda funkcja $x = g_n(k)$, gdzie k jest liczbą białych kul wśród n wyjętych, taktyką czystą Y -a jest dowolny wybór liczby y z przedziału $\langle 0, 1 \rangle$. Regulamin gry każe przyrodzie wypłacić X -owi kwotę $-(x-y)^2$. Procedura prowadząca do rozwiązania tego problemu jest dość zawiła. Podam tylko, że X ma czystą optymalną taktykę zdefiniowaną za pomocą funkcji

$$g_n(k) = \frac{k + \sqrt{n}/2}{n + \sqrt{n}}.$$

W ostatnich czasach cykl problemów związanych z zastosowaniem teorii gier do statystyki rozwinął się w całą teorię znaną pod nazwą teorii funkcji decyzyjnych. Zarys tej teorii czytelnik znajdzie w przystępnej książce D. Blackwella i M. A. Girshicka, *Theory of games and statistical decisions*, first edition, New York and London 1954.

Tytuł rozdziału czternastego przetłumaczyłem „Planowanie liniowe”. W ekonomii i technice często występuje zagadnienie minimizacji kosztów, które zależą liniowo od pewnych parametrów. Zwykle te parametry są związane jedną lub wieloma nierównościami liniowymi. Ogólnie problemat planowania liniowego żąda minimizacji wyrażenia $\sum_{k=1}^n c_k y_k$ pod m warunkami $\sum_{k=1}^n a_{ik} y_k \geq b_i$ ($i = 1, 2, \dots, m$), przez dobór układu $(y_1, y_2, \dots, y_n)$.

Skonstatujmy, że znane metody znajdowania ekstremów za pomocą rachunku różniczkowego w naszym przypadku zawodzą, gdyż na ogół w tego rodzaju zadaniach ekstremalne wartości są osiągane na brzegach obszaru ograniczonego danymi nierównościami, a nie w punktach, gdzie pochodne znikają. Autor pokazuje, że zagadnienie znalezienia optymalnych strategii i wartości gry w grze prostokątnej sprowadza się do wyżej zdefiniowanego problemu. Twierdzenie odwrotne jest prawdziwe przy pewnych ograniczeniach.

Rozdział piętnasty o n -osobowych grach zerowych zaczyna się, od uwagi, że matematycy nie są na ogół zadowoleni z wyników Neumanna i Morgensterna, jakkolwiek niemało, bo około 400 stron swego dzieła o grach poświęciła tej materii para autorska. Na pierwszy rzut oka sytuacja jest pomyślna: nawet gdy gra n -osobowa dopuszcza kilka ruchów naraz, ruchy losowe i informację niepełną, da się ona zredukować do równoczesnego wyboru elementów, przy czym każdy partner wybiera po jednym elemencie z pewnego raz na zawsze określonego, skończonego zbioru. Ale tu zaczyna się trudność związana z ekonomicznym sensem gry: jeżeli sumy wygrane (i przegrane) wyrażają się w wartościach tego samego rodzaju, np. w jednostkach określonej waluty, nie można zignorować faktu, że partnerzy mogą łączyć się w koalicje i asekurować się wzajemnie

przed złymi skutkami gry indywidualnej. Tak np. w grze „niezgodny traci” trzech partnerzy odkrywają równocześnie monety leżące na stole i zakryte dłońmi; jeżeli wszystkie trzy leżą jednakowo, mamy remis, jeżeli nie, to właściciel niezgodnej monety płaci po złotym zgodnym partnerom. Jasne jest, że wszyscy trzech mają równe szanse dopóty, dopóki się nie porozumiewają. Gdy jednak powstanie koalicja dwóch, którzy przed partią umówią się grać zawsze zgodnie i ustalą kolejność orłów i reszek, to wartość gry dla każdego z koalitionów będzie $\frac{1}{2}$, a dla outsidera -1 . Tym sposobem koalicja zmienia wartość gry. Stąd trudny problemat obliczenia wartości gier n -osobowych przy dopuszczeniu wszelkich koalicji. Mógłby ktoś zaproponować, żeby zmodyfikować gry przez zakaz koalicji. Taki zakaz przeciałby wprawdzie matematyczną kwestię, ale zwęziłby zakres zastosowań; przed chwilą stwierdziliśmy, że wiele sytuacji praktycznych redukuje się do gry przeciw przyrodzie; czy można zakazać ludziom wchodzenia w koalicje w tej grze bez nedorzecznego zmniejszania korzyści społecznej z postępowania? Ekonomisci znajdują zapewne lepsze przykłady.

