

MATERIAŁY I STUDIA

Zeszyt nr 165

Behawioralna ekonomia finansowa

Modyfikacja paradygmatów funkcjonujących
w nowoczesnej teorii finansów

Anna Cieślak

Projekt graficzny:

Oliwka s. c.

Skład i druk:

Drukarnia NBP

Wydął:

Narodowy Bank Polski
Departament Komunikacji Społecznej
00-919 Warszawa, ul. Świętokrzyska 11/21
tel. (22) 653 23 35, fax (22) 653 13 21

© Copyright Narodowy Bank Polski, 2003

Materiały i Studia rozprowadzane są bezpłatnie.

Dostępne są również na stronie internetowej NBP: <http://www.nbp.pl>

Spis treści

Spis rysunków i tabel	5
Streszczenie	6
Wstęp	7
1. Standardowa teoria finansów w konfrontacji z podejściem behawioralnym	11
1.1. Łamigłówki rynku kapitałowego	12
1.1.1. Premia za inwestycję w aktywa ryzykowne	12
1.1.2. Nadmierna zmienność kursów giełdowych	14
1.1.3. Zmienność zbyt zmienna	18
1.1.4. Nadmierne obroty giełdowe	19
1.1.5. Problem funduszy zamkniętych	21
1.2. Szum na rynku	24
1.2.1. Koncepcja inwestowania na podstawie szumu	24
1.2.2. Nadmierna i niedostateczna reakcja cen na informacje	28
1.3. Ograniczenia skuteczności arbitrażu	32
1.3.1. Ryzyko fundamentalne	33
1.3.2. Ryzyko szumu	34
1.3.3. Koszty implementacji strategii arbitrażowej	37
1.3.4. Koszty psychologiczne krótkiej sprzedaży	39
1.3.5. Ryzyko modelu	40
1.3.6. Warunki ograniczenia arbitrażu	40
2. Percepcja i przetwarzanie informacji	43
2.1. Wnioskowanie o nieznanym	43
2.1.1. Prawa prawdopodobieństwa i statystyki	44
2.1.2. Ograniczona racjonalność	45
2.1.3. Szkoła Tversky'ego i Kahnemana	46
2.1.4. Szkoła Gigerenzera	47
2.2. Heurystyki i systematyczne błędy wnioskowania	50
2.2.1. Dostępność	50
2.2.2. Reprezentatywność	55
2.2.3. Zakotwiczenie i dostosowanie	63
2.3. Rozszerzenia	64
2.3.1. Nadmierna pewność siebie	64
2.3.2. Selektywna percepcja i błąd afirmacji	66
2.3.3. Błąd post factum	67
3. Prospect theory: podejmowanie decyzji w warunkach niepewności	69
3.1. Normatywna teoria podejmowania decyzji w warunkach ryzyka	70
3.1.1. Wartość oczekiwana a oczekiwana użyteczność	70
3.1.2. Aksjomatyzacja preferencji	73
3.2. Anomalie w kształtowaniu się preferencji	76
3.2.1. Zaniedbywanie prawdopodobieństwa	77

3.2.2. Odwrócenie preferencji	77
3.2.3. Paradoks Allais	78
3.2.4. Problem Samuelsona	79
3.3. Deskryptywna teoria podejmowania decyzji w warunkach niepewności	81
3.3.1. Wartość oczekiwań	82
3.3.2. Funkcja wartości	82
3.3.3. Funkcja wag prawdopodobieństwa	87
3.3.4. Efekt odwrócenia	92
3.3.5. Kadrowanie	93
4. Praktyczne znaczenie <i>prospect theory</i>	97
4.1. Analiza anomalii w zachowaniach ekonomicznych	97
4.1.1. Krótkowzroczna awersja do strat	98
4.1.2. Efekt dyspozycji	101
4.1.3. Efekt utopionych kosztów	107
4.1.4. Efekt obdarowania oraz skłonność do zachowania <i>status quo</i>	112
4.2. Mentalna księgowość	116
4.2.1. Definicja i znaczenie	116
4.2.2. Hedonistyczne kadrowanie zysków i strat	117
4.2.3. Zastosowania	121
5. Wpływ czasu na podejmowanie decyzji	127
5.1. Percepcja przyszłości	127
5.1.1. Normatywny model wyboru międzyokresowego	127
5.1.2. Anomalie wyboru międzyokresowego	130
5.1.3. Deskryptywny model wyboru międzyokresowego	137
5.2. Pamiętanie przeszłości	142
5.2.1. Adaptacja	143
5.2.2. Retrospektywna ocena użyteczności	149
6. Zakończenie	154
7. Bibliografia	157
8. Literatura przedmiotu	164

Spis rysunków i tabel

1.1. Kurs akcji a wartość bieżąca dywidend	16
1.2. Odchylenia kursu Royal Dutch względem Shell od parytetu 60:40	36
3.1. Hipotetyczna funkcja wartości $v(x)$	83
3.2. Funkcja wag prawdopodobieństwa	88
4.1. Hedonistyczne kadrowanie zysków i strat w obrębie portfela	119
5.1. Dyskontowanie w warunkach niepewności	137
5.2. Dyskontowanie hiperboliczne	140
5.3. Zastosowanie funkcji wartości do analizy procesów adaptacyjnych	145
 Tabela 3.1. Paradoks Allais	 78
Tabela 3.2. Oczekiwane użyteczności w eksperymencie Allais	79
Tabela 3.3. Cztery wymiary stosunku do ryzyka	92
Tabela 5.1. Przykładowy przebieg adaptacji w sytuacji strat	147

Streszczenie

W niniejszej pracy zaprezentowano nurt ekonomii behawioralnej jako podejście alternatywne wobec paradygmatów wspartych racjonalnością i leżących u podstaw teorii finansów. Przeprowadzona tu analiza miała wykazać, czy wprowadzenie czynnika behawioralnego do nauki finansów pozwala na lepsze zrozumienie zachowań uczestników rynku.

Punkt wyjścia tych rozważań stanowi prezentacja anomalii rynkowych obserwowanych w makroskali i nie dających się wyjaśnić w ramach tradycyjnych koncepcji, jak np. CAPM czy APT. Dalej, już z mikroperspektywy, omówiony został proces decyzyjny jednostki ze szczególnym wskazaniem na pojawiające się w nim systematyczne odstępstwa od racjonalności. W pierwszym kroku skupiono się na zagadnieniu wnioskowania o nieznanym w oparciu o heurystyki oraz wskazano pojawiające się przy tym błędy postrzegania i przetwarzania informacji. Następnie zaprezentowano deskryptywny model podejmowania decyzji w warunkach niepewności – *prospect theory* Kahnemana i Tversky’ego. Wiele uwagi poświęcono zastosowaniom wspomnianej teorii w tłumaczeniu zachowań uznawanych za ekonomicznie nieracjonalne. Wreszcie, w ostatniej części opracowania, wcześniejsze rozważania wzbogacone zostały o element dynamiczny – wpływ czasu na dokonywanie wyborów ekonomicznych.

Słowa kluczowe: racjonalność, finanse behawioralne, ekonomia behawioralna, heurystyki, systematyczne błędy wnioskowania, prospect theory, teoria podejmowania decyzji w warunkach niepewności, anomalie wyboru międzyokresowego, dyskontowanie hiperboliczne, adaptacja

Klasyfikacja JEL: G11, G12, G14, D81, D91

Wstęp

Złoty wiek założenia o racjonalności jednostki w nauce ekonomii rozpoczął się w latach czterdziestych XX stulecia, stając się punktem wyjścia dla rozwoju nowoczesnej teorii finansów. Teoria ta opiera się na koncepcji *homo oeconomicus* – reprezentatywnej jednostki, która, uwzględnivszy wszelkie dostępne informacje, podejmuje decyzje w sposób maksymalizujący jej oczekiwaną użyteczność.

Ta ortodoksyjna koncepcja wyborów ekonomicznych stała się przesłanką powstania i ugruntowania się w latach siedemdziesiątych zunifikowanego podejścia do zachowań rynków finansowych na poziomie zagregowanym. I tak na mocy hipotezy efektywności rynku przyjęło się uważać, iż ceny aktywów finansowych ustalane są racjonalnie, dzięki czemu odzwierciedlają wszelkie dostępne informacje fundamentalne. Efektywność cen gwarantowana jest istnieniem racjonalnych inwestorów – arbitrażystów. Nie szkodzi jej również obecność jednostek nieracjonalnych, których działania są stochastyczne, przez co niwelują się wzajemnie.

Jednakże z upływem czasu dojrzało również pokolenie młodych ekonomistów, którzy chcieli odciąć się od schedy po propagatorach racjonalności. Standardowa teoria finansów przestała być dla nich satysfakcjonującym wyjaśnieniem obserwowanych ruchów cen giełdowych. Zjawiska te – niewytłumaczalne na gruncie tradycyjnych modeli równowagi rynkowej – nazwane zostały anomaliami. Stały się one motywacją do rozważań nad wpływem kognitywnych błędów popełnianych przez inwestorów na kształtowanie się cen aktywów. Tak oto w latach osiemdziesiątych XX wieku powstał agnostyczny, bowiem uznający ludzką nieracjonalność za regułę,

nurt analizy rynków – behawioralna ekonomia finansowa¹. Jego dociekania koncentrują się na systematycznych błędach popełnianych w procesach wnioskowania i podejmowania decyzji.

Jakkolwiek zwolennicy hipotezy efektywnego rynku wciąż poddają ostrej krytyce dokonania teoretyków behawioralnych, miejsce tego podejścia w nauce ekonomii i finansów zostało niedawno „uprawomocnione” poprzez uhonorowanie psychologa Daniela Kahnemana nagrodą Nobla w dziedzinie ekonomii. Jego prace, tworzone wspólnie ze zmarłym w 1996 roku Amosem Tversky’em, położyły podwaliny behawioralnego nurtu finansów.

Przedmiotem tego opracowania jest prezentacja alternatywy dla funkcjonujących w nauce finansów paradygmatów, które bazują na racjonalności. Owa alternatywna koncepcja opiera się na mikroekonomicznej analizie zachowań jednostek działających na rynkach finansowych. Celem pracy było stwierdzenie, czy wprowadzenie czynnika behawioralnego do nauki finansów pozwala lepiej zrozumieć działania uczestników rynku. Ów cel zrealizowano, poszukując odpowiedzi na następujące pytania:

- (1) Czy jednostki zdolne są wnioskować o nieznanym w sposób całkowicie obiektywny – zgodny z teorią rachunku prawdopodobieństwa i statystyki? Jeśli nie: czy ich sposób percepcji i przetwarzania informacji ulega systematycznym zniekształceniom?
- (2) Czy teoria oczekiwanej użyteczności, będąca powszechnie akceptowaną normatywną koncepcją podejmowania decyzji w warunkach ryzyka, stanowi również adekwatny opis rzeczywistych zachowań jednostek? Jeśli nie: czy zadaniu temu jest w stanie sprostać model deskryptywny – *prospect theory* – zaproponowany przez Kahnemana i Tversky’ego?
- (3) Czy decyzje jednostek wykazują dynamiczną konsekwencję, postulowaną przez teorię zdyskontowanej użyteczności? Jeśli nie: czy

¹ W literaturze brak zgodnego nazewnictwa określającego tę dziedzinę nauki. Mimo iż za najtrafniejsze należy uznać sformułowanie „behawioralna ekonomia finansowa”, w pracy tej posłużono się nim zamiennie z pojęciami takimi jak „behawioralny nurt finansów”, „behawioralne podejście do finansów” czy też prościej „finanse behawioralne”.

zniekształcenia międzyokresowego procesu decyzyjnego dają się wyjaśnić w ramach alternatywnego modelu deskryptywnego?

Odpowiedzi na postawione pytania badawcze udzielono w oparciu o przebadaną literaturę, jak również doświadczenia, jakie autorka zdobyła dzięki współpracy z firmą Cognitrend, zajmującą się behawioralnymi analizami rynków finansowych. Najcenniejszym źródłem wiedzy w zakresie behawioralnej ekonomii finansowej, z którego korzystano przy tworzeniu opracowania, były artykuły znanych ekonomistów i psychologów publikowane na łamach renomowanych czasopism naukowych. Za zasadne uznano bezpośrednio zapoznanie się nie tylko z najnowszymi opracowaniami z dziedziny, lecz także z pracami wcześniejszymi. Szczególnie ważne były powstałe już w latach siedemdziesiątych dzieła Daniela Kahnemana i Amosa Tversky'ego, których dorobek intelektualny należy postrzegać jako kluczowy dla zrozumienia podejścia behawioralnego.

Krótką historią badanej dziedziny narzuca pewne ograniczenia. Po pierwsze, dotychczas powstały jedynie nieliczne publikacje o charakterze kompendium; większość prac poświęcona jest specyficznym zagadnieniom. Po drugie, bardzo wąski zbiór literatury w języku polskim nie pozwala w sposób kompetentny odpowiedzieć na postawione tu problemy badawcze. Dlatego korzystano niemal wyłącznie z opracowań zagranicznych, w szczególności z pozycji anglojęzycznych oraz z nielicznych prac pisanych w języku niemieckim. Stąd też wynikły problemy terminologiczne: niektóre specjalistyczne pojęcia, jakimi się posłużono, są własnymi propozycjami nazewnictwa autorki.

* * *

Praca ta obejmuje pięć rozdziałów. Rozdział pierwszy traktuje o zachowaniach rynku obserwowanych z makroperspektywy. Omówiono tu anomalie rynkowe, których nie można wyjaśnić w ramach ugruntowanych modeli standardowej teorii finansów. Ponadto zaprezentowano dualistyczną koncepcję racjonalności inwestora i wiążące się z nią ograniczenia skuteczności

arbitrażu. Celem tych rozważań jest ukazanie potrzeby zrewidowania poglądów na temat *homo oeconomicus*. Stąd też kolejne rozdziały koncentrują się na mikroanalizie poczynąń inwestora.

W rozdziale drugim skupiono się na procesach wnioskowania o nieznanym. Omówiony tu etap „edycji” informacji, pomijany w podejściu tradycyjnym, przedstawiony został jako istotny determinant jakości podejmowanych następnie decyzji. Zwrócono przy tym szczególną uwagę na systematyczność błędów wnioskowania popełnianych przez jednostki.

W rozdziale trzecim skonfrontowano normatywną teorię oczekiwanej użyteczności z rzeczywistością. Wskazano na zjawiska będące pogwałceniem aksjomatów racjonalnego wyboru, a następnie omówiono deskryptywny model podejmowania decyzji w warunkach niepewności – *prospect theory* Kahnemana i Tversky’ego.

Czwarty rozdział prezentuje zastosowania *prospect theory* w wyjaśnieniu zachowań ekonomicznych niezgodnych z paradygmatami teorii wspartych racjonalnością. Została w nim przeprowadzona analiza diagnostyczna takich zachowań, a także zbadano ich konsekwencje.

Rozdział piąty, ostatni, wzbogaca dotychczasowe rozważania o element dynamiczny – wybór międzyokresowy. Zagadnienie to zostało rozpatrzone z dwóch perspektyw: przyszłości – w znaczeniu dyskontowania w czasie oraz przeszłości – w kontekście procesów adaptacyjnych i retrospektywnej oceny zdarzeń.

1

Standardowa teoria finansów w konfrontacji z podejściem behawioralnym

1

Celem niniejszego rozdziału jest wskazanie problematycznych punktów, w których tradycyjne, utrzymane w duchu nowoczesnej teorii finansów, myślenie o rynkach okazało się dalece niewystarczające. Omówione tu anomalie opisują zachowania zagregowanego rynku (makroskala). Trudność ich interpretacji stanowi motywację do mikroekonomicznej analizy zachowań jednostek działających na rynku, którą przeprowadzono w dalszych częściach tego opracowania.

W pierwszym kroku (podrozdział 1.1) przedyskutowane zostaną nierozwiązane – nie pozwalające się wyjaśnić na gruncie modeli równowagi rynkowej – „zagadki” rynków finansowych. Są to: premia za inwestycję w aktywa ryzykowne, nadmierna zmienność kursów, zbyt wysokie obroty giełdowe, a także problem wyceny funduszy zamkniętych. W trakcie prezentacji owych zagadnień zasygnalizowano wytłumaczenia behawioralne, które zostaną rozwinięte w kolejnych rozdziałach.

Następnie (w części 1.2) skupiono się na dualistycznej koncepcji rynku, wprowadzającej rozróżnienie na racjonalnych arbitrażystów oraz niepoinformowanych inwestorów, którzy podejmują decyzje na podstawie szumu informacyjnego.

Centralnym tematem części trzeciej (podrozdział 1.3) jest zagadnienie nieskuteczności arbitrażu. Analiza technicznych ograniczeń w stosowaniu arbitrażu pozwala częściowo odpowiedzieć na pytanie, dlaczego rynki bywają nieefektywne.

Warto zaznaczyć, że ostatecznym celem ekonomistów behawioralnych jest zrozumienie wpływu systematycznych błędów ludzkich na nieefektywność rynków finansowych. Zagadnienia poruszane w tym rozdziale, jakkolwiek istotne, należy rozpatrywać w kategoriach genezy podejścia behawioralnego. Odkryte anomalie i techniczne przyczyny nieefektywności cen stanowiły

bowiem pierwotną motywację badaczy do zainteresowania się indywidualnymi procesami wnioskowania i podejmowania decyzji w inwestowaniu.

1.1 Łamigłówki rynku kapitałowego

Ostatnie dwie dekady przyniosły szereg odkryć anomalii rynkowych, nie dających się wyjaśnić w ramach tradycyjnych koncepcji, jak np. model wyceny aktywów kapitałowych (CAPM – *capital asset pricing model*) czy teoria arbitrażu cenowego (APT – *arbitrage pricing theory*). Owe łamigłówki rynku kapitałowego zostaną obecnie omówione wraz ze wskazaniem potencjalnych ich wytłumaczeń w świetle teorii behawioralnych.

1.1.1 Premia za inwestycję w aktywa ryzykowne

Termin „zagadka premii za ryzyko” (*equity premium puzzle*) został wprowadzony do literatury finansów w 1985 roku przez Mehra i Prescottta (Mehra i Prescott, 1985). Stojące za tym określeniem zjawisko odnosi się do zdumiewająco wysokich stóp zwrotu z inwestycji w akcje (aktywa ryzykowne) w stosunku do obligacji, obserwowanych w toku całej niemal historii amerykańskiego rynku kapitałowego. Tzw. premia za inwestycję w aktywa ryzykowne, czyli różnica między stopami zwrotu z akcji i obligacji, była w przeszłości na tyle wysoka, że standardowe modele (np. CAPM) nie mogły dostarczyć jej satysfakcjonującego wytłumaczenia.

Obliczona dla rynku amerykańskiego premia wyniosła w latach 1926–1992 około 6 pkt. proc.² (Siegel, 1994)³. Również w innych krajach dowodzi się występowania tego efektu. Dla Szwecji i Australii premia plasuje się na poziomie około 8 pkt. proc., zaś dla Włoch, Francji i Japonii między 4 a 5 pkt. proc. (Dimson i in., 2000). Takie dane pozwalają odrzucić argument zwo-

² Przy stopie zwrotu z akcji równej około 7%, zaś z obligacji 1%.

³ Warto podkreślić, że niedawne spadki kursów giełdowych przyczyniły się do zredukowania premii za inwestycję w akcje. Zagadka zatem chwilowo rozwiązała się sama. Jednakże jej wcześniejsze permanentne występowanie oraz wynikające z tego problemy interpretacyjne warte są omówienia.

lenników efektywności rynku, że omawiane zjawisko wiąże się wyłącznie ze specyfiką szybko wzrastającego rynku kapitałowego w USA⁴.

W standardowym podejściu wysoka stopa zwrotu z akcji kojarzona jest z wynagrodzeniem inwestorów za ryzyko. W tle pojawia się ciche i powszechnie akceptowane założenie, że inwestor wykazuje do niego awersję, innymi słowy: jest ostrożny. Jednak zgodnie z CAPM już premia na poziomie 1 pkt. proc. stanowiłaby dostateczne wynagrodzenie za lokatę środków na rynku kapitałowym. Mehra i Prescott oszacowali ponadto, iż dla odnotowanej jej wielkości współczynnik względnej awersji do ryzyka musiałby wynosić około 30, przy czym za normalny uznaje się jego poziom równy 1⁵. Awersja do ryzyka odpowiadająca współczynnikowi 30 została zilustrowana przez Mankiw i Zeldes (1991). Rozważali oni loterię dającą równą szansę wygrania kwoty 100 000 j.p.⁶ lub 50 000 j.p. Tak silna awersja do ryzyka oznacza, że gracz musiałby być obojętny w wyborze między zaproponowaną loterią, a pewną wygraną w wysokości 51 209 j.p. Niewielu osobom ryzyko opisanej loterii może wydawać się aż tak znaczące. A zatem, w świetle dowodów dostarczonych przez badaczy, zarówno tradycyjny argument premii za dodatkowe ryzyko, jak i awersji do ryzyka należy uznać za niewystarczające wytłumaczenie zjawiska.

Obserwacja wysokich stóp zwrotu z akcji w porównaniu z obligacjami nasuwa pytanie o motywy inwestowania w mniej lukratywne obligacje, skoro akcje są tak dobrą lokatą. Wielu badaczy próbowało podać rozwiązanie tej zagadki, niemniej jednak ich wartość poznawcza bywała ograniczona⁷. Najbardziej przekonująca wydaje się argumentacja Benartzi'ego i Thaler (1995) czerpiąca inspirację z *prospect theory* Kahnemana i Tversky'ego

⁴ W całej pracy skrótem USA posłużono się dla oznaczenia Stanów Zjednoczonych Ameryki Północnej.

⁵ Mowa tu o mierze względnej awersji do ryzyka opracowanej przez Arrowa i Pratta (*Arrow-Pratt measure of relative risk aversion*) postaci: $RRA = -xu''(x)/u'(x)$, gdzie x – poziom bogactwa, $u(x)$ – funkcja użyteczności. Należy zauważyć, że współczynnik równy 1 jest charakterystyczny dla logarytmicznej postaci funkcji użyteczności.

⁶ W opracowaniu skrótem „j.p.” posłużono się w celu oznaczenia jednostek pieniężnych.

⁷ Prace tłumaczące premię za ryzyko to m.in. (Constantinides, 1990), (Epstein i Zin, 1990), (Fama i French, 2000).

(por. podrozdział 3.3). Łamiąc konwencjonalne zasady myślenia o decyzjach inwestorów, Benartzi i Thaler wysuwają argument krótkowzrocznej awersji do strat (*myopic loss aversion*). Zaproponowane przez nich rozwiązanie omówione zostanie w rozdziale 4, który poświęcono zastosowaniom *prospect theory* (por. punkt 4.1.1).

1.1.2 Nadmierna zmienność kursów giełdowych

Model efektywnego rynku zakłada, że aktualny kurs akcji P_t jest równy wartości bieżącej oczekiwanych strumieni pieniężnych (dywidend), jakie dana spółka wygeneruje. Tak zostaje wyznaczona wartość fundamentalna jej akcji. Prognoza przyszłych strumieni dywidend tworzona jest na podstawie wszystkich informacji dostępnych w danym momencie t . Na efektywnym rynku bieżący kurs akcji P_t stanowi optymalną prognozę wartości fundamentalnej P_t^* .

Rozumowanie to prowadzi do stwierdzenia, że kurs rynkowy „śledzi” wartość doskonałą (fundamentalną). Jeśli wartość P_t^* ulega zmianie, to kurs akcji P_t powinien zmieniać się w tym samym stopniu. Oczywiście, teoria nie wyklucza sporadycznych „pomyłek” rynku. Na poziomie zagregowanym nie mają one jednak znaczenia, jako że odchylenia kursu P_t *in plus* oraz *in minus* w stosunku do P_t^* znoszą się wzajemnie.

Prosty model kształtowania się ceny na rynku efektywnym można formalnie zapisać jako: $P_t = E_t(P_t^*)$, (1.1)

gdzie E_t określa prognozę uwarunkowaną wszystkimi publicznymi informacjami o spółce dostępnymi w momencie t . Z zapisu tego wynika, że każdy ruch kursu giełdowego musi mieć swoje źródło w nowej wiadomości wpływającej na wartość fundamentalną P_t^* danego papieru wartościowego. Jednocześnie zachodzi:

$$P_t^* = P_t + u_t, \quad (1.2)$$

gdzie u_t to błąd prognozy. Jeśli rynek jest efektywny, błąd prognozy u_t nie może być skorelowany z P_t , gdyż w przeciwnym razie prognoza nie byłaby

optymalna – istniałaby bowiem informacja, która nie została uwzględniona w cenie P_t . Mamy zatem:

$$\text{cov}(P_t, u_t) = 0, \quad (1.3)$$

$$\text{var}(P_t^*) = \text{var}(P_t) + \text{var}(u_t), \quad (1.4)$$

$$\text{var}(u_t) \geq 0 \Rightarrow \text{var}(P_t^*) \geq \text{var}(P_t), \quad (1.5)$$

gdzie:

$\text{cov}(P_t, u_t)$ – kowariancja kursu akcji P_t i błędu prognozy u_t ,

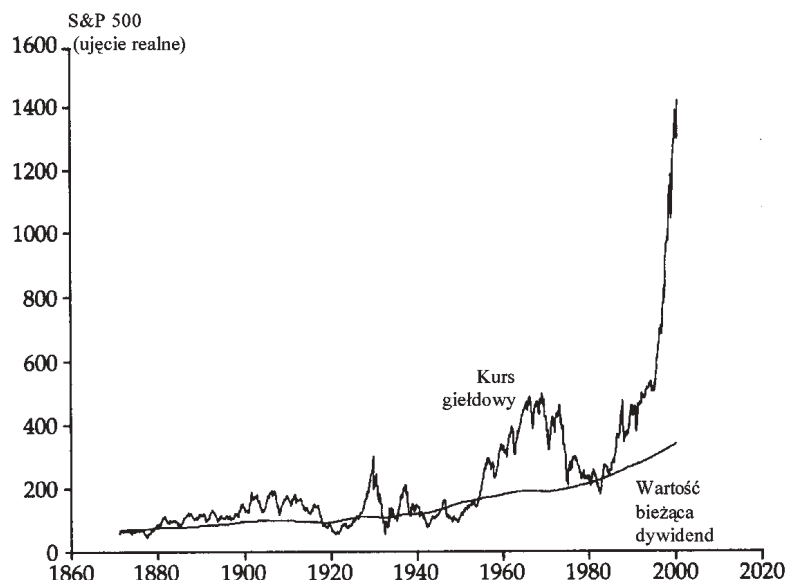
$\text{var}(\cdot)$ – symbol wariancji.

Z nierówności (1.5) wynika, że wariancja prognozy zmiennej P_t^* , czyli $\text{var}(P_t)$, może być co najwyżej równa wariancji tej zmiennej, czyli $\text{var}(P_t^*)$.

Jakie znaczenie dla analizy rynku kapitałowego ma powyższe stwierdzenie? Otóż przedstawiony model pokazuje teoretyczną zależność między kursem akcji a wartością fundamentalną: w warunkach efektywnego rynku wariancja kursów giełdowych (w długim okresie) powinna być wytłumaczalna zmiennością poziomu dywidend.

Badania przeprowadzone przez Shillera (1981, 1990a) dowodzą natomiast, że zmienność obserwowana na giełdzie wielokrotnie przewyższa zmienność wartości bieżącej oczekiwanych dywidend (por. rysunek 1.1)⁸. Oznacza to brak spełnienia warunku (1.5), który nakazuje, by zmienność kursu akcji była co najwyżej równa zmienności samej wartości fundamentalnej. Można stąd powątpiewać w deskryptywną moc powyższego modelu.

⁸ Pierwsze badania nad nadmierną zmiennością kursów prowadzone były równocześnie przez Shillera (1981) oraz LeRoy'a i Portera (1981). Badacze doszli do podobnych wniosków, stwierdzając nadmierną zmienność kursów akcji w stosunku do przewidywań hipotezy efektywnego rynku.

Rysunek 1.1: Kurs akcji a wartość bieżąca dywidend*

* Uwaga: Wartość bieżąca przyszłych dywidend dla danego miesiąca wyznaczona została na podstawie zsumowania wartości bieżących dywidend (w ujęciu realnym) wypłacanych w miesiącach następnych. Czynniki dyskontujące mają postać $(1+r)^{-t}$, gdzie r jest miesięczną realną stopą dyskontową, zaś t oznacza liczbę miesięcy między miesiącem danym a kolejnym. Do wyznaczenia przedstawionej na rysunku 1.1 krzywej wartości bieżącej dywidend zastosowano stałą stopę dyskontową $r=0,6\%$, czyli średnią geometryczną stopę zwrotu (ujęcie realne, skala miesięczna), jaka obowiązywała na rynku od stycznia 1871 roku do czerwca 1999 roku. Użycie stałej stopy r zgodne jest z założeniem, że na efektywnym rynku oczekiwane stopy zwrotu nie zmieniają się w czasie. Przyjęto ponadto, iż od grudnia 1999 roku (ostatni miesiąc, dla którego dostępne były dane) realna wartość dywidend wzrasta $0,1\%$ z miesiąca na miesiąc.

Źródło: (Shiller, 2000, s. 186).

Shiller oparł swoje badania na wartościach indeksu Standard and Poor's 500 (S&P 500) w latach 1871–2000. Posługiwał się metodologią polegającą na: (i) ustaleniu wartości fundamentalnej P_t^* równej sumie zdyskontowanych kolejnych dywidend poczynając od danego punktu w przeszłości przez wszystkie następne okresy⁹, (ii) porównaniu rzeczywistej zmienności indek-

⁹ Wielkość P_t^* należy tu rozumieć jako: $P_t^* = \sum_{\tau=t+1}^{\infty} \delta^{(\tau-t)} D_{\tau}$, gdzie: δ – stały czynnik dyskontujący, D_{τ} – dywidenda w okresie τ .

su ze zmiennością wartości fundamentalnej obliczoną na podstawie (i). Jak wynika z rysunku 1.1, w badanym okresie wartość bieżąca dywidend charakteryzowała się niewielką zmiennością oraz lekkim trendem wzrostowym, podczas gdy wahania kursów akcji były niezwykle gwałtowne. Zjawisko to nabrało szczególnego nasilenia w okresie fascynacji Nową Ekonomią, czyli w czasach „nieracjonalnej euforii” (*irrational exuberance*)¹⁰.

Ponadto, dla poparcia tezy o nadmiernej zmienności kursów giełdowych, Campbell i Shiller (1988) oszacowali, że jedynie 27% rocznej wariancji kursów na rynku amerykańskim można wytłumaczyć napływem informacji o przyszłych dywidendach¹¹.

Mimo iż dotychczas nie powstał spójny model tłumaczący to zjawisko, badacze nurtu behawioralnego wysunęli kilka hipotez, stanowiących potencjalną podstawę do stworzenia bardziej usystematyzowanego podejścia. Hipotezy te odwołują się do reguł, za pomocą których inwestorzy tworzą swoje przekonania o prawdopodobieństwie zdarzeń, tzw. heurystyk (por. podrozdział 2.2). Jednym z przykładów może być stosowanie „prawa małych liczb” (por. punkt 2.2.2.2), czyli wnioskowanie o populacji generalnej na podstawie małej próby danych. I tak inwestor, obserwując wzrost dywidend wypłacanych przez spółkę, zakłada, iż trend wzrostowy utrzyma się także w następnych okresach. Ta prosta ekstrapolacja prowadzi do przeszacowania ceny akcji w stosunku do ich wartości. Kurs wzrasta aż do momentu ogłoszenia faktycznych danych, niższych niż oczekiwane.

¹⁰ Określenie *irrational exuberance* pierwszy raz użyte zostało przez Alana Greenspana podczas jego wystąpienia w grudniu 1996 roku. Już wówczas przewodniczący amerykańskiej Rezerwy Federalnej sygnalizował swoje zaniepokojenie niezwykle wysokimi kursami giełdowymi. Opamiętanie inwestorów przyszło jednak z opóźnieniem.

¹¹ Publikacje o nadmiernych wahaniami kursów wzbudziły wiele kontrowersji dotyczących metodyki badań. Zwolennicy hipotezy efektywnego rynku, która została tu bezpośrednio zaatakowana, podważali np. słuszność założenia stałych realnych stóp dyskontowych, znajdującego się w tle rozważań Shillera. Uważali oni, że nadmierna zmienność rynku może zostać wytłumaczona „racjonalnymi zmianami stóp procentowych”. Jednakże zastosowanie alternatywnych podejść (np. zastąpienie stałej stopy – stopą międzyokresowej substytucji konsumpcji czy też zmienną stopą procentową papierów komercyjnych) wciąż nie daje podstaw do odrzucenia hipotezy o nadmiernej zmienności kursów giełdowych. Aby zapoznać się z dokładnym opisem procedury badawczej stosowanej przez Shillera z uwzględnieniem różnych stóp dyskontowych, por. (Shiller, 2002).

Inna hipoteza zakłada, że inwestorzy mają skłonność do przeceniania dostępnych im informacji prywatnych. Znajduje to odbicie w nadmiernej pewności siebie (*overconfidence* – por. punkt 2.3.1), która z kolei skłania do ignorowania ważnych informacji publicznych (np. dotyczących ogólnej sytuacji gospodarczej) i zbytnej wiary we własne analizy. W rezultacie optymistyczna informacja prywatna, choć obiektywnie niepotwierdzona, może wywołać nadmierny entuzjazm, zwiększając zmienność cen. Warunkiem koniecznym jest jednak systematyczny charakter powyższego wzorca zachowań, co zapewnia, że wpływy nastrojów optymistycznych i pesymistycznych nie zniosą się wzajemnie.

Procesy przetwarzania informacji oparte na heurystykach oraz powstające przy tym systematyczne błędy są przedmiotem rozważań w rozdziale 2.

1.1.3 *Zmienność zbyt zmienna*

Gdy rozważyć stabilność wariancji kursów akcji, zagadnienie ich nadmiernej zmienności wydaje się jeszcze bardziej frapujące. Na efektywnym rynku odchylenie standardowe zmian cen determinowane napływem nowych danych powinno wykazywać względną stabilność z okresu na okres. W rzeczywistości zaś sama zmienność kursów charakteryzuje się nadmierną wariancją; innymi słowy, odchylenie standardowe cen aktywów giełdowych jest niestabilne w czasie. Stwierdzenie to można poprzeć analizą reakcji rynku na zmiany w poziomie wahań kursów. Na przykład krachowi giełdowemu w 1987 roku towarzyszył siedmiokrotny wzrost zmienności kursów, mimo iż w tle nie pojawiła się żadna istotna informacja. Haugen, Talmor i Torous (1991), zainspirowani tym wydarzeniem, zbadali dzienne wahania indeksu Dow Jones Industrial Average (DJIA) w latach 1897–1988. Ich analiza potwierdziła występowanie niestabilności wariancji cen giełdowych w czasie. Badacze dowiedli ponadto, iż sam wzrost/spadek poziomu zmienności rynku indukuje zmiany cen. Inwestorzy uwzględniają zatem informacje o zmienności jako istotne przy wycenie aktywów. Jej nasilenie jest dla nich sygnałem, że zasady „gry rynkowej” uległy zmianie – wzrosło prawdopodobieństwo znacznych fluktuacji cen w krótkim okresie. Należy przy-

puszczać, że inwestor niechętny ryzyku zażąda wyższej rekompensaty, jeśli rynek jest zmienny w sposób niestabilny. Realizacja wyższych stóp zwrotu jest z kolei możliwa jedynie poprzez spadek cen aktywów (zakładając stały poziom dywidendy). Haugen, Talmor i Torous obserwują przy tym pozytywną i statystycznie istotną korelację między siłą negatywnej korekty cen a wzrostem zmienności w danym okresie¹².

Tak więc, wbrew nowoczesnej teorii finansów, rynek wykazuje wrażliwość na informacje pochodzące ze swojego wnętrza i reaguje na nie nawet w obliczu braku nowych danych fundamentalnych. Badania dotyczące poziomu zmienności, jak i stabilności wahań giełdowych stanowią motywację do zrewidowania poglądów na efektywność rynku oraz na racjonalność inwestorów, których to decyzje są ostatecznie uzewnętrzniane w cenach akcji.

1.1.4 Nadmierne obroty giełdowe

Jedną z konsekwencji efektywności rynku są niskie obroty giełdowe. Jeśli rynku nie można „pobić” poprzez konsekwentne osiąganie ponadprzeciętnej rentowności, a zmiany kursów są nieprzewidywalne, racjonalny inwestor powinien uprawiać pasywne strategie inwestycyjne, czyli lokować środki w portfelu odzwierciedlającym skład indeksu. Czynnikiem motywującym do aktywnego inwestowania winny być jedynie nowe informacje bądź względy strategiczne związane z działalnością gospodarczą inwestora.

Wolumen transakcji obserwowany codziennie na rynkach światowych¹³ okazuje się stanowczo wyższy od poziomu wynikającego z konieczności równoważenia portfela, zmiany jego profilu ryzyka, hedgingu, względów podatkowych czy też z potrzeby płynności. Mimo iż dotąd nie powstał mo-

¹² Wiarygodność tych wyników została dodatkowo wzmocniona przez Blacka, który stwierdził negatywną zależność między stopami zwrotu z akcji i zmianami wahań rynkowych; por. (Black, 1976). Podobne wnioski pojawiają się również w (French, Schwert i Stambaugh, 1987).

¹³ Stopa obrotu (całkowita liczba akcji sprzedanych w danym roku podzielona przez liczbę wszystkich akcji notowanych na giełdzie) na giełdzie w Nowym Yorku (NYSE) wzrosła z 42% do 78% w latach 1982–1999. Jeszcze wyraźniejszy wzrost obrotów – z 88% w 1990 roku do 221% w roku 1999 – miał miejsce na rynku nowych technologii NASDAQ (Shiller, 2000, s. 39).

del służący wyznaczeniu „słusznego” poziomu obrotów giełdowych, wielkość tę można aproksymować metodami pośrednimi.

Badania związane z nadmiernymi obrotami giełdowymi przeprowadzone zostały przez Odeana (1999). Odean skoncentrował się na specyficznej grupie uczestników rynku – inwestorach indywidualnych, będących klientami brokerów dyskontowych. Taki wybór grupy badanych uzasadniony jest dwoma czynnikami: (i) dostępnością danych o dokonywanych transakcjach oraz (ii) „suwerennością” decyzji podejmowanych przez indywidualnych inwestorów (brak zakłóceń wynikających z problemu pośrednictwa). Przyjęta metodologia polegała na sprawdzeniu, czy inwestorzy byłoby w stanie osiągać wyższe stopy zwrotu, gdyby rzadziej przeprowadzali transakcje.

Na podstawie analizy 10 000 rachunków brokerskich Odean pokazał, że inwestorzy redukuja osiągane wyniki poprzez swą nadmierną aktywność. Pozwoliło to na postawienie hipotezy o ich zbytnej pewności siebie, która skłania do zajmowania pozycji nawet wówczas, gdy oczekiwane zyski nie wystarczają na pokrycie kosztów transakcyjnych. Obserwacja ta umożliwia połączenie zjawiska nadmiernych obrotów giełdowych z przesłanką, iż uczestnicy rynku nie są w stanie poprawnie wyceniać wszystkich dostępnych na rynku aktywów. Dlatego też muszą uciekać się do przybliżonych metod wyboru portfela – wspomnianych heurystyk, na podstawie których alokują swoje środki. Co więcej, inwestorzy przeceniający swoje umiejętności przewidywania rynku, ulegają efektowi dyspozycji (*disposition effect*), czyli skłonności do przedwczesnej realizacji zysków i zwlekania z realizacją strat (por. punkt 4.1.2), przez co obniżają rentowność swojego portfela. Spowodowane nadmierną aktywnością niższe stopy zwrotu wiążą się też ze stawianiem na trend i stosowaniem strategii pozytywnej reakcji, czyli kupowaniem akcji, które ostatnio zyskały, i sprzedawaniem tych, które straciły na wartości (por. punkt 1.2.1.1).

Mimo iż Odean skupił się na testowaniu hipotezy o nadmiernej pewności siebie jedynie w wybranej grupie inwestorów, podobnych wyników można oczekiwać na poziomie całego rynku. Badania przeprowadzane przez psy-

chologów wykazują bowiem, że omawiana cecha znamionuje większość ludzi. Rynki finansowe zaś szczególnie predysponują do takich postaw. Podejmowanie decyzji inwestycyjnych wymaga pewności siebie, wiary w słuszność własnej interpretacji informacji oraz polegania na odważnych założeniach. Specyfika rynków finansowych może zatem powodować negatywną selekcję, przyciągając szczególnie te jednostki, które przesadnie wierzą w swoją zdolność do prognozowania zdarzeń.

O zbytnej pewności siebie i jej konsekwencjach będzie szerzej mowa w punkcie 2.3.1.

1.1.5 Problem funduszy zamkniętych

Fundusze zamknięte stanowią jedną z najbardziej zastanawiających zagadek w finansach (*closed-end fund puzzle*). Instytucje te inwestują w aktywa, będące przedmiotem obrotu giełdowego i sprzedają certyfikaty inwestycyjne, których liczba nie zmienia się w okresie między kolejnymi emisjami. W przeciwieństwie do funduszy otwartych umorzenie certyfikatu w zamian za wypłatę równowartości aktywów netto nie jest tu dopuszczalne. Inwestor może zlikwidować swoją pozycję jedynie poprzez odsprzedaż certyfikatu inwestycyjnego innej osobie.

W świetle takich obwarowań wartość fundamentalna funduszu zamkniętego wydaje się niezwykle transparentna. Jest to po prostu suma kapitalizacji rynkowych wszystkich aktywów, jakie dany fundusz posiada. Dlatego też prawidłowa wycena nie powinna stanowić najmniejszego problemu, tym bardziej że całkowita wartość rynkowa aktywów utrzymywanych przez fundusz jest regularnie publikowana w prasie.

Dziwi zatem fakt, że ceny certyfikatów funduszu znacznie odbiegają od wartości jego aktywów netto (WAN):

- (1) W momencie tworzenia funduszu certyfikaty wyceniane są zazwyczaj z premią (nawet 10%) w stosunku do WAN.
- (2) W trakcie działania funduszu ceny certyfikatów wskazują na dyskonto sięgające nawet 10% względem WAN.

- (3) Rozbieżność między kursem a wartością cechuje się znacznymi wahaniami w czasie istnienia funduszu i znika w momencie jego likwidacji.

1

Propagatorzy hipotezy efektywności rynku, jakkolwiek zdziwieni zjawiskiem, starali się wytłumaczyć ów fakt na gruncie racjonalnych teorii. Wyjaśnienia te bazują zasadniczo na trzech motywach: (i) kosztach zarządzania funduszem, (ii) jego zobowiązaniach podatkowych oraz (iii) niskiej płynności posiadanych przezeń aktywów.

Argumenty te nie są jednak w stanie rozwiązać zagadki. Oto, jak można je zbić (Haugen, 1999, s. 47). (i) W odpowiedzi na motyw kosztowy, wystarczy wskazać, że wynagrodzenia managerów funduszu – w przeciwieństwie do poziomu dyskonta – nie zmieniają się z dnia na dzień i są w porównaniu z nim niewspółmiernie niskie. (ii) Zdaniem zwolenników argumentu fiskalnego inwestorzy dyskontują w cenie certyfikatów zobowiązania podatkowe funduszu, wynikające ze wzrostu wartości posiadanych przezeń aktywów. Teza ta traci moc, jeśli zauważyć, że w rzeczywistości dyskonto obniża się przy wzrostach cen akcji i powiększa się przy ich spadkach. Gdyby względy podatkowe stanowiły faktyczną przyczynę obserwowanych anomalii, należałoby oczekiwać sytuacji odwrotnej. (iii) W interpretacji uwzględniającej płynność przyczyn błędnej wyceny upatruje się w tym, iż fundusze zamknięte inwestują w aktywa trudno zbywalne. Trzeba jednak stwierdzić, że problem dyskonta dotyka również funduszy utrzymujących aktywa o wysokiej płynności. Żadna z zaproponowanych koncepcji nie jest zatem wystarczająco przekonująca, by problem funduszy zamkniętych można było uznać za rozwiązany.

Zagadnienie funduszy zamkniętych wzbudziło szerokie zainteresowanie naukowców behawioralnego nurtu finansów. Warto tu powołać się na badania przeprowadzone przez Lee, Shleifera i Thaler (1991). Zaproponowali oni rozwiązanie w oparciu o koncepcję nastroju indywidualnych inwestorów.

rów, podejmujących decyzje na podstawie szumu (tzw. noise-traderzy¹⁴, por. podrozdział 1.2). Inwestorzy ci kierują się w swoich poczynaniach nieracjonalnymi zmianami własnych oczekiwań co do przyszłej rentowności funduszu. Ich nastroje balansują między optymizmem a pesymizmem, powodując odpowiednio zmiany w cenie certyfikatów inwestycyjnych. Tak też lokata środków w funduszu pociąga za sobą dwa rodzaje ryzyka: (i) ryzyko posiadanego przezeń portfela oraz (ii) ryzyko szumu, wynikające ze zmiany nastrojów indywidualnych uczestników. Z tego względu inwestycja w sam fundusz jest bardziej ryzykowna niż bezpośredni zakup ekwiwalentu jego portfela. Jeśli ryzyko szumu jest systematyczne – niedywersyfikowalne, inwestor będzie wymagał za nie dodatkowego wynagrodzenia. Stanowi to powód, dla którego fundusze są przeciętnie wyceniane z dyskontem wobec WAN.

W zaproponowanym modelu uchwycone zostaje również zjawisko premii w momencie emisji nowych certyfikatów: założyciele funduszu decydują się na rozpoczęcie działalności w okresach, kiedy na rynku przeważają nastroje optymizmu. Umożliwia to uplasowanie certyfikatów depozytowych po cenie przewyższającej WAN, czyli z premią. I rzeczywiście: liczba funduszy zamkniętych tworzonych w USA charakteryzuje się nieliniowością w czasie, ze szczególną kumulacją w okresach pozytywnego nastroju rynkowego.

Rozwiązanie Lee, Shleifera i Thaler'a pozwala ponadto zrozumieć przyczynę zaniku rozbieżności między ceną a WAN, kiedy fundusz kończy swoją działalność. W momencie jego likwidacji ryzyko szumu przestaje mieć znaczenie, stąd cena certyfikatów wykazuje konwergencję do wartości fundamentalnej.

Behawioralne podejście do problemu funduszy zamkniętych jest istotnym wyzwaniem dla hipotezy efektywnego rynku. Dowodzi ono bowiem, że nastroj inwestorów może prowadzić do zachowań niezgodnych z racjonalnym

¹⁴ Właścicielami certyfikatów depozytowych funduszy zamkniętych w USA są głównie inwestorzy indywidualni, stąd też do nich odnosi się określenie *noise trader* w modelu zaproponowanym przez Lee, Shleifera i Thaler'a.

modelem, a będących źródłem systematycznego ryzyka. O systematycznym charakterze ryzyka szumu będzie mowa w punkcie 1.2.1.3.

1.2 Szum na rynku

Ekonomiści przez dekady posługiwali się kategoriami efektywności rynku oraz racjonalności jego uczestników. Odkryte anomalie, jak te opisane w poprzedniej części rozdziału, poddały w wątpliwość dotychczasowe rozumienie rynku oraz zapoczątkowały rozwój koncepcji uwzględniających działanie dwóch rodzajów inwestorów: w pełni racjonalnych arbitrażystów oraz tzw. noise-traderów¹⁵. To dualistyczne podejście pozwoliło zidentyfikować efekty działania nieracjonalnych jednostek na racjonalnym rynku.

1.2.1 Koncepcja inwestowania na podstawie szumu¹⁶

Szum (*noise*) jest tym, co czyni nasze obserwacje niedoskonałymi (Black, 1986). Można go traktować jako antonim informacji – nieinformację. Inwestorzy nieracjonalnie sądzą, że шум, na podstawie którego działają, jest prawdziwą informacją i może zapewnić przewagę nad innymi uczestnikami rynku. Zgodnie z powyższą definicją, noise-traderzy reagują na pseudosygnały pochodzące z rynku, np. wskazówki brokerów, analityków technicznych, prognostyków rynkowych i guru inwestycyjnych. To determinuje ich dalsze działania.

Mimo popełnianych błędów noise-traderzy stanowią nieodzowny element rynku, zapewniając mu płynność. Ceną płynności jest jego nieefektywność informacyjna i niewłaściwa wycena aktywów.

¹⁵ Dla zachowania precyzji zdecydowano się na spolszczenie angielskiego wyrażenia *noise traders* jako „noise-traderzy”. Sugerowano się przy tym określeniem „market-makerzy” popularnie używanym w polskojęzycznej literaturze finansów.

¹⁶ Należy wspomnieć, iż w literaturze brak zgodnego podejścia do noise-traderów. W zależności od autora desygnatem tego określenia są niewyszukani inwestorzy indywidualni (De Bondt, 1998) tudzież spekulanci, niekiedy także osoby uprawiające strategię pozytywnej reakcji (*positive feedback traders*, Thaler, 1992). Punktem zbieżności powyższych ujęć jest wskazanie, że noise-traderzy działają w sposób nieracjonalny.

1.2.1.1 Strategie inwestycyjne

Nadmierna pewność siebie, niezdolność oceny prawdopodobieństw zdarzeń i wariancji kursów to kluczowe cechy nieracjonalnych uczestników rynku. Motywują one do takich zachowań inwestycyjnych, jak pogoń za trendem (*trend chasing*) czy też strategia pozytywnej reakcji (*positive feedback trading*). W pogoni za trendem inwestorzy dokonują ekstrapolacji przeszłych zmian cen w przyszłość; kupują też akcje po wzroście ich kursów i sprzedają po spadku (strategia pozytywnej reakcji). Popularne techniki to choćby zlecenia *stop-loss*, powodujące zamknięcie pozycji, jeśli cena dotknie uprzednio określonego poziomu bądź też strategię takie, jak ubezpieczenie portfela (*portfolio insurance*), polegające na zakupie akcji po wzroście ceny (celem zwiększenia ryzyka portfela) oraz sprzedaży po spadku (dla ograniczenia ryzyka). Zabiegi te należą do codzienności inwestycyjnej.

1.2.1.2 Przetrwanie na rynku

Standardowa teoria finansów ignoruje wpływ, jaki noise-traderzy wywierają na rynek. Wynika to z dwóch założeń zaproponowanych już przez Milтона Friedmana (1953):

- (1) Działania nieracjonalnych inwestorów uznaje się za stochastyczne, w związku z czym za pozbawione znaczenia dla zagregowanego popytu na aktywa.
- (2) Noise-traderzy – poprzez popełniane błędy – stopniowo tracą swój majątek na rzecz racjonalnych arbitrażystów i ostatecznie zostaną usunięci z rynku.

Pierwszy argument traci moc, jeśli dowieść, że nieracjonalni inwestorzy popełniają systematyczne błędy (por. podrozdziały 2.2 i 3.3). Wówczas zmiany popytu są między nimi powiązane i nie znoszą się wzajemnie. Liczne badania wskazują, że większość inwestorów posługuje się popularnymi modelami i wierzy w te same pseudosygnały (Shiller, 1990b). Stąd też założenie o korelacji ich działań jest uzasadnione.

W odpowiedzi na drugi argument warto przytoczyć słynne słowa Keyne-
sa¹⁷: „Rynek może pozostać nieracjonalny dłużej, niż ty możesz pozostać
wypłacalny”. Jako że utrata majątku noise-traderów na rzecz arbitrażystów
zachodzi niezwykle powoli (Figlewski, 1979), ich „upadku” nie należy brać
za pewnik. Wręcz przeciwnie, powodowani nadmierną pewnością siebie
i przesadnym optymizmem, noise-traderzy mogą działać bardziej agresyw-
nie i podejmować większe ryzyko niż arbitrażyści. Dopóki rynek to wynag-
radza, dopóty nieracjonalni inwestorzy będą zarabiać wyższe stopy zwrotu
mimo stosowania nieoptymalnych strategii.

Paradoksalnie, noise-traderzy mogą być rekompensowani za ryzyko, które
sami wnoszą. Stymulowani pewnością siebie podejmują działania nadmier-
nie ryzykowne, lecz równocześnie przynoszące ponadprzeciętną rentowność.
Oczywiście ryzykanctwo idzie w parze z podwyższoną wariancją osiąga-
nych wyników. Inaczej niż zakłada podejście Friedmanowskie, noise-
traderzy mogą przetrwać na rynku, a nawet zdobyć znaczny majątek. Muszą
się przy tym jednak liczyć z wyższym prawdopodobieństwem ponownej
utruty swego dorobku.

1.2.1.3 Systematyczny charakter ryzyka szumu

Genezą ryzyka szumu (*noise trader risk*) jest nieprzewidywalny nastrój
czy też sentyment inwestorów, oscylujący między nadmiernym optymizmem
a przesadnym pesymizmem. Korelacja nastrojów między poszczególnymi
osobami sprawia, że ryzyko szumu jest systematyczne, czyli nie daje się dy-
wersyfikować. Z tego względu stanowi ono element wyceniany na rynku.

Badacze dowodzą systematycznego charakteru ryzyka szumu poprzez ob-
serwację równoległych i jednokierunkowych zmian cen różnych aktywów.
W standardowym podejściu do finansów zakłada się, że łączne wahania kur-
sów spółek (*comovement*) mogą pojawić się jedynie wówczas, gdy istnieje
wspólny dla nich determinant wielkości fundamentalnych (przepływów pie-
niężnych). Teza ta została podważona za sprawą badań empirycznych. Wy-

¹⁷ Wszystkie tłumaczenia w tej pracy pochodzą od autorki.

kazują one współwystępowanie wahań także w przypadku akcji spółek niepowiązanych żadnymi czynnikami gospodarczymi (Lee, Shleifer i Thaler, 1991). Jednym z przykładów jest wysoka dodatnia korelacja między zmianami cen certyfikatów inwestycyjnych emitowanych przez fundusze zamknięte a kursami akcji małych firm w USA¹⁸. Jedyne, co łączy te aktywa, to typ posiadających je uczestników rynku – głównie inwestorzy indywidualni. Jeśli noise-traderzy są nieracjonalnie pesymistyczni, wówczas sprzedają swoje aktywa bez wyjątku, obniżając ich cenę niezależnie od napływających informacji. Podobnie kiedy kupują na fali ogólnego optymizmu giełdowego, ignorują dane o poszczególnych spółkach i zawyżają kursy. Zjawisko łącznych wahań pozwala zatem wnioskować o systematycznym charakterze ryzyka szumu.

1.2.1.4 Wpływ *noise-traderów* na zachowania innych inwestorów

Oczekiwania nieracjonalnych inwestorów są odzwierciedleniem ich nastroju i to one powodują niespodziewane zmiany popytu na aktywa. Systematyczny charakter tych zmian wpływa na zachowania innych uczestników rynku. Dla arbitrażysty przeciwdziałanie nieracjonalności *noise-traderów* może okazać się nieoptymalne, a nawet szkodliwe. Co więcej, lepsze ekonomicznie rezultaty ma szansa przynieść strategia arbitrażowa polegająca na uprzedzaniu *noise-traderów*, czyli zajmowaniu pozycji długiej, zanim tamci zaczną kupować i krótkiej, zanim zaczną sprzedawać. Takie interakcje między arbitrażystami i nieracjonalnymi uczestnikami rynku są katalizatorem „bąbli” cenowych oraz przyczyną autokorelacji stóp zwrotu. Chcąc wykorzystać błędy innych, arbitrażyści najpierw stawiają na trend, powielając popularne na rynku zachowania, a następnie dokonują kontry, spychając ceny z powrotem do wartości fundamentalnych. Prowadzi to do pozytywnych

¹⁸ Lee, Shleifer i Thaler (1991) dowodzą ponadto, że łączne wahania są charakterystyczne także dla zmian cen certyfikatów różnych funduszy zamkniętych. Korelacja poziomów dyskonta między funduszami wynosi około 0,53. Oznacza to, że fluktuacje cen certyfikatów są w dużej mierze tym powodowane samym nastrojem inwestorów.

autokorelacji stóp zwrotu w krótkim terminie, zaś negatywnych – w terminie średnim i długim.

Zjawisko autokorelacji kursów zostało szeroko udokumentowane w literaturze przedmiotu, dając początek hipotezom nadmiernej i niedostatecznej reakcji cen na rynku. Publikacje te wywołały gwałtowny sprzeciw ze strony zwolenników efektywności rynku, niezłomnie wierzących w błędzenie losowe kursów giełdowych (*random walk*) oraz w brak możliwości prognozowania rynku.

1.2.2 Nadmierna i niedostateczna reakcja cen na informacje

Nadmierna i niedostateczna reakcja (*over- and underreaction*) to koncepcje makroekonomiczne powstałe na gruncie analizy zachowań zagregowanego rynku kapitałowego. W finansach nadmierna i niedostateczna reakcja odnoszą się do ruchów cen aktywów, nie zaś do poczynań inwestorów. Oczywiście zmiany cen akcji są odzwierciedleniem działań ludzkich, które jednak nie dają się obserwować w sposób bezpośredni.

Określenia „nadmierna” i „niedostateczna” reakcja sygnalizują odchylenie kursów od pewnego poziomu uznawanego za „normalny”. Warto przypomnieć, iż na rynku efektywnym normalny jest ów stan, w którym wszelkie dostępne informacje są uwzględnione w cenie rynkowej.

1.2.2.1 Niedostateczna reakcja

Niedostateczna reakcja oznacza, że ceny akcji z niewystarczającą siłą reagują na nowe wiadomości, np. na ogłoszone przez spółkę zyski. Inaczej niż prognozuje hipoteza efektywnego rynku, informacje takie nie są natychmiast odzwierciedlane w cenie papierów wartościowych. Ze względu na opóźniony proces aktualizacji oczekiwań wobec spółki, wpływ świeżej informacji na kurs rozłożony jest na kilka okresów. Tak też po ujawnieniu danych pozytywnych kursy wchodzą w fazę powolnego trendu wzrostowego. I odwrotnie: po ukazaniu się wiadomości negatywnej kształtuje się trend spadkowy. Dowodzi to, iż początkowa reakcja ceny była zbyt słaba w stosunku do znaczenia informacji.

W związku ze spowolnionym przyswajaniem nowych wiadomości przez rynek akcje wykazują pozytywne autokorelacje stóp zwrotu w krótkim okresie (od 1 do 12 miesięcy). Efekt ten określany jest mianem *momentum*. Na jego mocy przedawnione informacje pozwalają na prognozowanie przyszłej rentowności.

Z psychologicznego punktu widzenia niedostateczna reakcja jest konsekwencją konserwatyizmu inwestorów, który należy rozumieć jako zbyt powolną aktualizację poglądów w stosunku do napływu nowych informacji (Edwards, 1968). Konserwatywny inwestor, opierając się na heurystyce zakotwiczenia i dostosowania (por. punkt 2.2.3), niedostatecznie dostosowuje swoje oczekiwania wobec przyszłych zysków spółki, przez co zaniża jej cenę.

Hipoteza niedostatecznej reakcji potwierdzona została w licznych przekrojowych analizach stóp zwrotu na rynku kapitałowym w USA. Jak pokazują Bernard i Thomas (1989), spółki, które w poprzednim okresie zaskoczyły inwestorów nadspodziewanie dobrymi wynikami finansowymi¹⁹, odnotowują następnie wyższe stopy zwrotu niż średnia na rynku. W sytuacji takiej przyjęło się mówić o dryfie kursów po ogłoszeniu wyników finansowych (*post-earnings announcement drift*). Zjawisko dryfu wskazuje, iż pierwotnie reakcja ceny musiała być zbyt słaba.

Podobnie inne badania w zakresie autokorelacji stóp zwrotu (*momentum*) potwierdzają niedostateczną reakcję cen akcji. Wyniki uzyskane przez Chana, Jegadeesha i Lakonishoka (1996) dla dużej próby amerykańskich firm (lata 1977–1993) pokazują, że portfel o najgorszej rentowności w ostatnich 6 miesiącach poprzedzających badanie w następnym półroczu przegrywał z portfelem najlepszym o około 9 pkt. proc.

¹⁹ Prosta (naiwna) metodą określania „zaskoczenia” inwestorów rezultatami spółki może być tzw. standaryzowany nieoczekiwany wynik finansowy (SUE – *standardised unexpected earnings*). Wielkość tę definiuje się jako różnicę w wynikach spółki między danym kwartałem a tym samym kwartałem zeszłego roku podzieloną przez odchylenie standardowe jej wyników.

Prace empiryczne nad zjawiskiem niedostatecznej reakcji wywołały spore zainteresowanie krótkoterminowymi strategiami inwestycyjnymi (głównie od 6 do 12 miesięcy), opierającymi się na *momentum*. Polegają one na zakupie akcji „zwycięskich” w przeszłości i sprzedaży przeszłych „przegranych”.

1.2.2.2 Nadmierna reakcja

Obserwuje się, że w okresie od 3 do 5 lat kursy nadmiernie reagują na serie jednokierunkowych informacji (tzn. stale pozytywnych lub negatywnych). Zjawisko to przyjęło się określać mianem nadmiernej reakcji. W konsekwencji akcje spółek dostarczających kolejno pozytywnych informacji stają się przeszacowane, zaś akcje spółek ujawniających dane negatywne stają się niedoszacowane. Takie błędy wyceny są funkcją nastroju inwestorów, którzy popadają w nadmierny optymizm wobec spółek o dobrych wynikach i przesadny pesymizm wobec tych, których historia wyników finansowych była rozczarowująca. Ta regularność pozwala przypuszczać, że inwestowanie w oparciu o analizę przeszłych serii informacji daje ponadprzeciętne stopy zwrotu.

Nadmierna reakcja bywa często przypisywana heurystyce reprezentatywności (por. punkt 2.2.2). Rewidując swoje przekonania o danej spółce, inwestorzy przeceniają informacje najświeższe i zaniedbują poprzednie (tzw. bazowe); przywiązują też zbyt duże znaczenie do wiadomości przykuwających uwagę, dramatycznych, przy jednoczesnym ignorowaniu tendencji długoterminowych.

Pierwszymi badaczami, którzy postawili sobie za cel sprawdzenie, czy rynek kapitałowy nadmiernie reaguje, byli De Bondt i Thaler (1985). Na podstawie analizy stóp zwrotu na giełdzie nowojorskiej w latach 1933–1980 postawili oni następującą hipotezę: jeśli kursy w regularny sposób nadmiernie reagują na informacje zarówno pozytywne, jak i negatywne, to przyszłe stopy zwrotu z akcji powinny być przewidywalne jedynie w oparciu o ich minioną rentowność. Dane fundamentalne można wówczas całkowicie zignorować. Dla wsparcia tego twierdzenia De Bondt i Thaler pokazali, że portfele przeszłych skrajnych „przegranych” są w stanie „pobić” wyniki

portfeli byłych największych „zwycięzców” nawet o 25 pkt. proc., już po korekcie ze względu na ryzyko²⁰. Obserwację tę nazwano efektem zwycięzcy i przegranego (*winner-loser effect*).

W obliczu wielu głosów krytyki ze strony zwolenników efektywności rynku hipoteza nadmiernej reakcji została poparta dalszymi badaniami²¹. Kolejne prace skupiły się na innych miarach wartości spółek, np. stosunek ceny do zysku (*P/E – price earnings ratio*) czy też wartości księgowej do wartości rynkowej (*book-to-market ratio*). Potwierdziły one wcześniejsze wyniki: akcje o potencjale wzrostu (*growth stocks*) – wyceniane przez rynek wysoko w stosunku do wartości fundamentalnej – w perspektywie kolejnych 3 do 5 lat charakteryzują się niższymi stopami zwrotu niż akcje o potencjale wartości (*value stocks*), które rynek wycenia nisko w odniesieniu do rzeczywistej ich wartości.

Publikacje dotyczące zjawiska nadmiernej reakcji wywołały żywe zainteresowanie strategią inwestycyjną zwaną kontrarianizmem (*contrarian strategy*). Kontrarianizm polega na kupowaniu akcji o niskich stopach zwrotu w przeszłości (byłych przegranych) i sprzedaży akcji o wysokich stopach zwrotu w przeszłości (minionych zwycięzców).

1.2.2.3 Słowa krytyki

Dowody nadmiernej i niedostatecznej reakcji rynku uznawane są przez wielu badaczy za rezultat eksploracji danych (*data mining*). Na mocy tego zarzutu behawioralny nurt finansów popadł w ogień ostrej krytyki ze strony zwolenników podejścia tradycyjnego. Oto słowa Famy, jednego z głównych twórców hipotezy efektywnego rynku: „Efektywność rynku przetrwała wyzwanie rzucone jej przez odkrywców (...) anomalii stóp zwrotu. Zgodnie z hipotezą efektywnego rynku, traktującą anomalie jako zrzędzenie losu,

²⁰ De Bondt i Thaler klasyfikowali portfele „przegranych” i „zwycięzców” na podstawie stóp zwrotu z aktywów w okresie 3 lat poprzedzających moment formacji portfela. Najwięksi „przegranii” to akcje, których rentowność znajdowała się w pierwszym (najniższym) decylnym, zaś ekstremalnie zwycięzcy to akcje o rentowności na poziomie ostatniego (najwyższego) decyla.

²¹ Por. np. (De Bondt i Thaler, 1987), (De Bondt i Thaler, 1990), (Chopra, Lakonishok i Ritter, 1992).

1

rzekoma nadmierna reakcja na informacje jest tak samo częsta jak reakcja niedostateczna, zaś kontynuacje ponadprzeciętnych stóp zwrotu, występujące po danym zdarzeniu, równie częste jak regresje do średniej (...)” (Fama, 1998).

Również badacze behawioralni, mimo statystycznej istotności zjawisk nadmiernej i niedostatecznej reakcji, rezygnują z przypisania im centralnego miejsca w nurcie behawioralnym. Anomalie te – wykryte w latach osiemdziesiątych i na początku lat dziewięćdziesiątych – powinny być raczej traktowane jako motywacja do dalszych poszukiwań. Ostatecznym zaś zadaniem behawioralnego podejścia do ekonomii i finansów jest zrozumienie „czynnika psychologicznego”, tak często cytowanego w obliczu niewytłumaczalnych zachowań rynków finansowych.

1.3 Ograniczenia skuteczności arbitrażu

Zgodnie z postulatem Friedmana mechanizm arbitrażu zapewnia utrzymanie efektywności rynku (por. punkt 1.2.1.2). W tradycyjnej teorii finansów arbitraż definiowany jest jako strategia inwestycyjna pozbawiona ryzyka polegająca na jednoczesnym zakupie i sprzedaży tego samego aktywu (lub jego doskonałego substytutu) na dwóch rynkach po korzystnej cenie (Shleifer, 2000, s. 28). W praktyce natomiast arbitraż przeprowadzany jest najczęściej między aktywami niedoskonale substytucyjnymi i w bardzo szerokim kontekście powinien być rozumiany jako działalność inwestycyjna, zmierzająca do maksymalizacji zysków poprzez wykorzystanie błędów nieracjonalnych uczestników rynku.

Wyobraźmy sobie, że grupa noise-traderów niesłusznie (bez oparcia w danych fundamentalnych) nabiera pesymistycznych przekonań o perspektywach rozwojowych firmy Ford. Wyrazem pesymizmu jest oczywiście zwiększona podaż akcji Forda, która obniża jego cenę do, załóżmy, 15 j.p. Arbitrażysta, wiedząc, że wartość wynosi 20 j.p., kupuje akcje Forda po atrakcyjnej cenie, jednocześnie zabezpieczając swoją pozycję poprzez krótką sprzedaż doskonałego substytutu, np. akcji General Motors (GM). Wzmó-

ny popyt na aktywa Forda przywraca ich właściwą wycenę (na poziomie wartości fundamentalnej), zaś arbitrażysta realizuje zysk bez ryzyka kosztem błędów niepoinformowanych inwestorów. W konsekwencji swych kolejnych niewłaściwych decyzji nieracjonalny uczestnik traci swój majątek i zostaje usunięty z rynku.

Tradycyjna linia obrony efektywności rynku, oparta o działalność racjonalnych jednostek, bazuje zatem na założeniu, że odchylenie ceny giełdowej od wartości fundamentalnej otwiera atrakcyjną możliwość arbitrażu. Poinformowani inwestorzy natychmiast wykorzystają tę szansę dla osiągnięcia zysku bez ryzyka. Dzięki nim cena i wartość zrównają się.

Przedstawiciele behawioralnego nurtu finansów nie podważają pierwszego argumentu. Faktycznie, lukratywne możliwości inwestowania są powszechnym zjawiskiem na rynkach finansowych. Jednak bardzo często wykorzystanie rozbieżności między ceną a wartością wiąże się z podjęciem wysokiego ryzyka. Sprawia to, że sytuacje oczywistego prze- lub niedoszacowania cen aktywów utrzymują się przez długi czas przy braku reakcji ze strony arbitrażystów. Fakt ten podważa tradycyjny pogląd, jakoby brak możliwości systematycznego osiągania ponadprzeciętnych zysków był dowodem efektywności rynku (w znaczeniu: cena równa się wartość). Brak szans na „pobicie” rynku niekoniecznie musi wynikać z tego, że kursy odzwierciedlają wartości fundamentalne. Równie dobrze może być on spowodowany ograniczeniami skuteczności arbitrażu. Rodzaje i źródła ryzyka w arbitrażu są tematem kolejnych punktów²².

1.3.1 Ryzyko fundamentalne

Ryzyko fundamentalne wiąże się z nieistnieniem doskonałych substytutów, czyli takich (co najmniej dwóch różnych) papierów wartościowych, które we wszystkich stanach natury dają inwestorowi identyczny dochód (dywidendę) z tytułu ich posiadania. Wróćmy do wcześniejszego przykładu.

²² Podobną systematykę źródeł ryzyka w arbitrażu można znaleźć w (Barberis i Thaler, 2001). Dalsza literatura w tej dziedzinie to (Shleifer, 2000).

1

Arbitrażysta kupuje niedoszacowane akcje firmy Ford po cenie 15 j.p. i jednocześnie zajmuje krótką pozycję w substytucyjnych aktywach GM. W ten sposób zabezpiecza się przed napływem negatywnych informacji dotyczących całego przemysłu samochodowego. Jeśli takie miałyby się bowiem pojawić, straty na długiej pozycji zostaną wyrównane przez zyski z zajętej pozycji krótkiej. Tak działa strategia w teorii. W praktyce jednak aktywa bardzo rzadko są doskonale substytucyjne. Zwykle akcje firm z tej samej branży charakteryzują się różną wrażliwością na nowe informacje. Ponadto zabezpieczenie się substytutem nie chroni arbitrażysty przed napływem negatywnych danych specyficznych jedynie dla akcji, które kupił (w tym wypadku akcje Forda). Stąd też wyeliminowanie całego ryzyka fundamentalnego jest niemożliwe ze względu na brak idealnych substytutów.

1.3.2 Ryzyko szumu

O ryzyku szumu pisano już w punkcie 1.2.1.3. Teraz warto się zastanowić, dlaczego działalność noise-traderów ogranicza aktywność arbitrażystów, czyniąc rynek permanentnie niedoskonałym.

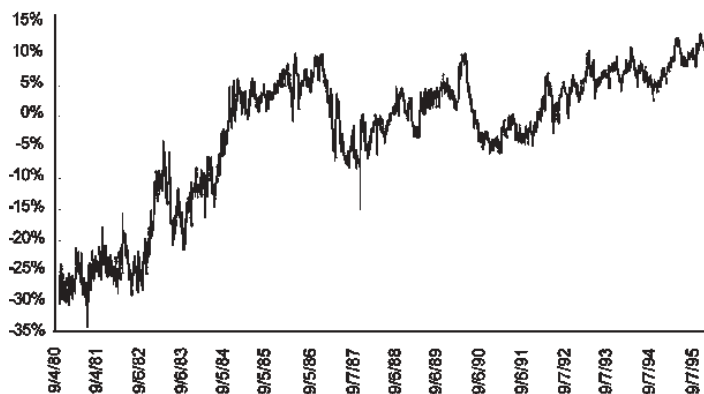
Założmy na chwilę, że GM jest doskonałym substytutem Forda. Poprzez kupno akcji Forda i sprzedaż GM ryzyko fundamentalne zostaje zatem wyeliminowane. Lecz nawet teraz arbitrażysta musi liczyć się z nieracjonalnymi zmianami nastroju noise-traderów, czyli z ryzykiem szumu. Objawia się ono nasileniem pesymistycznych nastrojów i dalszym spadkiem kursu Forda, przynosząc straty na pozycji długiej.

Oczywiście dla arbitrażysty z nieograniczonym horyzontem inwestycyjnym ryzyko szumu nie istnieje. W nieskończonej (lub bardzo długiej) perspektywie cena akcji wykazuje bowiem konwergencję do wartości fundamentalnej. Jednak w praktyce perspektywa inwestycyjna arbitrażysty jest zwykle ograniczona do krótkiego bądź średniego okresu. Wynika to z faktu, że arbitrażyści zarządzają cudzymi środkami, co nakłada na nich limity w ustalaniu horyzontu inwestycyjnego. Shleifer i Vishny (1997) nazywają ów problem pośrednictwa (*agency problem*) „oddzieleniem umysłu od kapitału”. Osoby oddające swe środki w zarządzanie są zwykle nieświadome

sposobu funkcjonowania rynku, stąd ich ocena skuteczności zarządzania portfelem opiera się często na błędnych przesłankach. Jeśli w krótkim okresie kursy rozwijają się niepomyślnie dla arbitrażysty (pośrednika), inwestorzy – laicy (klienci) będą postrzegali to jako stratę wynikającą z niekompetencji zarządzającego. Może to prowadzić do wycofywania przez nich środków. Arbitrażysta powodowany obawą przed takim rozwojem zdarzeń działa tak, jakby jego horyzont był faktycznie krótki. W rezultacie zamyka pozycję, zanim nastąpi odpowiednia korekta kursu w stosunku do wartości fundamentalnej. Dodatkowe ograniczenia mogą pojawić się ze strony instytucji, od których arbitrażyści pożyczają środki oraz aktywa potrzebne do implementacji swojej strategii. Widząc, że arbitrażysta ponosi straty (choćby niezrealizowane), kredytodawcy mogą domagać się uzupełnienia wartości zabezpieczenia lub wycofać swoje środki, wywołując tym samym przedwczesną likwidację pozycji arbitrażowej.

Ryzyko szumu jest zatem istotnym czynnikiem powstrzymującym racjonalnych uczestników rynku przed wykorzystaniem niepoprawnych cen akcji celem zysku. Najczęściej cytowanym przykładem tego zjawiska jest przypadek „aktywów bliźniaczych” Royal Dutch i Shell Transport Company (Froot i Dabora, 1999).

Od 1907 roku Royal Dutch i Shell działają jako jednostki połączone ekonomicznie w oparciu o parytet 60:40 (obliczony na podstawie przepływów pieniężnych). Mimo to ich akcje notowane są niezależnie od siebie: Shell w Wielkiej Brytanii, Royal Dutch w Holandii i USA. Ich aktywa stanowią 60% i 40% praw do przepływów pieniężnych odpowiednio dla Royal Dutch i Shell. Stąd też ich kapitalizacje rynkowe powinny odpowiadać relacji 60:40. Rysunek 1.2 pokazuje faktyczne odchylenia kursów od parytetu w latach 1980–1995.

Rysunek 1.2: Odchylenia kursu Royal Dutch względem Shell od parytetu 60:40

Źródło: (Froot i Dabora, 1999).

W omawianym okresie odchylenia od parytetu przybrały znaczące rozmiary: od 35% niedoszacowania do 10% przewartościowania aktywów Royal Dutch. Co charakterystyczne zatem, kursy odnotowały odchylenie zarówno *in plus*, jak i *in minus*. Jest to przykład trwale utrzymującego się błędu wyceny, który nie został skorygowany przez mechanizm arbitrażu. Badacze przypadku, Froot i Dabora (1999), tłumaczą to ryzykiem szumu. Wiadomo bowiem, że akcje obu firm dają prawa do tych samych przepływów pieniężnych (są doskonałymi substytutami), co wyklucza istnienie ryzyka fundamentalnego. Krótka sprzedaż nie jest tu ani ograniczona, ani obciążona ponadprzeciętnymi kosztami. Ponadto arbitrażysta może w obiektywny sposób – niezależnie od jakiegokolwiek modelu – stwierdzić autentyczność błędu wyceny. Jedyne źródła niepewności należy upatrywać w działaniach nieracjonalnych inwestorów, które pogłębiają rozbieżność ceny i wartości.

Wyobraźmy sobie, że w połowie 1983 arbitrażysta kupuje względnie niedoszacowane (-10%) aktywa Royal Dutch, sprzedając jednocześnie akcje Shell. Zajęta pozycja przynosi mu ciągłe straty przez sześć kolejnych miesięcy, kiedy to dyskonto kursu Royal Dutch osiąga poziom aż 25% . Równie ciekawa sytuacja ma miejsce we wrześniu 1980 roku, kiedy niedowartościowanie Royal Dutch sięga ponad 30% . Od tego momentu realizacja arbi-

trażowego zysku wymagałaby czterech lat oczekiwania na korektę ceny, na to zaś jedynie nieliczni mogą sobie pozwolić.

Przypadek akcji bliźniaczych Royal Dutch i Shell, ze względu na ewidentną rozbieżność wartości i ceny, kusił już wielu perspektywą zysku. Jednym z przykładów jest osławiony fundusz hedgingowy Long-Term Capital Management (LTCM), który próbował wykorzystać relatywne przeszacowanie Royal Dutch pod koniec lat dziewięćdziesiątych, zajmując długą pozycję w akcjach Shell, a Royal Dutch sprzedając krótko. Niestety aż do chwili upadku LTCM dyskonto kursu Shell, zamiast maleć, wzrastało (Shefrin, 2000, s. 7). Powrót do parytetu nastąpił dopiero w 2000 roku.

1.3.3 Koszty implementacji strategii arbitrażowej

Dotychczasowe rozważania skupiały się na dość wyszukanych źródłach ryzyka w arbitrażu. Jednak już sama konstrukcja strategii arbitrażowej, przez wzgląd na krótką sprzedaż, może powodować komplikacje. Celem krótkiej sprzedaży jest, jak wcześniej wspomniano, wyeliminowanie ryzyka fundamentalnego. Aby jej dokonać, inwestor musi pożyczyć odpowiednie aktywa od pośrednika finansowego. Trudności techniczne związane z tą procedurą obejmują kilka aspektów. Zaczynając od najbardziej banalnych: po pierwsze, arbitrażysta ponosi zwykłe koszty transakcyjne (np. *spread* kupna – sprzedaży, prowizje). Po drugie, na wielu rynkach krótka sprzedaż podlega ograniczeniom lub wręcz zakazom prawnym. Po trzecie, nawet jeśli ograniczenia prawne nie istnieją, sam fakt pożyczania aktywów od pośrednika finansowego ogranicza „suwerenność” arbitrażysty. Pośrednicy bowiem najczęściej zarządzają cudzymi portfelami i pożyczają arbitrażystom te aktywa, w których sami zajmują długą pozycję. Jeśli ich klienci zadecydują o wycofaniu środków, wówczas pośrednik żąda zwrotu pożyczki od arbitrażysty. Ten z kolei zmuszony jest do natychmiastowego odkupienia akcji, co szczególnie na mało płynnym rynku wiąże się z wysoką premią. Owa premia przypada inwestorom uprawiającym strategię *short-squeeze* – zajmującym dużą pozycję w aktywach sprzedanych krótko przez arbitrażystę. Ten zaś, postawiony przed koniecznością natychmiastowego zwrotu akcji pożycz-

1

nych od pośrednika, zmuszony jest do zapłacenia żądanej, zwykle bardzo wysokiej, premii – stąd nazwa *short-squeeze* (Shleifer, 2000, s. 47). Twierdzi się, że obawa przed takim scenariuszem zdarzeń stanowi czynnik ograniczający aktywność arbitrażową na wielu wschodzących rynkach (*emerging markets*).

Doskonałym przykładem sytuacji, w której koszty uniemożliwiły przeprowadzenie arbitrażu, jest przypadek firmy 3Com i jej spółki-córki – Palm (Lamont i Thaler, 2001). W marcu 2000 roku firma 3Com uplasowała na rynku 5% swoich udziałów w Palm w ramach pierwszej oferty publicznej (IPO). Jednocześnie ogłosiła zamiar sprzedaży pozostałych 95% w ciągu kolejnych 9 miesięcy. Dotychczasowi akcjonariusze 3Com w momencie IPO stali się pośrednio również właścicielami 1,5 udziału w Palm za każdy udział w 3Com. Formalnie akcje te miały zostać przydzielone w momencie pełnej sprzedaży spółki-córki na rynku. W ten sposób stosunek kursu 3Com do Palm powinien ukształtować się na poziomie 3:2, jako że inwestor planujący kupić akcje Palm mógł także pośrednio stać się ich właścicielem poprzez zakup udziałów w spółce-matce.

W pierwszym dniu po IPO cena Palm kształtowała się na zaskakująco wysokim poziomie 95 dolarów za akcję, podczas gdy kurs 3Com obniżył się do 81 dolarów. Zgodnie zaś z parytetem 3:2 akcje 3Com powinny być wyceniane na co najmniej 142 dolary. Tak więc rzeczywista cena wskazywała na niespotykane wręcz niedoszacowanie innych, poza Palm, aktywów w biznesowym portfelu 3Com. Paradoksalnie, ich implikowana wartość rynkowa ukształtowała się na ujemnym poziomie (– 61 dolarów na akcję). Mimo iż rozbieżność między ceną a wartością fundamentalną była dla wszystkich oczywista, arbitrażyści nie podchwycili jej jako doskonałej okazji do wzbogacenia się. Utrzymujące się przewartościowanie Palma urosło do rangi gorącej debaty na łamach „The Wall Street Journal” i „The New York Times”. Każdy bardziej świadomy inwestor musiał zadać sobie pytanie o słuszność hipotezy efektywnego rynku.

Badacze przypadku 3Com i Palm, Lamont i Thaler (2001), poszukują wytłumaczenia w niezwykle wysokich kosztach krótkiej sprzedaży. Inwestor, który chciał wykorzystać przeszacowany kurs Palm poprzez krótką sprzedaż jego akcji, napotykał dwie przeszkody: ograniczoną podaż akcji lub, jeśli akcje były dostępne, niezwykle wysokie koszty pożyczenia ich od pośrednika. W czerwcu 2000 roku żądana stopa procentowa wyniosła 35%, czyniąc strategię arbitrażową nieopłacalną. W tym więc przypadku długotrwała rozbieżność między ceną a wartością fundamentalną spowodowana była kosztami implementacji strategii arbitrażowej.

1.3.4 Koszty psychologiczne krótkiej sprzedaży

Należy zauważyć, że koszty krótkiej sprzedaży nie sprowadzają się jedynie do wysokości stopy procentowej. Shiller (2002) wskazuje dodatkowo na koszty psychologiczne, które mogą stanowić przeszkodę dla wielu potencjalnych arbitrażystów. W praktyce strategie arbitrażowe postrzegane są nie tylko jako ryzykowne, lecz także jako wymagające pewnej „wiedzy tajemnej”. Jak wiadomo, inwestor dokonujący krótkiej sprzedaży musi liczyć się z wymuszoną przedwcześnie likwidacją. Proceduralnie sytuacja taka nie powinna stanowić problemu, gdyż dane aktywa można zawsze pożyczyć od innego pośrednika. Konieczność nieoczekiwanej likwidacji jest jednak elementem wprowadzającym dodatkową niepewność i dlatego też istotnym z psychologicznego punktu widzenia.

Dalszym, ważniejszym jeszcze, aspektem krótkiej sprzedaży są potencjalnie nieograniczone straty z niej wynikające. Maksymalna strata inwestora zajmującego długą pozycję ogranicza się do równowartości początkowej lokaty kapitału. Natomiast strata na krótkiej sprzedaży aktywów może dalece przekroczyć inwestycję początkową (jest nielimitowana). Oczywiście w sytuacji pogłębiających się ujemnych wyników inwestor ma możliwość ich ucięcia poprzez natychmiastowe zamknięcie pozycji (tzw. pokrycie „shorta”). Jednakże, jak przewiduje *prospect theory* Kahnemana i Tversky’ego, po stronie strat jednostki wykazują skłonność do ryzyka, przez co realizują spotęgowane straty z opóźnieniem (por. punkty 3.3.4, 4.1.2 i 4.1.3).

1.3.5 Ryzyko modelu

Ostatnim wartym omówienia rodzajem ryzyka w arbitrażu jest tzw. ryzyko modelu. Aby stwierdzić, że cena danego aktywu jest błędna, arbitrażysta posługuje się modelem, na podstawie którego szacuje wartość fundamentalną. Obiektywnie istnieje tylko jedna prawdziwa wartość. W rzeczywistości jednak ustalenie słusznego jej poziomu nie jest zadaniem banalnym. Wyniki uzyskiwane na podstawie różnych metod są zwykle rozbieżne i silnie uzależnione od przyjętych założeń. Arbitrażysta niezwykle rzadko jest w stanie z pewnością stwierdzić, że dany błąd wyceny faktycznie istnieje (taka wyjątkowa sytuacja miała miejsce w wypadku opisanych wcześniej aktywów bliźniaczych). Zwykle obecność anomalii kursowych można wykazać jedynie na gruncie zastosowanego modelu wyceny.

1.3.6 Warunki ograniczenia arbitrażu

Dyskutowane rodzaje ryzyka wymagają dodatkowo uzupełnienia o warunki, przy których arbitraż jest faktycznie ograniczony.

Załóżmy, że arbitrażysta stwierdza rozbieżność między kursem a wartością fundamentalną danego papieru wartościowego. W pierwszym kroku należy się zastanowić nad sytuacją, gdy akcje spółki nie posiadają doskonałego substytutu. Oczywiście w takim wypadku arbitrażysta jest narażony na ryzyko fundamentalne. Barberis i Thaler (2001) podają dwa warunki dostateczne, by skuteczność arbitrażu była ograniczona:

- (1) Arbitrażysta wykazuje awersję do ryzyka.
- (2) Ryzyko fundamentalne ma charakter systematyczny, czyli nie może zostać wyeliminowane poprzez dywersyfikację.

Oba warunki bliskie są rzeczywistości. Pierwszy oznacza, że inwestor nie będzie samotnie dążył do wyeliminowania ryzyka przez zajęcie bardzo dużej pozycji w błędnie wycenianych akcjach spółki. Drugi warunek zapewnia, że cena nie zostanie skorygowana przez wielu inwestorów, z których każdy zajmie niewielką pozycję.

W drugim kroku należy uwzględnić sytuację, kiedy doskonały substytut jest dostępny. Mimo iż jego istnienie chroni inwestora zarówno przed ryzykiem fundamentalnym, jak i ryzykiem modelu, skuteczność arbitrażu może nadal podlegać ograniczeniom. Problemem pozostaje bowiem ryzyko szumu, czyli ryzyko nieprzewidywalnych zmian nastroju nieracjonalnych inwestorów. De Long i in. (1990) podają następujące warunki wystarczające do ograniczenia skuteczności arbitrażu:

- (1) Arbitrażysta wykazuje awersję do ryzyka i ma krótki horyzont inwestycyjny.
- (2) Ryzyko szumu ma charakter systematyczny, czyli nie zostanie usunięte poprzez dywersyfikację.

Tak jak poprzednio, warunek pierwszy zapewnia, że cena nie zostanie skorygowana działalnością jednostki. Warunek drugi implikuje korelację nastrojów – pesymistyczne lub optymistyczne oczekiwania wobec danej spółki udzielają się szerokiej rzeszy inwestorów. Jego spełnienie gwarantuje, że rozbieżność ceny i wartości nie zostanie usunięta przez wielu drobnych uczestników rynku.