Autor uważa za swój pierwszy obowiązek wykład teorii Neumanna i Morgensterna, chociaż ma ją za niekompletną. Za nimi definiuje funkcję charakterystyczną gry zerowej n -osobowej (dalej, aż do odwołania, termin *gra* będziemy rozumieli jako prostokątną grę n -osobową o sumie zero). Funkcję tę określa się na wszystkich podzbiorach T zbioru partnerów N . Jeżeli ten zbiór rozpadnie się na dwie koalicje T i $N-T$, to gra zamieni się na dwuosobową grę osoby prawnej T przeciw osobie prawnej $N-T$. Gra ta podpada pod teorię dwuosobowych gier prostokątnych i da się domknąć przez taktyki mieszane. Tak otrzymamy wartość $v(T)$ nowej gry i funkcja $v(T)$ jest właśnie funkcją charakterystyczną *ex definitione*. Spełnia ona następujące trzy związki: 1° $v(N) = 0$, 2° $v(N-T) = -v(T)$, 3° dla wszelkich dwóch rozłącznych podzbiorów R i T zbioru N zachodzi nierówność $v(R \cup T) \geq v(R) + v(T)$. Prawdziwe jest twierdzenie odwrotne: Jeżeli jakaś funkcja zbioru g spełnia te trzy warunki, to istnieje gra, dla której jest ona funkcją charakterystyczną. Kwestię tendencji do łączenia się w koalicje ujmuje się na bazie funkcji charakterystycznych. Dzieli się mianowicie wszystkie gry n -osobowe na istotne i nieistotne. Nieistotne to takie, że żadna koalicja nie może zapewnić większej łącznej wygranej, niż suma wygranych, które mogą sobie zapewnić jej członkowie bez wchodzenia w koalicje. Istotne, to pozostałe gry. Warunkiem koniecznym i dostatecznym na nieistotność jest addytywność funkcji charakterystycznej. W grach istotnych musi nastąpić tendencja do łączenia się w koalicje.

W paragrafie 2 (rozdział XV) rozpatruje się gry w zredukowanej postaci. Oznaczmy funkcję wypłaty dla gracza $\{i\}$ przez $M_i(x_1, x_2, \dots, x_n)$

($i = 1, 2, \dots, n$). Struktura gry się nie zmieni, jeżeli do każdej funkcji M_i dodamy stałą a_i tak, by $M_i + a_i = 0$, jak również, gdy wszystkie funkcje M_i pomnożymy przez dodatni współczynnik k . Niech M_i i M'_i oznaczają funkcje wypłaty w grach Γ i Γ' . Gry te nazywają się *strategicznie równoważne*, jeżeli istnieje wzajemnie jednoznaczne przekształcenie $f = (f_1(x_1), f_2(x_2), \dots, f_n(x_n))$ zbioru taktyk każdego gracza gry Γ na zbiór taktyk odpowiedniego gracza gry Γ' takie, że

$$(6) \quad M_i(x_1, x_2, \dots, x_n) = kM'_i(f_1(x_1), f_2(x_2), \dots, f_n(x_n)) + a_i$$

dla każdego i , gdzie $k > 0$ oraz $\sum_{i=1}^n a_i = 0$. Jeżeli dwie gry są strategicznie równoważne, to ich funkcje charakterystyczne v i v' spełniają związek

$$(7) \quad v(T) = kv'(T) + \sum_{i \in T} a_i,$$

gdzie liczby a_i i k są zdefiniowane przez (6). Można więc mówić o równości funkcji charakterystycznych — dwie funkcje v i v' są równoważne, jeśli spełniają równanie (7).

Autor interesuje się grami, których funkcje charakterystyczne spełniają związek

$$(8) \quad v(\{1\}) = v(\{2\}) = \dots = v(\{n\}) = s,$$

gdzie s jest 0 lub 1. Gry te nazywa *zredukowanymi o module s* . Każda funkcja charakterystyczna jest strategicznie równoważna dokładnie jednej funkcji charakterystycznej spełniającej związek (8), więc przy rozpatrywaniu kwestii koalicji wystarczy rozpatrywać gry zredukowane. Warunkiem koniecznym i dostatecznym na to, żeby gra była nieistotna, jest równoważność jej z grą zredukowaną o module zero.