* * *

Powyższe rozważania pokazują, że arbitrażysta podlega licznym ograniczeniom, nawet jeśli warunki jego działania zbliżone są do teoretycznych. Przedstawiona systematyka rodzajów ryzyka nie wyczerpuje bynajmniej wszystkich limitacji w praktyce arbitrażu. Pozwala jednak nabrać dystansu do klasycznego założenia o braku ryzyka w tej strategii. Należy uznać, iż ryzyko wynikające z nieprzewidywalności nastrojów rynkowych połączone z ograniczonym horyzontem inwestycyjnym jest jednym z podstawowych czynników zmniejszających atrakcyjność arbitrażu. Ponadto jednostki nieracjonalne – noise-traderzy – mogą być dodatkowo wynagradzane ze względu na systematyczny charakter ryzyka, jakie same tworzą. Prowadzi to do paradoksalnej sytuacji, kiedy inwestorzy ci nie tylko zdolni są przetrwać na rynku, ale i osiągają lepsze wyniki niż uczestnicy racjonalni.

W świetle tych wniosków niesłuszna wydaje się zatem wiara zwolenników klasycznego podejścia, że efektywność rynku jest gwarantowana samym faktem istnienia racjonalnych arbitrażystów.

Proces decyzyjny w inwestowaniu zaczyna się od percepcji i przetwarzania napływających informacji. Na tej podstawie jednostki formułują przekonania, a następnie podejmują decyzje. Tradycyjna teoria finansów w duchu markowitowskim uznaje ów pierwszy etap za dany, koncentrując się bezpośrednio na ostatniej fazie – wyborze portfela inwestycyjnego. W świetle udowodnionych zniekształceń informacji, powstających w wyniku jej „edycji”, to standardowe podejście wydaje się daleko idącym uproszczeniem.

Treścią niniejszego rozdziału są zagadnienia wnioskowania i tworzenia poglądów o nieznanym otoczeniu. W pierwszej części przedstawiono konkurujące ze sobą koncepcje związane z racjonalnością jednostki oraz z metodami, jakimi posługuje się ona w celu przetwarzania informacji. Następnie omówione zostało podejście oparte na heurystykach i systematycznych błędach wnioskowania, zapoczątkowane pracami Tversky’ego i Kahnemana, które rozszerzono o dodatkowe elementy w części trzeciej.

2.1 Wnioskowanie o nieznanym

Podłoże decyzji stanowią przekonania dotyczące prawdopodobieństwa zajścia określonego zdarzenia. Przekonania zaś to rezultat procesu wnioskowania o nieznanym, którego przesłankami są napływające informacje. Do tego miejsca teoretycy różnych obozów (ekonomiści i psychologowie) wydają się jednomyślni. Jednakże sam przebieg procesów postrzegania informacji i wnioskowania wciąż stanowi przedmiot ożywionej debaty. Zwolennicy podejścia racjonalnego obstają przy twierdzeniu, że rezultaty ludzkiego rozumowania wskazują na użycie reguł matematycznych. Propagatorzy nurtu psychologicznego zgodnie oddalają takie założenie, pozostając jednocześnie między sobą w głębokiej niezgodzie odnośnie faktycznego sposobu przetwarzania informacji i jego konsekwencji.

W bieżącym podrozdziale omówiono te konkurencyjne podejścia w porządku chronologicznym. Już teraz warto zasygnalizować, że dalsze rozważania bazują na koncepcji proponowanej przez Tversky'ego i Kahnemana (podrozdział 2.2).

2.1.1 Prawa prawdopodobieństwa i statystyki

W ekonomii i finansach definicja jednostki racjonalnej – *homo oeconomicus* – obejmuje jej dwie cechy: konsekwencję działania oraz dążenie do podnoszenia poziomu osobistego dobrobytu. W ujęciu tym racjonalność jest postrzegana przez pryzmat rezultatów dokonywanych wyborów. W oparciu o analizę konsekwencji ludzkich poczynań przyjmuje się, że (i) jednostki przetwarzają wszelkie możliwe informacje, odpowiednio uaktualniając swoje przekonania, (ii) na podstawie przekonań formują preferencje i podejmują decyzje w taki sposób, aby zmaksymalizować swoją oczekiwaną użyteczność. Obecnie przedyskutowane zostanie pierwsze z założeń.

Podejście tradycyjne zakłada, że w celu „obróbki” informacji i tworzenia swoich osądów jednostki posługują się narzędziami teorii prawdopodobieństwa i statystyki. Źródła takiego myślenia należy poszukiwać w historii, kiedy to klasyczni teoretycy – Condorcet, Poisson czy Laplace – stawiali teorię prawdopodobieństwa na równi ze zdrowym rozsądkiem wyedukowanych ludzi. Tak oto dostępne narzędzia matematyczne przyjęło się uznawać za opis właściwego ludziom procesu wnioskowania. Nie dziwi zatem przyjmowane w teorii ekonomii i finansów podstawowe założenie o racjonalnej jednostce rozumującej zgodnie z prawem Bayesa. Ludziom przypisuje się umiejętność automatycznej aktualizacji wyobrażeń o prawdopodobieństwie wraz z napływem wszelkich nowych informacji. Bayesowska racjonalność nakazuje, by ów proces aktualizacji przebiegał zgodnie z formułą:

$$p(H / D) = \frac{p(H) \cdot p(D / H)}{p(D)}, \quad (2.1)$$

gdzie:

$p(H/D)$ – prawdopodobieństwo *a posteriori* zdarzenia H po ukazaniu się informacji D,

$p(H)$ – prawdopodobieństwo *a priori* (bazowe) zdarzenia H przed pojawieniem się informacji D,

$p(D/H)$ – prawdopodobieństwo ujawnienia informacji D, jeśli H jest prawdą,

$p(D)$ – prawdopodobieństwo całkowite zdarzenia D.

Założenie bayesowskiego myślenia nie oznacza, iż każdy osąd poprzedzony jest mozolnymi obliczeniami. Chodzi raczej o przeciętny rezultat wnioskowania, który jest zgodny z zasadami prawdopodobieństwa. Również teoretycy, tacy jak Milton Friedman, akceptują to mniej restrykcyjne podejście; wystarczy, że konsekwencje wnioskowania są takie, jak gdyby ludzie stosowali prawa rachunku prawdopodobieństwa i statystyki.

2.1.2 Ograniczona racjonalność

Pierwszy głos sprzeciwu wobec założenia o bezgranicznej mocy przetwórczej ludzkiego umysłu padł ze strony Herberta Simona (1955). Simon zasugerował dwie podstawowe przyczyny, które uniemożliwiają osiągnięcie doskonałej racjonalności: ograniczenia czasowe i ograniczenia technologiczne. Po pierwsze, w czasie wyznaczonym na podjęcie decyzji jednostka albo nie jest w stanie uzyskać dostępu do wszystkich informacji istotnych dla danego problemu, albo nie jest zdolna do ich dokładnego przetworzenia. Po drugie, względy technologiczne nie pozwalają na analizę informacji w sposób umożliwiający dokładne wyznaczenie prawdopodobieństwa celem optymalizacji wyborów. Owe ułomności sprawiają, że ludzie nie kierują się prawami matematyki – są co najwyżej racjonalni w sposób ograniczony.

Zdolność formułowania i rozwiązywania przez człowieka złożonych zadań jest niezwykle mała w porównaniu ze skalą problemów spotykanych w rzeczywistości. Konfiguracja wielowymiarowej funkcji celów oraz wielu potencjalnych ścieżek działania sprawia, że jednostka cierpi z powodu „przeła-

dowania” informacją. Pod natłokiem danych nie jest w stanie rozpatrzeć wszystkich ewentualności tak, aby dokonać optymalnego wyboru. Stąd naturalną skłonnością człowieka jest upraszczanie: rozważanie jedynie tych możliwości, które wydają się istotne. Na podstawie tak wyodrębnionych opcji jednostka decyduje się na określone działanie, które jest „wystarczająco dobre”, lecz przypuszczalnie nieoptymalne. Ów algorytm postępowania nazwał Simon (1957) „strategią satysfakcjonowania” (*satisficing*). Jednym z najlepszych przykładów jej zastosowania jest początkowe ustalenie poziomu swych aspiracji i wybór pierwszej napotkanej możliwości, która pozwala osiągnąć ten poziom.

Reasumując, koncepcja ograniczonej racjonalności Simona miała spełnić trzy zadania: (i) zaproponować alternatywę wobec *homo oeconomicus*, (ii) która byłaby kompatybilna z posiadanym przez jednostki dostępem do informacji, a także z ich zdolnościami obliczeniowymi, (iii) i która działałaby w rzeczywistym środowisku, w jakim jednostki się poruszają (Simon, 1955).

2.1.3 Szkoła Tversky’ego i Kahnemana

Podejście Tversky’ego i Kahnemana powstało na podstawie obserwacji rzeczywistych zachowań, niezgodnych z tradycyjnym rozumieniem racjonalności (por. punkt 2.1.1). Tversky i Kahneman (1974) twierdzą, że ograniczona racjonalność determinowana jest dwiema charakterystykami świata zewnętrznego: presją czasu oraz złożonością informacji.

Aby przezwyciężyć swe ograniczone zdolności analizy danych, ludzie wykształcili zespół heurystyk. Heurystyki należy rozumieć jako strategie oparte na intuicyjnej ocenie rzeczywistości i stosowane (zarówno rozmyślnie, jak i nieświadomie) w celu zastąpienia złożonego procesu estymacji prawdopodobieństwa i prognozowania wartości operacjami prostego wnioskowania. Dzięki heurystykom jednostki zawężają zbiór możliwości i jednocześnie zwiększają swoje szanse na znalezienie rozwiązania w sytuacji decyzyjnej.

Tversky i Kahneman (1974) definiują trzy podstawowe rodzaje heurystyk:

- (1) dostępność (*availability*) – wnioskowanie na podstawie informacji łatwo dostępnych w pamięci,

- (2) reprezentatywność (*representativeness*) – wnioskowanie na podstawie podobieństwa,
- (3) zakotwiczenie i dostosowanie (*anchoring and adjustment*) – wnioskowanie na podstawie pierwotnie zasugerowanej wartości (tzw. kotwicy), która jest dostosowywana w celu oszacowania wartości rzeczywistej.

Na uwagę zasługują ambiwalentne konsekwencje polegania na tych regułach, często podkreślane przez autorów. Z jednej strony bowiem heurystyki stanowią wydajne narzędzie orientacji w natłoku informacji (funkcja pozytywna), z drugiej zaś dają źródło dotkliwym i systematycznym błędom wnioskowania (funkcja negatywna), które uniemożliwiają osiągnięcie optymalnych ekonomicznie wyników. Wprowadzenie podejścia heurystycznego²³ do nauki o finansach ma również dwojakie implikacje: ze względu na systematyczny charakter popełnianych błędów umożliwia ono prognozowanie zachowań jednostek, kładąc jednocześnie kres przyjmowaniu do modelowania rynków upraszczającego założenia o ich racjonalności.

Zaproponowana w podrozdziałach 2.2 i 2.3 klasyfikacja procesów postrzegania i przetwarzania informacji opiera się zasadniczo na podejściu Tversky'ego i Kahnemana, jako że koncepcja heurystyk i systematycznych błędów wnioskowania (*heuristics-and-biases approach*) znalazła szerokie zastosowanie w analizie zachowań inwestorów.

2.1.4 Szkoła Gigerenzera

Najmłodsza koncepcja racjonalności, postulowana przez grupę badaczy skupionych wokół profesora Gerda Gigerenzera, wyrosła na gruncie sprzeciwu wobec szkoły Tversky'ego i Kahnemana. Centralnym punktem owej krytyki jest fakt, iż w podejściu heurystycznym zachowania ludzi porówny-

²³ Określenia „podejście heurystyczne” używa się tu dla oznaczenia podejścia zaproponowanego przez Tversky'ego i Kahnemana, opartego na heurystykach i systematycznych błędach wnioskowania (*heuristics-and-biases approach*). Uściślenie to jest istotne, ponieważ także przedstawiciele innych szkół (por. punkt 2.1.4) posługują się określeniem „heurystyka”, aczkolwiek w odmiennym znaczeniu.

wane są z klasycznie pojmowaną racjonalnością. Gigerenzer zadaje pytanie o zasadność traktowania jako błędne poczynień, które oceniane są przez pryzmat tak nierealistycznego obrazu człowieka²⁴.

W zamian Gigerenzer proponuje spojrzeć na zachowania ludzi w oderwaniu od jakiegokolwiek modelu wartościującego. Co więcej, w jednym z głównych artykułów, napisanym wspólnie z Toddem (Gigerenzer i Todd, 1999), wygłasza pogląd, że ekonomiczna racjonalność może dawać gorsze rezultaty niż stosowanie tzw. heurystyk szybkich i oszczędnych (*fast and frugal heuristics*). Podejście Gigerenzera stanowi zatem optymistyczne wyobrażenie o ludzkich możliwościach wnioskowania. Używając heurystyk szybkich i oszczędnych, jednostki dają wyraz swej zdolności adaptacji do nowo zaistniałych warunków. Powołanie się na aspekt środowiskowy zbliża racjonalność w ujęciu Gigerenzera do koncepcji Simona.

Inspiracją dla heurystyk szybkich i oszczędnych była zresztą Simonowska strategia satysfakcjonowania (por. punkt 2.1.2). Na jej podstawie grupa badawcza Gigerenzera zdefiniowała zespół algorytmów – probabilistycznych modeli myślowych (*probabilistic mental models*) – stosowanych przez jednostki w procesie wnioskowania o nieznanym i podejmowania decyzji. Jednym z takich prostych, aczkolwiek wydajnych algorytmów jest „wybierz to, co najlepsze”. Jest on wart wspomnienia ze względu na próby jego implementacji w analizie zachowań na rynkach finansowych. Składa się nań pięć zasad, stosowanych kolejno w procesie decyzyjnym: (i) rozpoznawania, (ii) poszukiwania wskazówek, (iii) rozróżniania, (iv) zastępowania wskazówek, (v) maksymalizacji przy dokonywaniu wyborów. Jak widać, pierwsze cztery kroki składają się na wnioskowanie, zaś w ostatnim podejmowane są decyzje. Niekiedy proces zostaje ucięty już po pierwszym etapie. Heurystyka²⁵

²⁴ Warto wspomnieć, iż krytyka Gigerenzera znacznie upraszcza podejście heurystyczne. Tversky i Kahneman nie wartościują heurystyk jako reguł niewłaściwego postępowania. Badacze wielokrotnie podkreślają, że reguły te są przydatną metodą radzenia sobie ze złożonością informacji, lecz uwrażliwiają na częste błędy wynikające z ich zastosowania. Zarzuty Gigerenzera dały początek ożywionej debacie na łamach *Psychological Review* i zostały skutecznie odparte. Por. Kahneman i Tversky (1996) oraz Gigerenzer (1996).

²⁵ Określenie „heurystyka” pochodzi od autorów podejścia (por. uwagi w przypisie 23).

rozpoznawania (krok (i)) stanowi, że jednostka kończy proces (dokonuje wyboru) w momencie, gdy rozpozna dany obiekt. Zasada ta podpowiada, że jeśli inwestor ma do wyboru akcje dwóch spółek, z których jedna jest przez niego rozpoznawana, wówczas jego wybór padnie właśnie na nią²⁶. Jeśli żaden obiekt nie zostanie rozpoznany (albo zostanie rozpoznany więcej niż jeden), wówczas proces jest kontynuowany. Rozpoczyna się poszukiwanie wskazówek (krok (ii)) i sprawdzanie, czy pozwalają one na wartościujące rozróżnienie alternatyw (krok (iii)). Jeśli wskazówki są niewystarczające do podjęcia wyboru, wówczas jednostka poszukuje nowych informacji (krok (iv)), aby ostatecznie dokonać wyboru (krok (v)) obiektu o najwyższej wartości tudzież wyboru losowego, jeśli obiekt najlepszy nie został znaleziony.

Zaproponowane przez badaczy probabilistyczne modele myślowe mają poprawić ludzkie umiejętności szybkiego wnioskowania pod presją czasu i przy ograniczonych zdolnościach poznawczych. Konstrukcję tych modeli znamionują zarówno cechy podejścia deskryptywnego, jak i normatywnego²⁷.

Wstępne badania wskazują na potencjalnie większą skuteczność heurystyk szybkich i oszczędnych niż wnioskowania bayesowskiego, co podkreśla ich wartość normatywną. Inspiracją dla takich wniosków są dobre wyniki osiągnięte przez niewyrafinowanych inwestorów, którzy podejmują decyzje bez dostępu do wiedzy eksperckiej oraz bez dogłębnego zrozumienia zasad funkcjonowania rynku. Gigerenzer i Goldstein (1999) upatrują źródeł ich powodzenia właśnie w stosowaniu heurystyk szybkich i oszczędnych. Dodatkowym potwierdzeniem tej tezy ma być, zdaniem badaczy, ponadprzeciętna rentowność ich eksperymentalnego portfela, zbudowanego jedynie

²⁶ Oczywiście „rozpoznawanie” rozumiane jest w tym wypadku jako przywołanie w pamięci pozytywnych cech obiektu. Heurystyka rozpoznawania działa w następujący sposób: jeśli jeden obiekt jest rozpoznawany, zaś drugi nie, wówczas wnioskujemy, że obiekt rozpoznawany ma wyższą wartość od nierozpoznanego przy zastrzeżeniu, że wynik rozpoznania jest pozytywnie skorelowany z kryterium rozpoznawania. Jeśli korelacja jest negatywna, wówczas wnioskujemy o wyższości wartości obiektu nierozpoznanego nad rozpoznany (Gigerenzer i Todd, 1999).

²⁷ Definicję modelu normatywnego i deskryptywnego podano w rozdziale 3 (por. punkt 3.1.2 oraz uwagi wprowadzające w podrozdziale 3.2).

w oparciu o heurystykę rozpoznawania testowaną na laikach – bez zaangażowania jakiejkolwiek wiedzy eksperckiej²⁸ (Borges i in., 1999).

Dodatkowo należy zastanowić się nad wartością deskryptywną wspomnianych heurystyk, czyli nad tym, czy stanowią one właściwy opis faktycznych zachowań uczestników rynku. Aktywni inwestorzy, szczególnie o krótkim horyzoncie inwestycyjnym, niewątpliwie myślą algorytmami, opierając swe decyzje na ograniczonej liczbie wskazówek. Intuicyjnie można zatem przypuszczać, że omawiana koncepcja dość dobrze odzwierciedla sposoby wnioskowania stosowane w inwestowaniu. Jednakże na potwierdzenie tej hipotezy za pomocą rygorystycznych testów należy jeszcze poczekać.

2.2 Heurystyki i systematyczne błędy wnioskowania

Jak wcześniej wspomniano, koncepcja wnioskowania o prawdopodobieństwie oparta na heurystykach i systematycznych błędach znajduje swój początek w pracach Kahnemana i Tversky'ego, publikowanych już w latach siedemdziesiątych. Wypadnie podkreślić, iż autorzy nie uznają heurystyk *per se* za niewłaściwe. Wręcz przeciwnie – twierdzą, że jednostka działająca pod presją czasu i przytłoczona dużą ilością złożonych informacji nie może się bez nich obejść. Uwrażliwiają jednak na systematyczność wynikających przy tym błędów.

2.2.1 Dostępność

Heurystyka dostępności stanowi, że jednostka tworzy przekonania o prawdopodobieństwie lub częstotliwości zdarzeń kierując się łatwością przywoływania w myślach stosownych doświadczeń bądź konotacji (Tversky i Kahneman, 1974). Mentalna dostępność przykładów jest przydatną wskazówką w szacowaniu prawdopodobieństw, albowiem po ilustracji zdarzeń częstych sięga

²⁸ Badanie to jest pionierskim eksperymentem w zakresie stosowania heurystyki rozpoznawania. W obliczu braku głębszych dowodów w niniejszym opracowaniu zasygnalizowano jedynie, iż pierwsze próby implementacji heurystyk szybkich i oszczędnych miały już miejsce. Autorka rezygnuje jednak z omówienia eksperymentu ze względu na ograniczoną wartość poznawczą (badania zostały przeprowadzone na rynku silnie wzrostowym, brak natomiast ich potwierdzenia w innych warunkach).

się zwykle łatwiej i szybciej, niż ma to miejsce w przypadku zajęć rzadkich. Dostępność jest jednak determinowana nie tylko faktyczną częstością sytuacji, lecz także innymi czynnikami, nieistotnymi w kontekście określanego prawdopodobieństwa. Stąd stosowanie tejże heurystyki otwiera szerokie możliwości popełniania błędów wnioskowania, czego przykłady zostaną obecnie omówione.

2.2.1.1 Doświadczenia

Gdy ocena prawdopodobieństwa zależy od łatwości przypominania sobie odpowiednich zdarzeń, wówczas liczebność owych wspomnień wpływa na poziom szacowanego prawdopodobieństwa. Klasy zdarzeń, dla których jednostka jest w stanie podać wiele przykładów, uznawane są za bardziej prawdopodobne niż te o ograniczonej egzemplifikacji, mimo iż faktyczna częstość występowania obu klas może być identyczna.

Modelowy przykład takiego błędu podają Tversky i Kahneman (1973). Uczestnikom eksperymentu odczytano listę nazwisk znanych osób obu płci i poproszono ich o osądzenie, czy lista zawiera więcej nazwisk kobiet, czy mężczyzn. Różnym grupom respondentów zaprezentowano różne listy nazwisk, wśród których raz relatywnie bardziej znani byli mężczyźni, innym zaś razem – kobiety. W każdej próbie uczestnicy błędnie uznawali za bardziej liczne te klasy (płci), w których nazwiska były bardziej znane, a zatem lepiej rozpoznawalne, choć ich faktyczna liczebność była niższa niż w wypadku klas gorzej rozpoznawalnych.

Podobną metodologię badawczą, bliższą jednak rzeczywistości inwestora, przyjął Stephan (1999). Respondenci otrzymali raport z wyników giełdowych dla akcji 51 niemieckich spółek, z których 25 było dobrze znanych, zaś 26 mało znanych²⁹. Gdy raporty wskazywały, że wśród akcji znanych firm przeważają spadki kursów, zaś wśród akcji mniej znanych – wzrosty, wówczas większość respondentów reprezentowała mylny pogląd o przewa-

²⁹ Zaszeregowanie danej spółki do odpowiedniej grupy zostało dokonane na podstawie wstępnych testów.

dze spadków w kontekście całego rynku. Natomiast gdy w grupie dobrze znanych spółek przeważały wzrosty, a w grupie mniej znanych – spadki, wówczas respondenci błędnie przekonani byli, że kursy większości spółek na rynku wzrosły. Za każdym razem dobrze znane spółki okazywały się bardziej dostępne, przez co dominowały w ocenie całej „populacji”. Będąc świadomym tego efektu, Peter Lynch, były manager największego na świecie funduszu inwestycyjnego – Magellan Fund, postulował unikanie inwestowania w akcje spółek percepcyjnie łatwo dostępnych dla większości inwestorów. Dostępność uznawał bowiem za główną przyczynę przewartościowania aktywów giełdowych (Bazerman, 1998, s. 7).

Także wyjątkowość i dramatyzm sytuacji czy też aktualność zdarzeń mogą być bodźcem do stosowania omawianej heurystyki. Zjawiskiem powszechnym jest przecenianie prawdopodobieństwa wypadku samochodowego tuż po jego ujrzaniu (aktualność). Podobnie na rynku kapitałowym inwestorzy, którzy doświadczyli krachu giełdowego, przeszacowują szanse jego ponownego wystąpienia w danym odcinku czasu (dramatyzm).

2.2.1.2 Wyobrażenia

Nawet jeśli jednostka nie dysponuje odpowiednimi doświadczeniami, by na ich podstawie określić prawdopodobieństwo, heurystyka dostępności może mimo to działać. W przypadku braku własnych doświadczeń ludzie posługują się wyobrażnią. Z jej to pomocą konstruują przykłady potencjalnych zdarzeń, zaś sprawność, z jaką są w stanie to czynić, determinuje ich przekonania o prawdopodobieństwie.

Założmy, że inwestor ma do wyboru akcje dziesięciu spółek, spośród których chce utworzyć portfel. Ile istnieje możliwości zbudowania portfela złożonego z dwóch aktywów, a ile – portfela ośmioelementowego? W analogicznym eksperymencie przeprowadzonym przez Tversky’ego i Kahnemana (1973) respondenci podają przeciętnie 70 kombinacji dla portfela dwuelementowego oraz 20 dla ośmioelementowego. Oczywiście w obu przypadkach liczba możliwych konstrukcji jest identyczna (45), bowiem kombinacje dwuelementowe z dziesięciu w sposób jednoznaczny określają kombinacje

obejmujące osiem elementów. Dlaczego zatem podawane odpowiedzi są tak odległe od właściwej? Otóż uwaga respondentów skupia się przypuszczalnie na pierwszej dostępnej w myślach możliwości stworzenia pięciu rozłącznych zbiorów dwuelementowych i tylko jednego ośmioelementowego. Badacze tłumaczą wyniki eksperymentu faktem, iż zbiory dwu elementów są łatwiej wyobrażalne niż zbiory liczniejsze, na przykład ośmioelementowe.

Wyobrażenia odgrywa także ważną rolę w ocenie prawdopodobieństw w codziennym życiu. Subiektywne ryzyko związane z uruchomieniem nowego przedsięwzięcia może być przykładowo oceniane poprzez mnożenie przykładów potencjalnych trudności, które zadecydują o ostatecznym fiasku inwestycji. Żywotność owych wyobrażeń jest w stanie wpłynąć na przeszacowanie prawdopodobieństwa niepowodzenia. I odwrotnie: ryzyko przedsięwzięcia może być niedoceniane, jeśli ewentualne zagrożenia trudno „mieszczą się w głowie” bądź też w ogóle nie przychodzą na myśl.

Heurystyka dostępności wpływa również na decyzje inwestycyjne. Niektórzy badacze powołują się na nią w tłumaczeniu nadmiernej skłonności inwestorów do lokowania środków w akcjach krajowych w porównaniu z zagranicznymi (*home bias*). Przypuszczalnie inwestorzy nie szacują ryzyka zagranicznych akcji jedynie w oparciu o historyczną wariancję ich stopy zwrotu, lecz kierują się także przesłankami takimi, jak ograniczona znajomość funkcjonowania zagranicznych rynków, instytucji i firm. Właśnie owa zawężona dostępność wyobrażeń o obcym rynku może implikować przecenianie ryzyka inwestycji za granicą. W tym wypadku bezpośrednią negatywną konsekwencją polegania na heurystyce dostępności jest brak międzynarodowej dywersyfikacji portfela (French i Poterba, 1991)³⁰.

³⁰ French i Poterba (1991) powołują się na badania Tversky'ego i Heatha (1991), którzy pokazują, że ludzie wyżej szacują ryzyko gier, które nie są im znane niż tych, z którymi mieli okazję się zapoznać. Dzieje się tak nawet wówczas, gdy jednostki są świadome idetycznego rozkładu prawdopodobieństwa w obu grach.

2.2.1.3 Iluzoryczne korelacje

Posługiwanie się heurystyką dostępności w naturalny sposób zachęca do dostrzegania iluzorycznych korelacji. Wiąż skojarzeniowa istniejąca między zdarzeniami rzutuje na ocenę częstości ich współwystępowania (Chapman i Chapman, 1969; Tversky i Kahneman, 1974). I tak, gdy skojarzenia są mocne, zdarzenia zdają się „chodzić parami”. Weźmy przykład firmy, która ma być wkrótce notowana na giełdzie NASDAQ. Informacja taka wzbudza w obserwatorach rynku natychmiastową asocjację owej spółki z rynkiem nowoczesnych technologii, motywując do przypisania jej akcjom dużego potencjału wzrostu (*growth stock*).

Warto w tym miejscu wspomnieć o wpływie tzw. schematów przyczynowego myślenia (*causal schemata*) na procesy wnioskowania i percepcję prawdopodobieństwa (Tversky i Kahneman, 1982). Myślenie przyczynowo-skutkowe umożliwia spójną interpretację zdarzeń, która jest naturalnym dążeniem człowieka. Dlatego jednostki cechuje skłonność do postrzegania sekwencji zdarzeń w kontekście relacji kauzalnych nawet wówczas, gdy związek jest czysto przypadkowy, zaś korelacja pozorna. Liczne badania wykazały ponadto, że z większą pewnością siebie wnioskuje się o skutkach na podstawie przyczyn, niż odwrotnie: diagnozuje przyczyny na podstawie obserwowanych skutków. Jeśli dane są dwa zdarzenia X i Y , dla których (i) prawdopodobieństwa brzegowe są identyczne, $p(X) = p(Y)$, a także $p(Y/X) = p(X/Y)$ oraz (ii) zajście X jest postrzegane jako naturalna przyczyna Y , to ludzie wykazują skłonność do postrzegania silniej prawdopodobieństwa, że Y zostało wywołane przez X , czyli $p(Y/X) > p(X/Y)$.

Tezę tę potwierdzają wyniki wielu eksperymentów, z których przykładowy podano poniżej (Tversky i Kahneman, 1982).

Które ze stwierdzeń jest bardziej sensowne? (W nawiasach kwadratowych umieszczono odsetek osób wybierających daną opcję; liczba respondentów $N=70$.)

- | | |
|--|-------|
| (i) Adam jest ciężki, ponieważ jest wysoki. | [90%] |
| (ii) Adam jest wysoki, ponieważ jest ciężki. | [10%] |

Mimo iż wzrostu i ciężaru ciała nie uznaje się za połączone związkiem przyczynowo- skutkowym, respondenci upatrują przyczynę wagi we wzroście i dlatego przypisują pierwszemu stwierdzeniu większe prawdopodobieństwo. Choć tego rodzaju eksperymenty przeprowadzone wśród inwestorów najprawdopodobniej nie są jeszcze dostępne, iluzoryczne związki kauzalne wydają się potencjalnie istotne w analizie informacji ekonomicznych, a następnie we wnioskowaniu (np. ocena jakości spółki uzależniona od prestiżu giełdy, na której spółka jest notowana).

Inna forma błędnego wnioskowania przyczynowo-skutkowego opiera się na chronologii zdarzeń. Zajście X , które miało miejsce przed zdarzeniem Y , jest w naturalny sposób uznawane za wysoce prawdopodobną przyczynę Y , szczególnie jeśli pojawieniu się X nie towarzyszyły inne zdarzenia. Przykładowo, gdy Fed (Bank Rezerwy Federalnej USA) zapowiada obniżkę stóp procentowych, a następnie wartość dolara ulega deprecjacji, wówczas ogłoszona obniżka będzie postrzegana jako naturalna przyczyna deprecjacji. Jeśli jednak informacja o obcięciu stóp procentowych pojawi się równolegle np. z ujawnieniem wiadomości o rezygnacji kluczowego członka rządu, wówczas związek obniżenia stóp i deprecjacji wyda się słabszy, mimo iż faktyczna korelacja zdarzeń nie uległa zmianie.

2.2.2 *Reprezentatywność*

Heurystyka reprezentatywności polega na ocenie stopnia, w jakim próba odpowiada całej populacji, dane zdarzenie – klasie zdarzeń, określone zachowanie – typowi osobowości lub, bardziej ogólnie, wynik – określonemu modelowi. W poszukiwaniu odpowiedzi na pytanie o prawdopodobieństwo tego, że obiekt X należy do klasy Y , że zdarzenie X zostało wygenerowane przez proces Y , czy też, że proces Y wygeneruje zdarzenie X , jednostki często polegają na reprezentatywności. Ich estymacja prawdopodobieństwa determinowana jest podobieństwem X do Y tudzież stopniem, w jakim X wydaje się charakterystyczne dla Y .

2.2.2.1 Prawdopodobieństwo bazowe

Jak pisano w punkcie 2.1.1, osoba racjonalna aktualizuje swoje przekonania o prawdopodobieństwie zdarzeń zgodnie z prawem Bayesa. Wnioskowanie za pomocą reprezentatywności implikuje bezpośrednio pogwałcenie tego prawa poprzez zaniedbywanie prawdopodobieństwa bazowego (*a priori*).

W klasycznym eksperymencie testującym powyższą hipotezę respondenci mają przewidzieć profesję osoby (prawnik czy inżynier) po zapoznaniu się z jej krótkim opisem (Kahneman i Tversky, 1973). Opis wskazuje na cechy charakteru zgodne ze stereotypem danego zawodu. Dodatkowo informacje o częstości występowania obu profesji są różne w różnych grupach respondentów (w jednym przypadku stosunek liczby prawników do liczby inżynierów kształtuje się jak 7:3, zaś w drugim jak 3:7). Zgodnie z hipotezą o zaniedbywaniu prawdopodobieństwa bazowego, uczestnicy eksperymentu przewidują zawód osoby, kierując się jedynie podobieństwem jej opisu do stereotypu; prawdopodobieństwa bazowe (częstość występowania danego zawodu w całej populacji) nie grają natomiast roli. Wyniki takie powielane były w licznych kolejnych eksperymentach.

De Bondt i Thaler (1985) jako pierwsi zastosowali heurystykę reprezentatywności (w sensie zaniedbywania prawdopodobieństwa bazowego) do wytłumaczenia zjawiska nadmiernej reakcji kursów na informacje i związanego z tym efektu zwycięzcy i przegranego (por. punkt 1.2.2.2). Badacze sugerowali, że inwestorzy przywiązują nadmierną wagę do informacji świeżych i przykuwających uwagę, ignorując dane bazowe, np. o wielkościach fundamentalnych spółki, sytuacji w branży czy też o ogólnym położeniu gospodarczym. Przykładem to potwierdzającym jest powszechne zjawisko wzrostu kursu spółki po ogłoszeniu przez nią informacji o nowym kontrakcie handlowym. Kurs wzrasta nawet wtedy, gdy wielkość kontraktu jest znikoma w relacji do całkowitego wolumenu sprzedaży spółki (Wärneryd, 2001, s. 128).

2.2.2.2 Percepcja losowości

Ludzie dzielą przekonanie, że sekwencja zdarzeń wygenerowana przez proces losowy odzwierciedla charakterystyczne cechy owego procesu niezależnie od swojej długości. Rozważając wyniki kolejnych rzutów monetą (O – orzeł, R – reszka), seria O-R-O-R-R-O zdaje się bardziej prawdopodobna niż O-O-O-R-R-R, która nie sprawia wrażenia losowej, podobnie zresztą jak O-O-O-O-R-O, która postrzegana jest jako niereprezentatywna dla rzutów symetryczną monetą (Kahneman i Tversky, 1972). Jeśli w kolejnych pięciu próbach stale wypada reszka, jednostki przeszacowują prawdopodobieństwo otrzymania orła za szóstym razem, mimo iż – na podstawie niezależności zdarzeń – prawdopodobieństwo wyrzucenia orła bądź reszki jest w każdym rzucie identyczne i wynosi 0,5. Losowość jest w naturalny sposób rozumiana jako proces samokorygujący – w celu utrzymania równowagi odchylenie w jednym kierunku powinno znosić się z odchyleniem w kierunku przeciwnym. Takie myślenie przyjęło się nazywać „błędem hazardzisty” (*gambler's fallacy*). Hazardzista, odnotowawszy na tarczy ruletki osiem czarnych liczb, jest niemalże pewien, że w dziewiątym wypaść musi czerwona, choć w rzeczywistości szanse wyrzucenia liczby czerwonej bądź czarnej są identyczne.

Ów błąd wynika ze złej interpretacji jednego z podstawowych praw statystycznych – prawa wielkich liczb. Zaniedbując rozmiar próby, jednostki sądzą, iż prawo to stosuje się także do krótkich serii, w domyśle: że małe próby są reprezentatywne dla całej populacji czy też dla generującego je procesu. Tę ludzką przypadłość Tversky i Kahneman (1971) określili dla kontrastu mianem „prawa małych liczb”.

Wiara w „prawo małych liczb” skłania również do przedwczesnego dostrzegania regularności w sekwencjach czysto losowych. W ślad za Gilovichem, Vallone i Tversky’em mówi się tu często o efekcie „gorącej ręki” (*hot hand fallacy*). Gilovich i in. (1985) jako pierwsi udokumentowali to zjawisko, obserwując zawodowych graczy w koszykówkę. Zarówno fani, jak i sami sportowcy wierzą w „gorącą rękę” czy też „dobrą passę” gracza. Dla

przykładu: jeśli jego cztery ostatnie rzuty były trafne, większość obserwatorów ocenia prawdopodobieństwo trafienia w piątym rzucie powyżej długoterminowej średniej skuteczności zawodnika. Obszerna analiza Gilovicha i in. wykazała jednak, iż intuicja ta jest mylna. Badacze udowodnili, że wyniki poprzednich rzutów nie podnoszą szans sukcesu w następnym, oraz że odchylenia od długookresowej przeciętnej skuteczności gracza są jedynie dziełem przypadku. Co więcej, wiara w „dobrą passę” zawodnika – poprzez nieoptymalną „alokację” piłki – okazuje się rzutować negatywnie na wyniki osiągane przez drużynę.

Zjawisko „gorącej ręki” jest wszechobecne również w świecie finansów, gdzie objawia się choćby pod postacią wiary w umiejętności managerów funduszy w oparciu o krótkie serie ich sukcesów (Kahneman i Riepe, 1998). Interesujące badania w tym zakresie przeprowadzili Stephan i Kiell (2000), którzy analizowali zachowania licznej grupy profesjonalnych inwestorów³¹. Pierwszej grupie respondentów przedstawiony został następujący scenariusz³²:

W prestiżowym konkursie na najlepszego eksperta walutowego uczestnicy mają przewidzieć dla kolejnych 10 dni handlowych wzrost lub spadek kursu dolara wobec jena. Tylko jedna osoba z 500 – pani X – dokonuje poprawnej prognozy. Jak ocenilibyś jej wynik?

A. Wybitny / B. Bardzo dobry / C. Ponad średnią / D. Średni / E. Czysty przypadek

Drugiej grupie zaprezentowano ten sam scenariusz, zmieniając jednak informację o liczbie uczestników eksperymentu z 500 na 90.

Przy założeniu, że fluktuacje kursów walutowych generowane są przez proces stochastyczny, prawdopodobieństwo, że co najmniej jedna osoba z 500 poprawnie odgadnie sekwencję, posługując się zwykłym rzutem mo-

³¹ W eksperymencie Stephana i Kiella brało udział 142 inwestorów, zarządzających łącznym kapitałem o wartości około 350 miliardów euro. Por. (Stephan i Kiell, 2000).

³² Wersja Stephana i Kiella została tu uproszczona w sposób nie zmieniający jej sensu.

netą, jest stosunkowo wysokie i wynosi 0,386³³. Prawdopodobieństwo tego samego zdarzenia w grupie 90 osób jest znacznie niższe i równe 0,084. Przyjrzyjmy się teraz odpowiedziom uczestników eksperymentu. W pierwszej grupie (scenariusz z 500 uczestnikami), 44% respondentów uznało wynik osiągnięty przez panią X za wybitny lub bardzo dobry (odpowiedzi A lub B), podczas gdy w drugiej grupie (scenariusz z 90 uczestnikami) zdanie to podzieliło jedynie 27% zapytanych osób. Jak widać, umiejętności pani X zostały lepiej ocenione w sytuacji, gdy prawdopodobieństwo odgadnięcia wyniku (bez wiedzy eksperckiej) było wyższe. Otrzymane rezultaty wskazują na posługiwanie się heurystyką reprezentatywności: zwycięstwo w gronie 500 specjalistów wydaje się mocniejszym dowodem fachowości niż pokonanie 89 przeciwników.

Jakkolwiek wiedza na temat funkcjonowania rynków finansowych jest istotnym czynnikiem w inwestowaniu, wielu „ekspertów” może odnosić sukcesy jedynie za sprawą przypadku. Ludzki umysł skłonny jest jednak przypisywać im ponadprzeciętną sprawność w prognozowaniu zdarzeń. Piśma finansowe bezustannie przysparzają klientów funduszom, reklamując ich ponadprzeciętne wyniki w poprzednim roku. Należy jednak zauważyć, że duże instytucje, jak np. Fidelity czy Dreyfus, zawsze będą mogły pochwalić się doskonałymi wynikami jednego z wielu swoich funduszy. Ten jako jedyny zostanie zareklamowany szerokiej rzeszy inwestorów; inne, z osiągnięciami poniżej przeciętnej, nie zostaną wymienione.

2.2.2.3 Regresja do średniej

Regresja do średniej jest statystycznym zjawiskiem charakteryzującym proces losowy. W jej wyniku po wartościach skrajnie wysokich lub niskich następują realizacje mniej ekstremalne. Prawidłowość tę można zaobserwować w życiu codziennym, porównując wzrost rodziców i dzieci, zdolności

³³ Wyznaczone na podstawie schematu Bernoulli’ego prawdopodobieństwo zdarzenia, że co najmniej jedna osoba z 500 odgadnie 10 kolejnych zmian kursu; prawdopodobieństwo sukcesu $p \approx 0,000977$, liczba prób $n=500$, liczba sukcesów $k \geq 1$.

intelektualne pomiędzy rodzeństwem, osiągnięcia studentów na kolejnych egzaminach czy też wyniki finansowe firm w poszczególnych latach. Podstawą występowania regresji do średniej jest to, iż wartości ekstremalne cechuje ograniczona wariancja: wartości wysokie wykazują tendencję do spadków, z kolei wartości niskie do wzrostów, ciężąc w obu przypadkach ku długookresowej średniej.

Mimo powszechności występowania tego statystycznego faktu jednostki nie rozwinęły odpowiedniej intuicji w jego rozpoznawaniu (Tversky i Kahneman, 1974), a to z powodu posługiwania się przeszłymi wynikami (np. danymi o sprzedaży firmy z roku 2002) do prognozowania przyszłych realizacji (np. sprzedaży w roku 2003) poprzez prostą ich ekstrapolację.

Także w ocenie trendów ekonomicznych ludzie ulegają tej formie reprezentatywności. Dwie podstawowe koncepcje, które sterują prognozami rynku, to kontynuacja trendu lub też jego odwrócenie. Warto przypomnieć (por. punkt 1.2.2), że badania empiryczne dowodzą kontynuacji trendu – pozytywnej autokorelacji stóp zwrotu – w krótkim okresie (od 6 do 12 miesięcy) oraz zjawiska odwrócenia trendu – negatywnej autokorelacji – w średnim i długim terminie (od 3 do 5 lat). Optymalne inwestowanie powinno zatem łączyć strategię opartą na *momentum* (por. punkt 1.2.2.1) w perspektywie krótkoterminowej z kontrarianizmem (por. punkt 1.2.2.2) w dłuższym okresie. I faktycznie: zgodnie ze swoimi przekonaniem inwestorzy zaliczają się do zwolenników *momentum* bądź kontrarianizmu. Trzeba jednak zadać następujące pytanie: (i) czy dokonywany przez nich wybór strategii jest wrażliwy na długość przeszłego trendu oraz, (ii) jaki odsetek inwestorów łączy strategię *momentum* z kontrarianizmem w sposób optymalny, uwzględniający wspomniane autokorelacje.

Znalezienie odpowiedzi na te pytania było jednym z celów badawczych opisanego w poprzednim punkcie eksperymentu Stephana i Kiella (2000). (i) W badanej grupie profesjonalnych inwestorów wybór strategii był w dużej mierze niezależny od czasu trwania trendu. Pokazują to następujące wyniki: w krótkiej perspektywie (6 miesięcy) odsetek zwolenników kontrarianizmu

plasował się na poziomie 36% i niewiele różnił się od wyników dla horyzontu długiego (3 lata), w którym strategię tę wybierało 40% uczestników. W świetle występowania pozytywnej autokorelacji stóp zwrotu w krótkim okresie, wybór strategii kontrariańskiej dla okresu sześciomiesięcznego jest błędem. (ii) Ze względu na strukturę danych odpowiedź na drugie pytanie została udzielona poprzez zdefiniowanie odpowiednich przedziałów. I tak: odsetek inwestorów wybierających optymalną kombinację obu strategii – *momentum* w krótkim terminie, kontrarianizm w długim – kształtował się między 14% a 34%. Jednocześnie najmniej korzystnego połączenia obu strategii można było oczekiwać od 11% inwestorów przy optymistycznym szacunku aż do 30% w wersji najbardziej pesymistycznej.

Należy jeszcze raz podkreślić, że uczestnikami eksperymentu Stephana i Kiella byli zawodowi uczestnicy rynku, zajmujący duże pozycje inwestycyjne. Uzyskane rezultaty potwierdzają postawioną przez Tversky'ego i Kahnemana tezę o braku zrozumienia regresji do średniej. Co więcej, uprzytomniają, iż nawet profesjonalizm nie czyni wolnym od systematycznych błędów wnioskowania.

2.2.2.4 Błąd koniunkcji

Błąd koniunkcji jest potencjalnie najważniejszym przykładem mylnego wnioskowania o prawdopodobieństwie. Wyobraźmy sobie, że dany jest ogólny rys osobowości danej jednostki, który na podstawie własnych spekulacji uzupełniamy o dodatkowe atrybuty (np. zawód, poglądy polityczne). Jak stanowi jedno z podstawowych praw rachunku prawdopodobieństwa, uściślenie zdarzenia nie zwiększa szansy jego wystąpienia. Zatem prawdopodobieństwo tego, że osoba podziela poglądy lewicowe (zdarzenie X) i jest jednocześnie artystą (zdarzenie Y), musi być z definicji niższe niż prawdopodobieństwo, że osoba ta jest artystą (zdarzenie Y). Formalnie zasadę koniunkcji można zapisać jako $p(X \cap Y) \leq p(Y)$.

Błąd koniunkcji polega na przypisywaniu większego prawdopodobieństwa koniunkcji zdarzeń niż pojedynczemu zdarzeniu (np. elementowi opisu).

Dzieje się tak wówczas, gdy koniunkcja sprawia wrażenie bardziej reprezentatywnej niż pojedynczy jej komponent. Powszechność występowania tego efektu i jego siłę potwierdziły obszerne badania zapoczątkowane przez Tversky'ego i Kahnemana. Warto tu powołać się na jedno z nich przeprowadzone w czerwcu 1982 roku (Tversky i Kahnemana, 1983), w którym zapytani eksperci ocenili prawdopodobieństwo (i) całkowitego zawieszenia stosunków dyplomatycznych między USA a Związkiem Radzieckim jako niższe niż prawdopodobieństwo (ii) radzieckiej inwazji na Polskę oraz całkowitego zawieszenia stosunków dyplomatycznych między USA a Związkiem Radzieckim. Eksperci postrzegali zatem koniunkcję wydarzeń (zawieszenie stosunków plus inwazja) jako bardziej możliwą niż pojedyncze wydarzenie (zawieszenie stosunków).

Występowanie błędu koniunkcji w procesach wnioskowania u zawodowych uczestników rynku walutowego potwierdzili także Stephan i Kiell w swoim wcześniejszym eksperymencie³⁴ (Stephan i Kiell, 1998). Oto jedno z zadanych respondentom pytań:

Które zdarzenie jest bardziej prawdopodobne? (W nawiasach kwadratowych podano średnie oszacowane przez respondentów prawdopodobieństwo.)

(i) W wyniku zwiększenia tempa wzrostu inflacji amerykańska gospodarka wykazuje pierwsze oznaki przegrzania. Fed decyduje się na podwyżkę stóp procentowych. [35%]

(ii) Fed decyduje się na podwyżkę stóp procentowych. [27%]

Ponownie koniunkcja zdarzeń, jako bardziej reprezentatywna, oszacowana została na bardziej prawdopodobną. W nowszym eksperymencie Stephan i Kiell (2000) dostarczyli dalszych dowodów błędu koniunkcji popełnianego przez ekspertów rynku kapitałowego. Przykładowo respondenci oceniali szansę zajścia koniunkcji zdarzeń: „spadek Nikkei i spadek Dow Jones” wyżej niż pojedynczego zdarzenia: „spadek Nikkei”. Ponownie wiedza bądź też

³⁴ Badana grupa obejmowała 36 zawodowych dealerów walutowych – pracowników jednego z największych na świecie banków.

intuicja ekspercka okazała się niezgodna z racjonalnymi zasadami teorii prawdopodobieństwa.

2.2.3 Zakotwiczenie i dostosowanie

Studia empiryczne wykazały, że ludzie estymują wielkości na podstawie początkowo zasugerowanego poziomu – tzw. kotwicy, którą dostosowują w celu „wyłuskania” ostatecznej oceny (Slovic i Lichtenstein, 1971; Tversky i Kahneman, 1974). Wartość początkowa może zostać narzucona przez samo sformułowanie problemu czy też poprzez pobieżne obliczenie, dokonywane w celach orientacyjnych. Charakterystyczne, że w obu przypadkach dostosowanie początkowej wielkości jest zwykle niedostateczne. Oznacza to, iż różne kotwice pociągają za sobą odmienny stopień dostosowania. Stąd wartości oszacowane ciążą w kierunku poziomu wskazanego na początku. Zjawisko to przyjęło się określać jako zakotwiczenie.

Istnieje wiele doskonałych dowodów na to, że jednostki ulegają efektowi zakotwiczenia. Jednym z klasycznych przykładów jest eksperyment polegający na oszacowaniu wartości $8!$ w dwóch „wariantach”: (i) $1 \times 2 \times 3 \times 4 \times 5 \times 6 \times 7 \times 8$ oraz (ii) $8 \times 7 \times 6 \times 5 \times 4 \times 3 \times 2 \times 1$. Średnie szacunki respondentów to 512 w pierwszym przypadku oraz 2250 w drugim, przy prawidłowej odpowiedzi 40 320. Tversky i Kahneman (1974) interpretują te wyniki w kontekście „zakotwiczenia” na początkowych czynnikach mnożenia (np. 1,2,3 lub 8,7,6), a następnie niedostatecznego dostosowania, które jest tym znaczniejsze, im niższa była wartość początkowej kotwicy.

Późniejsze prace potwierdziły siłę efektu zakotwiczenia nie tylko wśród laików, lecz także w gronie ekspertów. Znany przykładem jest badanie przeprowadzone wśród grupy agentów nieruchomości przez Northcrafta i Neale (1987). Uczestnicy eksperymentu (agenci nieruchomości oraz studenci) poproszeni zostali o dokonanie wyceny pewnego domu na podstawie jego oględzin i po zapoznaniu się z pakietem informacyjnym o nieruchomości. Informatory rozdano w czterech wersjach, które różniły się ceną wywoławczą. Oszacowana przez uczestników rzeczywista wartość domu wskazywała na wpływ zasugerowanej ceny (kotwicy) na wycenę zarówno

profesjonalistów (agentów nieruchomości), jak i laików (studentów). Co istotne, studenci otwarcie przyznawali, że w swoich szacunkach kierowali się podaną ceną, czemu zaprzeczali agenci nieruchomości. Pozwala to konkludować, że eksperci są podatni na błąd kotwiczenia i być może nie są tego świadomi.

Mówiąc o działaniu opisywanego efektu wśród praktyków rynku kapitałowego, ponownie wypadnie powołać się na badania Stephana i Kiella (2000). Podzieleni na grupy inwestorzy mieli oszacować, czy za dwanaście miesięcy niemiecki indeks giełdowy DAX uplasuje się poniżej bądź powyżej pewnego zasugerowanego poziomu. W grupie pierwszej poziom ów określono na 4500 punktów (niska kotwica), w drugiej zaś na 6500 punktów (wysoka kotwica). Oczywiście określona w arbitralny sposób wartość początkowa nie powinna mieć jakiegokolwiek znaczenia diagnostycznego dla prognozy DAX. Wbrew temu oczekiwaniu, a zgodnie z hipotezą zakotwiczenia, prognozy respondentów zależały od nieistotnej wskazówki o poziomie indeksu. W pierwszej grupie (kotwica 4500) szacowany średni poziom DAX wyniósł 5765, w drugiej zaś (kotwica 6500) 5930. We wcześniejszych doświadczeniach podobne wyniki uzyskano dla cen złota, kursu funta brytyjskiego do marki niemieckiej oraz indeksu Dow Jones (Stephan i Kiell, 1998).

2.3 Rozszerzenia

Zniekształcenia procesu wnioskowania mogą być – i często są – powodowane nakładającym się działaniem więcej niż jednej heurystyki. Odpowiednie tego przykłady zostaną pokrótce omówione.

2.3.1 *Nadmierna pewność siebie*

Wiele obserwacji codziennych zachowań świadczy o przesadnej wierze jednostek we własne możliwości. Ów nadmierny optymizm może dotyczyć oceny własnego charakteru, tężyzny fizycznej, zdolności intelektualnych, i co najważniejsze – trafności własnych osądów. Nadmierna pewność siebie właściwa jest zarówno laikom, jak i osobom dysponującym wiedzą ekspercką.

Analizy dokonywane przez profesjonalistów polegają zwykle na wyrażeniu przekonania co do takich wielkości, jak na przykład poziom indeksu Dow Jones w danym okresie. Istota tych zadań polega na zaproponowaniu subiektywnego rozkładu prawdopodobieństwa. Obrazuje to następujący problem (Kahneman i Riepe, 1998):

Podaj swój najlepszy szacunek poziomemu, jaki Dow Jones osiągnie od dziś za miesiąc. Następnie (i) wybierz najwyższą wartość indeksu taką, że na poziomie ufności 99% Dow Jones ukształtuje się poniżej tej wartości; (ii) dodatkowo wybierz taką najniższą wartość, że Dow Jones przekroczy ją z prawdopodobieństwem 99%.

Analityk finansowy, określając górną i dolną wartość indeksu, ustala jednocześnie subiektywny przedział wartości z ufnością 98%. Oznacza to, że z prawdopodobieństwem 98% faktyczny poziom Dow Jones zmieści się w wyznaczonych granicach, zaś z prawdopodobieństwem 2% wykróczy poza oszacowany przedział. Innymi słowy, jeśli analityk właściwie kalibruje swoje prognozy, wówczas średnio 98 ze 100 analiz powinno być trafnych.

Niestety jedynie niewielki odsetek osób wykazuje taką sprawność. Literatura przedmiotu dostarcza licznych przykładów systematycznych błędów, polegających na zbyt wąskiej ocenie przedziałów dla szacowanych wartości (De Bondt, 1998; Kahneman i Riepe, 1998). Rzeczywiste realizacje wykraczają poza oszacowany przedział w 15%–20% przypadków, podczas gdy poprawna ocena dopuszcza maksymalnie 2% błędów. Zbyt wąskie ustalanie przedziałów ufności to dowód nadmiernej pewności siebie i przeceniania własnych możliwości. Te własności ludzkiego wnioskowania można łącznie określić mianem „błędnej kalibracji”.

Nadmierna pewność siebie wśród wykwalifikowanych uczestników rynku była jednym z problemów badawczych Stephana i Kiella (1998). Wyniki tu otrzymane dobitnie potwierdzają błędną kalibrację cechującą inwestorów. Trafność szacunków dokonywanych przez respondentów (przy zakładanym poziomie ufności 90%) okazała się bowiem niezwykle niska zarówno w pytaniach z wiedzy ogólnej, jak i w prognozowaniu wielkości na rynkach fi-

nansowych. Odsetek trafień kształtował się następująco: wiedza ogólna – 21%, prognozy kursów akcji – 17%, prognozy kursów walutowych – 29%.

Przecenianie siebie jest dodatnio skorelowane z czynnikami takimi, jak stopień trudności zagadnienia czy złożoności zjawiska (Fischhoff, Slovic, Lichtenstein, 1977). Griffin i Tversky (1992) twierdzą ponadto, że ograniczone możliwości przewidywania zdarzeń – w tym także prognozowania cen aktywów na rynkach finansowych – sprawiają, że eksperci są bardziej pewni siebie niż niedoświadczeni inwestorzy.

2.3.2 *Selektywna percepcja i błąd afirmacji*

Jednym z istotnych czynników zniekształcających procesy wnioskowania jest selektywna percepcja informacji. Jej genezy należy poszukiwać w ograniczonych zdolnościach przyswajania danych czy też w celowym unikaniu tych informacji, które mogłyby zaburzyć ukształtowane poglądy jednostki.

Selektywna percepcja rzutuje na interpretację wiadomości, skąd bierze się jej znaczenie dla uczestników rynków finansowych. Analitycy rynkowi, działając pod presją czasu, wychwytyują tylko te informacje, które pasują do ich scenariusza zdarzeń. Najlepiej, gdy nowo wyszukane wiadomości potwierdzają wcześniejsze prognozy (Goldberg i von Nitzsch, 2000, s. 61). Mechanizm ten, w połączeniu z konkretnymi oczekiwaniami jednostek, pozwala wytłumaczyć rozbieżne komentarze dotyczące tych samych wydarzeń na rynku. Jeśli analityk spodziewa się dalszego wzrostu kursów akcji, wówczas ogłoszenie danych o wzroście bezrobocia będzie odbierane pozytywnie w kontekście ograniczenia inflacji. I odwrotnie: w sytuacji spodziewanego spadku kursów, zwiększenie liczby bezrobotnych skojarzone zostanie ze spowolnieniem wzrostu gospodarczego i w konsekwencji z osłabieniem rynku.

Selektywna percepcja w naturalny sposób wiąże się zatem z pragnieniem potwierdzenia własnych sądów. Ów tzw. błąd afirmacji (*confirmation bias*) polega na poszukiwaniu informacji spójnych z przekonaniami przy jednoczesnym ignorowaniu faktów podważających słuszność wnioskowania. Dzieje się tak, mimo iż stwierdzenie, czy coś jest prawdą, wymaga rozważenia potencjalnych dowodów temu przeczących (Einhorn i Hogarth, 1978).

Inwestor decydujący się na zakup akcji danej spółki koncentruje się zazwyczaj na wyszukiwaniu argumentów przemawiających za tą decyzją, chociaż analiza kontrargumentów prowadziaby często do lepszych rezultatów. Konsekwencją może być optymizm motywujący do nadmiernej aktywności inwestycyjnej, której skutki omówiono w punkcie 1.1.4.

2.3.3 *Błąd post factum*

Ludzka skłonność do przeceniania swej przeszłej wiedzy – nazywana tu błędem *post factum* (*hindsight bias*) – została szeroko udokumentowana w literaturze psychologii³⁵. Po fakcie jednostki rzadko są zdolne zrekonstruować to, jak oceniały prawdopodobieństwo danego zdarzenia, zanim miało ono miejsce. W większości przypadków wstecznie przeszacowują dokładność swoich prognoz. Wystarczy zauważyć, jak często słychać wykrzyknienia „wiedziałam, że tak będzie” jako wyraz żalu, że faktyczny rozwój sytuacji nie sprzyjał podjętej uprzednio decyzji. Wiedza po fakcie opiera się na pamięci, która z kolei może być zawodna. Wspomnienia wydarzeń ulegają zniekształceniom poprzez subiektywne wrażenia, ograniczoną zdolność zapamiętywania oraz zakłócenia w procesie przypominania. W ten sposób pamiętanie polega nie tylko na przechowywaniu informacji, lecz także na konstruowaniu ich od nowa (Wärneryd, 2001, s. 134).

Dlatego też zdarzenia, jakich nie byli w stanie przewidzieć nawet eksperci czy *insiderzy*, *ex post* zdają się nieuniknione i oczywiste. Świat rynków finansowych dostarcza wielu takich przykładów: w godzinę po zamknięciu giełdy słyszy się często głosy analityków – guru tłumaczące, dlaczego rynek zachował się tak, a nie inaczej. Na podstawie tych komentarzy i logicznych argumentacji trudno nie dojść do błędnego przekonania, że rynki łatwo prognozować. Oczywiście, jeśli zachowania rynku mogły być zostać przewidziane, wówczas większość uczestników zachowałaby się była inaczej. Inne byłyby również implikacje tych zachowań. Mimo iż niemal każdy inwestor dobrze zna ów tok rozumowania, fascynacja interpretacjami przeszłości trwa.

³⁵ Wczesne prace z tego zakresu pochodzą od Fischhoffa (1975, 1980) oraz Wooda (1978).

Większość badaczy błędu *post factum* upatruje w nim źródło potencjalnych zagrożeń³⁶. Po pierwsze, skłonność do przeceniania swej przeszłej wiedzy stymuluje do nadmiernej wiary w możliwość przewidywania zdarzeń. To z kolei powstrzymuje proces uczenia się na błędach. Po drugie, efekt ten zniekształca percepcję przeszłych decyzji, które – mimo iż obiektywnie zasadne w momencie ich podejmowania – po fakcie uznawane są za nierozsądne. W tym kontekście błąd *post factum* może mieć nietrywialne znaczenie dla doradców finansowych. Jeśli bowiem spadek kursu akcji wydaje się nieunikniony, po tym jak już nastąpił, inwestorzy – laicy mogą zarzucać specjalistom brak umiejętności prognozowania kursów i rekomendowania odpowiednich działań inwestycyjnych.

* * *

Opierające się na intuicji heurystyki to efektywny sposób radzenia sobie ze złożonością informacji pod presją czasu. Ceną, jaką jednostki płacą za tę zdolność sprawnego wnioskowania, są systematyczne pomyłki. Określenie „systematyczny” oznacza, iż błędy te stanowią niejako domyślny element działania ludzi, przez co też determinują funkcjonowanie całego rynku.

Krytycy podejścia heurystycznego oskarżają je o brak związku z rzeczywistym środowiskiem, w jakim jednostki funkcjonują. Zarzuty te są skutecznie odpierane w nowych eksperymentach prowadzonych poza salą laboratoryjną, czego przykład stanowią choćby licznie tu cytowane badania Stephana i Kiella (1998, 2000) czy też De Bondta (1998). Przytoczone dowody nakazują uznać etap „edycji” informacji i tworzenia osądów za proces kreatywny, wpływający na podejmowane następnie decyzje. Stąd podejście markowit-zowskie do selekcji portfela, ignorujące sposób, w jaki jednostki kształtują przekonania, zdaje się dalece niewystarczającym wytłumaczeniem wyborów inwestycyjnych.

³⁶ Jednakże grupa naukowców skupionych wokół Gigerenzera odżegnuje się od uznawania tej skłonności za błąd przetwarzania informacji. Badacze ci twierdzą, że błąd *post factum* jest „produktem ubocznym” dwóch procesów: (i) aktualizowania wiedzy po otrzymaniu nowej informacji, (ii) wyciągania wniosków na podstawie uaktualnionej wiedzy w sposób szybki i oszczędny. Ów „produkt uboczny” to niska cena, jaką jednostka płaci za sprawną pamięć i aktualną wiedzę (Hoffrage i Hertwig, 1999).

3

Prospect theory: podejmowanie decyzji w warunkach niepewności

Decyzje są końcowym etapem złożonego procesu postrzegania, przetwarzania i oceny nowych informacji. Omówiony w poprzednim rozdziale początkowy etap „edycji” stanowi preludium do fazy oceny, w której formują się preferencje i zapadają decyzje. Wysilek przeznaczony na przetwarzanie informacji (np. śledzenie poziomu indeksu giełdowego) prowadzi niejako automatycznie do tworzenia osądów („rynek znajduje się w fazie byka”) i decyzji („kupuję akcje”).

W celu analizy zachowań jednostek działających na rynkach finansowych szczególnie istotne jest zrozumienie, jak decydują one w warunkach niepewności. Niepewność bowiem jest jedną z podstawowych przyczyn złożoności rynku. Często to ona właśnie katalizuje zachowania nieracjonalne.

Inwestor, podobnie jak uczestnik loterii, stawia na niepewny wynik w przyszłości. W tym sensie inwestycja w aktywa rynku kapitałowego przypomina grę losową. Należy jednak zwrócić uwagę na pewien niuans: rozróżnienie między prawdopodobieństwem obiektywnym (ryzykiem) a subiektywnym (niepewnością)³⁷. Uczestnik gry losowej, w przeciwieństwie do inwestora, jest w stanie określić dokładne prawdopodobieństwo wystąpienia danego wyniku. Powtarzając grę wielokrotnie, może spodziewać się wyniku równego wartości oczekiwanej rozkładu prawdopodobieństwa stojącego w tle jego gry. Stawia go to w sytuacji uprzywilejowanej w stosunku do inwestora, który poznaje prawdopodobieństwo poprzez aproksymację (np. na podstawie danych historycznych).

Rozdział ten przedstawia deskryptywny model podejmowania decyzji w warunkach niepewności – *prospect theory* Kahnemana i Tversky’ego. Podejście to okazuje się szczególnie adekwatne w interpretacji zachowań wykraczających poza ramy tradycyjnej teorii oczekiwanej użyteczności, a mi-

mo to poruszających rynki finansowe. *Prospect theory* koncentruje się na preferencjach jednostek. Chociaż transformacja na poziom zagregowanego rynku nie została jeszcze zaproponowana, model Kahnemana i Tversky'ego pozwala postawić hipotezy tłumaczące anomalie rynkowe.

Pierwsza część niniejszego rozdziału przedstawia normatywną teorię podejmowania decyzji, która w dużej mierze zdeterminowała kształt nowoczesnej nauki finansów. W drugiej części omówiono podstawowe koncepcje składające się na *prospect theory* – funkcję wartości oraz funkcję wag prawdopodobieństw – wraz ze wskazaniem ich wartości poznawczej.

3.1 Normatywna teoria podejmowania decyzji w warunkach ryzyka

Powszechnie akceptowanym modelem podejmowania decyzji w warunkach ryzyka jest teoria oczekiwanej użyteczności von Neumanna i Morgensterna. Jako podejście normatywne stanowi ona swoisty „przepis” na to, jak jednostki powinny decydować, by ich wybory można było uznać za racjonalne.

3.1.1 Wartość oczekiwana a oczekiwana użyteczność

Otóż zgodnie z teorią oczekiwanej użyteczności w pełni racjonalny uczestnik rynku (*homo oeconomicus*) zawsze wybierze alternatywę, która przyniesie mu największą korzyść. Jak podpowiada intuicja, wybór powinien być determinowany wartością oczekiwaną³⁷ definiowaną jako:

³⁷ Dokładne rozróżnienie między prawdopodobieństwem obiektywnym a subiektywnym zob. Savage (1954) oraz ryzykiem a niepewnością zob. Knight (1921).

³⁸ Rozważane są tu zdarzenia, których realizacja wiąże się z wypłatą pieniężną, co odpowiada rzeczywistości, w jakiej porusza się uczestnik rynku.

$$EV = \sum_i p_i \cdot x_i, \quad (3.1)$$

gdzie:

EV – wartość oczekiwana,

p_i – prawdopodobieństwo wystąpienia i -tego zdarzenia (wypłaty x_i),

x_i – wartość wypłaty w i -tym zdarzeniu.

Za hipotezą tą kryje się założenie, że użyteczność pieniądza równa jest jego monetarnej wartości.

Dowód obalający tę hipotezę przedstawił Daniel Bernoulli w swoim wy tłumaczeniu „paradoksu petersburskiego” już w 1738 roku. Gra petersburska polega na rzucie monetą. Reguły są następujące: gra trwa do momentu wyrzucenia reszki; jeśli reszka wypada w pierwszym rzucie, uczestnik otrzymuje 2 j.p. i gra dobiega końca; jeśli reszka pojawi się dopiero w n -tym rzucie, gra zostaje zakończona, zaś gracz zabiera ze sobą wygraną 2^n j.p. W ten sposób gracz zawsze odchodzi z nagrodą, z tym jednak zastrzeżeniem, że jej wysokość jest nieznana. Łatwo można policzyć tu wartość oczekiwaną EV :

$$EV = 1/2 \cdot 2 + 1/4 \cdot 4 + 1/8 \cdot 8 + 1/16 \cdot 16 + \dots = 1 + 1 + 1 + 1 + \dots = \infty$$

Bernoulli pisze: „Mimo iż standardowe obliczenia pokazują, że wartość oczekiwana (...) jest nieskończenie duża, należy uznać, że każdy w miarę rozsądny człowiek z wielkim zadowoleniem sprzedałby tę szansę za 20 dukatów” (Bernoulli, 1738/1954; za: Baron, 2000, s. 239).

W wytłumaczeniu tego zjawiska Bernoulli sugerował, że użyteczność³⁹ bogactwa nie jest równa jego wartości pieniężnej, lecz co najwyżej proporcjonalna do jej logarytmu⁴⁰. Stwierdzenie to jest równoznaczne z malejącą krańcową użytecznością bogactwa.

³⁹ Użyteczność nie powinna być tu rozumiana jako „przyjemność”, „satysfakcja” czy „szczęście”. Koncepcja użyteczności odnosi się raczej do złożonego zestawu celów stawianych sobie przez jednostkę; w tym sensie reprezentuje ona wszystko to, co jednostka pragnie osiągnąć (Baron, 2000, s. 224).

⁴⁰ Bernoulli uznawał za wysoce prawdopodobne, że każdy wzrost bogactwa, nieważne jak nieznaczny, przyniesie wzrost użyteczności, który jest odwrotnie proporcjonalny do ilości dóbr już posiadanych. Stwierdzenie to implikuje funkcję logarytmiczną, jednakże Bernoulli nie podał uzasadnienia dla takiego kształtu krzywej użyteczności. Teza Bernoulli’ego oznacza, że przyrost użyteczności odczuwany przy wzroście całkowitego bogactwa z 1000 j.p. do 10 000 j.p. jest taki sam jak przy wzroście z 10 000 j.p. do 100 000 j.p.

Malejąca krańcowa użyteczność wyjaśnia, dlaczego ludzie niechętnie przyjmują tzw. sprawiedliwe zakłady, w których cena uczestnictwa równa jest wartości oczekiwanej gry. Niewielu znalazłoby się chętnych do zapłaty 10 j.p. za udział w grze, w której istnieją równe szanse wygrania 16 j.p. lub 4 j.p. Odrzucenie takiego zakładu oznacza, że dana osoba wykazuje awersję do ryzyka: przedkłada pewną wypłatę równą wartości oczekiwanej gry nad samą grę. W literaturze ekonomii awersję do ryzyka uznaje się za wyraz racjonalności. Przyjmuje się również, iż jest to zachowanie charakteryzujące większość uczestników rynków finansowych. Skłonność do ryzyka postrzegana jest natomiast jako swoista patologia.

Wkład Bernoulli'ego do teorii podejmowania decyzji w warunkach ryzyka można by streścić następująco: racjonalność nakazuje maksymalizować oczekiwaną użyteczność, nie zaś wartość oczekiwaną. Wybory w warunkach ryzyka determinowane są bowiem jednocześnie przez malejącą krańcową użyteczność bogactwa oraz indywidualny stosunek do ryzyka. Maksymalizacja wartości oczekiwanej mogłaby stanowić właściwe kryterium wyboru jedynie dla osoby charakteryzującej się neutralnością wobec ryzyka (czyli postawą niemalże niespotykaną w rzeczywistości).

Na gruncie odkryć Bernoulli'ego von Neumann i Morgenstern stworzyli najbardziej znany i akceptowany model racjonalnych preferencji i podejmowania decyzji w warunkach ryzyka. Teoria oczekiwanej użyteczności przedstawia relację preferencji za pomocą funkcji, która przypisuje każdemu poziomowi całkowitego bogactwa odpowiednią użyteczność. Jednostka postawiona przed wyborem w warunkach ryzyka określa użyteczność danego wyniku, ważąc ją prawdopodobieństwem jego wystąpienia; wyznacza zatem jego oczekiwaną użyteczność EU :

$$EU = \sum_i p_i \cdot u(x_i), \quad (3.2)$$

gdzie:

EU – oczekiwana użyteczność,

p_i – prawdopodobieństwo wystąpienia i-tego wyniku x_i ,

$u(x_i)$ – użyteczność bogactwa osiąganego przy i-tym wyniku.

Wybór pada na tę opcję, która wiąże się z najwyższą użytecznością. Niezwykle istotny jest przy tym fakt, że racjonalna jednostka rozważa poszczególne alternatywy zawsze w kontekście całego swojego majątku – w kategoriach absolutnych⁴¹.

O funkcji użyteczności można myśleć jako o danej – permanentnie przypisanej konkretnej osobie. Oznacza to stabilność jej preferencji w czasie. Przez całe życie człowiek sięga po tę samą „miarę” swojej użyteczności i na jej podstawie dokonuje wyborów. Kolejne decyzje wpływają na poziom jego całkowitego bogactwa, co wywołuje przesunięcia wzdłuż krzywej użyteczności, ale nie zmienia jej kształtu.

3.1.2 Aksjomatyzacja preferencji

Jak wcześniej wspomniano, teoria oczekiwanej użyteczności jest normatywnym modelem podejmowania decyzji w warunkach ryzyka. Model normatywny należy rozumieć jako pewien standard, który definiuje najlepszy możliwy sposób myślenia, pozwalający jednostce osiągać swoje cele najefektywniej. Postępowanie zgodne z takim modelem jest zatem równoznaczne z zachowaniem w pełni racjonalnym⁴².

⁴¹ Zaproponowany opis oczekiwanej użyteczności jest ogromnym uproszczeniem teorii von Neumanna i Morgensterna. Pominęto tu ważną procedurę wyprowadzenia funkcji użyteczności (tzw. twierdzenie reprezentacji, *representation theorem*), obejmującą zagadnienia aksjomatyzacji preferencji, relacji preferencji oraz ciągłej funkcji użyteczności możliwej do przedstawienia w przestrzeni euklidesowej.

⁴² Model normatywny należy traktować jako ideał myślenia, który prowadzi do optymalnych rezultatów. Dodatkowo definiuje się także model preskryptywny jako konstrukcję myślową, której celem jest zbliżenie faktycznego ludzkiego sposobu myślenia do modelu normatywnego. Najlepszym modelem preskryptywnym jest zatem staranie się myśleć w kategoriach normatywnych (Baron, 2000, s. 32).

W tym znaczeniu teoria oczekiwanej użyteczności jest najlepszym sposobem podejmowania decyzji. O jej normatywnym charakterze decydują leżące u podstaw aksjomaty przypisane racjonalnemu myśleniu i podejmowaniu decyzji. Aksjomatyzacja, wspólne dzieło von Neumanna i Morgensterna (1944), zakłada cztery kluczowe⁴³ właściwości preferencji: eliminację (*cancellation*), przechodniość (*transitivity*), dominację (*dominance*) oraz niezmienność (*invariance*). Zostaną one pokrótce omówione⁴⁴.

3.1.2.1 Eliminacja

Własność eliminacji bywa także nazywana „zasadą rzeczy pewnej” (*sure-thing principle*). Ustanawia ona, że jeśli istnieje taki stan natury, który prowadzi zawsze do jednego i tego samego wyniku niezależnie od dokonywanego wyboru, wówczas wybór nie powinien zależeć od tego stanu natury. Inaczej mówiąc, stan natury, który daje zawsze ten sam wynik, niezależnie od decyzji, nie powinien mieć wpływu na preferencje.

Załóżmy, że osoba ma do wyboru dwie loterie. Pierwsza możliwość to otrzymanie A , jeśli jutro będzie padał deszcz (oraz wypłata zerowa, jeśli jutro nie będzie padało); druga to wygranie B , jutro będzie padał deszcz (i brak wypłaty w przeciwnym razie). Jeśli gracz wyżej ceni A niż B ($A \succ B \Rightarrow u(A) > u(B)$), wybierze pierwszą loterię, obie dają bowiem ten sam wynik (brak wygranej), jeśli jutro nie pada deszcz. Własność eliminacji argumentuje się tym, że tylko jeden stan zostanie w rzeczywistości zrealizowany (albo pada deszcz, albo nie), co nakazuje rozważać wyniki alternatyw osobno dla każdego z przyszłych stanów natury. Racjonalny wybór między opcjami zależy jedynie od stanów, które prowadzą do różnych wyników.

⁴³ Von Neumann i Morgenstern wskazują na dwa dodatkowe założenia o charakterze bardziej „technicznym”: porównywalność (*comparability*) i ciągłość (*continuity*). W pracy tej zrezygnowano z ich omówienia ze względu na małe znaczenie dla dalszego wywodu.

⁴⁴ Aksjomatyzacja preferencji w teorii oczekiwanej użyteczności przedstawiona została za Kahnemanem i Tversky’em (1986).

3.1.2.2 Przechodność

Przechodność powinna charakteryzować preferencje zarówno w świecie, gdzie ryzyko istnieje, jak i w takim, gdzie go brak. Przechodność jest warunkiem koniecznym i wystarczającym, by preferencje (np. $A \succ B$) można było zaprezentować za pomocą uporządkowanej skali użyteczności (w przestrzeni euklidesowej), tak by $u(A) > u(B)$. Wymóg przechodności jest spełniony, gdy możliwe jest przypisanie użyteczności do każdej opcji niezależnie od użyteczności innych dostępnych opcji.

3.1.2.3 Dominacja

Dominacja wydaje się najbardziej intuicyjną zasadą racjonalnych preferencji. Jeśli dana alternatywa jest lepsza niż inna przynajmniej w jednym stanie natury i przynajmniej tak samo dobra we wszystkich innych stanach natury, wówczas opcja ta jest dominująca i powinna zostać wybrana. Nieco szersze pojmowanie tej własności pozwala rozważać ją w znaczeniu dominacji stochastycznej – preferencji rozkładu prawdopodobieństwa związanego z daną opcją. Opcja A jest lepsza od opcji B ($A \succ B$), gdy dystrybuenta jej rozkładu prawdopodobieństwa przesunięta jest w prawo w stosunku do dystrybuenty rozkładu charakteryzującego opcję B .

3.1.2.4 Niezmiennność

Mimo iż niezmiennność preferencji nie jest aksjomatem opisanym *explicite* przez von Neumanna i Morgensterna, stanowi ona podstawowy determinant normatywnego charakteru teorii oczekiwanej użyteczności. Zapewnia bowiem konsekwencję i stabilność preferencji. Fakt ten czyni ją najbardziej ogólną zasadą racjonalnego wyboru.

Zasada niezmienności stanowi, że różne sposoby prezentacji tego samego problemu decyzyjnego nie wpływają na preferencje. Osoba podejmująca decyzję musi być w stanie przeniknąć otoczkę słowną, w którą „zapakowany” jest dany problem. Własność niezmienności odpowiada intuicyjnemu rozumieniu racjonalności: wybory powinny być niewrażliwe na zmiany samej formy przedstawienia alternatyw.

3.2 Anomalie w kształtowaniu się preferencji

Teoria oczekiwanej użyteczności miała być w zamyśle jej twórców usankcjonowaniem zasad racjonalnego wyboru w warunkach ryzyka – niczym ponad model normatywny. Paradoksalnie jednak zaczęło się uważać, iż jest ona także właściwym opisem rzeczywistych zachowań – modelem deskryptywnym, zaś odstępstwa od niego, choć możliwe, mają charakter stochastyczny. Pozwalało to przyjąć, że znoszą się one poprzez agregację.

Badania preferencji wskazują jednak na systematyczne odchylenia zachowań jednostek od zaksjomatyzowanego ideału. Owa systematyczność sprawia, że zabieg agregacji przestaje działać – model traci wartość deskryptywną.

Kahneman i Tversky (1986) systematyzują omówione wyżej aksjomaty, uznając niezmienność i dominację za kluczowe wyznaczniki racjonalnych preferencji, przechodniość za zasadę wątpliwą, zaś eliminację za pozbawioną mocy⁴⁵. Rzeczywiste wybory niezgodne z aksjomatami przechodniości i eliminacji nie stanowią bezpośredniego zagrożenia dla deskryptywnego charakteru teorii oczekiwanej użyteczności – w przeciwieństwie do zachowań gwałcących aksjomaty niezmienności i dominacji, które nie mogą zostać opisane przez żaden model racjonalnego myślenia. Jedynym znanym podejściem tłumaczącym taki rodzaj aberracji jest *prospect theory*.

W kontekście genezy teorii Kahnemana i Tversky'ego warto zasygnalizować obszary rozbieżności między ujęciem modelowym a rzeczywistymi zachowaniami. Zaobserwowane anomalie stanowiły bowiem motywację do lepszego zrozumienia nieoptymalnych wyborów dokonywanych przez ludzi.

⁴⁵ Najbardziej znane i najwcześniejsze przykłady rzeczywistych zachowań odbiegających od zasady eliminacji pochodzą od Allais (1953) i Ellsberga (1961). Stały się one argumentem za jej odrzuceniem na rzecz zasad bardziej ogólnych.

3.2.1 Zaniedbywanie prawdopodobieństwa

Zaniedbywanie prawdopodobieństwa jest dość banalnym przykładem zniekształceń preferencji w warunkach niepewności. Wiąże się ono najczęściej z nieświadomością osoby podejmującej decyzję lub z niskim stopniem zrozumienia rzeczywistości. Warto jednak przypomnieć, że rzeczywistość rozgrywa się – inaczej niż loteria – w warunkach niepewności, gdzie prawdopodobieństwo może być jedynie szacowane. Trudność ta często sprawia, że wybory obciążone niepewnością bywają dokonywane z jej nieświadomym pominięciem.

3.2.2 Odwrócenie preferencji

Nawet jeśli prawdopodobieństwo jest uwzględniane, istnieje możliwość przypisania mu subiektywnej wartości w zależności od rodzaju problemu decyzyjnego. Przykładem takiego wzorca zachowań jest następujący eksperyment (Lichtenstein i Slovic, 1973):

Dane są dwie loterie: P oraz \$. Loteria P charakteryzuje się wysokim prawdopodobieństwem wygranej, lecz niską możliwą wygraną. Odwrotna zależność występuje w loterii \$.

Loteria P: wygrana 2 j.p. z prawdopodobieństwem 29/36

Loteria \$: wygrana 9 j.p. z prawdopodobieństwem 7/36

Większość respondentów wybiera loterię P.

W kolejnym kroku uczestnicy doświadczenia mają zdecydować, jaką wartość prezentuje dla nich każda z loterii. Co ciekawe, większość osób wyżej wycenia loterię poprzednio odrzuconą – \$. Jest to oczywisty przykład odwracania się preferencji (loteria odrzucona ma dla respondentów wyższą wartość). Podobne rezultaty powielane były także w eksperymentach z wypłatami pieniężnymi.

Jak można wytłumaczyć fakt, że jednostki przeceniają loterię \$? Badacze twierdzą, że w zależności od rodzaju zadania, respondenci koncentrują się bardziej na ryzyku (zadania polegające na wyborze loterii) albo na bezwzględnej wysokości wypłaty (zadania polegające na wycenie). W świetle

omówionych aksjomatów takie odwrócenie preferencji podważa zasadność zastosowania teorii oczekiwanej użyteczności do opisu rzeczywistych decyzji.

3.2.3 Paradoks Allais

Deskryptywna moc teorii oczekiwanej użyteczności została podważona już w 1953 roku przez Maurice'a Allais⁴⁶. Zaproponował on następujący hipotetyczny problem decyzyjny (tabela 3.1):

Tabela 3.1: Paradoks Allais

<i>Loteria X</i>	<i>Kwota (j.p.)</i>	<i>Prawd*.</i>
<i>Opcja 1</i>	1000	1,00
<i>Opcja 2</i>	1000	0,89
	5000	0,10
	0	0,01

<i>Loteria Y</i>	<i>Kwota (j.p.)</i>	<i>Prawd.</i>
<i>Opcja 3</i>	1000	0,11
	0	0,89
<i>Opcja 4</i>	5000	0,10
	0	0,90

* Prawdopodobieństwo

Źródło: Przykład zaadaptowany z (Baron, 2000, s. 249).

Allais zaobserwował, że respondenci decydują się na opcję 1 w loterii X oraz opcję 4 w loterii Y. W sytuacji X możliwość otrzymania 5000 j.p. nie jest dostatecznie kusząca w porównaniu z pewną wygraną 1000 j.p. Wybór opcji 2 pociągnąłby za sobą ryzyko odejścia z pustymi rękoma. Inaczej rzecz się ma w loterii Y. Tu różnica prawdopodobieństwa wygranej między dwoma opcjami jest niewielka, co skłania uczestników eksperymentu do „spróbowania szczęścia” i wyboru opcji 4 z wyższą potencjalną wypłatą.

⁴⁶ Allais przeprowadził także wiele innych eksperymentów wskazujących na preferencje niezgodne z teorią oczekiwanej użyteczności. Omawiany eksperyment jest dowodem pogwałcenia aksjomatu eliminacji.

Spróbujmy teraz przeanalizować te wybory w kategoriach teorii oczekiwanej użyteczności (tabela 3.2).

Tabela 3.2: Oczekiwanie użyteczności w eksperymencie Allais

<u>Loteria X</u>			
Opcja 1	0,01 u(1000)	+ 0,10 u(1000)	+ 0,89 u(1000)
Opcja 2	0,01 u(0)	+ 0,10 u(5000)	+ 0,89 u(1000)
<u>Loteria Y</u>			
Opcja 3	0,01 u(1000)	+ 0,10 u(1000)	+ 0,89 u(0)
Opcja 4	0,01 u(0)	+ 0,10 u(5000)	+ 0,89 u(0)

Źródło: Przykład zaadaptowany z (Baron, 2000, s. 250).

Zgodnie z aksjomatem eliminacji wybory powinny być niezależne od wyników, które zostaną zrealizowane bez względu na podjętą decyzję. Dlatego od osoby racjonalnej należałoby oczekiwać wyboru parami opcji 1 i 3 lub opcji 2 i 4. Zasada eliminacji pozwala na pominięcie identycznych wyników w każdej z loterii. Rzeczywiste wybory wskazują na następujące relacje preferencji:

Opcja 1 > Opcja 2

$$0,01 u(1000) + 0,10 u(1000) > 0,01 u(0) + 0,10 u(5000)$$

Opcja 3 < Opcja 4

$$0,01 u(1000) + 0,10 u(1000) < 0,01 u(0) + 0,10 u(5000)$$

Ze względu na ciągłość i monotoniczność funkcji użyteczności jedna z powyższych nierówności musi być fałszywa. Wybór opcji 1 i 4 jest pogwałceniem zasady eliminacji. Teoria oczekiwanej użyteczności nie jest w stanie wytłumaczyć preferencji tego rodzaju.

3.2.4 Problem Samuelsona

Rozważmy loterię polegającą na rzucie monetą: wygrana oznacza wypłatę 200 dolarów, przegrana – stratę 100 dolarów. Przed takim to wyborem postawił Paul Samuelson swojego znajomego – ekonomistę. Znajomy odrzucił zakład, zaznaczając jednocześnie, że przyjąłby 100 takich loterii. Uzasadnienie: „Nie założę się, ponieważ strata 100 dolarów ciążyłaby mi bardziej niż cieszyłaby mnie wygrana 200” (Thaler, 1999). Skłoniło to Samuelsona (1963) do przeprowadzania dowodu nieracjonalności takiego zachowania.

W sensie normatywnym jednostka powinna maksymalizować oczekiwaną użyteczność wraz z każdym podejmowanym wyborem. Konsekwentne przestrzeganie zasad wynikających z teorii oczekiwanej użyteczności prowadzi bowiem do najlepszych wyników w perspektywie całego życia. Na tej podstawie Samuelson twierdzi, że jednostka skłonna uczestniczyć w danej loterii stukrotnie powinna zdecydować się także na jednorazową próbę.

Taki sposób myślenia pomija jednak fakt, że rzeczywiste wybory dokonywane są w warunkach niepewności (bez znajomości rozkładu prawdopodobieństwa). Jednostka podejmując decyzję nigdy nie ma pewności, czy dana okazja powtórzy się w trakcie jej życia. Odpowiedź znajomego Samuelsona pozwala na jeszcze jedną niezwykle istotną konkluzję: ludzi znamionuje awersja do strat – cecha ignorowana przez podejście racjonalne.

* * *

Powyższa lista zachowań odbiegających od paradygmatu *homo oeconomicus* nie jest w żadnym razie wyczerpująca. Ma ona jednak pokazać potrzebę zbudowania modelu tłumaczącego nieoptymalne działania ludzi. Podane przykłady, mimo iż przedstawione w konwencji anegdot, nabierają szczególnego znaczenia poprzez swój systematyczny charakter. Można o nich myśleć w kategoriach niewykorzystanego przez ludzi potencjału. Co więcej, zachowania te właściwe są także profesjonalistom (zawodowym inwestorom, managerom funduszy, market-makerom, dealerom walutowym)⁴⁷, których codzienne działania determinowane są niepewnością. Jedynie zrozumienie przyczyn leżących u podstaw nieracjonalnych decyzji pozwala zrewidować nieoptymalne zachowania. Stąd też zagadnienie podejmowania decyzji w warunkach niepewności stanowi jeden z najważniejszych problemów poruszanych w niniejszej pracy.

⁴⁷ Istnieją liczne badania dotyczące nieracjonalnych zachowań profesjonalnych uczestników rynku: dealerów walutowych (Stephan i Kiell, 1998), managerów funduszy (Stephan i Kiell, 2000), prywatnych inwestorów (Odean, 1998a, 1999), zawodowych market-makerów (Coval i Shumway, 2001), doradców inwestycyjnych (Weber i Camerer, 1998). Są one cytowane w wielu miejscach tego opracowania.