Zagadnieniu łączenia się w koalicje i określenia optymalnych taktyk w n -osobowej grze poświęca McKinsey następne dwa rozdziały (XVI i XVII). W tym celu wprowadza pojęcie imputacji. Niech $(u_1, \dots, u_n)$ będzie wektorem o składowych $u_1, \dots, u_n$, gdzie u_i oznacza sumę, jaką otrzyma gracz $\{i\}$ po zakończeniu partii. Składowe tego wektora zależą od wyniku partii, od podziału na koalicje oraz od sposobów rozliczania się wewnątrz tych koalicji. Oczywiście (gdyż jak założyliśmy, rozpatrujemy gry o sumie zero) $\sum_{i=1}^n u_i = 0$. Poza tym tylko taki wektor może być zrealizowany w konkretnej partii, dla którego $x_i \geq v(\{i\})$ (gdyż gracz $\{i\}$ *ex definitione* ma taktykę gwarantującą mu wartość $v(\{i\})$ bez wchodzenia w koalicje). Wektor o tych własnościach autor nazywa *imputacją*. Między imputacjami wprowadza pojęcie dominacji: Imputacja $(u_1, \dots, u_n)$ *dominuje* imputację $(w_1, \dots, w_n)$, jeżeli istnieje taki niepusty zbiór graczy T , że $v(T) \geq \sum_{i \in T} u_i$ i $u_i > w_i$ dla każdego $i \in T$. Mogą więc istnieć

imputacje, które się wzajemnie dominują! Intuicyjnie jest dosyć oczywiste, że w aktualnej partii mogłaby być realizowana każda taka imputacja, której nie dominuje żadna inna imputacja, gdyż żaden zbiór graczy nie miałby ekonomicznych motywów do zamiany jej na inną; jeżeli w trakcie partii taka imputacja została zaproponowana, to każdy z graczy powinien się na nią zgodzić, gdyż w jakiejkolwiek koalicji by grał i chociażby zawarł najkorzystniejszą umowę o sposobie rozliczania się wewnątrz tej koalicji (ale oczywiście taką, która by każdemu z jego kontrahentów $\{i_j\}$ gwarantowała wartość $v(\{i_j\})$), nie otrzyma nic więcej, jeśli przeciwnicy będą grać dobrze. Jeżeli jednak gra jest istotna, to takich imputacji nie ma, może jednak istnieć zbiór A imputacji o tej własności, że żadna imputacja z A nie jest dominowana przez inną imputację z A i każda imputacja spoza A jest dominowana przez pewną imputację z A . Taki zbiór autor nazywa *rozwiązaniem*.

Zauważmy zasadniczą różnicę w intuicyjnym usprawiedliwieniu użycia słowa „rozwiązanie” w przypadku gier dwuosobowych a n -osobowych ($n > 2$). W pierwszym przypadku rozwiązanie oznacza zbiór prawdopodobieństw, z jakimi gracz stosuje swoje czyste strategie w celu maksymalizacji oczekiwanej wygranej, w drugim przypadku rozwiązanie podaje zbiór możliwych sposobów podziału wygranej na końcu partii. Można by czuć zawód z takiego postawienia kwestii; mogłoby się nasunąć pytanie, czy gracz znając rozwiązanie danej gry n -osobowej mógłby grać z większą oczekiwaną korzyścią, niż gdyby wcale nie znał tej teorii. Dla tego też niektórzy uważają, że teoria Neumanna i Morgensterna nie odpowiada wymogom praktyki. Poza tym nie wiadomo, czy dla $n \geq 5$ każda n -osobowa gra ma rozwiązanie w wyżej zdefiniowanym sensie.

W dalszej części rozdziału XVI autor wprowadza pojęcie izomorfizmu. Dwie gry są izomorficzne, jeżeli istnieje wzajemnie jednoznaczna odpowiedniość między imputacjami jednej i drugiej gry zachowująca relację dominacji. Każde dwie gry strategicznie równoważne są izomorficzne. Operując pojęciem izomorfizmu autor pokazuje, że jeżeli każda gra w zredukowanej formie ma rozwiązanie, to każda gra ma rozwiązanie.

W rozdziale siedemnastym rozpatruje się gry, dla których suma wypłat niekoniecznie jest równa zeru. Zauważmy, że gry takie można traktować jako gry $(n+1)$ -osobowe o sumie 0. Można mianowicie dołączyć fikcyjnego $(n+1)$ -go gracza i założyć, że funkcja wypłaty dla niego jest równa sumie funkcji wypłaty dla pozostałych graczy pomnożonej przez -1 . Trzeba byłoby jednak wprowadzić dość sztuczne ograniczenia dla takich gier, jak zakaz wchodzenia tego gracza w koalicje itd. Ta konstrukcja posłuży nam jednak do rozszerzenia pojęcia funkcji charakterystycznej na przypadek ogólny: Niech Γ będzie pewną grą n -osobową, a Γ' jej $(n+1)$ -osobowym przedłużeniem zdefiniowanym wyżej. Niech v' będzie

funkcją charakterystyczną gry Γ' . Funkcję charakterystyczną v gry Γ definiuje się jako identyczną z funkcją v' na zbiorach nie zawierających $(n+1)$ -go gracza. Inne pojęcia, jak np. gry istotnej, gry zredukowanej, imputacji, rozwiązania, definiuje się tak samo, jak poprzednio.