3.3 Deskryptywna teoria podejmowania decyzji w warunkach niepewności

Prospect theory jako deskryptywna teoria podejmowania decyzji w warunkach niepewności⁴⁸ została zaproponowana przez Kahnemana i Tversky'ego w 1979 roku. Celem przyświecającym jej twórcom było stworzenie modelu opisującego faktyczne wybory dokonywane przez ludzi. W tym sensie *prospect theory* miała wypełnić lukę między podejściem normatywnym a rzeczywistością.

Błędem jest więc postrzeganie *prospect theory* jako teorii konkurencyjnej wobec oczekiwanej użyteczności. Jest ona raczej rozszerzeniem i modyfikacją tradycyjnego myślenia o wyborach w warunkach niepewności. Jednocześnie teorię tę należy uznać za niezwykle ważną dla przyszłego rozwoju ekonomii i nauki finansów, gdyż jako jedyna tłumaczy w sposób kompleksowy systematyczne zachowania gwałcące zasady racjonalnego myślenia i działania.

Dwie nowatorskie koncepcje leżące u podstaw *prospect theory* to:

- (1) S-kształtna funkcja wartości $v(x)$, która modyfikuje tradycyjną funkcję użyteczności w następujących aspektach: zyski i straty, nie zaś absolutny poziom majątku, są nośnikami subiektywnej wartości (użyteczności); funkcja wartości jest wklęsła w dziedzinie zysków, wypukła w dziedzinie strat oraz bardziej stroma dla strat niż dla zysków,
- (2) funkcja wag prawdopodobieństwa $\pi(p)$ jako nieliniowa transformacja prawdopodobieństw, która przeszacowuje niskie prawdopodobieństwa, zaś zaniża oceny prawdopodobieństw średnich i wysokich.

Koncepcje te wraz z ich implikacjami zostaną dokładnie omówione w dalszej części rozdziału.

⁴⁸ Określenie „teoria podejmowania decyzji w warunkach niepewności” stosowane jest w literaturze źródłowej, mimo iż pierwotna wersja *prospect theory* opiera się na loteriach o dokładnie określonym prawdopodobieństwie.

3.3.1 Wartość oczekiwań

W pierwotnej wersji modelu Kahneman i Tversky operują prostymi regularnymi loteriami $(x,p;y,q)$ – zwanymi oczekiwaniami (*prospects*) – z dwiema niezerowymi realizacjami występującymi z zadaniem prawdopodobieństwem⁴⁹. Loterię postaci $(x,p;y,q)$ należy czytać jako: otrzymanie wypłaty x z prawdopodobieństwem p , otrzymanie y z prawdopodobieństwem q lub otrzymanie wyniku zerowego z prawdopodobieństwem $(1-p-q)$, gdzie $p+q \leq 1$.

Dana loteria jest ściśle dodatnia, jeśli $x, y > 0$ i $p+q=1$; ściśle ujemna, gdy $x, y < 0$ i $p+q=1$ oraz regularna w pozostałych przypadkach, czyli jeśli $p+q < 1$ lub $x \leq 0 \leq y$ albo $x \geq 0 \geq y$.

Jeśli $(x,p;y,q)$ jest regularną loterią, to osoba podejmująca decyzję przypisze jej wartość:

$$V(x, p; y, q) = \pi(p) \cdot v(x) + \pi(q) \cdot v(y), \quad (3.3)$$

gdzie:

$\pi(\cdot)$ – waga decyzyjna,

$v(\cdot)$ – postrzegana wartość danej realizacji x ,

oraz $v(0) = 0$, $\pi(0) = 0$, $\pi(1) = 1$ ⁵⁰.

3.3.2 Funkcja wartości

Zaproponowana w *prospect theory* funkcja wartości $v(x)$ (*value function*) przedstawia sposób, w jaki jednostki doświadczają zdarzeń i oceniają ich wyniki (zyski lub straty). Można o niej myśleć jako o reprezentacji elementów składających się na subiektywną wartość (przyjemność/użyteczność/szczęście) postrzeganą przez jednostki.

O nowatorstwie funkcji wartości stanowią jej trzy podstawowe własności:

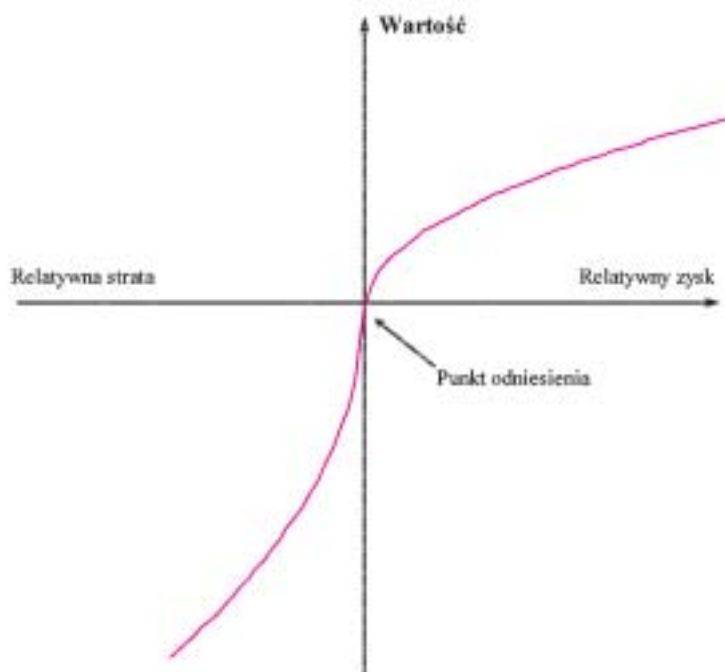
⁴⁹ Późniejsze prace autorów przyniosły rozszerzenia na wiele realizacji. Por. (Kahneman i Tversky, 1992).

⁵⁰ Należy zauważyć, że V jest zdefiniowane na przestrzeni loterii L , zaś v jest zdefiniowane na realizacjach loterii. Obie reprezentacje są równoważne dla loterii pewnych, gdzie $V(x,1) = V(x) = v(x)$.

- (1) Relatywna ocena – funkcja wartości jest zdefiniowana w dziedzinie strat oraz w dziedzinie zysków, które są względne wobec punktu odniesienia (*reference point*).
- (2) Malejąca wrażliwość – kształt funkcji (wklęsłość dla zysków, wypukłość dla strat) reprezentuje malejącą wrażliwość na dodatkowe zyski i straty, czyli krańcowa subiektywna wartość zysków i strat maleje wraz z ich wzrostem.
- (3) Awersja do strat – bardziej strome nachylenie funkcji w dziedzinie strat oznacza, że straty odczuwane są silniej niż zyski o tych samych wartościach bezwzględnych.

Własności te reprezentuje kształt funkcji zbliżony do litery S (rysunek 3.1).

Rysunek 3.1: Hipotetyczna funkcja wartości $v(x)$



Źródło: (Kahneman i Tversky, 1979).

Kahneman i Tversky (1992) proponują następującą postać funkcyjną:

$$v(x) = \begin{cases} x^\alpha & \text{dla } x \geq 0 \\ -\lambda(-x)^\beta & \text{dla } x < 0 \end{cases} \quad (3.4)$$

gdzie:

$v(x)$ – postrzegana wartość,

x – relatywny zysk/relatywna strata,

λ – współczynnik awersji do strat ($\lambda > 1$),

α, β – parametry determinujące malejącą wrażliwość ($0 < \alpha, \beta \leq 1$).

Na podstawie eksperymentów laboratoryjnych Kahneman i Tversky wyznaczyli następujące poziomy parametrów funkcji wartości: $\lambda = 2,25$, $\alpha = \beta = 0,88$.

3.3.2.1 Relatywna ocena

Punktem centralnym *prospect theory* jest stwierdzenie, że nośnikami postrzeganej i odczuwanej wartości są zmiany w poziomie bogactwa (zyski bądź straty), nie zaś, jak zakłada teoria oczekiwanej użyteczności, absolutny poziom bogactwa. Przyjęcie takiego założenia zgodne jest z pryncypiami postrzegania i oceny zjawisk przez człowieka.

„Nasz aparat percepcji nastawiony jest na ocenę zmian lub różnic, nie zaś wielkości absolutnych. Reakcje na czynniki takie jak jasność, głośność lub temperatura zależą od przeszłego i bieżącego ich poziomu, który definiuje nasz punkt odniesienia⁵¹. Świeżo napływające bodźce postrzegane są w stosunku do tego punktu. W ten sposób jednostka może postrzegać temperaturę danego przedmiotu jako wysoką bądź niską w zależności od poziomu, do jakiego się zaadaptowała. Ta sama zasada obowiązuje w kwestiach niezwiązanych z postrzeganiem zmysłowym takich, jak zdrowie, prestiż czy zamożność. Przykładowo ten sam poziom majątku może oznaczać dla dwóch róż-

⁵¹ Koncepcja relatywnej oceny została zaczerpnięta z psychofizyki. Wytlumaczenie relatywnej percepcji za pomocą poziomów adaptacji zaproponował Harry Helson (1964). Jego prace stanowiły po części inspirację dla twórców *prospect theory*.

nych osób skrajne ubóstwo lub niezmierzone bogactwo w zależności od ich faktycznego stanu posiadania [przyp. tłum.: poziomu, do którego się zaadaptowały]” (Kahneman i Tversky, 1979).

Następstwem relatywnego postrzegania zjawisk jest także relatywna ich ocena – w stosunku do punktu odniesienia. Wyniki leżące „w prawo” od punktu odniesienia, uznawane są za zyski (ocena pozytywna), wyniki leżące „po lewej stronie” punktu odniesienia traktowane są jako straty (ocena negatywna). Sam punkt odniesienia jest neutralny, odpowiada bowiem poziomowi, do którego jednostka się przystosowała. Poziom ten stanowi najczęściej *status quo*, np. aktualny stan majątkowy⁵².

Warto uświadomić sobie znaczenie punktu odniesienia dla osoby aktywnie działającej na rynku kapitałowym. Wyobraźmy sobie inwestora, który kupuje akcje spółki za 100 j.p. Cena zakupu stanowi dlań naturalny punkt odniesienia w tej transakcji. Jeśli następnie kurs akcji wzrośnie do 110 j.p., inwestor będzie cieszył się z uzyskania dodatkowych 10 j.p. Jeśli natomiast kurs spadnie do 90 j.p., inwestor będzie cierpiał z powodu straty 10 j.p. Mimo iż obecnie konstatacja ta wydaje się banalna, ma ona niezwykle duże znaczenie dla zrozumienia zachowań poruszających rynki finansowe (por. punkt 4.1.2).

3.3.2.2 Malejąca wrażliwość na zmiany

Reakcje ludzkiej percepcji na zmiany w poziomie bodźca zewnętrznego dają się opisać za pomocą wklęsłej funkcji, której argumentem jest natężenie owego bodźca. Przykładem tej własności może być postrzeganie temperatury otoczenia: różnica między 3°C a 6°C jest łatwiej odczuwalna niż między 13°C a 16°C. Kahneman i Tversky (1979) twierdzą, że zasada ta obowiązuje także dla wypłat pieniężnych; subiektywna różnica wartości między zyskiem 100 j.p. a 200 j.p. wydaje się większa niż między 1100 j.p. a 1200 j.p. Podobnie subiektywna różnica między stratą 100 j.p. a 200 j.p. wydaje się

⁵² Kahneman i Tversky nie precyzują, co może być punktem odniesienia. Punkt odniesienia zmienia się w zależności od rodzaju aktywności. Można przyjąć, że jednostka posługuje się wieloma funkcjami wartości, z których każda ma swój punkt odniesienia.

większa niż między 1100 j.p. a 1200 j.p., o ile większa ze strat nie okaże się nie do zniesienia (np. nie przekroczy ograniczenia budżetowego danej jednostki lub innego limitu administracyjnego).

Powyższy przykład opisuje wcześniej wspomnianą własność funkcji wartości: malejącą wrażliwość na zmiany w poziomie bodźca (zysk/strata). Jest ona implikowana wklęsłością funkcji powyżej punktu odniesienia ($v''(x) < 0$ dla $x > 0$) i wypukłością poniżej punktu odniesienia ($v''(x) > 0$ dla $x < 0$).

Dla inwestora malejąca wrażliwość oznacza, że bardziej cieszy się on z pierwszego zarobionego złotego niż z drugiego, z drugiego bardziej niż z trzeciego, itd. W dziedzinie zysków, obserwacja taka nie jest bynajmniej rewolucyjna i daje się wytłumaczyć awersją do ryzyka. W dziedzinie strat, podobnie, strata pierwszego złotego jest bardziej bolesna niż drugiego, drugiego bardziej niż trzeciego, itd. Tym razem jednak malejąca wrażliwość stanowi bodziec do zachowań ryzykownych (zob. punkt 3.3.4).

3.3.2.3 Awersja do strat

Znamienną własnością ludzkiej percepcji relatywnych zmian majątku jest fakt, że straty bołą bardziej niż cieszą zyski tych samych bezwzględnych rozmiarów (Kahneman i Tversky, 1979). Dowodem może być często obserwowana niechęć jednostek do przyjmowania symetrycznych zakładów, np. $[(x; 0,5); (-x; 0,5)]$. Co więcej, awersja do takich loterii wzrasta wraz ze zwiększaniem licytowanej kwoty x . Jeśli $x > y \geq 0$, wówczas spośród par gier $[(x; 0,5); (-x; 0,5)]$ oraz $[(y; 0,5); (-y; 0,5)]$ preferowana jest ta druga.

Na podstawie równania (3.3) mamy zatem:

$$v(y) + v(-y) > v(x) + v(-x) \Rightarrow v(-y) - v(-x) > v(x) - v(y). \quad (3.5)$$

Podstawiając $v(y=0)=0$ otrzymujemy:

$$v(x) < -v(-x). \quad (3.6)$$

Gdy y zmierza do x , przy założeniu, że funkcja $v(x)$ jest różniczkowalna dla wszystkich $x \neq 0$:

$$v'(x) < v'(-x). \quad (3.7)$$

Jak widać, funkcja wartości jest bardziej stroma w dziedzinie strat niż w dziedzinie zysków. Własność tę przyjęło się określać mianem awersji do strat (współczynnik λ w postaci funkcyjnej (3.4)). Straty odczuwane są dużo mocniej (od $\lambda = 2$ do $\lambda = 2,5$ ⁵³ razy) niż odpowiadające im kwotowo zyski.

Oczywiście awersja do strat uzależniona jest także od relacji między wielkością straty a poziomem całego majątku. Podczas gdy dla dużego inwestora ból z powodu straty poniesionej na inwestycji wartej 1000 j.p. jest niewiele większy niż radość z zysku o podobnym rozmiarze, dla osoby o przeciętnych dochodach straty w takiej pozycji mogą być bardzo bolesne. Nie oznacza to jednak, że indywidualne współczynniki awersji do strat są niższe dla dużych inwestorów, a wyższe dla małych. Można przyjąć, że uczestnicy rynku charakteryzują się podobną awersją do strat w relacji do ich całkowitego potencjału inwestycyjnego (do właściwego im punktu odniesienia).

* * *

Jeśli zaakceptować tradycyjne podejście, w którym użyteczności poszczególnych zdarzeń ważone są numerycznym (obiektywnym) prawdopodobieństwem, wówczas funkcja wartości $v(x)$ może służyć jako samodzielna koncepcja, tłumacząca zachowania w warunkach ryzyka. Jednakże Kahneman i Tversky idą o krok dalej, uznając matematycznie wyznaczone prawdopodobieństwa zdarzeń za ułomną reprezentację percepcji ryzyka. W kolejnym etapie rozważań nad *prospect theory* przybliżona zostanie zaproponowana przez nich funkcja wag prawdopodobieństwa.

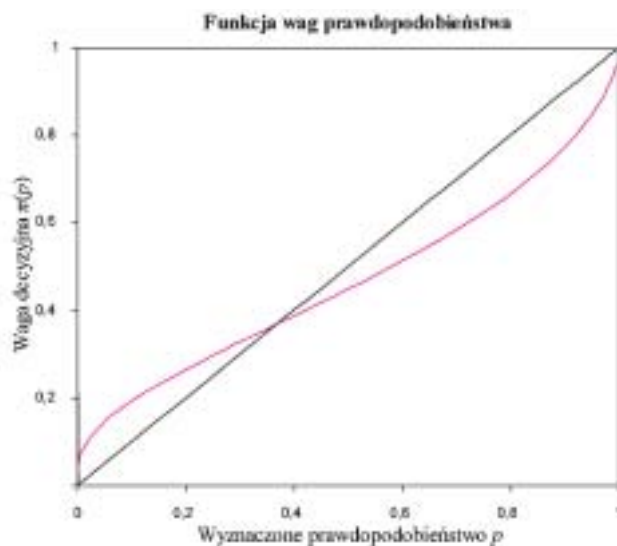
3.3.3 Funkcja wag prawdopodobieństwa

Tradycyjne podejście do decyzji w warunkach ryzyka zakłada, że prawdopodobieństwa pojmowane są zgodnie z ich matematycznie wyznaczonym poziomem. Wiele eksperymentalnych sytuacji decyzyjnych pokazuje jednak, że ludzie nie przestrzegają tej zasady (Kahneman i Tversky, 1979, 1986,

⁵³ Parametr λ wyznaczony został na podstawie badań empirycznych. Por. (Kahneman i Tversky, 1992).

1992). Ich percepcja okazuje się zniekształcona nawet wówczas, gdy numeryczne prawdopodobieństwa są jawnie podane. *Prospect theory* daje wyraz temu zjawisku za pomocą ich nieliniowej transformacji – funkcji wag prawdopodobieństwa $\pi(p)$ (*probability weighting function*) (rysunek 3.2).

Rysunek 3.2: Funkcja wag prawdopodobieństwa $\pi(p)$



Źródło: (Kahneman i Tversky, 1979).

Funkcja $\pi(p)$ wskazuje na malejącą wrażliwość na zmiany prawdopodobieństw wraz z oddalaniem się od punktu odniesienia, którym może być zarówno $p=0$ (niemożliwość), jak i $p=1$ (pewność). Co ciekawe, kształt funkcji wyróżnia się charakterystyczną asymetrią: obszar wypukły jest około dwukrotnie większy niż obszar wklęsły. Taki przebieg funkcji $\pi(p)$ nie jest w żadnym razie intuicyjny, jednakże został dowiedziony zarówno w wielu eksperymentach laboratoryjnych, jak i poprzez obserwację rzeczywistości. O tym – w dalszej części paragrafu.

Oto jak o omawianym zjawisku piszą autorzy (Kahneman i Tversky, 1992): „Przy ocenie niepewności rolę odgrywają dwie naturalne granice – pewność i niemożliwość – zlokalizowane na końcach skali prawdopodobieństwa. Malejąca wrażliwość oznacza, że wpływ danej zmiany prawdopodobieństwa zmniejsza się wraz z oddalaniem się od tych granic. Na przykład

dziesięcioprocentowy wzrost szansy wygranej ma silniejszy wpływ [przyp. tłum.: jest silniej postrzegany], gdy oznacza on zmianę prawdopodobieństwa z 90% na 100% lub z 0% na 10% niż zmianę z 30% na 40% lub 60% na 70%”.

Konsekwencją malejącej wrażliwości jest wklęsłość funkcji $\pi(p)$ dla prawdopodobieństw bliskich 0 i wypukłość dla prawdopodobieństw bliskich 1. Należy jednak zaznaczyć, że w otoczeniu pewności i niemożliwości funkcja nie jest ściśle określona: bardzo małe prawdopodobieństwa mogą być zarówno znacznie przewartościowane, jak i całkowicie pomijane.

Omówionym własnościom odpowiada następująca postać funkcji $\pi(p)$ zaproponowana przez Preleca (2000):

$$\pi(p) = \exp\{-\beta(-\ln p)^\alpha\},$$

(3.8)

gdzie:

$\pi(p)$ – waga decyzyjna (transformacja prawdopodobieństwa),

p – prawdopodobieństwo (wyznaczone zgodnie z teorią prawdopodobieństwa),

α, β – parametry funkcji, $\alpha, \beta > 0$.

Warto jeszcze raz podkreślić, iż kształt funkcji $\pi(p)$ opiera się na rzeczywistym sposobie oceny niepewności przez jednostki. Pozwala on na następującą klasyfikację systematycznych zachowań ludzkich w warunkach ryzyka:

- (1) efekt pewności (*certainty effect*) – przecenianie zdarzeń pewnych w stosunku do wysoce prawdopodobnych,
- (2) niedoszacowanie prawdopodobieństw wysokich i średnich,
- (3) przewartościowanie niskich prawdopodobieństw.

3.3.3.1 Efekt pewności i niedoszacowanie wysokich i średnich prawdopodobieństw

Dla zobrazowania tego zjawiska rozważmy następujące loterie (Kahneman i Tversky, 1979):

Loteria 1⁵⁴: $A(4000;0,80)$ [20%] $B(3000;1,0)$ [80%]

Loteria 2: $C(4000;0,20)$ [65%] $D(3000;0,25)$ [35%]

Wybierając opcję B w loterii 1 i opcję C w loterii 2, ponad połowa respondentów zachowuje się w sposób niezgodny z modelem racjonalnym. W kategoriach teorii oczekiwanej użyteczności decyzję tę można przedstawić jako:

Loteria 1: $u(3000) > 0,80 u(4000)$

Loteria 2: $0,25 u(3000) < 0,20 u(4000)$ ⁵⁵ $\Leftrightarrow u(3000) < 0,80 u(4000)$

Oczywistym jest, że obie nierówności nie mogą być prawdziwe.

Zachowanie takie daje się wytłumaczyć przy użyciu funkcji $\pi(p)$. Z jej kształtu widać, że wysokie prawdopodobieństwa zdarzeń są niedoceniane w stosunku do zdarzeń pewnych:

$$\frac{\pi(0,20)}{\pi(0,25)} > \frac{\pi(0,80)}{\pi(1,0)}.$$

Kahneman i Tversky trafnie charakteryzują ten sposób myślenia: „Stuprocentowe prawdopodobieństwo jest nadal gorsze niż pewność”. I rzeczywiście, mając przed sobą loterię wypłacającą 100 j.p. z prawdopodobieństwem 99% każdy byłby gotowy zapłacić więcej za otrzymanie wygranej z pewnością (zmiana prawdopodobieństwa z 99% na 100%) niż w przypadku loterii, w której wygrywa się z prawdopodobieństwem 46% z opcją zmiany na 47%.

⁵⁴ Loterie zostały przedstawione w konwencji używanej przez Kahnemana i Tversky'ego. Przykładowo loterię 1 należy rozumieć jako „realizacja A: wypłata 4000 j.p. z 80% prawdopodobieństwem; realizacja B: pewna wypłata 3000 j.p.”. W nawiasach kwadratowych podany jest odsetek respondentów wybierających daną opcję.

⁵⁵ Każda ze stron w tej nierówności wynika z obustronnego pomnożenia nierówności w loterii 1 przez czynnik 0,25.

3.3.3.2 Przewartościowanie niskich prawdopodobieństw

Wyniki następujących eksperymentów służą zademonstrowaniu zjawiska przeszacowania niskich prawdopodobieństw (Kahneman i Tversky, 1979):

Loteria 1: **A(5000;0,001) [72%]** **B(5;1,0) [28%]**

Loteria 2: **C(-5000;0,001) [17%]** **D(-5;1,0) [83%]**

W loterii 1 większość respondentów wybiera alternatywę ryzykowną (A), co dowodzi skłonności do ryzyka. Analogiczne zjawisko można zaobserwować codziennie, kiedy miliony osób kupują kupony lotto, płacąc cenę dalece przewyższającą wartość oczekiwaną wygranej. Podobnie ofiarą przeszacowania niskich prawdopodobieństw padają inwestorzy, kupujący tanie opcje *out-of-the-money*. Mimo iż prawdopodobieństwo wykonania takiej opcji jest niezwykle niskie, inwestorzy dają się skusić znikomą szansą zarobienia znacznej kwoty.

Loteria 2 polega na wyborze „mniejszego zła”, czyli preferowanej straty. W tym wypadku, większość respondentów znamionuje awersja do ryzyka: wybierają stratę pewną, lecz małą. Z tej własności czerpie swoje zyski przemysł ubezpieczeniowy, gdzie klient decyduje się na niewielką pewną stratę (koszt polisy), zabezpieczając się jednocześnie przed stratą znaczną, lecz niezwykle mało prawdopodobną.

* * *

Przedstawione własności funkcji $\pi(p)$ są wyrazem charakterystycznego stosunku do ryzyka, który można przedstawić w czterech wymiarach⁵⁶ (tabela 3.3

⁵⁶ Warto tu przypomnieć, że stosunek jednostki do ryzyka można określać na podstawie jej wyborów. O awersji do ryzyka świadczy przedkładanie wypłaty pewnej ponad loterię z wartością oczekiwaną równą wypłacie pewnej lub od niej wyższą; skłonność do ryzyka cechuje jednostki, preferujące loterię w stosunku do wypłaty pewnej wyższej od lub równej wartości oczekiwanej loterii. Tak też skonstruowany został test przedstawiony w tabeli.

Tabela 3.3: Cztery wymiary stosunku do ryzyka

<i>Prawdopodobieństwo</i>	<i>Zysk</i>	<i>Strata</i>
<i>Niskie</i>	<i>Skłonność do ryzyka</i> $C(100;0,05) = 14^*$	<i>Awersja do ryzyka</i> $C(-100;0,05) = -8$
<i>Wysokie</i>	<i>Awersja do ryzyka</i> $C(100;0,95) = 78$	<i>Skłonność do ryzyka</i> $C(-100;0,95) = -84$

* Uwaga: W eksperymencie wartość oczekiwana (EV) loterii porównywana była z kwotą, jaką respondenci byliby w stanie zapłacić za udział w loterii. Np. za udział w loterii (100;0,05) o wartości oczekiwanej $EV = 5$ j.p. respondenci byli skłonni zapłacić 14 j.p. Oznacza to ich skłonność do podejmowania ryzyka.

Źródło: Opracowanie własne na podstawie (Tversky i Fox, 1995) oraz (Prelec, 2000).

Przecenianie niskich prawdopodobieństw generuje jednocześnie awersję do ryzyka dla mało prawdopodobnych strat i upodobanie do ryzyka dla mało prawdopodobnych zysków. Natomiast niedocenywanie prawdopodobieństw średnich i wysokich powoduje awersję do ryzyka w dziedzinie zysków i skłonność do ryzyka w dziedzinie strat. Skłonność do ryzyka dla mało prawdopodobnych zysków tłumaczy popularność gier hazardowych, zaś awersja do ryzyka dla mało prawdopodobnych strat – chęć ubezpieczania się. Skłonność do ryzyka dla wysoce prawdopodobnych strat stymuluje do podejmowania dodatkowego ryzyka w celu uniknięcia ich realizacji. Omówione wzorce zachowań zostały potwierdzone przez liczne badania.

W tym miejscu trzeba przypomnieć, iż podejście normatywne wyklucza zachowania świadczące o skłonności do ryzyka, uznając je za nieracjonalne, zaś na rynkach finansowych wręcz za patologiczne.

3.3.4 Efekt odwrócenia

Efekt odwrócenia (*reflection effect*) jest jedną z najważniejszych cech znamionujących preferencje zgodne z *prospect theory*. Polega on na zmianie kolejności preferencji przy przejściu od loterii zdefiniowanych na zyskach do ich lustrzanego odbicia w dziedzinie strat⁵⁷. Efekt ten jest tłumaczony

⁵⁷ Należy zaznaczyć, iż efekt odwrócenia został niejawnie przedstawiony w tabeli 3.3. Kahneman i Tversky (1979) testują zmiany preferencji dla różnych poziomów prawdopodobieństwa. Jednakże o efekcie odwrócenia przyjęło się mówić w sytuacji wyboru między loteriami (zdefiniowanymi zarówno dla zysków, jak i strat), z których jedna jest loterią pewną. Taka sytuacja decyzyjna została przedstawiona w przytoczonym przykładzie.

malejącą wrażliwością, wynikającą z kształtu funkcji wartości (wklęsłość funkcji dla zysków oraz wypukłość dla strat).

Rozważmy następujący przykład:

Loteria 1: $A(5000;1,0)$ $B(10\ 000;0,50)$

Loteria 2: $C(-5000;1,0)$ $D(-10\ 000;0,50)$

W pierwszej sytuacji (loteria 1) należy wybrać między pewną wygraną A , 5000 j.p., a ryzykowną opcją B , dającą 50% szansę otrzymania 10 000 j.p. Awersja do ryzyka intuicyjnie skłania większość respondentów do wyboru opcji pewnej.

W drugiej sytuacji (loteria 2) większość osób decyduje się na alternatywę ryzykowną (D), odrzucając pewną stratę (C), co dowodzi skłonności do ryzyka. Zachowanie to można także zinterpretować za pomocą malejącej wrażliwości na dodatkowe straty: jednorazowa strata 10 000 j.p. doświadczana jest słabiej niż dwukrotność straty 5000 j.p.

Zachowania zgodne z efektem odwrócenia – ostrożność w dziedzinie zysków i ryzykanctwo w dziedzinie strat – mają niebanalne znaczenie dla rentowności inwestycji dokonywanych na rynkach finansowych. Konsekwencje tego zjawiska omówione zostaną w punkcie 4.1.2.

3.3.5 Kadrowanie

Zgodnie z klasyfikacją aksjomatów zaproponowaną przez Kahnemana i Tversky'ego (por. punkt 3.1.2) niezmiennosc jest najważniejszą własnością racjonalnych preferencji. Wskazanie na preferencje gwałcące tę zasadę kładzie kres użyciu teorii oczekiwanej użyteczności jako adekwatnego opisu ludzkich preferencji.

Jak wcześniej wspomniano, niezmiennosc preferencji oznacza, iż różne formy prezentacji problemu decyzyjnego nie mają wpływu na dokonywane wybory. Odkrycie rzeczywistych zachowań, które systematycznie łamią tę zasadę, stanowiło istotną motywację do stworzenia lepszego opisu ludzkich preferencji. Tak oto *prospect theory*, jako jedyna teoria podejmowania decyzji w warunkach ryzyka, jest w stanie wytłumaczyć zależność wyborów od

sposobu ich przedstawienia. W niniejszej pracy zjawisko to (*framing*) zdecydowano się nazywać kadrowaniem.

Ten sam problem decyzyjny może być przedstawiony na różne sposoby (np. w kontekście zysków bądź strat) – różnie skadrowany – względem punktu odniesienia. Zależność decyzji od formy doskonale przedstawia opisany poniżej eksperyment Kahnemana i Tversky'ego (1984). Odsetek respondentów wybierających daną opcję został podany w nawiasach.

Sytuacja 1: Wyobraźmy sobie, że USA przygotowuje się do wybuchu epidemii, która ma pochłonąć 600 ofiar. Istnieją dwa alternatywne programy walki z zarazą. Szacowane konsekwencje tych programów są następujące:

Jeśli uruchomiony zostanie program A, można być pewnym, że 200 żyć ludzkich zostanie uratowanych. [72%]

Jeśli uruchomiony zostanie program B, z prawdopodobieństwem 1/3 uratowanych zostanie 600 osób, zaś z prawdopodobieństwem 2/3 wszyscy zginą. [28%]

Sformułowanie użyte w sytuacji 1 nakazuje podświadomie przyjąć za punkt odniesienia potencjalną liczbę ofiar (600). Opcje A i B przedstawione są w kategoriach względnych zysków (liczba uratowanych osób). Jak można oczekiwać na podstawie *prospect theory*, respondenci, powodowani awersją do ryzyka, wybierają scenariusz pewny (A), w którym przeżyje 200 osób.

Rozważmy obecnie sytuację 2 (zdanie wprowadzające nie ulega zmianie).

Sytuacja 2:

Jeśli uruchomiony zostanie program C, 400 osób poniesie pewną śmierć. [22%]

Jeśli uruchomiony zostanie program D, z prawdopodobieństwem 1/3 nie zginie nikt, zaś z prawdopodobieństwem 2/3 śmierć poniesie 600 osób. [78%]

Jak łatwo zauważyć, opcje A i C oraz B i D są odpowiednio tożsame pod względem konsekwencji (liczba ofiar/liczba ocalonych). Jednakże sytuacja 2 niejawnie definiuje stan „zero ofiar” jako punkt odniesienia. Ponownie odpowiedzi respondentów zgodne są z preferencjami opisywanymi w *prospect*

theory. Sformułowanie problemu w kategoriach strat (liczba zgonów) skłania do ryzykanctwa – odrzucenia opcji pewnej (program C) na rzecz loterii (program D).

Jest to oczywisty przykład pogwałcenia aksjomatu niezmienności preferencji. Okazuje się bowiem, że wybory determinowane są sposobem ich prezentacji, nie zaś rzeczywistą treścią problemu. Co więcej, nawet świadomość i dojrzałość intelektualna respondenta nie zapobiega tej niekonsekwencji⁵⁸.

Można by przypuszczać, że podobne „występki” przeciw racjonalności zdarzają się jedynie w sytuacjach nie angażujących wypłat pieniężnych. Nic bardziej mylnego: kadrowanie właściwe jest również wyborom istotnym ekonomicznie. Dowodzi tego kolejny eksperyment (Kahneman i Tversky, 1979).

Sytuacja 1: Niezależnie od tego, co posiadasz, otrzymałeś dodatkowo 1000 j.p. Dokonaj wyboru między:

A: wygraną 1000 j.p. z prawdopodobieństwem 50% [16%]

B: pewną wygraną 500 j.p. [84%]

Sytuacja 2: Niezależnie od tego, co posiadasz, otrzymałeś dodatkowo 2000 j.p. Dokonaj wyboru między:

C: stratą 1000 j.p. z prawdopodobieństwem 50% [69%]

D: pewną stratą 500 j.p. [31%]

Także w tej sytuacji wybory większości uczestników badania dają dowód ostrożności w dziedzinie zysków i ryzykanctwa w dziedzinie strat. Należy zauważyć, że sformułowanie A jest tożsame z C i podobnie B z D. Problemy te można sprowadzić do identycznych postaci:

$$A = [(2000; 0, 50); (1000; 0, 50)] = C \text{ oraz } B = (1500) = D.$$

Opcje A i C, po uwzględnieniu początkowej wypłaty pewnej, dają 50% szansę otrzymania 2000 j.p. oraz 50% szansę otrzymania 1000 j.p. Opcje B

⁵⁸ W eksperymencie Kahnemana i Tversky’ego respondenci odpowiadają na pytania jedno po drugim. Niekonsekwencji wyborów nie można zatem tłumaczyć ich oddaleniem w czasie. Nawet po kilkukrotnej konfrontacji z problemem, respondenci nadal decydują się na opcję pewną przy sformułowaniu „życia ocalone” oraz opcję ryzykowną przy sformułowaniu „życia utracone”. Por. (Kahneman i Tversky, 1984).

oraz D natomiast oferują pewną wygraną 1500 j.p. Problemy przedstawione w sytuacji 1 i 2 są niczym innym jak dwuetapową loterią. Wypłata pewna (1000 j.p. w sytuacji 1, 2000 j.p. w sytuacji 2) otrzymana w kroku pierwszym staje się automatycznie punktem odniesienia dla decyzji podejmowanych w kroku drugim. Tu z kolei działa omówiony wcześniej efekt odwrócenia (por. punkt 3.3.4); wybór pada na opcję pewną po stronie zysków i ryzykowną po stronie strat.

* * *

Prospect theory tłumaczy preferencje zwykłego człowieka. Określenie „zwykły” nie przeczy profesjonalizmowi czy wyrafinowaniu jednostki, ma jednak pokazać, że nawet takie cechy nie chronią przed nieracjonalnością. W rzeczywistości trudno znaleźć reprezentantów podgatunku *homo oeconomicus*, stąd *prospect theory* okazuje się tak dokładnie odzwierciedlać mechanizmy rządzące ludzkimi wyborami. Codzienne decyzje determinowane są przez punkty odniesienia, ryzykanctwo w sytuacji strat i ostrożność w sytuacji zysków, preferencję zdarzeń pewnych, przecenianie zdarzeń mało prawdopodobnych, kadrowanie... Jednostki działające na rynkach finansowych – profesjoniści – podlegają takim zjawiskom ze zwielokrotnioną siłą. Katalizatorami są tu bowiem presja czasu i niepewność – siły wrogie racjonalnemu działaniu. Dowodów istotności *prospect theory* w analizie zachowań ekonomicznych dostarczono w następnym rozdziale.

Krytycy behawioralnego podejścia do rynków finansowych i wyborów ekonomicznych utrzymują, iż loterie analizowane przez Kahnemana i Tversky'ego rzadko występują w rzeczywistości. Badania przeprowadzane wśród profesjonalnych inwestorów dowodzą jednak, że zachowania „zawodowców” nie różnią się od tych, jakie właściwe są uczestnikom eksperymentów omówionych w poprzednim rozdziale. Czy zrealizować już teraz zysk, czy może liczyć na większy, lecz niepewny? Czy zlikwidować pozycję z wynikiem ujemnym, czy może odczekać z nadzieją na „uśmiech losu”, ryzykując jednocześnie jeszcze większą stratę?

Ekonomiczne konsekwencje preferencji opisywanych przez *prospect theory* są przedmiotem dyskusji w niniejszym rozdziale. W pierwszej części omówiono zachowania ekonomiczne sprzeczne z tradycyjnymi paradygmatami, lecz pozwalające się wyjaśnić na gruncie modelu Kahnemana i Tversky'ego. Część druga przybliża koncepcję mentalnej księgowości oraz jej zastosowania.

4.1 Analiza anomalii w zachowaniach ekonomicznych

W nauce finansów anomaliami przyjęło się nazywać te obserwowalne zjawiska, które niezgodne są z teorią oczekiwaną użyteczności i modelami pochodnymi (por. podrozdział 1.1). Ekonomisci od dawna deklarowali gotowość wzbogacenia podejścia normatywnego o czynnik psychologiczny, o ile pozwoli to na lepszy opis rzeczywistości. Tak oto *prospect theory* została stworzona i rozumiana w odpowiednim momencie, by sprostać zaistniałemu zapotrzebowaniu. Za sprawą elementów ją tworzących – awersji do strat, efektu odwrócenia i nieliniowej funkcji wag prawdopodobieństwa – możliwe stało się wytłumaczenie „mistycznego” dziania się na rynkach finansowych.

4.1.1 Krótkowzroczna awersja do strat

Opisywana w punkcie 1.1.1 premia za inwestycję na rynku kapitałowym nurtowała świat nauki przez niemal dwie dekady. Pytanie o przyczynę niezmiernie wysokich stóp zwrotu z akcji, przewyższających w ujęciu historycznym rentowność obligacji o 6 pkt. proc., pozostało jednak otwartą kwestią⁵⁹.

W roku 1995 Benartzi i Thaler zaproponowali rozwiązanie zagadki w duchu teorii Kahnemana i Tversky'ego (Benartzi i Thaler, 1995). Punkt wyjścia ich analizy to założenie awersji do strat właściwej inwestorom oraz nadmiernej (w odniesieniu do faktycznego horyzontu inwestycyjnego) częstotliwości oceny stopy zwrotu z portfela. Kombinacja obu tych charakterystyk – zwana krótkowzroczną awersją do strat (*myopic loss aversion*) – w modelu Benartzi'ego i Thaler'a stanowi wytłumaczenie zagadkowo wysokiej premii za ryzyko.

4.1.1.1 Badania empiryczne

Benartzi i Thaler poszukują takiego połączenia awersji do strat z częstotliwością ewaluacji portfela, które tłumaczyłoby ów niezmiernie wysoki w przeszłości poziom premii za ryzyko. W tym celu stawiają następujące pytania: (i) Jak często inwestor musiałby szacować wartość portfela, by być obojętnym w wyborze między inwestycją w akcje bądź w obligacje? (ii) Załóżmy, że inwestor wycenia swój portfel w interwałach wyznaczonych w pytaniu (i). Jakie proporcje akcji i obligacji zawierałby wówczas jego portfel?

Przeprowadzając symulację dla miesięcznych stóp zwrotu z lat 1926–1990⁶⁰, Benartzi i Thaler wykazali, że: (i) średnia częstotliwość⁶¹, z jaką inwestor musiałby wyceniać zwrot z portfela, wynosi około 12 miesięcy; (ii)

⁵⁹ Warto przypomnieć, że zgodnie z CAPM premia w wysokości 1 pkt. proc. byłaby dostatecznym wynagrodzeniem za inwestowanie na rynku akcji.

⁶⁰ Dane obejmowały miesięczne stopy zwrotu odnotowane na rynku amerykańskim dla akcji, obligacji oraz bonów skarbowych i zostały dostarczone przez CRSP (*Center for Research in Security Prices*).

⁶¹ Należy odróżniać częstotliwość ewaluacji od planowanej długości inwestycji. Częstotliwość, z jaką inwestor wycenia swój portfel powinna mieć dla niego jedynie charakter informacyjny.

musiałby wyceniać zwrot z portfela, wynosi około 12 miesięcy; (ii) inwestor wyceniający portfel średnio raz do roku maksymalizowałby swoją użyteczność, inwestując około połowę środków w akcje i połowę w obligacje. Odpowiedzi te są wzajemnie jednoznaczne: przy dwunastomiesięcznym horyzoncie inwestycyjnym osoba o preferencjach zgodnych z *prospect theory* żądałaby około 6 pkt. proc. premii za ryzyko. Takie wynagrodzenie czyniłoby obojętnym wybór między inwestycją w akcje a zakupem obligacji.

Benartzi i Thaler idą w swoich rozważaniach o krok dalej, zadając pytanie, przy jakim okresie ewaluacji premia za inwestycję w akcje przyjęłaby poziom odpowiadający przewidywaniom klasycznego podejścia (np. CAPM – około 1 pkt. proc.). Zgodnie z oczekiwaniem, premia za ryzyko maleje wraz z wydłużaniem interwałów, w jakich portfel jest wyceniany. Dla okresów 5, 10 i 20 lat implikowana premia – wyznaczona na podstawie symulacji autorów – wynosi odpowiednio: 3 pkt. proc., 2 pkt. proc. i 1,4 pkt. proc. (przy jej faktycznej wielkości wynikającej z analizowanych danych równej 6,5 pkt. proc.). Osoba szacująca portfel raz do roku, przy 20-letnim horyzoncie inwestycyjnym, wymaga wynagrodzenia o 5,1 pkt. proc. przewyższającego poziom premii wynikający z rzeczywistego ryzyka⁶².

4.1.1.2 Interpretacja

Odniesienie dwunastomiesięcznego okresu wyceny portfela do rzeczywistości pozwala na przekonującą interpretację. Oczywiście trudno przypuszczać, że każdy dokonuje takiego rachunku dokładnie z tą samą częstością. Termin jednego roku wydaje się jednak dobrym uśrednieniem. Otóż inwestorzy dokonują rozliczeń podatkowych raz do roku, otrzymują także raporty roczne od swoich biur maklerskich, funduszy inwestycyjnych czy emerytalnych. Również inwestorzy instytucjonalni myślą o swoich wynikach w kategoriach roku, z taką bowiem częstotliwością przebiega proces raportowania i oceny sprawności ich działania.

⁶² Obserwowana premia minus premia uzasadniona ryzykiem 20-letniej inwestycji, czyli 6,5 pkt. proc. – 1,4 pkt. proc. = 5,1 pkt. proc.

Weźmy przykład młodego inwestora, który lokuje środki na rynku kapitałowym jako zabezpieczenie swojej starości (horyzont inwestycyjny około 30 lat). Inwestor ten, niezależnie od długiej perspektywy inwestycyjnej, przypuszczalnie wycenia swój portfel dużo częściej (np. corocznie). Użyteczność, jaką kojarzy z jego rentownością, stanowi odbicie wyników osiągniętych w krótkiej perspektywie. Należy zauważyć, iż krótkoterminowe wahania kursów akcji mogą być bardzo znaczące, co zniechęca do inwestowania w te aktywa. Im częściej inwestorzy kalkulują wyniki księgowe („papierowe”) swoich inwestycji na giełdzie, tym mniej atrakcyjne będą się one wydawały. Wobec tego premię za ryzyko można zinterpretować w kontekście kosztów psychicznych (postrzeganych przez jednostkę), wywołanych dużą częstotliwością wyceny portfela.