Cała teoria v. Neumanna dla gier wieloosobowych opiera się na pojęciu funkcji charakterystycznej. Jeżeli dwie gry mają identyczne funkcje charakterystyczne, to mają identyczne rozwiązania w sensie v. Neumanna. Autor uważa, że takie postawienie sprawy jest rzeczą co najmniej wymagającą dyskusji. Rozpatrzmy np. dwuosobową grę o macierzach wypłaty dla gracza X i Y zdefiniowanych następująco:

$$\begin{array}{cc} M_1 & M_2 \\ \begin{array}{c} \parallel 0 \ 10 \parallel, \\ \parallel -1000 \ 0 \parallel. \end{array} \end{array}$$

Oczywiście jest to gra niezerowa. Gracz X ma tu jedną taktykę, a gracz Y dwie (wybrać pierwszą lub drugą kolumnę). Jeżeli Y wybierze pierwszą kolumnę, to X otrzyma 0 zł, a Y straci 1000 zł; jeżeli Y wybierze drugą kolumnę, to X otrzyma 10 zł, a Y 0 zł. Intuicyjnie jest jasne, że X jest tutaj w lepszej pozycji niż Y . Jeżeli bowiem Y będzie maksymizował własny zysk, to wygra 0, a X wygra 10 zł. Można by naturalnie pomyśleć, że Y wymusi na X -ie pewną część wygranej za pomocą groźby, że w innym przypadku wybierze pierwszą kolumnę, ale ponieważ sam na tym straciłby 1000 zł, więc należy przypuszczać, że X nie potraktuje tej groźby poważnie. Zresztą gdyby Y otrzymał nawet pewną część wygranej, to trzeba się zgodzić, że znajduje się w gorszej sytuacji i każde rozwiązanie powinno uwzględniać tę asymetrię. To nie jest jednak prawdą dla rozwiązania w sensie v. Neumanna; definiuje on rozwiązanie na podstawie funkcji charakterystycznej, a ta dla naszego przykładu jest symetryczna: $v(X) = 0$, $v(Y) = 0$, $v(X, Y) = 10$.

W rozdziale osiemnastym jest mowa o otwartych problematach w teorii gier. Większość z nich są to problemy conceptualne, to jest takie, które wymagają wprowadzenia nowych pojęć, a nie stosowania znanych. Wyobraźmy sobie np. grę, w której każdy z partnerów wybiera funkcję z L^2 o normie 1, a wypłata jest równa odległości wybranych funkcji. Trudność polega na braku właściwej miary w L^2 . Jak określać mieszane taktyki, gdy czystych taktyk jest więcej niż kontinuum? Innym zagadnieniem jest znalezienie racjonalnego sposobu postępowania, gdy wartości funkcji wypłaty mają różną użyteczność dla różnych graczy: Tak na przykład 1 zł ma inne znaczenie dla ubogiej osoby i dla bogatej.

Najaktualniejszym dezyderatem jest zbudowanie bardziej zadowalającej teorii gier wieloosobowych. Jedną z takich prób jest teoria J. F. Nasha. Dzieli on mianowicie gry na kooperatywne i niekooperatywne.

W grach niekooperatywnych gracze nie mogą się porozumiewać, a więc zawierać umów itd. W grach kooperatywnych komunikacja jest dozwolona. Nash uważa niekooperatywne gry za podstawowe i redukuje do nich kooperatywne gry w następujący sposób: Pertraktacje dotyczące zawierania transakcji włącza on jako formalne posunięcia do niekooperatywnych gier (zatem posunięciami mogą być propozycje utworzenia koalicji, podziału zysków itp.). Chociaż teoria Nasha reprezentuje pewien postęp w tej dziedzinie, ma ona poważne niedociągnięcia. W praktyce np. jest niesłychanie trudno wprowadzić do niekooperatywnych gier posunięcia odpowiadające pertraktacjom między poszczególnymi graczami i to w sposób sprawiedliwy dla wszystkich graczy. Poza tym już sama teoria niekooperatywnych gier ma poważne braki.