4.1.1.3 Znaczenie dla rynku

Można by spekulować, że krótkowzroczna awersja do strat to przypadek indywidualnych inwestorów, zatem rozwiązanie zaproponowane przez Bernartzi’ego i Thaler’a ma niewielką wartość poznawczą. Większość aktywów rynku kapitałowego znajduje się bowiem w rękach organizacji, w szczególności funduszy emerytalnych i powierniczych oraz donacji. Czy krótkowzroczną awersję do strat przejawiają także inwestorzy instytucjonalni?

Istnieją dowody, że profesjonalizm inwestorów nie stanowi osłony przed stwierdzonymi błędami. Na przykład korporacyjne fundusze emerytalne w USA inwestują przeciętnie w akcje i obligacje w stosunku 60:40. Tak niski udział aktywów ryzykownych jest niezrozumiały w świetle bardzo długiego horyzontu inwestycyjnego, który umożliwia funduszom wykorzystanie premii za inwestycję w akcje. Ponownie przekonującym wytłumaczeniem tego faktu jest krótkowzroczna awersja do strat połączona z problemem pośrednictwa (*agency problem*). Perspektywę samego funduszu można uznać za nieskończoną, w przeciwieństwie do perspektywy osoby nim zarządzającej, najczęściej menedżera finansowego korporacji. Zarządzający funduszem z pewnością nie zakładają pozostania na tym stanowisku nieskończenie długo. Również system wynagradzania i raportowania skłania ich do myślenia o

wynikach swej pracy w rytmie rocznym. Ów krótki horyzont stanowi źródło konfliktu między celami managerów a uczestnikami funduszu. Optymalizowanie krótkoterminowych wyników (motywowane krótkowzroczną awersją do strat), jakkolwiek zgodne z interesami managerów, może prowadzić do nieoptymalnej alokacji środków klientów w długiej perspektywie.

Reasumując, interpretację Benartzi’ego i Thaler’a należy uznać za przekonującą, choć jedynie częściowe, wytłumaczenie premii za ryzyko. Model w pełni satysfakcjonujący powinien dodatkowo uwzględniać działanie innych czynników, np. czasu i konsumpcji. Rozwiązanie opierające się na awersji do strat zasługuje jednak na szczególną uwagę. Jest to bowiem pierwsza tak udana próba przeniesienia dokonań psychologów Kahnemana i Tversky’ego w dziedzinę finansów. Model preferencji wyjaśnianych przez *prospect theory* zastosowany do analizy zjawisk rynku finansowego okazuje się odpowiadać na pytania, z jakimi od dawna nie radziła sobie teoria oczekiwanej użyteczności.

4.1.2 Efekt dyspozycji

W 1985 roku Shefrin i Statman opublikowali artykuł analizujący jedno z najbardziej typowych i jednocześnie frapujących zachowań inwestorów: skłonność do przedwczesnej realizacji zysków i do zwlekania z realizacją strat (Shefrin i Statman, 1985). Ów tzw. efekt dyspozycji⁶³ (*disposition effect*) przypisany został awersji do strat oraz malejącej wrażliwości, odpowiedzialnym za zmianę postawy wobec ryzyka: z ostrożności w sytuacji zysków na zachowania ryzykowne w sytuacji strat⁶⁴. Efekt ten sprawia, że inwestorzy przetrzymują akcje spółek, które straciły na wartości (w porównaniu z ceną zakupu), zaś sprzedają te, których kursy wzrosły. Nawet

⁶³ Określenie „efekt dyspozycji”, zaproponowane przez Shefrina i Statmana, jest niefortunnym skrótem myślowym opisu: „dyspozycji do sprzedaży akcji wygrywających (tzw. zwycięzców) i przetrzymywania przegranych”. Zjawisko dyspozycji w psychologii jest bardzo ogólną koncepcją związaną ze skłonnością do określonego działania. Stąd pojęcie „efekt dyspozycji” powinno zostać uzupełnione opisem zachowania, do którego się odnosi. Por. (Wärneryd, 2001, s. 249).

⁶⁴ O zjawisku tym pisano w rozdziale poświęconym *prospect theory* i efektowi odwrócenia (por. punkt 3.3.4).

profesjonaliści ulegają chorobie „wyjścia na zero” (*get-evenitis*), czyli nadziei, że straty da się „odrobić”. To błędne przypuszczenie stanowi potencjalnie jedną z najważniejszych przyczyn obniżania rentowności inwestycji giełdowej (Shefrin, 2000, s. 24).

4.1.2.1 Podstawy teoretyczne

Warto tu przywołać kształt funkcji wartości (por. rysunek 3.1) i prześledzić zachowania inwestora. Inwestor, kupując akcje, wyraża jednocześnie oczekiwanie, że przyszła stopa zwrotu z tej inwestycji wynagrodzi poniesione ryzyko. W tej sytuacji cena zakupu pełni rolę punktu odniesienia. Wzrost kursu przesuwają inwestora w dziedzinę zysków, gdzie jego funkcja wartości jest wklęsła. Implikuje to awersję do ryzyka i malejącą wrażliwość na kolejne wyższe ceny. Naturalne jest przy tym dążenie do zagwarantowania, że papierowe zyski nie zostaną utracone. Pewność taką daje jedynie zamknięcie pozycji. Inaczej dzieje się, kiedy ceny akcji kształtują się wbrew oczekiwaniom. Wówczas inwestor oddala się na swojej funkcji wartości w strefę strat, gdzie malejąca wrażliwość (wypukły kształt krzywej) skłania go do ryzykowania. Powodowany awersją do strat gotów jest trzymać akcje, nawet jeśli działanie to jest ekonomicznie pozbawione sensu. Zachowanie takie przewidywają także Kahneman i Tversky (1979), konkludując, że „osoba, która nie pogodziła się ze swoimi stratami, skłonna jest przyjąć ryzyko, którego nie byłaby w stanie zaakceptować w innej sytuacji”.

Dlaczego efekt dyspozycji należy uznać za anomalię? Otóż rozumując racjonalnie, cena nabycia nie powinna mieć wpływu na decyzję o sprzedaży akcji. Czynnikiem rozstrzygającym o kupnie bądź sprzedaży akcji winno być natomiast przewidywanie odpowiednio wzrostu lub spadku jej kursu. Od uczestnika rynku należałoby oczekiwać myślenia zgodnego z następującym algorytmem: liczę na wzrost – zajmuję pozycję długą, spodziewam się spadku – zajmuję pozycję krótką.

Efekt dyspozycji jest ponadto sprzeczny z optymalizacją zobowiązań podatkowych z inwestycji kapitałowych⁶⁵, jako że straty z inwestycji mogą zostać wykorzystane do zmniejszenia zobowiązań fiskalnych. Względy podatkowe sugerują zatem, iż straty, w przeciwieństwie do zysków, powinny być realizowane w krótkim terminie. Przy zerowych kosztach transakcyjnych oraz braku rozróżnienia między krótko- a długoterminową stopą opodatkowania, w celu optymalizacji zobowiązań wobec fiskusa inwestorzy powinni od stycznia do grudnia stopniowo zwiększać sprzedaż akcji, na których ponieśli stratę, zaś akcje zyskujące na wartości utrzymywać w portfelu możliwie najdłużej (Constantinides, 1984).

4.1.2.2 Badania empiryczne

Efekt dyspozycji potwierdzony został w licznych analizach empirycznych, z których najpowszechniej uznawaną przeprowadził Odean (1998a)⁶⁶. Odean przetestował występowanie efektu dyspozycji dla 10 000 rachunków inwestycyjnych z danymi o transakcjach zakupu i sprzedaży dokonywanych przez indywidualnych inwestorów od stycznia 1987 roku do grudnia 1993 roku. Jak pokazała jego analiza, inwestorzy trzymali akcje „przegrywające” oraz „zwycięskie” odpowiednio przez 124 oraz 104 dni (ujęcie średnie). Często przytaczanym argumentem za przetrzymywaniem akcji tracących na wartości jest nadzieja, że straty zostaną odrobione. O iluzoryczności takiego myślenia świadczą dalsze wyniki badania Odeana. Zatrzymane w portfelu akcje przegrywające „odbiły się” przeciętnie o 5% w następnym roku, podczas gdy sprzedane akcje wygrywające zyskały w tym samym okresie 11,6%. Na podstawie prostej symulacji Odean konkludował, że efekt dyspozycji pomniejsza stopę zwrotu inwestora w badanej grupie o około 4,4 pkt. proc.

⁶⁵ Omawiany efekt analizowany był na rynku amerykańskim. Stąd wspomniane kwestie podatkowe stosują się do prawa amerykańskiego, które rozróżnia między zyskami i stratami krótkoterminowymi (opodatkowane jak zwykły dochód) a długoterminowymi (opodatkowane niższą stopą). Por. (Odean, 1998a).

⁶⁶ Inne prace to: (Gerke i Bienert, 1993), (Lakonishok i Smidt, 1986), (Weber i Camerer, 1998).

Co ciekawe, efekt dyspozycji zanika w grudniu, kiedy inwestorzy częściej realizują straty niż zyski. Znajduje to uzasadnienie we wspomnianych wcześniej względach podatkowych. W grudniu inwestorzy mają bowiem ostatnią okazję do sprzedaży akcji ze stratą, o którą obniżona zostaje podstawa opodatkowania.

Badacze zgodnie dowodzą, że przedwczesna realizacja zysków i zwlekanie z realizacją strat prowadzą do nieoptymalnych wyników ekonomicznych. Stąd zachowania takie należy uznać za nieracjonalne.

4.1.2.3 Interpretacja

Aktywni uczestnicy rynku od dawna byli świadomi istnienia owego tajemniczego efektu. Nikt jednak nie starał się zgłębić jego przyczyn. *Prospect theory* zdaje się uzupełniać tę lukę, wytyczając kierunek dalszych badań nad źródłem awersji do strat. Badacze upatrują przyczyn efektu dyspozycji w zjawiskach takich jak: **(i)** mentalna księgowość (np. strata papierowa mniej „boli” niż strata zrealizowana i zadeklarowana), **(ii)** pragnienie uczucia dumy i podziwu (szybka realizacja zysków) oraz unikanie żalu (zwlekanie z realizacją strat), a także **(iii)** problem samokontroli, czyli intrapersonalny konflikt między jednostką planującą (ograniczanie strat) a działającą (unikanie realizacji strat)⁶⁷.

Oto cytat z wypowiedzi doświadczonego profesjonalisty (Shefrin i Statman, 1985): „Mam sztywną i mocną zasadę: nigdy nie pozwalam stratom w jednej transakcji przekroczyć 10% [ceny zakupu]. (...) [Inwestorzy] odmawiają przyjęcia strat... Kiedy przełamujesz lody jako młody makler, najtrudniejszą rzeczą, jakiej przyjdzie ci się nauczyć, jest przyznanie się do błędu. (...) Musisz być twardy, żeby przyznać się przed innymi (...). Tylko wtedy przeżyjesz i będziesz grał w tę grę następnego dnia”.

Mechanizmami broniącymi inwestora przed efektem dyspozycji są arbitralnie ustanowione granice wymuszające zamknięcie pozycji (np. *stop-loss*

⁶⁷ W tym miejscu zagadnienia te zostają jedynie zasygnalizowane. Dokładniejsza ich analiza przeprowadzona zostanie w odpowiednich częściach tej pracy. Por. podrozdział 4.3 dotyczący mentalnej księgowości oraz punkt 5.1.3.3 o problemie samokontroli.

dla strat, *on-stop/market-if-touched* dla zysków). Podobną funkcję pełni data 31 grudnia (ostatni dzień roku obrachunkowego), która czyni realizację strat mniej bolesną.

4.1.2.4 Znaczenie dla rynku

Aktywność uczestników rynku – w szczególności profesjonalnych i instytucjonalnych o krótkim horyzoncie inwestycyjnym – określa siłę oddziaływania efektu dyspozycji na kursy. Punkt odniesienia tych inwestorów związany jest bowiem z danym rynkiem⁶⁸ (z kursami aktualnie obowiązującymi).

Na poziomie zagregowanym efekt dyspozycji może być odpowiedzialny za pozytywną zależność między zmianami kursów a obrotami giełdowymi, zidentyfikowaną przez Lakonishoka i Smidta (1986). Przez wpływ na podaż aktywów, efekt dyspozycji może także powodować stabilizowanie się cen w okolicach poziomu, przy którym aktywność rynku była wysoka w poprzednim okresie. Jeśli szerokie grono inwestorów nabywa akcje przy danej cenie,

⁶⁸ Punkt odniesienia aktywnych krótkoterminowych uczestników znajduje się na rynku, na którym działają. Jest nim cena, po której kupili dane aktywa. Grupa inwestorów o krótkim horyzoncie inwestycyjnym (od kilku minut do paru dni) obejmuje głównie dealerów/maklerów zatrudnionych w bankach i domach maklerskich. U tych jednostek efekt dyspozycji jest najsilniejszy ze względu na presję osiągania jak najlepszych wyników. Por. np. (Coval i Shumway, 2001). Nieco dłuższa perspektywa (od kilku dni do około roku) cechuje inwestorów średnioterminowych – głównie managerów funduszy inwestycyjnych oraz niektórych inwestorów indywidualnych. Również oni podlegają efektowi dyspozycji (Odean, 1998a), oceniając wartość swoich pozycji względem ceny zakupu, która stanowi dla nich punkt odniesienia (niekiedy rolę punktu odniesienia dla managerów funduszy pełni *benchmark*). Trzecia grupa to inwestorzy długoterminowi (horyzont powyżej roku), których aktywność rynkowa ma charakter sporadyczny i determinowana jest względami strategicznymi. Przykładem takich uczestników rynku są banki centralne, firmy dokonujące fuzji lub przejęć czy też osoby, które wykonały swoje opcje pracownicze. Inwestorzy ci, jakkolwiek ich motywacja jest różna, mają jedną wspólną cechę: ich punkt odniesienia nie znajduje się na rynku, w który się zaangażowali. Pierwszoplanowym celem banku centralnego interweniującego na rynku walutowym nie jest zysk, lecz stabilizacja kursu walutowego. Dla europejskiej spółki dokonującej przejęcia w USA punktem odniesienia może być np. cena firmy przejmowanej wyrażona w dolarach. Spółka ta, kupując dolary za euro, nie jest zainteresowana osiągnięciem zysku w transakcji na rynku walutowym, lecz utrzymaniem skalkulowanej ceny przejmowanych aktywów. Dla pracownika punktem odniesienia jest cena rozliczenia opcji (*strike price*), powyżej której sprzedaż akcji na rynku kapitałowym jest zawsze opłacalna. Dzięki procesowi adaptacji (por. punkt 5.2.1) inwestorzy ci nie oceniają wartości swych inwestycji w relacji do ceny zakupu aktywów. Efekt dyspozycji jest u nich najsłabszy. Wynika stąd ogólna zależność między horyzontem inwestycyjnym a siłą efektu dyspozycji: im krótszy horyzont, tym większa zależność oceny wartości aktywów od ceny ich zakupu i tym silniejszy efekt dyspozycji. Innymi słowy, efekt dyspozycji słabnie wraz z wydłużaniem się perspektywy inwestycyjnej.

należy przypuszczać, że wszyscy oni będą mieli ten sam punkt odniesienia (kurs, po którym kupili). Spadek ceny poniżej punktu odniesienia wywołuje niechęć inwestorów do sprzedaży ze stratą, co redukuje podaż akcji i jednocześnie spowalnia dalszy spadek kursów. I odwrotnie: kiedy cena wzrasta względem punktu odniesienia, wówczas skłonność do szybkiej sprzedaży zwiększa podaż akcji i ogranicza dalszy wzrost ich kursu.

Należy zaznaczyć, iż badania nad wpływem efektu dyspozycji na zachowania zagregowanego rynku znajdują się w bardzo wczesnej fazie. Można jednak przypuszczać, iż efekt ten jest jednym z najsilniejszych mechanizmów determinujących ruchy cen giełdowych⁶⁹. Świadomość jego występowania w skali mikro pozwala na postawienie kilku hipotez o „makrokompozycji” rynku. I tak, fazy wysokich obrotów giełdowych wskazują, iż wielu inwestorów, otwierających w tym momencie pozycję, ma ten sam punkt odniesienia (cena, po której kupują bądź sprzedają). Ich percepcja zysków i strat jest zatem jednorodna. Rozważmy rynek obniżających się kursów. Przy długotrwałym trendzie spadkowym na rynku wzrasta liczba inwestorów tracących na pozycji długiej – „uwięzionych” za stratą, której nie chcą zrealizować. Natomiast inwestorzy wygrywający opuszczają rynek, zamykając pozycję (krótką), aby zagwarantować sobie zyski. Prawdopodobieństwo ich ponownego „wejścia” zmniejsza się wraz z oddalaniem się ceny od poziomu, przy którym poprzednio zamknęli pozycję (wydłużaniem się danego trendu)⁷⁰. Jednocześnie inwestorzy ponoszący straty nie reagują. Skłonność do ryzyka motywuje ich do utrzymywania otwartej pozycji w oczekiwaniu na zwrot tendencji rynkowej. Oznaką takiej sytuacji może być niski poziom obrotów połączony z niewielką zmiennością kursów (rynek nie reaguje).

⁶⁹ Pominęto tu procesy adaptacyjne, jakim podlega percepcja punktu odniesienia inwestora. Adaptacja stanie się przedmiotem dyskusji w rozdziale 5. Już teraz warto jednak wspomnieć, iż w połączeniu z efektem dyspozycji stanowi ona niezwykle istotny mechanizm warunkujący aktywność na rynkach finansowych.

⁷⁰ Związane jest to z powszechną wiarą w regresję kursów do długoterminowej średniej: „co spadło musi kiedyś wzrosnąć”. Inwestor, który realizuje zyski (zamyka pozycję) liczy, że tendencja kształtująca się obecnie na rynku ulegnie niebawem odwróceniu. Por. rozważania na temat losowości i regresji do średniej w punktach 2.2.2.2 oraz 2.2.2.3.

Wyobraźmy sobie jednak, że trend rzeczywiście się odwraca, zaś kurs giełdowy zrównuje się z przeciętnym rynkowym punktem odniesienia właściwym większości uczestników. Jak wiadomo, zrównanie się ceny z punktem odniesienia implikuje neutralność wobec ryzyka. W kontekście *prospect theory* inwestor znajduje się na funkcji $v(x)$ w punkcie, dla którego wartość wynosi zero (por. rysunek 3.1). Teraz może zatem bez szwanku zamknąć pozycję. Oznacza to wzrost podaży, której towarzyszą większe obroty i wyższa zmienność cen.

Powyższe rozważania mają charakter ćwiczenia myślowego i nie zostały dotychczas wsparte dogłębnymi badaniami. Cenna uwaga kryje się jednak w stwierdzeniu, że dla lepszego przewidywania zagregowanych zachowań rynku konieczna jest znajomość rozkładu (funkcji gęstości) punktów odniesienia inwestorów. Wiedza taka pozwala na zdefiniowanie kierunku i siły działania efektu dyspozycji jako jednego z podstawowych wyznaczników aktywności rynku.

4.1.3 *Efekt utopionych kosztów*

Efekt utopionych kosztów (*sunk-cost effect*) jest pochodną omówionego powyżej efektu dyspozycji. Można o nim myśleć jako o uogólnieniu efektu dyspozycji dla sekwencji decyzji inwestycyjnych. Jednakże utopione koszty to także zjawisko o poważnych konsekwencjach mikro- i makroekonomicznych oraz społecznych i politycznych. Charakterystyczne jest to, że jego siła bywa doceniana dopiero *ex post*. Ze względu na nietrywialne implikacje efekt ten wymaga szerokiego omówienia.

Zgodnie z teorią ekonomii jedynie analiza teraźniejszych oraz przyszłych kosztów i korzyści powinna mieć wpływ na decyzje. Koszty historyczne nie odgrywają roli⁷¹. Należy jednak zadać pytanie, czy dla przeciętnego zjadacza chleba poniesione koszty również popadają w niepamięć. W decyzjach zarówno jednostek, jak i firm przeszłe koszty są raczej odpowiednikiem opóźnionej zmiennej z następstwami w teraźniejszości. Pod pojęciem efektu uto-

⁷¹ Pomija się tu znaczenie analizy finansowej, która bazuje na danych historycznych.

pionych kosztów należy zatem rozumieć zjawisko wpływu kosztów poniesionych na kolejne decyzje inwestycyjne. Zależność tę obrazuje następujący prosty przykład:

Firma Omega pracuje nad projektem wartym 10 milionów złotych, z czego 5 milionów złotych zostało już zainwestowanych. Jego uwieńczeniem ma być nowoczesny silnik hybrydowy dla pojazdów elektrycznych. Prototypowy model przechodzi właśnie pierwsze testy. W tym samym czasie manager projektu odkrywa, ku swemu zaskoczeniu, że firma Alfa jest już gotowa do wprowadzenia podobnego produktu na rynek. Konkurencyjna konstrukcja Alfy jest lżejsza i mniejsza – predestynowana do sukcesu komercyjnego. Manager Omegi stoi przed dylematem: przerwać inwestycję i rozsądniej zainwestować pozostałe 5 milionów złotych, czy kontynuować projekt, którego szanse sukcesu są niewielkie.

Ludzka skłonność do przeprowadzania spraw „od początku do końca” nakazuje, by projekt taki został ukończony. Eksperymenty dwojga psychologów Arkesa i Blumer (1985) pokazują, że około 85% respondentów decyduje się na kontynuację projektu analogicznego do powyższego przykładu. Porzucenie inwestycji teraz oznaczałoby spisanie 5 milionów złotych na pewną stratę. Zrealizowane straty bolą; stąd zgodnie z *prospect theory* naturalne jest dążenie do ich (tymczasowego) uniknięcia.

Badania nad efektem utopionych kosztów⁷² pozwalają na postawienie następujących hipotez (Thaler, 1999):

- (1) Wysilek i czas poświęcone na „ratowanie” projektu są tym większe, im większe koszty zostały już poniesione.
- (2) Za opcję mniej bolesną niż całkowite zarzucenie projektu uznaje się jego zawieszenie (zamrożenie środków).
- (3) Wraz z upływem czasu, niezależnie od utopionych kosztów, projekt zostaje uznany za w pełni zamortyzowany i ostatecznie anulowany.

⁷² Inne prace związane z tym efektem to np. (Thaler, 1980), (Gourville i Soman, 1998).

Dzięki procesowi adaptacji⁷³ jednostki stopniowo godzą się ze stratami i ostatecznie je ignorują.

4.1.3.1 Z życia wzięte

Anegdoty podobne do przytoczonej wyżej opowieści można traktować z pobłażaniem, sądząc, że analogiczne sytuacje nie zdarzą się nigdy w świecie wielkich pieniędzy i prawdziwych profesjonalistów. O namacalnych skutkach efektu utopionych kosztów przekonują jednak zdarzenia z życia wzięte.

Jako przykład Shefrin (2000, s. 24) przedstawia w swojej książce historię komputerowego potentata – firmy Apple. Firma ta jako pierwszy producent sprzętu informatycznego starała się wylansować urządzenie typu *handheld*. Ów produkt miał pełnić funkcję „osobistego cyfrowego asystenta”; nadano mu dumną nazwę „Newton”. Prezes Apple John Sculley był szczególnie oddany Newtonowi, wierząc, iż nowy produkt zrewolucjonizuje podejście do komputeryzacji. Projekt wystartował w 1987 roku, zaś w roku 1993 na rynku ukazał się pierwszy Newton. Niestety zarówno cena urządzenia (1000 dolarów), jak i jego funkcjonalność (błędy w rozpoznawaniu pisma) nie przyczyniły się do entuzjastycznego zainteresowania konsumentów, na które liczone. Po roku rozczarowujących danych o sprzedaży oczywiste stało się, że projekt przynosi straty. Dodatkowo na rynek agresywnie zaczął wchodzić konkurencyjny Palm – produkt firmy 3Com. Praca nad Newtonem nie została mimo to przerwana. Wręcz przeciwnie, projekt z roku na rok angażował kolejne ogromne inwestycje w rozbudowę funkcjonalności, zmianę organizacji produkcji, utworzenie specjalnego działu... Bezskutecznie – Newton nie zarobił ani centa. Nowi prezesowie przychodzili i odchodzili. W roku 1998, ponad 10 lat po uruchomieniu projektu, nowy prezes zdecydował kategorycznie o zarzuceniu idei Newtona.

Apple nie jest bynajmniej wyjątkowym przypadkiem. Utopione koszty w sposób jeszcze bardziej drastyczny dotknęły bank Barings PLC, kładąc kres

⁷³ O znaczeniu adaptacji będzie mowa w rozdziale 5, poświęconym problematyce międzyokresowej i czynnikowi czasu.

jego 232-letniej historii. Sprawcą była jedna osoba – Nicholas Leeson. W roku 1992 Leeson zaczął angażować się w nieautoryzowany handel, zajmując coraz większe niezrównoważone pozycje na rynku terminowym. Stopniowo zwiększał obroty i ryzyko swego portfela, uzależniając jego rentowność od ruchów Nikkei 225. Trzęsienie ziemi w Kobe, powodując znaczny spadek Nikkei, pociągnęło za sobą ogromne straty dla Leesona. Ten jednak, zamiast ograniczyć swoją pozycję, zdecydował się na zwiększenie inwestycji w nadziei na rychłe ożywienie rynku. Ożywienie jednak nie przyszło, zaś działalność Leesona zaowocowała stratą 1,4 miliarda dolarów w 1995 roku. Barings zbankrutował i został sprzedany holenderskiemu ING za symbolicznego funta.

Przytoczone historie mają jeden wspólny element: początkowe straty spotęgowane zostają przez serię kolejnych bezsensownych inwestycji. Przykłady te pozwalają na wyciągnięcie pouczających wniosków o sile efektu utopionych kosztów oraz braku samokontroli na poziomie jednostek i organizacji. Warto zauważyć, że kontynuowanie inwestycji skazanych na niepowodzenie jest równoznaczne z destrukcją wartości firmy dla jej właścicieli (inwestorów – w przypadku spółek giełdowych). Stąd efekt utopionych kosztów w lawinowy sposób rzutuje na sytuację ekonomiczną firm i jednostek.

4.1.3.2 Znaczenie dla rynku

Efekt utopionych kosztów nie oszczędza aktywnych uczestników rynków finansowych, szczególnie tych o krótkim horyzoncie inwestycyjnym.

W swojej pracy z 1979 roku Kahneman i Tversky wskazują na interesujące zachowanie osób uczestniczących w grach losowych: jeśli gracz najpierw stracił, obstawiana przezeń kwota wzrasta wraz ze zbliżaniem się końca rozgrywek. Jego ryzykanctwo powodowane jest brakiem adaptacji do strat lub utratą poprzednich niezrealizowanych zysków.

Jak powyższa obserwacja ma się do zachowań osób działających na rynkach finansowych? Doświadczeni uczestnicy i obserwatorzy rynku twier-

dza⁷⁴, że aktywny inwestor (szczególnie o krótkim horyzoncie inwestycyjnym) mentalnie identyfikuje się z hazardystą. Stąd niezwykle częstym posunięciem jest podejmowanie dodatkowo ryzykownych inwestycji pod koniec dnia handlowego z zamiarem „nadgonienia” strat. Potwierdza to najnowsze badanie zachowań market-makerów rynku futures z CBOT (*Chicago Board of Trade*) przeprowadzone w 2001 roku przez profesorów Covala i Shumway’a z Uniwersytetu w Michigan. Coval i Shumway (2001) obserwują tzw. „efekt popołudniowy”: inwestorzy, którzy ponieśli straty w początkowych godzinach handlu, decydują się na podjęcie dodatkowego ryzyka w drugiej połowie dnia. Zachowanie takie znacznie rzadziej występuje w przypadku inwestorów wygrywających. W wyniku porannych porażek przegrywający ponadprzeciętnie zwiększają ryzyko, wartość oraz częstotliwość dokonywanych transakcji⁷⁵. Oczywiście wyniki badania wskazują na działanie efektu utopionych kosztów przy „mikrohoryzoncie” inwestycyjnym. Należy jednak zauważyć, że rynek futures na CBOT to jeden z najbardziej płynnych i konkurencyjnych rynków na świecie. Stąd rozsądnym wydaje się przypuszczenie, iż analiza rynków o niższej częstotliwości transakcji oraz dłuższym horyzoncie inwestycyjnym dałaby podobne rezultaty.

Podsumowując, warto zastanowić się nad taktyką inwestowania, która leży u podstaw owego ryzykanctwa. Wy tłumaczeniem może być tak zwane „uśrednianie” (*averaging*). Polega ono na zwiększaniu zajmowanej pozycji poprzez dokupowanie taniejących akcji danej spółki⁷⁶. W ten sposób średnia cena zakupu (stanowiąca nowy punkt odniesienia) zostaje obniżona w kie-

⁷⁴ Autorka miała okazję przedyskutować ten problem z zawodowymi inwestorami rynku walutowego i kapitałowego.

⁷⁵ Coval i Shumway analizowali dane z 1998 roku, obejmujące 5 milionów transakcji. Ich wyniki pozwalają stwierdzić, że prawdopodobieństwo podjęcia ponadprzeciętnie wysokiego ryzyka przez inwestora przegrywającego jest o 16 pkt. proc. wyższe niż w przypadku inwestora wygrywającego. Podobnie rozpatrując zwiększanie częstotliwości transakcji i ich wartości, prawdopodobieństwo to jest wyższe o odpowiednio 28,6 pkt. proc. oraz 9,5 pkt. proc. dla inwestora, który odnotował straty (Coval i Shumway, 2001).

⁷⁶ Obrazuje to następujący przykład: inwestor kupuje 100 akcji danej spółki po 100 j.p. Następnie cena akcji spada do 80 j.p. Inwestor, wierząc, że redukcja ta jest krótkotrwała, decyduje się na zakup dodatkowo 100 akcji. Pozwala to na obniżenie średniej ceny zakupu do 90 j.p., gdyż $(100 \text{ j.p.} \cdot 100 + 80 \text{ j.p.} \cdot 100) / 200 = 90 \text{ j.p.}$

runku aktualnego kursu, inwestor zaś ulega złudzeniu, że straty łatwo będzie teraz odrobić. Strategia ta, jakkolwiek bywa skuteczna, podnosi ryzyko portfela i poprzez zwiększenie pozycji naraża inwestora na wielokrotnione straty.

4.1.4 Efekt obdarowania oraz skłonność do zachowania *status quo*

Swoim artykułem z 1980 roku Richard Thaler poruszył świat ekonomistów. Do owego czasu panowało przekonanie, że koszty utraconych możliwości (*opportunity costs*) oraz faktycznie poniesione (*out-of-pocket costs*) są ekwiwalentne, podobnie jak skłonność do zapłaty danej ceny za dobro równa jest skłonności do zaakceptowania tej ceny przez sprzedającego. Thaler (1980) stwierdził, że tak nie jest, dając początek rozważaniom nad efektem obdarowania.

Należy wspomnieć, że obserwacja Thaler, zgłębiona w następnych latach we współpracy z Kahnemanem i Knetschem (Kahneman, Knetsch i Thaler, 1991), odnosi się zasadniczo do dóbr użytkowych, pozbawionych ryzyka i nieprzeznaczonych do szybkiego obrotu. Niektórzy autorzy błędnie starają się rozszerzyć ją na rynek aktywów ryzykownych, co jest niezgodne z omawianą koncepcją. Zjawiskiem pokrewnym, obserwowanym także dla dóbr ryzykownych, jest natomiast skłonność do zachowania *status quo* opisana przez Samuelsona i Zeckhausera (1988). Stanie się ona przedmiotem dyskusji w punkcie 4.1.4.2.

4.1.4.1 Efekt obdarowania

Zjawisko polegające na tym, że ludzie żądają więcej za rezygnację z posiadania dobra, niż skłonni byliby za nie zapłacić, nazwane zostało efektem obdarowania (*endowment effect*).

Jeden z pierwszych eksperymentów potwierdzających efekt obdarowania przeprowadzony został przez Knetscha i Sindena (1984). Uczestnicy badania, podzieleni na dwie grupy, obdarowani zostali odpowiednio kuponem

loteryjnym⁷⁷ lub kwotą 2 dolarów. Następnie pierwszą grupę zachęcono do sprzedaży kuponu za 2 dolary, zaś drugą do jego zakupu po tej samej cenie. Oto rezultaty: 82% członków pierwszej grupy zdecydowało się zatrzymać kupon, zaś jedynie 28% reprezentantów drugiej grupy wyraziło chęć zakupu. Dalsze eksperymenty potwierdziły istnienie niezwykle dużych rozbieżności między wartością postrzeganą przez posiadacza (skłonność do przyjęcia ceny, SPC) a wyceną potencjalnego nabywcy (skłonność do zapłaty ceny, SZC)⁷⁸. Kahneman, Knetsch i Thaler (1991) udokumentowali rozbieżność między SPC a SZC rzędu od 2:1 do 3:1⁷⁹.

Ta niezgodność wydaje się ujawniać natychmiast po tym, jak dana osoba wejdzie w posiadanie artykułu. Nie jest ona bynajmniej ograniczona do przedmiotów cennych, lecz rozszerza się także na mało wartościowe towary codziennego użytku. Co więcej, kiedy nabywający i sprzedawcy zamieniają się rolami, uprzedni sprzedawcy nie są skłonni zapłacić sumy, jakiej wcześniej sami żądali.

Efekt obdarowania bywa przypisywany awersji do strat znanej z *prospect theory*: ujemna użyteczność płynąca z rezygnacji z przedmiotu jest wyższa (co do wartości bezwzględnej) od użyteczności jego nabycia. Wynikające stąd rozbieżności między ceną zakupu a ceną sprzedaży mogą przedstawiać znaczny problem dla praktyków dokonujących analizy kosztów i korzyści (*cost-benefit analysis*). Przed nimi stoi bowiem wyzwanie przypisania wartości pieniężnej dobrom nie będącym przedmiotem obrotu na rynku.

4.1.4.2 Skłonność do zachowania *status quo*

W literaturze przedmiotu efekt obdarowania uznaje się za część bardziej ogólnego zjawiska – skłonności do zachowania *status quo* (Schwartz, 1998,

⁷⁷ Sic! Kupon loteryjny jest, co prawda, dobrem ryzykownym, lecz nieprzeznaczonym do szybkiego obrotu. Ponadto w omawianym eksperymencie respondenci otrzymali kupony gratis. Wspomniana w uwagach wstępnych niekonsekwencja niektórych badaczy nie stosuje się zatem do tego przypadku.

⁷⁸ Dotyczy to szczególnie artykułów zakupionych przez jednostkę dla jej własnych celów.

⁷⁹ Badacze wyeliminowali wpływ kosztów transakcyjnych, efektu dochodowego i inklinacje sprzedających do negocjowania ceny.

s. 90). W przeciwieństwie do efektu obdarowania skłonność do *status quo* obejmuje problemy wymiany dóbr przeznaczonych do szybkiego obrotu, a także wszelkie transakcje o charakterze czysto finansowym.

Podobnie jak efekt obdarowania skłonność do zachowania *status quo* stanowi bezpośredni rezultat awersji do strat. Potencjalne niekorzyści związane z utratą stanu bieżącego przeważają nad przewidywanymi zaletami owej zmiany.

Działanie tego efektu zademonstrowali Samuelson i Zeckhauser⁸⁰ (1988). W przeprowadzonym przez nich badaniu dwie grupy respondentów postawione zostały przed różnymi wersjami hipotetycznej sytuacji decyzyjnej:

Wersja 1 (Status quo nie zdefiniowany)

Jesteś wiernym czytelnikiem gazety finansowej, jednak do niedawna nie miałeś środków, które mógłbyś zainwestować. Nagle otrzymałeś w spadku po wuju znaczną sumę pieniędzy. Rozważasz inwestycję w różne portfele: (i) akcje spółki o umiarkowanym ryzyku, (ii) akcje spółki o wysokim poziomie ryzyka, (iii) bony skarbowe, (iv) obligacje komunalne.

Wersja 2 (Jedna z opcji inwestycyjnych zdefiniowana jako status quo; zdanie wprowadzające nie ulega zmianie)

...Właśnie niedawno otrzymałeś w spadku po wuju portfel zawierający gotówkę i aktywa ryzykowne. Znaczna część portfela jest ulokowana w akcjach spółki o umiarkowanym ryzyku... (W przypadku zmiany składu portfela pomijamy podatki oraz prowizję brokerską.)

Opierając się na powyższej metodologii, Samuelson i Zeckhauser przeanalizowali liczne scenariusze. Dało im to możliwość oszacowania prawdopodobieństwa wyboru danej konstrukcji portfela w zależności od sposobu sformułowania problemu decyzyjnego. Uzyskane przez nich wyniki wskazują na znacznie większą popularność wariantów zdefiniowanych za pomocą *status quo* w porównaniu z innymi opcjami inwestycyjnymi. W wersji 2

⁸⁰ Eksperyment Samuelsona i Zeckhausera opiera się na hipotetycznej sytuacji. Badania skłonności do *status quo* w rzeczywistych wyborach ekonomicznych przeprowadzili Hartman, Doane i Woo (1991) a także Hershey, Johnson, Meszaros i Robinson (1990). Ich wyniki potwierdzają działanie opisywanego efektu.

(*status quo* zdefiniowany) połowa respondentów zdecydowała się na pozostawienie kapitału w akcjach o umiarkowanym ryzyku. Natomiast wybory portfeli w wersji 1 (brak *status quo*) charakteryzowały się dużą różnorodnością; akcje spółki o umiarkowanym ryzyku wybrała niespełna trzecia część respondentów.

Wnioski płynące z obserwowanego efektu są niezwykle ważne w czasach, kiedy ciężar zapewnienia sobie spokojnej starości coraz częściej spoczywa na jednostce. Przejście firm od systemów emerytalnych o zdefiniowanych wypłatach (*defined benefit plans*) do systemów o zdefiniowanej składce (*defined contribution plans*) nakłada na uczestnika funduszu obowiązek podjęcia szeregu decyzji: (i) o przystąpieniu do funduszu, (ii) o ilości środków ulokowanych w funduszu, (iii) o strukturze indywidualnego portfela. Świadomość występowania efektu *status quo* pozwala decydentom na skonstruowanie systemu emerytalnego zapewniającego optymalną kwotę oszczędności. Niezwykle pouczającą obserwację poczynił w tej dziedzinie Nofsinger (2001). Jedna z dużych firm amerykańskich zreformowała system uczestnictwa w korporacyjnym funduszu emerytalnym. W starym systemie pracownik sam decydował o przystąpieniu do funduszu i o sposobie alokacji wpłat. Zmiana polegała na wprowadzeniu: (i) domyślnego uczestnictwa w funduszu, (ii) zdefiniowanej wielkości wpłat oraz (iii) określonego sposobu ich alokacji. Uczestnik mógł zadeklarować rezygnację z uczestnictwa w funduszu, bądź też samodzielnie zmodyfikować wersję domyślną. Konsekwencje reformy są zgodne z hipotezą o działaniu efektu *status quo*. W starym systemie jedynie 37% pracowników decydowało się na przystąpienie do funduszu; wprowadzenie nowego systemu zaowocowało wzrostem tego odsetka do 86% (jedynie 14% pracowników zrezygnowało z opcji domyślnej). Skłonność do zachowania *status quo* wpłynęła również na sposób konstrukcji portfela; w domyśle – na jego rentowność.

Warto zauważyć, iż wpływ *status quo* na decyzje jednostki jest wyrazem zależności preferencji od formy (w rozdziale 3 nazwano to kadrowaniem,

por. punkt 3.3.5). Jak przewiduje *prospect theory*, rzeczywiste wybory ludzi dowodzą pogwałcenia aksjomatu niezmienności preferencji.

4.2 Mentalna księgowość

Wprowadzenie pojęcia mentalnej księgowości (*mental accounting*) do literatury ekonomii i finansów jest wspólną zasługą Thaler (1980, 1985) oraz Kahnemana i Tversky'ego (1984). Koncepcję tę należy uznać za wyraz odwagi jej autorów; podważa ona bowiem jedną z najświętszych zasad w ekonomii: zasadę substytucyjności pieniądza, mówiącą o tym, że „pieniądz nie ma metki” (*money has no label*) – źródło, z którego pochodzi, nie ma wpływu na sposób jego alokacji i wydatkowania.

4.2.1 Definicja i znaczenie

„Mentalna księgowość stanowi zespół myślowych operacji dokonywanych przez jednostki i gospodarstwa domowe w celu organizowania, oceny i analizy transakcji finansowych” (Thaler, 1999).

O księgowości mentalnej (myślowej) można myśleć przez odniesienie jej do rachunkowości stosowanej w firmach w celu rejestrowania transakcji gospodarczych. W przeciwieństwie jednak do rachunkowości finansowej, księgowość mentalna nie rządzi się normatywnymi zasadami. Określające ją reguły można poznać jedynie przez obserwację zachowań ludzkich.

Centralnym elementem koncepcji jest funkcja wartości opisana w *prospect theory* wraz z integralnymi jej własnościami: punktem odniesienia, malejącą wrażliwością na dodatkowe zmiany oraz awersją do strat. Stanowi ona odpowiednik mentalnego konta, otwieranego dla każdej zawieranej transakcji. Zastosowanie funkcji wartości pozwala opisać sposób postrzegania, kodowania i oceny zdarzeń ekonomicznych przez jednostki. Co więcej, połączenie myślowych kont z własnościami funkcji wartości implikuje, że stosunek danej jednostki do ryzyka zmienia się w zależności od transakcji (dokładniej: w zależności od mentalnego konta, na którym transakcja została zarejestrowana). W przypadku inwestora oznacza to dyspozycję do zachowań ryzykownych w transakcjach przynoszących straty przy jednoczesnej awersji

do ryzyka w transakcjach zyskowych. W ten sposób podważony zostaje paradygmat, jakoby stosunek do ryzyka był stałą cechą osobowości.

4.2.2 *Hedonistyczne kadrowanie zysków i strat*

Kadrowanie to jeden z kluczowych procesów, jakim podlega osoba formująca preferencje i podejmująca decyzje w myśl *prospect theory*.

4.2.2.1 Zakres percepcji zysków i strat

Posługując się funkcją wartości w analizie zachowań inwestycyjnych, należy sprecyzować zakres zmian majątkowych, wobec których inwestor wykazuje awersję do strat. Za szeroko zdefiniowane uznaje się zyski i straty rozważane w kategoriach całego majątku. Wąska definicja obejmuje natomiast zmiany wartości wyizolowanych komponentów majątku, jak na przykład portfela akcji czy też pojedynczego papieru wartościowego posiadanego przez inwestora.

Liczne badania⁸¹ dowodzą, że proces mentalnej księgowości inwestora przebiega zgodnie z wąską definicją zysków i strat (Barberis i Huang, 2001). Innymi słowy, nośnikami użyteczności są dlań zmiany wartości poszczególnych aktywów w portfelu. Taka percepcja wyników finansowych zwana jest wąskim kadrowaniem (*narrow framing*).

W ramach wąskiego kadrowania inwestor prowadzi dwa rodzaje księgowości: (i) księgowość poszczególnych transakcji oraz (ii) księgowość portfela. Zajęcie pozycji na rynku równoznaczne jest z otwarciem indywidualnego konta. Punktem odniesienia jest tu cena zakupu akcji; kolejne jej zmiany określają odpowiednio położenie inwestora na funkcji wartości, determinując jednocześnie jego stosunek do ryzyka. Zamknięcie pozycji oznacza przeniesienie zysku/straty z (myślowego) konta indywidualnej inwestycji na nadrzędne konto wyniku z portfela.

⁸¹ Dowody posługiwania się różnymi formami wąskiego kadrowania zysków i strat można znaleźć w (Redelmeier i Tversky, 1992), (Kahneman i Lovallo, 1993), (Benartzi i Thaler, 1999).

Przedstawiony powyżej scenariusz dotyczy najczęściej inwestorów indywidualnych. Swoistą osobliwość stanowi natomiast mentalna księgowość prowadzona przez profesjonalnych uczestników rynku czy managerów funduszy, których wyniki odnoszone są do kryterium zewnętrznego (*benchmark*). Otwierając nowy okres rozliczeniowy, profesjonalisci z założenia znajdują się na funkcji wartości w dziedzinie strat, postrzegając poziom zadany poprzez *benchmark* za punkt odniesienia. Stanowi to bodziec do zwiększania ryzyka zajmowanych pozycji.

4.2.2.2 Umysłowa arytmetyka: integracja – segregacja

Portfel inwestora, składający się z różnych aktywów, daje szerokie pole do uprawiania „kreatywnej księgowości”. Niektóre transakcje okazują się „strzałem w dziesiątkę”, do innych wstyd się przyznać. Powstaje zatem naturalne pytanie: jak połączyć różne wyniki finansowe w obrębie jednego portfela?

Słusznym wydaje się założenie, że jednostki kodują kombinacje wyników finansowych w sposób, jaki gwarantuje im największe zadowolenie tudzież ogranicza odczucia negatywne. Proces ten przyjęło się nazywać hedonistycznym kadrowaniem (*hedonic framing*).

Arytmetyka umysłowa polega na łączeniu zdarzeń (x,y) tak, aby rezultat przedstawiał się jak najkorzystniej. Wyniki finansowe mogą być oceniane łącznie – integrowane: $v(x+y)$ bądź osobno – segregowane: $v(x)+v(y)$. Zależnie od przyjętej konwencji, jednostka jest zdolna manipulować odczuwaną wartością v .

Kształt funkcji wartości pozwala wydedukować zasady hedonistycznego kadrowania zysków i strat. Niech $x > 0$ i $y > 0$. Rozważmy cztery możliwe kombinacje wyników:

- (1) dwa zyski (x,y) ; ze względu na wklęsły kształt funkcji wartości w dziedzinie zysków: $v(x) + v(y) > v(x+y)$,
- (2) dwie straty $(-x,-y)$; ze względu na wypukłość funkcji wartości w dziedzinie strat: $v(-x) + v(-y) < v(-x-y)$,

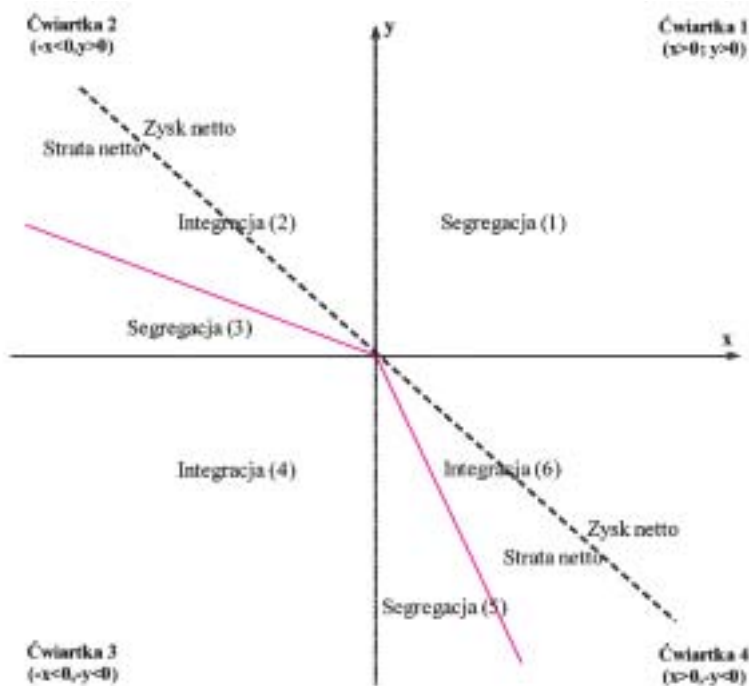
(3) zysk netto $(x, -y)$, gdzie $x > y$; otrzymujemy:

$$v(x) + v(-y) < v(x - y),$$

(4) strata netto $(x, -y)$, gdzie $x < y$; wówczas nie można określić zwrotu nierówności bez dodatkowych informacji. Im mniejsze x w porównaniu z y , tym bardziej prawdopodobne jest, że $v(x) + v(-y) > v(x - y)$ ze względu na malejącą wrażliwość na dodatkowe zmiany dla głębokich strat (obrazowo: x jest traktowane jako „dobra strona” w nieszczęściu y).

Przedstawione zasady hedonistycznego kadrowania pozwalają sformułować zestaw rekomendacji (rysunek 4.1): (i) segreguj zyski, (ii) integruj straty, (iii) niweluj mniejsze straty większymi zyskami, (iv) segreguj małe zyski z dużymi stratami.

Rysunek 4.1: Hedonistyczne kadrowanie zysków i strat w obrębie portfela



Źródło: Opracowanie własne na podstawie (Shefrin i Statman, 1984).

4.2.2.3 Dowody empiryczne

W kontekście przedstawionych zasad powstaje pytanie, czy intuicja zaczerpnięta z *prospect theory* ma zastosowanie w rzeczywistości. Pierwszy eksperyment laboratoryjny w dziedzinie preferowanych metod kadrowania zysków i strat pochodzi od Thaler (1985). Metodologia badawcza oraz uzyskane wyniki zostaną teraz pokrótce zaprezentowane⁸².

Uczestnicy eksperymentu (87 studentów Uniwersytetu Cornella w USA) poproszeni zostali o porównanie sytuacji, jakie przydarzyły się dwóm osobom: A oraz B i wydanie werdyktu, która z osób jest bardziej zadowolona.

*Sytuacja 1: Osoba A wygrywa dwie loterie z wypłatami odpowiednio 50 j.p. i 25 j.p. B wygrywa pojedynczą loterię wypłacającą 75 j.p. Która z osób czuje się **lepiej**?*

A: **64,4%** B: 18,4% Bez różnicy: 17,2%

*Sytuacja 2: Tego samego dnia osoba A najpierw ponosi stratę 100 j.p., a następnie 50 j.p. B za jednym zamachem traci 150 j.p. Która z osób czuje się **gorzej**?*

A: **75,9%** B: 16,1% Bez różnicy: 8,0%

*Sytuacja 3: Osoba A tego samego dnia traci w wyniku złego zrządzenia losu 200 j.p. i wygrywa w zakładach totalizatora sportowego 25 j.p. B traci w wyniku złego zrządzenia losu 175 j.p. Która z osób czuje się **gorzej**?*

A: 21,8% B: **72,4%** Bez różnicy: 5,8%

*Sytuacja 4: Osoba A wygrywa 100 j.p. na loterii i tego samego dnia traci przez złe zrządzenie losu 80 j.p. B wygrywa na loterii 20 j.p. Która z osób czuje się **lepiej**?*

A: 25,3% B: **70,0%** Bez różnicy: 4,7%

Rezultaty badania zgodne są z zasadami hedonistycznego kadrowania. Należy jednak zauważyć, że sposób prezentacji problemu został narzucony z zewnątrz. Respondenci nie musieli poddawać go samodzielnej „edycji”; wyrażali jedynie swoje preferencje co do optymalnego sposobu kadrowania.

⁸² Przytoczone poniżej sytuacje przedstawiają logikę badawczą stosowaną przez Thaler, nie są jednak wiernym tłumaczeniem treści przedstawionej respondentom.

Thaler słusznie podkreśla, że istnieją granice „samooszustwa”. Trudno przypuszczać, iż jednostki zdolne są do segregowania jednorazowego zysku 50 j.p. jako 2 razy po 25 j.p. lub w wersji jeszcze bardziej ekstremalnej jako 50 razy po 1 j.p. Tak doskonałej umiejętności hedonistycznego edytowania ludzi nie posiadają. Natomiast hipoteza segregacji tudzież integracji zysków i strat ma jak najbardziej zastosowanie w rzeczywistości. Inwestor, który poniósł stratę 10 000 j.p. na inwestycji w akcje, lecz jednocześnie zarobił 15 000 j.p. na rynku złota, z dużym prawdopodobieństwem przeprowadzi prosty zabieg integracji obu wyników, z satysfakcją stwierdzając zysk 5000 j.p.

4.2.3 Zastosowania

4.2.3.1 Zyski papierowe – papierowe straty

W systemie mentalnej księgowości decyzja o zamknięciu konta lub pozostawieniu go otwartym ma charakter uznaniowy. Rozważmy sytuację zakupu 100 akcji po cenie jednostkowej 10 j.p. Początkowa wartość inwestycji (punkt odniesienia) ulega zmianom *in plus* oraz *in minus* wraz z fluktuacjami kursu, generując papierowe zyski bądź straty, które są z kolei rejestrowane na odpowiednich kontach. Trudno zaprzeczyć intuicji, iż papierowe straty odczuwane są mniej boleśnie niż straty zrealizowane. Zamknięcie konta „na minusie” pociąga za sobą niedogodności proceduralne: straty stają się jawne i zostają udokumentowane; dowiadują się o nich władze podatkowe, znajomi inwestorzy i rodzina. Nie sposób zatem dziwić się chęci ich uniknięcia. Preferencja papierowych strat w stosunku do zrealizowanych to nic innego jak uwolniony efekt dyspozycji (por. punkt 4.1.2).

Wyobraźmy sobie zawodowego inwestora stojącego przed koniecznością zgromadzenia pewnej kwoty w gotówce. Dysponuje on dwiema otwartymi inwestycjami: strata na jednej z nich wynosi 10%, na drugiej zaś nieznaczne 2%. Którą pozycję zamknąć? W odpowiedzi na to pytanie profesjonalni inwestorzy niemalże jednogłośnie decydują o likwidacji pozycji z dwuprocentową stratą, dając wyraz swemu ryzykanctwu.

Rynek miał okazję poznać konsekwencje takich decyzji w latach 2001–2002, kiedy powszechne stało się wyprzedawanie akcji dobrych spółek (*blue chips*) w celu utrzymania inwestycji w aktywach firm wysokiej technologii (*high tech*) mimo ich jak najgorszych rokowań. Jak przypuszczają niektórzy badacze, owo przywiązanie do papierowych strat wywołało lawinowe redukcje kursów w kolejnych okresach⁸³.

4.2.3.2 Kształtowanie portfela

Zapewne nikogo nie zdziwi stwierdzenie, że zarządzanie osobistymi środkami pieniężnymi to nic innego jak alokowanie ich na kontach mentalnych przyporządkowanych różnym celom. Pieniądze w budżetach gospodarstw domowych mają swoje „metki”: pieniądze emerytalne, przeznaczone na cele edukacyjne, na rozrywkę, na wakacje. Co zaskakujące, podobna idea „metkowania” przyświeca inwestorowi konstruującemu swój portfel.

Laureat nagrody Nobla Harry Markowitz (1952) starał się wpoić inwestorom, że kluczem do zbudowania dobrego portfela inwestycyjnego jest uwzględnienie kowariancji między poszczególnymi aktywami (mentalnymi kontami). Inwestor markowiczowski optymalizuje zależność między średnią stopą zwrotu a wariancją portfela; działa racjonalnie i ma jednorodny stosunek do ryzyka (ostrożność).

Okazuje się jednak, że zalecenia noblisty wywołały niewielki rezonans. Wielu inwestorów, za radą profesjonalnych doradców finansowych, decyduje się na warstwową konstrukcję portfela na wzór piramidy (Shefrin, 2000, s. 120). Każda warstwa stanowi oddzielne konto mentalne, dające wyraz potrzebie bezpieczeństwa, chęci zapewnienia potencjału wzrostu czy też aspiracjom. Każda warstwa determinowana jest też odmiennym stosunkiem do ryzyka. Na dole piramidy znajdują się aktywa mające zapewnić bezpieczeństwo (fundusze rynku pieniężnego, certyfikaty depozytowe); nieco wyżej – obligacje. Dalej umiejscowione są akcje oraz aktywa rynku nieruchomości.

⁸³ Z powyższą opinią autorka zetknęła się podczas rozmowy z psychologiem giełdowym Joachimem Goldbergem w lutym 2003 roku.

mości, które postrzegane są w kategoriach potencjału wzrostowego. Na szczycie piramidy sytuują się inwestycje spekulacyjne, takie jak opcje *out-of-the-money* czy też kupony loteryjne, które celują w aspiracje do bogactwa.

Konsekwencją uprawiania mentalnej księgowości jest niedostateczna dywersyfikacja. Alokacja inwestycji między różnymi kontami o z góry określonych celach powoduje, że kowariancje idą w niepamięć. Zgodnie z teorią Markowitza portfele takie leżą poza granicą efektywności (*efficient frontier*).

Należy jednak zauważyć, iż konstrukcje nieefektywne w rozumieniu tradycyjnym stanowią domenę świata profesjonalistów, którzy wręcz zalecają warstwową budowę portfela. Struktura portfeli rekomendowanych przez doradców finansowych odbiega od klasycznej zasady średniej – wariancji. Fisher i Statman (1997) wskazują na tę niezgodność w zaleceniach managerów funduszy inwestycyjnych, sugerujących utrzymywanie różnych proporcji akcji do obligacji w zależności od „typu” inwestora. Rekomendacja tego rodzaju w oczywisty sposób podważa podstawowe twierdzenie CAPM o podziale funduszy (*two-fund separation theorem*) między część ryzykowną portfela (zawierającą stałą proporcję akcji i obligacji) oraz część pozbawioną ryzyka.

4.2.3.3 Problem dywidend

Dlaczego inwestorzy wykazują tak silną preferencję dywidend gotówkowych? Pytanie o celowość wypłacania dywidend przez firmy od długiego czasu nurtuje teoretyków nauki finansów. Powszechnie przyjmuje się, że w świecie bez podatków i kosztów transakcyjnych dywidendy są doskonałym substytutem zysków kapitałowych (wynikających ze wzrostu wartości rynkowej spółki). Rozumowanie stojące w tle jest następujące: wypłata jednostki pieniężnej dywidendy prowadzi do obniżki kursu akcji danej firmy dokładnie o kwotę wypłaconą. Stąd inwestor powinien jednakowo cenić jednostkę dywidendy i jednostkę uzyskaną przez odsprzedaż posiadanych akcji, czyli dywidendę „samodzielnej produkcji” (*homemade dividend*, Shefrin i Statman, 1984). Jeśli uwzględnić wpływ podatków, argument przeciwko dywidendom staje się oczywisty, jako że stopa ich opodatkowania prze-

wyższa stopę opodatkowania zysków kapitałowych. Dlatego racjonalny inwestor powinien zrezygnować z dywidend, jeśli spółka dysponuje możliwościami inwestycyjnymi, które generują zwrot wyższy niż jej koszty kapitału: w domyśle – kreuje dodatnią wartość dla akcjonariuszy. Z normatywnego punktu widzenia nie sposób zanegować tej zasady. Mimo to preferencje inwestorów stale i niezmiennie wskazują na wyższość dywidend nad zyskami kapitałowymi. Tak więc podważeniu ulega tradycyjnie przyjmowane założenie o doskonałej substytucyjności pomiędzy pieniądzem w formie dywidendy a przyrostem wartości rynkowej spółki (pieniądz w formie zysku kapitałowego).

Istnieje wiele koncepcji stanowiących próbę wyjaśnienia owego fenomenu⁸⁴. Ciesząc się dużym uznaniem wytłumaczenie behawioralne pochodzi od profesorów Shefrina i Statmana (1984). Punktem wyjścia dla ich rozważań jest opisywana w *prospect theory* zależność preferencji od formy. Shefrin i Statman zastanawiają się nad czynnikami skłaniającymi spółki do kadowania części swoich zysków jako dywidendy, mimo iż rozwiązanie to nie jest optymalne w sensie Pareto⁸⁵. Zgodnie z hipotezą mentalnej księgowości jednostki nadają pieniądzwowi z różnych źródeł metki: „pieniądz z dywidendy”, „pieniądz z kapitału”. Idąc tym śladem badacze wskazują, że poprzez wypłacanie dywidend firmy ułatwiają inwestorom segregowanie zysków i strat w sposób zapewniający najwyższą użyteczność. Przybliży to następujący przykład:

W ciągu roku wartość rynkowa firmy wzrosła o 10 j.p. na akcję. Firma ma do wyboru dwie opcje: (i) integrację: nie wypłacać dywidendy, traktując wzrost wartości rynkowej jako zysk kapitałowy 10 j.p. albo (ii) segregację: wypłacić dywidendę w wysokości 2 j.p., pozostawiając 8 j.p. zysku kapitałowego.

⁸⁴ Pominęto tu koncepcje związane ze standardową teorią finansów, jako że wykraczają one poza zakres opracowania.

⁸⁵ Rozwiązaniem optymalnym w sensie Pareto byłoby natomiast zastąpienie dywidend skupem własnych akcji przez firmy. Wówczas obie strony – zarówno inwestorzy (przez wzgląd na niższe opodatkowanie), jak i firmy (dzięki korzystnej reinwestycji środków) – odniosłyby większe korzyści.

Ze względu na wklęsły kształt funkcji wartości użyteczność inwestora wygenerowana przez zabieg segregacji jest wyższa niż w przypadku integracji: $v(2) + v(8) > v(10)$.

Podobny zabieg możliwy jest także dla strat. Firma, której wartość rynkowa obniżyła się o 10 j.p. na akcję, może (i) zdecydować się na wypłatę dywidendy 2 j.p., dając inwestorowi użyteczność $v(2) + v(-12)$ bądź też (ii) nie wypłacać dywidendy, pozostawiając inwestora ze stratą kapitałową o wartości $v(-10)$. Ponownie segregacja jest lepszą metodą hedonistycznego kadrowania. Ze względu na kształt funkcji wartości otrzymujemy $v(2) + v(-12) > v(-10)$.

Preferencję dywidend gotówkowych można zatem wytłumaczyć hedonistycznym kadrowaniem (por. punkt 4.2.2.2). Zjawisko to objaśniają dwa regiony segregacji przedstawione na rysunku 4.1. Aby zrozumieć logikę kryjącą się za zabiegiem segregacji dywidend i zysków/strat kapitałowych, warto przytoczyć cytaty ze znanego podręcznika dla maklerów giełdowych (Shefrin i Statman, 1984):

„Kupując akcje, które wypłacają sowiłą dywidendę, większość inwestorów przekonana jest o swojej rozwadze, polegającej na zapewnieniu sobie przyszłych dochodów. Inwestorzy ci mają wrażenie, że dodatkowy potencjał zysku [przyp. tłum.: z kapitału] to ‘ekstra’ korzyść. Inaczej, jeśli papier traci na wartości w porównaniu z ceną zakupu – wówczas pocieszają się, że dywidenda zwraca im część kosztów”.

Skrajnie prawy segment segregacji (na rysunku 4.1 oznaczony numerem (1)) związany jest z zyskiem kapitałowym; strategia segregacji dywidend i zysków pozwala inwestorowi cieszyć się z „ekstra korzyści”, czyli ze wzrostu wartości rynkowej spółki. W skrajnie lewym segmencie (pole (3)) segregacja pełni zupełnie inną funkcję: dywidenda traktowana jest jako „pocieszenie” po znacznej stracie kapitałowej. W regionie strat umiarkowanych inwestor dokonuje z kolei integracji (pole (2)), równoważąc małe straty niewielkimi kwotowo dywidendami.

4

Mentalna księgowość to jedynie część rozwiązania „zagadki” dywidend, jakie zaproponowali Shefrin i Statman. W kolejnym kroku badacze dokonują zręcznej kompilacji tej koncepcji z problemem samokontroli. Zasada podziału dochodów na dywidendy i zyski kapitałowe pozwala inwestorowi z mniejszym wysiłkiem powściągnąć potrzebę natychmiastowej konsumpcji. W walce z problemem samokontroli ludzie często ustalają sztywne reguły, które pozwalają na utrzymanie dyscypliny, np. „oszczędzaj pensję żony, wydawaj tylko dochody męża” albo podobnie „finansuj konsumpcję pieniędzmi z dywidend, nie ruszaj kapitału”. Ostatnie zdanie dobitnie potwierdza wielokrotnie wspomnianą metodę „metkowania” pieniądza pochodzącego z różnych źródeł. Dywidendy księgowane są, jak widać, na koncie bieżącego dochodu, zaś zyski ze wzrostu wartości rynkowej akcji – na koncie kapitału. Jeśli firmy zdecydowałyby się na rozwiązanie optymalne w sensie Pareto (skup własnych akcji w miejsce wypłat dywidend), wówczas ów zapewniający poczucie bezpieczeństwa misterny mechanizm samokontroli przestałby działać.

O samokontroli i znaczeniu czynnika międzyokresowego będzie mowa w kolejnym rozdziale.

Dotychczasowe rozważania koncentrowały się na procesach podejmowania decyzji w warunkach ryzyka w teraźniejszości. Analiza ta wymaga rozszerzenia o czynnik czasu.

Celem niniejszego rozdziału jest pokazanie wpływu czasu zarówno na przyszły stan jednostek determinowany dzisiejszymi decyzjami, jak i na percepcję wyborów dokonywanych w przeszłości. W kontekście przyszłości wpływ ten ujawnia się pod postacią indywidualnej stopy dyskontowej, odniesiony do przeszłości określa natomiast przebieg procesów adaptacyjnych.

5.1 Percepcja przyszłości

Uczestnicy rynku rzadko podejmują decyzje istotne jedynie dziś. Zwykle konsekwencje wyborów dokonanych w teraźniejszości rozciągają się na przyszłe okresy. Oczwistym przykładem jest rynek kapitałowy, gdzie alokacja środków dokonywana w chwili obecnej implikuje stopę zwrotu w przyszłości.

Proces podejmowania decyzji i percepcja czasu pozostają w ścisłym związku zarówno w codziennym życiu jednostek, jak i na poziomie zagregowanym – gospodarki i społeczeństwa. Ile czasu zainwestować we własną edukację? Ile oszczędzać „na stare lata”? Jak inwestować? Czy karczować lasy dziś kosztem efektu cieplarnianego w następnych dekadach? Każde z tych pytań zawiera w sobie kombinację czasu i problemu decyzyjnego.

Wpływ czasu na decyzje uczestników rynku, jakkolwiek zgłębiany przez ekonomistów, wymaga świeżego podejścia. Liczne badania wskazują bowiem na systematyczne rozbieżności między obserwowanymi zachowaniami ludzi a ugruntowanym modelem wyborów międzyokresowych.

5.1.1 Normatywny model wyboru międzyokresowego

W 1937 roku Paul Samuelson opublikował pięciostronicowy artykuł, w którym zaproponował uogólniony model wyboru międzyokresowego –

teorię zdyskontowanej użyteczności⁸⁶ (Samuelson, 1937). Teoria ta stała się powszechnie akceptowanym podejściem normatywnym do decyzji dokonywanych w czasie⁸⁷.

Do wyznaczenia międzyokresowej użyteczności używa się następującej postaci funkcyjnej:

$$V(X) = \sum_t \delta^t \cdot u(x_t), \quad (5.1)$$

gdzie:

$V(X)$ – wartość bieżąca przyszłych wypłat,

$X = (x_1, \dots, x_n)$ – wartości (np. konsumpcja) występujące w kolejnych okresach,

$\delta \equiv (1 + \rho)^{-1}$ – indywidualny czynnik dyskontujący dla jednego okresu,

ρ – indywidualna stopa dyskontowa dla jednego okresu,

$u(x_t)$ – użyteczność w okresie t .

Na podstawie sformułowania (5.1) jednostka jest w stanie określić swoje preferencje wobec sekwencji wypłat (np. konsumpcji w kolejnych okresach) $X = (x_1, \dots, x_n)$ poprzez porównanie jej subiektywnej wartości z użytecznością innych opcji.

Model Samuelsona usiłuje zagregować wszelkie psychologiczne uwarunkowania wyborów międzyokresowych za pomocą jednego parametru – stopy dyskontowej, której zadaniem jest uchwycenie preferencji czasowych jednostki. Wielkość tę można zdefiniować jako miarę gratyfikacji („cenę czasu”), która czyni obojętnym wybór między wartością bieżącą a przyszłą. Przyjmuje się przy tym następujące założenia:

⁸⁶ Wcześniej dominowała koncepcja Irvinga Fishera, stworzona dla dwóch okresów i opierająca się na krzywych obojętności. Por. (Fisher, 1930). Podejście Samuelsona umożliwiło rozszerzenie analizy na przypadek wielookresowy.

⁸⁷ Mimo iż sam autor wykazywał duży sceptycyzm wobec normatywnego znaczenia swego dzieła, zdyskontowana użyteczność uzyskała status jednej z centralnych koncepcji w nowoczesnej nauce finansów. Popularność modelu wynika głównie z jego matematycznej elegancji i łatwości zastosowania.

- (1) Jednostki mają dodatnie preferencje czasowe⁸⁸ (dodatnią stopę dyskontową), co oznacza, iż przedkładają dzisiejsze wpływy gotówki (zyski) nad jutrzejsze, przy jednoczesnym dążeniu do opóźnienia wydatków (strat).
- (2) Indywidualna stopa dyskontowa jest identyczna w każdej jednostce czasu (cecha rozpadu wykładniczego) i niezależna od dyskontowanej wartości.
- (3) Preferencje czasowe spełniają warunek rozdzielności względem czasu (*time separability*) stanowiący, że wartość bieżąca szeregu wypłat jest równa sumie zdyskontowanych wartości poszczególnych jego elementów.

Powyższe podejście gwarantuje stabilność preferencji w czasie. Od osoby racjonalnej oczekuje się konsekwencji w dokonywaniu wyborów. Jeśli przedkłada ona 100 j.p. dziś ponad 110 j.p. za miesiąc, należy zakładać, że postawiona przed wyborem 100 j.p. za 10 lat wobec 110 j.p. za 10 lat i jeden miesiąc, zadecyduje tak samo.

Ponadto przyjmuje się, iż indywidualne preferencje czasowe są ściśle powiązane ze stopą procentową obowiązującą na rynku. Istnienie rynków kapitałowych pozwala na dokonywanie tzw. wewnętrznego arbitrażu (*internal arbitrage*, Loewenstein i Thaler, 1989). Osoba o wyższej od rynkowej indywidualnej stopie dyskontowej winna zadłużyć się na rynku kapitałowym w celu zaspokojenia potrzeby natychmiastowej konsumpcji. I odwrotnie: jednostka charakteryzująca się niższą indywidualną stopą dyskontową niż rynkowa powinna stać się pożyczkodawcą. Tak więc możliwość wewnętrznego arbitrażu istnieje wówczas, gdy indywidualne preferencje czasowe są odmienne od poziomu stóp procentowych. Jednostki pożyczają bądź zadłużają się do momentu, gdy ich krańcowa stopa substytucji między konsumpcją dziś i jutro zrówna się ze stopą rynkową.

⁸⁸ Określenie „dodatnie preferencje czasowe”, mimo braku precyzji, jest często stosowane przez specjalistów w dziedzinie wyboru międzyokresowego. Por. (Frederick, Loewenstein i O'Donoghue, 2002). Bezpośrednią konsekwencją dodatnich preferencji czasowych jest dodatnia indywidualna stopa dyskontowa, innymi słowy preferencja konsumpcji dzisiejszej nad jutrzejszą.

Warto wspomnieć, iż Samuelson odżegnywał się od uznania modelu zdyskontowanej użyteczności za deskryptywnie właściwy, podkreślając arbitralność założenia, jakoby jednostki maksymalizowały wartość wyrażenia postaci (5.1). Mimo tych zastrzeżeń jego koncepcja została szybko zaadaptowana jako podstawowe narzędzie analizy decyzji międzyokresowych. Zakres jej obecnych zastosowań rozciąga się od problemów oszczędzania w gospodarce, poprzez zagadnienia podaży pracy, wyceny aktywów finansowych, decyzji edukacyjnych, aż po analizę przestępczości. Jednakże konfrontacja Samuelsonowskiej teorii z rzeczywistością nie pozwala na przypisanie jej roli deskryptywnego modelu ludzkich preferencji. Najnowsze prace podważają założenie dodatnich stóp dyskontowych i stabilności preferencji w czasie. Odkryte anomalie stanowią przedmiot dalszych rozważań.

5.1.2 Anomalie wyboru międzyokresowego

Badania nad dynamiczną niekonsekwencją decyzji zapoczątkowane zostały publikacją Thaler (1981), prezentującą wyniki prostego eksperymentu. Jego uczestników zapytano, jakiego żądałoby wynagrodzenia za odroczenie otrzymania pewnej kwoty, by uczynić ją równie atrakcyjną, jak wypłata dzisiaj. W ćwiczeniu manipulowano trzema zmiennymi: długością okresu oczekiwania, wartością oraz znakiem (zysk, strata)⁸⁹. Na podstawie odpowiedzi zaobserwowano trzy silne wzorce zachowań, będące pogwałceniem teorii zdyskontowanej użyteczności:

- (1) Implikowana wyborami jednostek stopa dyskontowa jest ujemnie skorelowana z długością okresu oczekiwania na wypłatę.
- (2) Stopa dyskontowa maleje wraz ze wzrostem wartości wypłaty.
- (3) Zyski są dyskontowane zdecydowanie silniej niż straty⁹⁰.

⁸⁹ Dokładniej należałoby zastąpić określenia „zysk”, „strata” wyrażeniami „zapłata dokonana na rzecz respondenta” oraz „zapłata dokonana przez respondenta na rzecz osoby trzeciej”. W literaturze źródłowej pojęcia „zysk”, „strata” używane są bowiem w takim właśnie znaczeniu.

⁹⁰ Sformułowanie „silniejsze dyskontowanie” jest skrótem myślowym, służącym wyrażeniu, że zyski dyskontowane są przy użyciu wyższej stopy dyskontowej.

Dalsze eksperymenty stwierdziły kolejne anomalie. Wskazanie decyzji międzyokresowych niezgodnych z modelem normatywnym, a jednak stanowiących część codzienności, nie jest trudne. Jednostki – wbrew predykcjom klasycznej teorii – nie dyskontują za pomocą stałego dodatniego parametru (por. punkt 5.1.2.1). Stopy dyskontowe, a wraz z nimi dokonywane wybory, zależą od kwoty (por. punkt 5.1.2.2), jej znaku – wpływ pieniężny bądź wydatek (por. punkt 5.1.2.3), od źródła, z jakiego pochodzi pieniądz (por. punkt 5.1.2.4), czy też od ryzyka związanego z przyszłą wypłatą (por. punkt 5.1.2.5). Najważniejsze z dotychczasowych odkryć omówione zostaną poniżej.

5.1.2.1 Negatywne preferencje czasowe

Podstawowe założenie modelu normatywnego stanowią dodatnie preferencje czasowe, które są odzwierciedleniem niecierpliwości – wyrazem silniejszej chęci konsumowania dziś niż w przyszłości. Dla pojedynczych wypłat w przyszłości założenie to ma również moc deskryptywną, traci ją natomiast dla sekwencji. Dodatnia stopa dyskontowa implikuje bowiem, że jednostki czerpią większą satysfakcję z malejących sekwencji wypłat (np. konsumpcji, dochodów, które obniżają się z okresu na okres) niż z rosnących. Rozszerzając ów tok rozumowania można by przewrotnie wnioskować, jakoby podupadający z czasem standard życiowy albo pogarszający się stan zdrowia były lepsze niż profil stopniowej poprawy ich jakości.

Prowadzone w ostatnich latach obserwacje wyborów między sekwencjami zdarzeń podważają – zgodnie zresztą z intuicją – deskryptywną wartość założenia stałe dodatnich stóp dyskontowych. Jednostki znamionuje silna skłonność do utrzymywania wzrastającego profilu bogactwa czy konsumpcji, co dowodzi istnienia ujemnych stóp dyskontowych (Loewenstein i Prelec, 1993).

Jako przykład często wskazuje się zachowania amerykańskich podatników, którzy – dokonując nadpłat podatków – oferują rządowi nieoprocentowaną pożyczkę, mimo iż ekonomicznie lepszym dla nich rozwiązaniem byłoby dostosowanie stopy zachowania dochodu w czasie (Loewenstein i Thaler, 1989). O to jednak niewielu się dotychczas pokusiło.

Dalszą ilustracją problemu są trzynaste pensje wypłacane pracownikom pod koniec roku. Tzw. trzynastka nie stanowi bynajmniej dodatkowego dochodu i wynika jedynie z podziału rocznego uposażenia na 13 części. W związku z tym powinna być raczej traktowana jako utrata odsetek od zatrzymanego przez pracodawcę kapitału, nie zaś jako dodatkowy wpływ pieniężny. Nawet jednak świadomi pracownicy opowiadają się za utrzymaniem trzynastek. Co więcej, zdają się przedkładać rosnący profil płacowy nad opadający bądź równo rozłożony w czasie, mimo iż ta druga opcja daje w sumie lepszy wynik (Hsee, Abelson i Salovey, 1991).

Powyższe przykłady koncentrują się na negatywnych preferencjach czasowych dla wpływów pieniężnych, czyli dla bodźców o dodatnim ładunku emocjonalnym. Dążenie do stopniowej poprawy dotyczy również zdarzeń negatywnych (np. strat lub bólu). Varey i Kahneman (1992) dowodzą, że jednostki lepiej oceniają strumienie malejącego dyskomfortu niż strumienie o nasilających się wrażeniach negatywnych, nawet jeśli obie sekwencje dają identyczny wynik po skumulowaniu. Podobne wzorce zachowań stwierdzone zostały dla serii pieniężnych.

Tłumacząc występowanie negatywnych preferencji czasowych, badacze powołują się na zjawisko zależności wyborów od formy ich prezentacji, co w pracy tej nazwano kadrowaniem. Skadrowanie wyboru międzyokresowego jako sekwencji zdarzeń motywuje jednostki do dalekowzroczności, wyrażającej się chęcią przesunięcia chwili otrzymania najlepszych wyników na później. Inaczej działa segregacja, polegająca na przedstawieniu decyzji w oderwaniu od siebie wzajemnie. Zabieg ów promuje wysokie dodatnie stopy dyskontowe zgodne z zasadą „lepiej dziś niż jutro”.

5.1.2.2 Efekt skali

Efekt skali (*magnitude effect*) mówi o tym, że wysokość wypłaty ma wpływ na poziom stopy dyskontowej. W szczególności małe kwoty dyskontowane są znacznie silniej niż kwoty duże⁹¹. W eksperymencie Thaler

⁹¹ Efekt ten stwierdzony został także przez innych badaczy. Por. (Holcomb i Nelson, 1992).

(1981) respondenci nie odczuwali różnicy między otrzymaniem 15 j.p. teraz i 60 j.p. za rok, 250 j.p. teraz i 350 j.p. za rok, 3000 j.p. teraz i 4000 j.p. za rok. Indywidualne stopy dyskontowe ρ ⁹² obliczone na podstawie tych decyzji wynoszą odpowiednio: 139%, 34% oraz 29%.

W celu wytłumaczenia efektu można odnieść się do problemu samokontroli. Oczekiwanie na wypłatę (nagrodę pieniężną) wymaga wysiłku, którego niewielka suma 15 j.p. może nie być warta. Dlatego też wynagrodzenie, jakiego jednostki wymagają za opóźnienie, jest w tym przypadku bardzo wysokie (139%). Warto wskazać w tym miejscu na wnioski płynące z psychofizyki: różnica między kwotą 15 j.p. a 21 j.p. postrzegana jest jako mniejsza niż między 250 j.p. a 350 j.p., choć proporcje w obu parach są identyczne.

5.1.2.3 Asymetria w dyskontowaniu zysków i strat

Kolejnym punktem rozbieżności między teorią normatywną a rzeczywistymi wyborami jednostek jest asymetria w dyskontowaniu zysków i strat, zwana również efektem znaku (*sign effect*). Istnieje wiele empirycznych potwierdzeń faktu, że straty dyskontowane są słabiej niż zyski. Loewenstein (1988) udokumentował, iż respondenci nie odczuwali różnicy między zyskiem 10 j.p. teraz a 21 j.p. za rok ($\rho = 74\%$) oraz 100 j.p. teraz a 157 j.p. za rok ($\rho = 45\%$). Po stronie wyników ujemnych relacje te prezentowały się następująco: natychmiastowa strata 10 j.p. wobec 15 j.p. za rok ($\rho = 41\%$) oraz strata 100 j.p. wobec 133 j.p. za rok ($\rho = 29\%$). Dowody jeszcze bardziej znaczącej asymetrii uzyskał Thaler (1981), wskazując na stopy dyskontowe dla zysków od 3 do 10 razy przewyższające odpowiednie parametry dla strat.

Warte wspomnienia są, obserwowane u niektórych osób, negatywne preferencje czasowe dla wyników ujemnych. Wyrażają się one większą chęcią

⁹² Thaler wyznacza indywidualną stopę dyskontową, używając formuły dla ciągłej kapitalizacji odsetek.

realizacji straty natychmiast niż jej opóźnienia⁹³. Efekt ten dotyczy szczególnie małych kwot i daje się wytłumaczyć awersją do długu (Loewenstein i Thaler, 1989). Pozwala on zrozumieć, dlaczego jednostki przedwcześnie spłacają kredyty (szczególnie tanie – np. kredyt studencki), mimo iż optymalne rozwiązanie polegałoby na bezpiecznym ulokowaniu środków na rynku w celu uzyskania dodatkowego procentu i spłacie zobowiązania kredytowego w zwykłym terminie.

5.1.2.4 Efekt źródła

Wbrew przewidywaniom modelu zdyskontowanej użyteczności, ludzie stosują różne stopy dyskontowe w zależności od źródła dochodów. Ów tzw. efekt źródła (*source effect*) jest negacją powszechnie akceptowanego w ekonomii faktu, że pieniądź nie ma „metki” – jednostki monetarne są wzajemnie doskonałymi substytutami (por. podrozdział 4.2). Przekonujących przykładów pogwałcenia paradygmatu substytucyjności pieniądza w wyborach międzyokresowych dostarczyli Shefrin i Thaler (1988). Wykazali oni, że stosunek wydatków na konsumpcję do oszczędności zależy od formy, w jakiej jednostki otrzymują dochody. Jeśli dochody stanowią dodatkową część comiesięcznych pensji (np. premia wypłacana co miesiąc), wówczas zostaną przeznaczone na codzienną konsumpcję. Inaczej dzieje się w sytuacji jednorazowych przychodów (np. premia otrzymana w jednej transzy), które postrzegane są zazwyczaj w kategoriach oszczędności i przeznaczane na konkretny cel. Ciekawa obserwacja wiąże się z percepcją przyszłych dochodów (np. darowizna otrzymana dziś przy zastrzeżeniu jej nienaruszenia przez pięć kolejnych lat). Okazuje się, iż przyszłe dochody, jakkolwiek pewne, niemal zawsze traktowane są jako oszczędności. Taki sposób klasyfikacji charakterystyczny jest nawet dla osób dysponujących dużymi zasobami płynnych

⁹³ Jak wiadomo, zgodnie z klasycznym podejściem do wyboru międzyokresowego jednostka racjonalna powinna dążyć do jak najwcześniejszych wpływów gotówki i możliwie najbardziej opóźniać jej wypływy (np. spłata kredytu). Takie preferencje są wówczas oznaką podatniej indywidualnej stopy dyskontowej.

środków, które mogłyby zostać wydane dziś, skoro istnieje pewność przyszłego wpływu gotówki (np. wspomnianej darowizny).

Preferencje czasowe zależne od źródła dochodu wspierają hipotezę mentalnej księgowości prowadzonej celem samokontroli (Shefrin i Thaler, 1988). W przeciwieństwie do kwot znacznych, interpretowanych w kategoriach oszczędności i ostrożniej konsumowanych, małe sumy rejestrowane są przez jednostki na kontach natychmiastowej konsumpcji. Można się stąd spodziewać, iż koszt czekania na małe wypłaty pieniężne będzie pojmowany jako niezrealizowana konsumpcja, podczas gdy cena oczekiwania na większe kwoty potraktowana zostanie jako niezrealizowany procent od oszczędności. Pozwala to zrozumieć stwierdzony empirycznie fakt dyskontowania wyższych wartości za pomocą stóp bliższych realiom rynkowym.

5.1.2.5 Wybór międzyokresowy w warunkach ryzyka

Sformułowanie „wybór międzyokresowy w warunkach ryzyka” w najlepszy sposób opisuje warunki, w jakich działa inwestor. Jego decyzje wymagają określenia wartości oczekiwanej zysków oraz zdefiniowania ich w kategoriach środków, jakie byłby skłonny dziś zainwestować. Stopy zwrotu z inwestycji w aktywa ryzykowne zależą od przyszłych ruchów cen, które z dzisiejszej perspektywy mogą być jedynie ocenione w przybliżeniu. Metody ilościowe, jakkolwiek pomocne, nie są w stanie usunąć niepewności związanej z przyszłymi stopami zwrotu. Wobec tego każde posunięcie uczestnika rynku łączy w sobie dwa komponenty: preferencje czasowe i stosunek do ryzyka.

Tymczasem okazuje się, iż zagadnieniu temu poświęcono dotąd niezwykle mało uwagi. Częściowym usprawiedliwieniem tego uchybienia mogą być problemy „techniczne” związane z badaniem zachowań przy jednoczesnym uwzględnieniu czasu i niepewności. Nieliczne dotychczas prace w tej dziedzinie pozwalają jednak na wartościowe spostrzeżenia.

Normatywny model zakłada, iż dyskontowanie w czasie jest niezależne od dyskontowania ze względu na ryzyko. Podejściu temu przeczą badania empiryczne, dowodząc, że przyszłe wypłaty obciążone ryzykiem dyskontowane

są za pomocą niższej stopy niż wyniki pewne w przyszłości. Słabsze dyskontowanie przyszłej niepewności wskazuje *implicite* na silną preferencję zdarzeń, jakie zajądą z pewnością.

Występowanie niższej stopy dyskontowej w warunkach ryzyka dokumentuje choćby praca Ahlbrechta i Webera (1997). Jak przedstawia rysunek 5.1, porównując wielkość niepewną w przyszłości, np. przyszłą loterię L_t , z jej pewnym ekwiwalentem dziś $CE_0(L_t)$ (ścieżka B) respondenci używają niższej stopy dyskontowej, niż w sytuacji, gdy najpierw uwzględniony zostaje czynnik ryzyka $CE_t(L_t)$, a następnie czas $CE_0(CE_t)$ (droga A+C). Dwie metody dyskontowania tej samej przyszłej kwoty obciążonej ryzykiem dają różną ocenę jej wartości bieżącej, taką, że: $CE_0(L_t) > CE_0(CE_t)$.

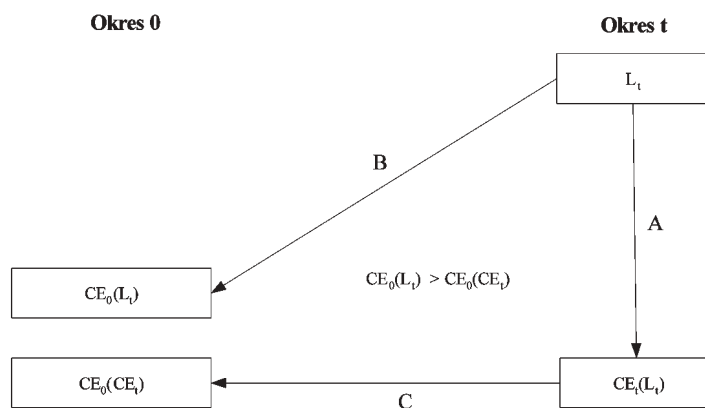
Obserwacja ta godzi w tradycyjne założenie mówiące, iż ponoszenie ryzyka wymaga dodatkowego wynagrodzenia (premii). Stąd inwestycje ryzykowne powinny być dyskontowane przy użyciu wyższej stopy niż wypłaty pewne. Sprzeczne z ekonomiczną intuicją rezultaty badań pozwalają na postawienie następującej hipotezy: jeśli jednostki dyskontują jednocześnie ze względu na czas i ryzyko (wyznaczenie $CE_0(L_t)$), wówczas element ryzyka skupia na sobie większą uwagę niż samo przesunięcie w czasie (Stevenson, 1992). Dwuetapowe sformułowanie (skadrowanie) tego samego problemu (wyznaczanie $CE_0(CE_t)$) skłania jednostki do uwzględnienia ceny obu czynników.