Bibliografia obejmuje około 120 tytułów. Nie ma w niej ani jednego nazwiska matematyków polskich. Niech mi będzie wolno wobec tego przypomnieć nieliczne pozycje naszych autorów. Zacznę od felietonu w „Myśli Akademickiej” jednodniówce studenckiej, która ukazała się we Lwowie w 1925 r., *Definicje potrzebne do teorii gry i pościgu*. Chronologicznie wyprzedza ona wszystkie pozycje cytowane przez McKinsey’a. Można w niej znaleźć definicję minimaksowych taktyk dla takiej gry jak pościg. Przez wiele lat nikt nie zajął się pościgiem. Dopiero prof. Choquet, który po wojnie dowiedział się we Wrocławiu o zagadnieniu pościgu, napisał rozprawę, której niestety nie możemy odnaleźć w bibliografii. Zagadnienie pościgu przy brzegu zbadał M. Warmus, a A. Zięba rozwiązał zagadnienie pościgu jednego statku przez dwa na pełnym morzu. Wcześniej (1943) powstało zagadnienie podziału pragmatycznego, który jest grą — rozwiązałem je dla trzech partnerów, a S. Banach i B. Knaster dla dowolnej liczby. Częściową informację o tych rezultatach daje *Kalejdoskop matematyczny*. Niedawno ogłoszono przedwojenną grę Banacha i Mazura, która polega na pisaniu kolejnym przez dwóch partnerów cyfr rozwinięcia decymalnego liczby rzeczywistej — jeżeli padnie ona z góry w dany zbiór A wygrywa X , jeżeli w CA , wygrywa Y .⁽²⁾ Z teorią gier wiążą się rezultaty uzyskane przez J. Schreiera przed wojną i dotyczące najmniejszej liczby rozgrywek potrzebnych do klasyfikacji graczy, gdy każda rozgrywka decyduje o lepszości jednego z dwóch partnerów; ta teoria ma zastosowanie przy mierzeniu przez porównywanie. C. Ryll-Nardzewski ulepszył teorię pościgu wpro-

⁽²⁾ S. Zubrzycki, *On the game of Banach and Mazur*, Coll. Math. IV (1957), str. 227-229; A. Zięba, *An example of the game of Banach and Mazur*, Coll. Math. IV (1957), str. 230-231.

Treść odsyłaacza ¹⁾ na str. 227 oraz data przy tytule [3] na str. 229 wynikały z nieporozumienia: cytowana praca była niedostępna od roku 1941 aż do jesieni roku 1957 i cytowano błędnie z pamięci jej tytuł i datę wydania.

wadząc pod jeden z dwóch symboli „min”, „max” zamiast zbioru taktyk zbiór rejsów, co logicznie upraszcza teorię a zapewne nie tylko w tej grze jest możliwe. Tenże uogólnił twierdzenie o kategoryczności gier skończonych; niestety znaczna część wspomnianych wyników nie została ogłoszona.

Na zakończenie podam definicję „warszawskiej szubienicy uogólnionej”. W punktach kratowych na płaszczyźnie X i Y stawiają na przemian swoje symbole. Wygrywa ten, który pierwszy ustawi obok siebie n swoich symboli pionowo lub poziomo. Dla $n \geq 5$ gra przestaje być trywialna. Wspomnieć również należy o „grze wrocławskiej”. Każdy partner wpłaca dowolną kwotę (co najmniej 1 zł), jej wysokość notuje arbiter, ale inni partnerzy jej nie znają; gdy wszyscy wpłacili, oblicza się średnią wpłatę i cała pula przypada temu, którego wpłata najmniej się różni od średniej.

Od daty ukazania się książki McKinseya uzyskano w Polsce nowe rezultaty; jeden z młodych adeptów teorii gier, S. Trybuła, przyczynił się do uzupełnienia niniejszej recenzji, która w pierwotnej formie czekała parę lat na ogłoszenie (nie z winy autora). W chwili oddania jej do druku ukazał się już trzeci tom publikacji ciągłej *Contributions to the Theory of Games* (Princeton Univ. Press), w którym J. C. Oxtoby pisze o grze Banacha i Mazura, cytując m. i. pracę Mycielskiego, Świerczkowskiego i Zięby ogłoszoną w Bull. Acad. Pol. Sci., cl. III, 4(1956), str. 485-488. W Colloquium Mathematicum IV 2 (1957) podaje A. Zięba elementarny dowód minimaxowego twierdzenia v. Neumanna.