Jeśli wskazana asymetria jest silna u wielu osób, należy się zastanowić nad jej konsekwencjami dla rynku kapitałowego⁹⁴ oraz dla instytucji finansowych, oferujących klientom produkty o niepewnym w przyszłości profilu wypłat. Inwestorzy, których działania zgodne są z omawianym wzorcem, będą koncentrować się bardziej na ryzyku niż na horyzoncie czasowym. Może to prowadzić do wyboru aktywów o zbyt umiarkowanym ryzyku, ge-

⁹⁴ Dotychczasowe wnioski na temat asymetrii między dyskontowaniem w warunkach ryzyka i pewności opierają się na eksperymentach laboratoryjnych. Dlatego też dyskusja nad konsekwencjami tego zjawiska dla rynku kapitałowego polega jedynie na stawianiu hipotez.

nerujących oczekiwaną stopę zwrotu jedynie w długiej perspektywie (Ahlbrecht i Weber, 1997).

Rysunek 5.1: Dyskontowanie w warunkach niepewności



Źródło: Opracowanie własne na podstawie (Ahlbrecht i Weber, 1997).

5.1.3 Deskryptywny model wyboru międzyokresowego

W normatywnym modelu wyboru międzyokresowego przyjmuje się założenie dyskontowania wykładniczego (rozpadu wykładniczego). Oznacza to, że przesunięcie wypłaty o jednostkę czasu w przyszłość powoduje spadek jej wartości bieżącej o daną liczbę punktów procentowych (stałą z okresu na okres). Dyskontowanie wykładnicze bazuje na aksjomacie stacjonarności (*stationarity*). Ustanawia on, że jeśli osoba woli otrzymać kwotę A_s w momencie s od kwoty B_t w momencie t , to identyczne opóźnienie obu wypłat o czas d nie zmieni tych preferencji, czyli $A_s \phi B_t \Leftrightarrow A_{s+d} \phi B_{t+d}$. Czynniki dyskontujące między momentami s i t oraz $(s + d)$ i $(t + d)$ jest identyczny.

Brak spełnienia tej zależności w rzeczywistych wyborach zaobserwował Thaler⁹⁵ (1981), pytając respondentów o podanie przyszłej kwoty równoważnej z otrzymaniem 15 j.p. dziś. W eksperymencie manipulowano czasem

⁹⁵ Choć powołano się w tym miejscu na pracę Thaler (1981), warto nadmienić, że w literaturze cytuje się również wcześniejsze badania Ainslie (1975). Inni badacze zagadnienia to również Benzion, Rapoport i Yagil (1989), Chapman i Elstein (1995), Redelmeier i Heller (1993).

odroczenia wypłaty⁹⁶: o miesiąc, rok oraz 10 lat. Przeciętne kwoty uznawane za przyszły ekwiwalent dzisiejszych 15 j.p. wyniosły: 20 j.p. za miesiąc, 50 j.p. za rok oraz 100 j.p. za 10 lat, co pozwala wyznaczyć indywidualne stopy dyskontowe ρ na poziomie odpowiednio: 345%, 120% i 19%⁹⁷.

Obniżanie się stóp dyskontowych wraz z wydłużaniem okresu oczekiwania na wypłatę jest najważniejszym dowodem empirycznym podważającym deskryptywną moc podejścia tradycyjnego. Bezpośrednim rezultatem tego zjawiska jest bowiem niekonsekwencja wyborów międzyokresowych.

5.1.3.1 Dyskontowanie hiperboliczne

Terminem „dyskontowanie hiperboliczne” (w odróżnieniu od powszechnie zakładanego dyskontowania wykładniczego) przyjęło się określać malejące wraz z oddaleniem w czasie stopy dyskontowe. Opisana równaniem hiperboli funkcja dyskontująca najlepiej oddaje bowiem wyniki badań empirycznych. Loewenstein i Prelec (1992) proponują następującą jej postać:

$$\phi(t) = (1 + \alpha t)^{-\frac{\beta}{\alpha}}, \quad \alpha, \beta > 0, \quad (5.2)$$

gdzie:

α – współczynnik określający stopień odchylenia funkcji $\phi(t)$ od dyskontowania za pomocą stałego czynnika; dla $\alpha \rightarrow 0$ otrzymujemy funkcję dyskontowania wykładniczego zgodną z podejściem tradycyjnym $\phi(t) = e^{-\beta \cdot t}$,⁹⁸

⁹⁶ Nie należy mylić tego z efektem skali, gdzie manipulowano wielkością wypłaty, zaś opóźnienie w czasie było niezmiennie.

⁹⁷ Oznacza to, że $15 = 20 \cdot e^{-(3,45)(1/12)} = 50 \cdot e^{-(1,20)(1)} = 100 \cdot e^{-(0,19)(10)}$. Wyniki badań można jednak zinterpretować w znaczeniu średniej stopy dyskontowej dla danego okresu, tzn. 345% to średnia (w skali roku) stopa stosowana dla okresu od dziś do za miesiąc; 100% to średnia stopa dla okresu od miesiąca do roku, licząc od dnia dzisiejszego; 7,7% to średnia stopa dla okresu od roku do dziesięciu lat, patrząc od teraz. To jest: $15 = 20 \cdot e^{-(3,45)(1/12)} = 50 \cdot e^{-(3,45)(1/12)} \cdot e^{-(1,00)(11/12)} = 100 \cdot e^{-(3,45)(1/12)} \cdot e^{-(1,00)(11/12)} \cdot e^{-(0,077)(9)}$.

⁹⁸ Niech $\alpha = \frac{1}{k}$, $k \neq 0$, wówczas mamy

$$\lim_{\alpha \rightarrow 0} (1 + \alpha \cdot t)^{-\frac{\beta}{\alpha}} = \lim_{k \rightarrow \infty} \left(1 + \frac{t}{k} \right)^{-\beta \cdot k} = \lim_{k \rightarrow \infty} \left(\left(1 + \frac{1}{\frac{k}{t}} \right)^{\frac{k}{t}} \right)^{-\beta \cdot t} = e^{-\beta \cdot t}.$$

β – parametr określający preferencje czasowe jednostki,

t – czas.

Wartość bieżąca użyteczności sekwencji wypłat $X = (x_1, \dots, x_n)$ otrzymywanych w kolejnych okresach przyjmuje zatem postać:

$$U(x_1, t_1; \dots; x_n, t_n) = \sum_i v(x_i) \cdot \phi(t_i), \quad (5.3)$$

gdzie:

$U(\cdot)$ – wartość bieżąca użyteczności przyszłych wypłat,

$v(\cdot)$ – funkcja wartości,

$\phi(\cdot)$ – funkcja dyskontująca hiperbolicznie.

Podstawą wyprowadzenia (5.3) jest aksjomat rozciągłości (*stretching*), zaproponowany przez Harvey'a (1986). Jednostka przedkładająca wypłatę A_s w momencie s nad B_t w momencie t , dokona takiego samego wyboru dla okresów odpowiednio $d(s+1)-1$ oraz $d(t+1)-1$, gdzie d oznacza opóźnienie. Stosowne relacje preferencji można zatem zapisać jako:

$$A_s \phi B_t \Leftrightarrow A_{d(s+1)-1} \phi B_{d(t+1)-1}.$$

Odmienne niż w przypadku aksjomatu stacjonarności, rozciągłość jest wrażliwa na przesunięcie w czasie. Wyobraźmy sobie osobę, która woli 100 j.p. dziś od 110 j.p. jutro. Rozciągłość oznacza, że preferencje te nie ulegną zmianie w sytuacji wyboru między 100 j.p. za 10 dni a 110 j.p. za 21 dni (parametr $d = 11$). Jednak ta sama jednostka, postawiona przed alternatywą 100 j.p. za 10 dni lub 110 za 11 dni, może zmienić swe upodobania na korzyść wypłaty wyższej i późniejszej. Jakkolwiek zaprezentowane rozumowanie odpowiada intuicji, stanowi ono pogwałcenie aksjomatu stacjonarności.

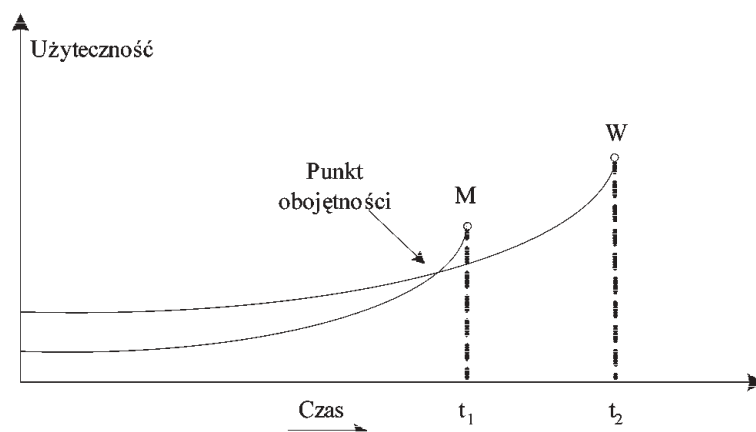
Chociaż model hiperboliczny nie oferuje jedynej alternatywy wobec dyskontowania wykładniczego, jest to ujęcie wsparte najliczniejszymi badania-

mi, a zatem i powszechnie przyjmowane jako satysfakcjonujący opis rzeczywistych wyborów⁹⁹.

5.1.3.2 Dynamiczna niekonsekwencja wyborów

Występujący w podejściu hiperbolicznym negatywny związek między odaleniem wypłaty w czasie a stopami dyskontowymi prowadzi do odwracania się preferencji. Problem ten ilustruje rysunek 5.2.

Rysunek 5.2: Dyskontowanie hiperboliczne



Źródło: (Loewenstein i Thaler, 1989).

Załóżmy, że jednostka musi wybrać między dwiema nagrodami pieniężnymi: mniejszą, lecz wcześniejszą (M w momencie t_1) oraz większą, lecz późniejszą (W w momencie $t_2 > t_1$). Oś y reprezentuje użyteczność obu wypłat dla każdego momentu w czasie; oś x mierzy czas, jaki pozostał do otrzymania M lub W .

W modelu hiperbolicznym siła dyskontowania jest funkcją czasu, co wskazuje na możliwość przecinania się krzywych użyteczności¹⁰⁰. Jeśli moment wypłaty jest odległy, wówczas jednostki zdecydują się na wypłatę

⁹⁹ Inne postaci funkcyjne dla dyskontowania hiperbolicznego to: $f(t) = 1/t$ (Ainslie, 1975), $f(t) = 1/(1 + \alpha t)$ (Mazur, 1987). Obecna dyskusja w literaturze wskazuje na możliwość dyskontowania subaddytywnego (Read, 2001), które oznacza, że siła dyskontowania (wysokość stopy dyskontowej) w danym interwale czasu wzrasta wraz ze stopniem jego podziału na podokresy. Zagadnienie to pozostaje jednak w dużym stopniu niezbadane, dlatego też w pracy tej nie poświęcono mu uwagi.

¹⁰⁰ Sytuacja taka jest wykluczona przy zastosowaniu podejścia rozpadu wykładniczego, w którym czynnik dyskontujący jest niezależny od czasu.

wyższą. Jednak wraz z kurczeniem się dystansu czasowego, bliższa, aczkolwiek mniejsza nagroda wydaje się bardziej kusząca i zostaje wybrana. Z ekonomicznego punktu widzenia oznacza to odwrócenie preferencji, które jest najważniejszą konsekwencją dyskontowania hiperbolicznego.

5.1.3.3 Problem samokontroli

Przecinające się na rysunku 5.2 krzywe użyteczności, symbolizujące dynamiczną niekonsekwencję zachowań, to wprowadzenie w tematykę samokontroli jednostek. Problem ten zręcznie obrazują Loewenstein i Thaler (1989): „Rano, kiedy pokusa jest jeszcze odległa, ślubujemy sobie, że dziś uda nam się pójść wcześniej spać, wytrwać przy postanowionej diecie i nie spożywać alkoholu. Tej samej nocy przesiadujemy do trzeciej nad ranem, zażywamy porcję czekoladowej dekadencji i degustujemy [likieri wszelkiego rodzaju]”.

Wielu badaczy wskazuje w tym miejscu na wynikający z dyskontowania hiperbolicznego problem konfliktu intrapersonalnego (Thaler i Shefrin, 1981; Shefrin i Thaler, 1988). Zmiana preferencji jest rezultatem działalności dwóch wewnętrznych „ja”: racjonalnie planującego – organizatora – oraz niecierpliwego i krótkowzrocznego – wykonawcy. Zadaniem organizatora jest maksymalizacja przyszłej użyteczności, podczas gdy wykonawca dba o natychmiastową satysfakcję. Ponieważ dyskontowanie hiperboliczne rodzi antagonizm między teraźniejszymi i przyszłymi preferencjami, „hiperboliczny konsument” (*hyperbolic consumer*, Harris i Laibson, 2001) odczuwa dysonans między swoim faktycznym stanem oszczędności a wyższym zaplanowanym ich poziomem – jego dążenie do natychmiastowej gratyfikacji podkopuje długoterminowe plany oszczędzania.

Jak sugerują Harris i Laibson (2001), jednostki doświadczające problemów z samokontrolą wykazują tendencję do ograniczania inwestycji w środki łatwo zbywalne (akcje i obligacje) oraz do utrzymywania niewielkich (w relacji do majątku bądź dochodów) zasobów gotówkowych. Idzie to w parze ze strategią agresywnego zadłużania się na rynku szybko odnawialnych kredytów (np. wykorzystywanie limitu kart kredytowych). Dalsze konsekwen-

cje działań jednostek dyskontujących hiperbolicznie, na jakie zwraca uwagę Palacios-Huerta (2002), to niższa w porównaniu z przypadkiem wykładniczym stopa oszczędności, inwestowanie mniejszej części majątku, a zatem ograniczony popyt na aktywa finansowe, implikujący ich niższe ceny i wyższe stopy zwrotu.

Należy przyjąć, iż jednostki są zasadniczo świadome problemów z samokontrolą. Świadczą o tym podejmowane przez nie codzienne zabiegi ograniczania siebie. Ochroną przed przyszłym ubóstwem są mechanizmy, które zabezpieczają przed niecierpliwością „ja” wykonawcy. Thaler i Shefrin (1981) sygnalizują możliwość posługiwania się dwiema technikami. Pierwsza, trudniejsza, opiera się na silnej woli i jest rzadko stosowana ze względu na ubytek odczuwanej użyteczności, do jakiego prowadzi. Druga polega na uciekaniu się do zewnętrznych ograniczeń celem zapanowania nad poczynaniami krótkowzrocznego „ja”. Dowodem popularności tej formy samokontroli są chociażby systemy przymusowego oszczędzania (np. automatyczne odprowadzanie składek na rzecz funduszy emerytalnych), czy też omawiana w poprzednim rozdziale preferencja dywidend.

* * *

Kończąc rozważania nad percepcją przyszłości, warto zastanowić się nad możliwością uchwycenia wyborów międzyokresowych na poziomie makroekonomicznym. Powszechnie przypisuje się tę rolę rynkowej stopie procentowej skorygowanej o wpływ podatków, która ma stanowić agregację indywidualnych preferencji czasowych. Zaobserwowane anomalie, ze szczególnym wskazaniem na dyskontowanie hiperboliczne, poddają jednak w wątpliwość istnienie jednej społecznej stopy dyskontowej.

5.2 Pamiętanie przeszłości

We wcześniejszych rozważaniach wielokrotnie podkreślano, że ludzka percepcja wartości zależy od punktu odniesienia. Satysfakcja płynąca z posiadanego majątku jest funkcją różnicy między stanem bieżącym a pozio-

mem, do którego jednostka się zaadaptowała. Poziom ów nie jest bynajmniej statyczny – zmienia się wraz z upływem czasu oraz z przeszłymi doświadczeniami.

W tym podrozdziale przedyskutowano dynamiczne aspekty odczuwania użyteczności, rozważając kolejno znaczenie procesów adaptacyjnych oraz retrospektywnej oceny zdarzeń.

5.2.1 *Adaptacja*

Wiele codziennych obserwacji dowodzi ludzkiej zdolności przystosowania się do nieoptymalnych bądź wręcz szkodliwych warunków. Adaptacja to mechanizm, który redukuje efekty (percepcyjne, fizjologiczne czy emocjonalne) działania stałego lub powtarzającego się bodźca – np. straty na inwestycji (Frederick i Loewenstein, 1992). W szczególności proces określany mianem „hedonistycznej adaptacji” wiąże się ze zjawiskami wyzwalającymi reakcje o charakterze bardziej kognitywnym niż zmysłowym, prowadząc do zmian zainteresowań, systemu wartości, przekonań czy celów.

Adaptacja spełnia dwie podstawowe funkcje: (i) chroni jednostkę, ograniczając wpływ działania bodźca zewnętrznego oraz (ii) potęguje percepcję tych zmian, które w sposób istotny odbiegają od poziomu uznawanego za normalny. Umiejętność przystosowania się jest sposobem oszczędzania energii. Pozwala ona zaakceptować rzeczy, których nie da się zmienić oraz daje odwagę zmieniać to, co możliwe.

O znaczeniu adaptacji dla inwestorów można wnioskować chociażby na podstawie obserwacji niedawnych wydarzeń na rynkach kapitałowych. Praca ta powstała w momencie, kiedy indeksy na giełdach światowych oscylowały wokół swych najniższych poziomów odnotowanych na przestrzeni dekady. W okresie największej fascynacji rynkiem nowych technologii trudno było wyobrazić sobie sytuację, kiedy portfele aktywów odnotowują spadek wartości o kilkadziesiąt punktów procentowych. Ten czarny scenariusz zrealizował się jednak w przypadku bardzo wielu osób i stał się częścią codzienności giełdowej. Dzięki modyfikacji swych celów (np. zmiana perspektywy

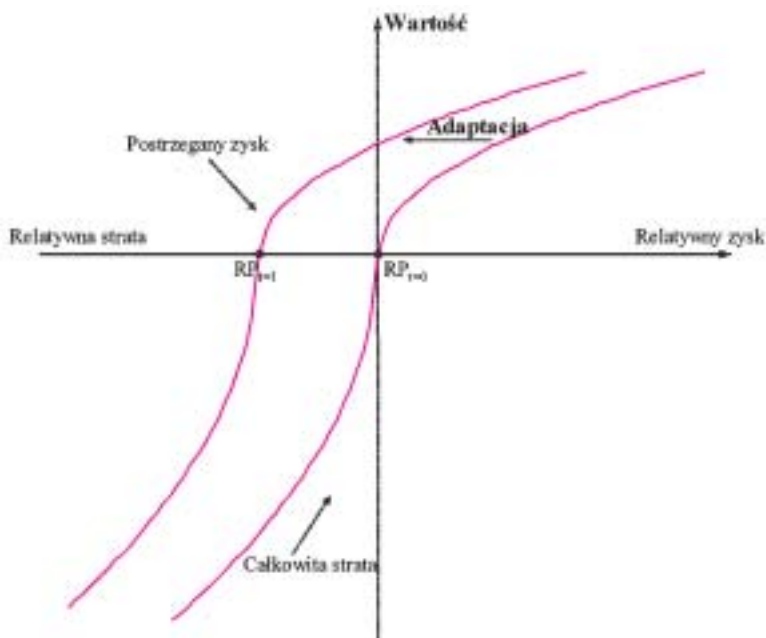
inwestycyjnej, skupienie uwagi na innych aktywach) inwestorzy z czasem przywykli do poniesionych strat.

5.2.1.1 Adaptacja a *prospect theory*

Dotychczasowe rozważania nad *prospect theory* ograniczały się do przypadku statycznego, w którym punkt odniesienia nie ulega przesunięciom. Analiza zjawiska adaptacji wymaga uwzględnienia dynamicznych zmian w położeniu krzywej wartości.

Upływ czasu jako nośnik procesów przystosowawczych zmienia sposób postrzegania wartości. Stąd też inwestor, odnotowawszy spadek ceny posiadanych aktywów, stopniowo przyzwyczaja się do straty. Owo „przyzwyczajanie się” można zobrazować jako przesunięcie na funkcji wartości ku (zaadaptowanemu) punktowi odniesienia.

Rozważmy przykład inwestora, który odnotowuje stratę 10 000 j.p. (por. rysunek 5.3). Wydarzenie to przesunęło go z punktu odniesienia $RP_{t=0}$ w dziedzinę ujemnego wyniku finansowego. Załóżmy dalej, że w następnych miesiącach sytuacja na rynku jest stabilna. Sam upływ czasu, mimo braku jakichkolwiek ruchów cen, sprawia jednak, że odczuwana wartość utraczonych 10 000 j.p. ulega zmianie. Inwestor „godzi się” z takim stanem rzeczy i stopniowo uznaje go za nowy punkt odniesienia ($RP_{t=1}$). Teraz pomyślna, nawet niewielka, zmiana ceny postrzegana będzie jako zysk, co na mocy efektu dyspozycji skłania do jego szybkiej realizacji. Skutek: mimo iż zagregowany wynik pozostaje ujemny, inwestor zamyka pozycję, oceniając ją w relacji do nowego punktu odniesienia.

Rysunek 5.3: Zastosowanie funkcji wartości do analizy procesów adaptacyjnych

Źródło: Opracowanie własne.

Zaprezentowany tok myślenia pozwala na istotne spostrzeżenie – mianowicie, że upływ czasu wywołuje zmianę stosunku do ryzyka. W świetle *prospect theory* początkowa strata (bodziec negatywny) stymuluje do zachowań ryzykownych, co w przypadku inwestora z przytoczonego przykładu oznacza utrzymywanie otwartej pozycji mimo ujemnego wyniku finansowego. Z biegiem czasu, dzięki procesom przystosowawczym, inwestor wraca do punktu odniesienia, przyjmując neutralny stosunek do ryzyka. Adaptacja przyczynia się zatem do stopniowego słabnięcia efektów wynikających z kształtu funkcji wartości, np. efektu dyspozycji, utopionych kosztów (por. podrozdział 4.1). Osiągnięcie nowego punktu odniesienia jest równoznaczne z pełnym dostosowaniem. Teraz wrażliwość zarówno na relatywne zyski, jak i na straty jest ponownie najwyższa (największe nachylenie funkcji wartości).

5.2.1.2 Pomiar adaptacji

Pierwszy ilościowy model adaptacji stworzony został przez Helsona (1964). Usiłował on uchwycić wpływ przeszłych bodźców (np. strat w przeszłości) na sposób doświadczania bodźców teraźniejszych (np. strat bieżą-

cych). Helson wprowadził pojęcie poziomu adaptacji, czyli takiego nasilenia bodźca, które jest postrzegane jako neutralne i nie wywołuje dalszej reakcji¹⁰¹. O poziomie adaptacji można zatem myśleć jako o dynamicznym punkcie odniesienia na funkcji wartości z *prospect theory*.

Równanie powszechnie używane do modelowania procesów adaptacyjnych przyjmuje postać (Hardie, Johnson i Fader, 1993):

$$AL_t = \alpha X_{t-1} + (1 - \alpha)AL_{t-1}, \quad (5.4)$$

gdzie:

AL_t – poziom adaptacji w okresie t ,

X_t – poziom bodźca (np. strata) w okresie t ,

α – parametr określający tempo adaptacji.

Jeśli $\alpha = 1$, wówczas poziom adaptacji AL_t odzwierciedla jedynie poziom bodźca z poprzedniego okresu X_{t-1} , a odczuwana wartość określona jest przez różnicę między bodźcami w okresach bieżącym i zeszłym. Jeśli zaś $\alpha = 0$, wówczas AL_t nie zależy od przeszłych bodźców.

Równanie (5.4) pokazuje, że poziom adaptacji w danym momencie jest średnią ważoną przeszłych bodźców, z których ostatnie otrzymują wagę wyższą niż poprzednie. Użyteczność odczuwana w momencie t , czyli u_t daje się przedstawić jako funkcja różnicy między stanem teraźniejszym a poziomem, do którego jednostka się zaadaptowała:

$$u_t = f(X_t - AL_t), \quad (5.5)$$

Uproszczony przebieg procesów adaptacyjnych inwestora w kolejnych okresach t pokazuje przykład w tabeli 5.1.

¹⁰¹ Helson za poziom adaptacji przyjmował średnią arytmetyczną przeszłych poziomów bodźca. Ujęcie to pomija rolę czasu i fakt silniejszego odczuwania zdarzeń niedawnych niż bardziej odległych w czasie. Stało się to głównym przedmiotem krytyki omawianego podejścia.

Tabela 5.1: Przykładowy przebieg adaptacji w sytuacji strat*

	t_1	t_2	t_3	t_4	t_5	t_6	t_7
X	-3	-5	-6	-4	-5	-5	-5
AL	-3**	-3	-4	-5	-4,5	-4,75	-4,875
u	0	-2	-2	1	-0,5	-0,25	-0,125

* Uwagi: Założono, że $\alpha = 0,5$; X – poziom bodźca, np. strata ponoszona na inwestycji w danym okresie t ; AL – postrzegany jako neutralny poziom adaptacji; u – postrzegana wartość straty, $u_t = X_t - AL_t$.

** Dla uproszczenia założono, że poziom straty w przeszłych okresach był stabilny i wynosił -3 j.p.

Źródło: Opracowanie własne; przykład zaadaptowany z (Frederick i Loewenstein, 2002).

Przykład ten ilustruje dwie podstawowe cechy modelowego podejścia do adaptacji:

- (1) Postrzegana wartość straty zależy od wyników w przeszłości (warto zwrócić uwagę, że w okresie t_2 strata postrzegana jest silniej niż w okresie t_5 , mimo iż obiektywne warunki są w obu momentach identyczne).
- (2) Intensywność bodźca o stałym poziomie maleje wraz z czasem (odczucia związane ze stratą 5 j.p. słabną w okresach od t_5 do t_7).

5.2.1.3 Znaczenie adaptacji

„Adaptacja jest jedną z podstawowych własności ludzkiej natury. Nawet w przypadku najważniejszych w życiu zdarzeń, osiągnąwszy nowy stan stabilności, mamy tendencję do powracania do naszego uprzedniego hedonistycznego poziomu [przyp. tłum.: poziomu satysfakcji]. Tak więc proces *stawania się* bogatym, nie zaś *bycia* bogatym, może stanowić źródło największej satysfakcji; i kiedy przyzwyczaimy się już do nowego standardu życia, możemy być tak samo szczęśliwi, jak wówczas, gdy byliśmy biedni” (Rabin, 2002).

Badania zachowań uczestników gier losowych dowodzą, że systematycznie przeszacowują oni swoją radość z potencjalnej wygranej. W momencie zaś, gdy ją faktycznie realizują, odnotowany przez nich wzrost użyteczności jest niewspółmierny do oczekiwań (Brickman, Coates, Janoff-Bulman, 1978). Owo „rozczarowanie” wynika bezpośrednio z przesunięcia punktu

odniesienia. Ludzie, przewidując przyszłą wygraną, zdają się jednak nieświadomi, jak szybko to dostosowanie następuje.

Także dla funkcjonowania rynku kapitałowego adaptacja ma nietrywialne znaczenie. Inwestorzy, pytani o motywby ich działalności na rynku, najczęściej podają chęć zarabiania pieniędzy i pomnażania majątku. Po stronie zysków adaptacja może zatem sprawiać, iż dążenie do bogactwa jest celem samym w sobie. Pozwala to domniemywać, że wśród uczestników rynku, podobnie jak w przypadku uczestników gier losowych, stopniowe przesunięcia punktu odniesienia powodują uczucie „nienasycenia”, które z kolei prowadzi do nadmiernej aktywności inwestycyjnej.

W sytuacji strat działanie adaptacji daje się pokazać na wspomnianym już przykładzie rynku nowych technologii. Trudno przypuszczać, iż którykolwiek z inwestorów kupujących akcje w roku 1999 liczył się z redukcją ich wartości, jaka niedawno miała miejsce. Co więcej, straty rzędu 50%, 70% czy 95% wartości były na początku „ery nowej ekonomii” zgoła niewyobrażalne. A jednak stały się faktem i zostały zaakceptowane. Dlaczego? Otóż spadki kursów następowały powoli, dając możliwość przyzwyczajania się do nich. Z tej perspektywy powiedzenie „czas goi rany” nabrało znaczenia także dla rynku kapitałowego.

Dotychczas zagadnienie adaptacji zajmowało badaczy w bardzo ograniczonym stopniu, jednak nieliczne przeprowadzone eksperymenty pozwalają na interesujące wnioski. Wpływ czasu na siłę efektu dyspozycji (por. punkt 4.1.2) wśród inwestorów australijskich obserwował Brown i in. (2002). Na podstawie analizy pierwszych ofert publicznych (IPO) na giełdzie australijskiej¹⁰² stwierdzono, że skłonność uczestników subskrypcji do szybszej realizacji zysków niż strat z czasem maleje. Natężenie efektu dyspozycji okazało się największe w okresie pierwszych dwóch tygodni, kiedy to częstotliwość realizacji była dwukrotnie wyższa dla zysków niż dla strat. Zjawisko uznane zostało za statystycznie istotne w okresie 90 dni od daty

oferty publicznej; natomiast po 264 dniach hipoteza o jego występowaniu musiała zostać odrzucona. Oznacza to, iż upływ 264 dni wystarczył na dokonanie się procesów przystosowawczych i zmianę postawy inwestorów ze skłonności do ryzyka na neutralność.

Proces adaptacji daje się także zaobserwować na wyższych poziomach agregacji, pociągając za sobą konsekwencje gospodarcze i społeczne. Badania potwierdziły bardzo niską korelację między wzrostem dochodu narodowego w danym kraju a poprawą zadowolenia społecznego. Dobrym przykładem jest Japonia: między 1958 a 1987 rokiem dochód narodowy *per capita* zwiększył się spektakularnie, bo aż pięciokrotnie. Jednocześnie, zgodnie z hipotezą postawioną w przytoczonej wypowiedzi Rabina, subiektywny poziom zadowolenia nie uległ zmianie (Easterlin, 1995).

5.2.2 *Retrospektywna ocena użyteczności*

Postrzegana atrakcyjność planowanej inwestycji zależy od przeszłych zdarzeń oraz doświadczanej w związku z nimi użyteczności. Pamiętana użyteczność ulega jednak systematycznym zniekształceniom i jest wrażliwa na działanie obiektywnie nieistotnych czynników. Błędna ocena przeszłości określa jakość przyszłych decyzji, stąd warta jest omówienia.

Dyskusja tego zagadnienia wymaga dygresji związanej z terminologią (Kahneman, 2000).

- (1) Chwilowa użyteczność (*moment utility*) mówi o sposobie odczuwania bodźca w danym momencie.
- (2) Profil użyteczności (*utility profile*) reprezentuje sekwencję chwilowych użyteczności postrzeganych w trakcie danego epizodu bądź dla serii niezależnych zdarzeń.
- (3) Pamiętana użyteczność (*remembered utility*) to dokonana przez jednostkę ogólna subiektywna ocena przeszłego zdarzenia.

¹⁰² Brown i in. badali zachowania uczestników subskrypcji w 480 pierwszych emisjach publicznych, jakie miały miejsce między grudniem 1995 roku a grudniem 2000 roku na rynku australijskim.

Większość studiów empirycznych nad retrospektywną ewaluacją skupia się na ocenie zdarzeń zmysłowych (wywołujących ból bądź przyjemność). Z ekonomicznego punktu widzenia istotne są jednak wnioski dotyczące oceny przeszłych zysków i strat. Taka zmiana perspektywy badawczej nie pozwala na bezkrytyczną inkorporację wyników dotychczasowych eksperymentów. Nieliczne dotąd prace związane z retrospektywną oceną wydarzeń ekonomicznych są w dużej mierze zasługą Langer, Sarina i Webera (2002) oraz Stephana i Kiella (2000).

5.2.2.1 Zasada ekstremum i końca

Zasada ekstremum i końca (*peak-end rule*) po raz pierwszy zaproponowana została przez Fredricksona i Kahnemana (1993) jako model służący pobieżnej ocenie pamiętanej użyteczności.

Zasada ta oznacza, iż sposób, w jaki jednostki dokonują oceny przeszłych zdarzeń odpowiada w przybliżeniu średniej arytmetycznej z dwóch wielkości: ekstremalnej i końcowej, które wystąpiły w czasie trwania epizodu. Tak wyznaczona wartość uznawana jest za reprezentatywną dla całej sekwencji, stąd też determinuje pamiętaną użyteczność.

Zastanówmy się nad działaniem zasady ekstremum i końca w przypadku inwestora. Przeszłym zdarzeniem może tu być obserwowana sekwencja wyników z danej inwestycji w kolejnych dniach, wartością ekstremalną – najwyższy poziom straty w rozważanym okresie, zaś końcową – ostatni obserwowany rezultat.

Odpowiedź na pytanie, czy inwestorzy oceniają sekwencje strat z przeszłości zgodnie z zasadą ekstremum i końca, była celem badawczym w eksperymencie Stephana i Kiella (2000). Respondentom (zawodowym uczestnikom rynku kapitałowego) przedstawiony został prosty scenariusz:

Od 1994 do 1996 roku trzech indywidualni inwestorzy A, B i C posiadali w swoim portfelu papiery wartościowe przynoszące straty. Każdy z nich zainwestował sumę 20 000 j.p. i ostatecznie poniósł stratę 5000 j.p. Rozkład strat w czasie trzech lat był jednak różny dla każdej z osób i przedstawiał się

następująco: $A(-2000, -2000, -1000)$, $B(-500, -3000, -1500)$,
 $C(-1000, -3000, -1000)$.

Uczestnicy doświadczenia poproszeni zostali o uszeregowanie osób A , B i C pod względem odczuwanego przez nie niezadowolenia. Wyznaczone na podstawie zasady ekstremum i końca wartości $A: -1500$, $B: -2250$, $C: -2000$ pozwalają przypuszczać, iż ocena niezadowolenia powinna się kształtować następująco: $B > C > A$. 71% badanych osób wyraziło pogląd, że inwestorzy A , B oraz C oceniają swoje straty w sposób niejednorodny; aż 53% z tych odpowiedzi szeregowało negatywne odczucia inwestorów zgodnie z predykcją omawianej zasady.

Możliwość zastosowania tej reguły do ewaluacji sekwencji wypłat pieniężnych potwierdził również eksperyment Langer, Sarina i Webera (2002). Ważną cechą jego konstrukcji było uzależnienie osiąganych zysków bądź strat od osobistego zaangażowania respondenta, przez co warunki laboratoryjne zbliżone zostały do rzeczywistości, w jakiej działają inwestorzy i managerowie finansowi. Badacze zaobserwowali dwa znamienne efekty:

- (1) Sekwencje z większą liczbą ekstremalnych strat postrzegane były jako bardziej przykre (choć sumaryczny wynik pozostawał identyczny, jak w sekwencjach bardziej „wygładzonych”) – efekt ekstremum.
- (2) Respondenci przedkładali dłuższe serie mniejszych strat, lecz charakteryzujące się wyższą całkowitą stratą, ponad ciągi krótsze, ale z wyższą stratą na końcu – efekt końca.

5.2.2.2 Ignorowanie długości sekwencji

Zasada ekstremum i końca wskazuje, iż ocena przeszłej użyteczności jest niezależna od długości sekwencji (*duration neglect*). Prowadzi to do pogwałcenia standardowej teorii oczekiwanej użyteczności, bowiem pamiętana użyteczność jest niewrażliwa na czas trwania doświadczenia nacechowanego ujemnie. Co więcej, odczucia negatywne mogą ulec osłabieniu poprzez wydłużanie sekwencji o kolejne bodźce znamionujące się zmniejszonym natężeniem.

Występowanie tego zjawiska w sytuacjach ekonomicznych udokumentowane zostało w badaniu Langer, Sarina i Webera. Respondenci oceniali lepiej serie długie, ale o mniejszych pojedynczych stratach, mimo iż całkowity ich wynik był gorszy niż w innych dostępnych opcjach o mniejszej liczbie pojedynczych realizacji. Badacze widzieli też możliwość, iż zwiększenie wariancji poziomu strat może powstrzymać przed ignorowaniem długości sekwencji¹⁰³.

Oczywiście normatywne podejście nie dopuszcza zjawisk takich, jak to obecnie omawiane czy też, jak zasada ekstremum i końca. Jednostka postępująca racjonalnie powinna oceniać sekwencje strat na podstawie prostego ich sumowania (dyskontowanego). Pionierskie prace wykazały, że dzieje się tak jedynie wówczas, gdy ewaluacja nie jest obciążona czynnikiem zaangażowania emocjonalnego bądź nie wiąże się z wysiłkiem respondenta. Zaangażowanie emocjonalne jest jednak zjawiskiem powszechnym na giełdzie i w wielu innych sytuacjach ekonomicznych. Mimo, iż badania naukowe znajdują się jeszcze w fazie początkowej, dotychczasowe rezultaty pozwalają na wyciągnięcie pierwszych wniosków.

Wydaje się, że adaptacja przebiega łatwiej w sytuacji małych kolejnych strat. Przy małych fluktuacjach rynku (niskiej zmienności), wrażliwość na straty może ulec obniżeniu, co z kolei pociąga za sobą brak reakcji w przypadku długich sekwencji negatywnych wyników. Może to prowadzić do utrzymywania otwartych pozycji przy pogłębiających się stratach. Jednocześnie nawet krótkotrwałe i nieznaczne odwrócenie niekorzystnego trendu jest w stanie wywołać pośpieszną likwidację inwestycji i dać podstawę do lepszej oceny globalnego rezultatu.

* * *

Przedstawione w tym rozdziale zagadnienia wskazują na niezwykle istotną rolę czasu w wyborach ekonomicznych. Przyszłe decyzje determinowane są

¹⁰³ Langer, Sarin i Weber powołują się w tym miejscu na obserwację Ariely'ego (1998), że jednostki zwracają większą uwagę na czas trwania epizodu wraz ze wzrostem wariancji w poziomie bodźca. Zjawisko to nie zostało jeszcze przeanalizowane dla wypłat pieniężnych.

bowiem zarówno sposobem dyskontowania przyszłych zdarzeń, procesami adaptacyjnymi, jak i oceną przeszłych doświadczeń oraz płynącą z nich pamiętaną użytecznością.

6

Zakończenie

Zaprezentowane dotychczas rozważania pozwalają udzielić odpowiedzi na pytania, jakie postawiono we wstępie do tej pracy.

- (1) Czy jednostki zdolne są wnioskować o nieznanym w sposób całkowicie obiektywny – zgodny z teorią rachunku prawdopodobieństwa i statystyki? Jeśli nie: czy ich sposób percepcji i przetwarzania informacji ulega systematycznym zniekształceniom?

Liczne badania empiryczne, a także obserwacja rzeczywistych zachowań uczestników rynku pozwalają uznać, iż jednostki, poprzez ograniczoną zdolność przetwarzania informacji, nie budują swych przekonań o prawdopodobieństwie zdarzeń zgodnie z obiektywnymi prawami matematyki. Przytłoczone wielością informacji muszą wspierać się heurystykami, których ceną są systematyczne błędy wnioskowania. Niepewność i presja czasu, dwa najistotniejsze determinanty działań uczestników rynków finansowych, szczególnie predysponują do tego typu zniekształceń. Tak zdeformowana informacja rzutuje na dokonywane następnie wybory inwestycyjne.

- (2) Czy teoria oczekiwanej użyteczności, będąca powszechnie akceptowaną normatywną koncepcją podejmowania decyzji w warunkach ryzyka, stanowi również adekwatny opis rzeczywistych zachowań jednostek? Jeśli nie: czy zadaniu temu jest w stanie sprostać model deskryptywny – *prospect theory* – zaproponowany przez Kahnemana i Tversky'ego?

Wartość teorii oczekiwanej użyteczności jako opisu rzeczywistych decyzji zostaje najdobitniej zanegowana poprzez wskazanie preferencji gwałcących aksjomat niezmienności, czyli niezależności preferencji od formy prezentacji problemu decyzyjnego. Jedynym zaproponowanym dotychczas podejściem wyjaśniającym takie wybory jest *prospect theory*, co zapewnia jej status właściwego deskryptywnego modelu podejmowania decyzji w warunkach

ryzyka. Elementy tworzące tę teorię umożliwiają wyjaśnienie takich frapujących (i nieracjonalnych z ekonomicznego punktu widzenia) zachowań inwestorów, jak np. krótkowzroczna awersja do strat, uleganie efektowi dyspozycji czy podatność na wpływ utopionych kosztów; pozwalają także uchwycić zmienność stosunku jednostek do ryzyka, której nie dopuszcza teoria normatywna.

- (3) Czy decyzje jednostek wykazują dynamiczną konsekwencję, postulowaną przez teorię zdyskontowanej użyteczności? Jeśli nie: czy zniekształcenia międzyokresowego procesu decyzyjnego dają się wyjaśnić w ramach alternatywnego modelu deskryptywnego?

Badania empiryczne obalają paradygmat dodatnich preferencji czasowych i dynamicznej konsekwencji wyborów międzyokresowych. Lepszym ich opisem jest natomiast model dyskontowania hiperbolicznego, który daje wyraz doświadczanemu przez ludzi problemowi samokontroli i wynikającemu stąd zjawisku odwracania się preferencji. Wpływ czasu na decyzje nie ogranicza się jedynie do percepcji przyszłości. Czas jest także nośnikiem procesów adaptacyjnych oraz czynnikiem determinującym sposób oceny użyteczności w retrospekcji. Sam upływ czasu może implikować zmianę postawy wobec ryzyka, co z kolei określa jakość decyzji podejmowanych w przyszłości. Koncepcja adaptacji stanowi zatem dopełnienie deskryptywnego modelu wyjaśniającego wpływ czasu na wybory ekonomiczne.

Celem badawczym tej pracy było stwierdzenie, czy wprowadzenie czynnika behawioralnego do nauki finansów pozwala na lepsze zrozumienie zachowań uczestników rynku i modyfikację dotychczasowych paradygmatów. W oparciu o zaprezentowane w rozdziałach 1 – 5 dowody można z całym przekonaniem udzielić pozytywnej odpowiedzi na to pytanie. Należy stwierdzić, iż uwzględnienie elementów psychologicznych pozwala na lepsze zrozumienie wyborów ekonomicznych jednostek, zaś behawioralna ekonomia finansowa jest w stanie uzupełnić lukę interpretacyjną między tradycyjną teorią a rzeczywistym funkcjonowaniem rynków finansowych. Trzeba przy tym zaznaczyć, iż behawioralny nurt finansów nie dostarczył jeszcze spójnej

hipotezy dającej się weryfikować na poziomie zagregowanym. Jednakże jego dynamiczny rozwój w ostatnich latach pozwala wierzyć, że takie zunifikowane podejście zostanie niebawem stworzone.

7 Bibliografia*

Pozycje książkowe

1. Baron, J. (2000). *Thinking and Deciding*. Cambridge: Cambridge University Press.
2. Bazerman, M. (1998). *Judgment in Managerial Decision Making*. Wydanie IV, New York: John Wiley & Sons.
3. Goldberg, J., von Nitzsch, R. (2000). *Behavioral Finance: Gewinnen mit Kompetenz*. Wydanie III, München: Finanzbuch Verlag.
4. Haugen, R. (1999). *Beast on the Wall Street. How Stock Volatility Devours Our Wealth*. Upper Saddle River: Prentice Hall.
5. Schwartz, H. (1998). *Rationality Gone Awry? Decision Making Inconsistent with Economic and Financial Theory*. Westport, Connecticut, London: Praeger.
6. Shefrin, H. (2000). *Beyond Greed and Fear. Understanding Behavioral Finance and the Psychology of Investing*. Boston: Harvard Business School Press.
7. Shiller, R.J. (2000). *Irrational Exuberance*. Princeton: Princeton University Press.
8. Shleifer, A. (2000). *Inefficient Markets. An Introduction to Behavioral Finance*. New York: Oxford University Press.
9. Thaler, R. (1993). *Advances in Behavioral Finance*. New York: Russell Sage Foundation.
10. Thaler, R. (1994). *Quasi Rational Economics*. New York: Russell Sage Foundation.

* W części oznaczonej jako „Bibliografia“ zawarto publikacje, z którymi zapoznano się w sposób bezpośredni przy pisaniu tej pracy. Por. uwagi w przypisie 105 dot. pozycji „Literatura przedmiotu”.

11. Wörneryd, K.-E. (2001). *Stock-Market Psychology. How People Value and Trade Stocks*. Cheltenham (UK), Northampton (MA, USA): Edward Elgar.
12. Zaleśkiewicz, T. (2003). *Psychologia inwestora giełdowego*. Gdańsk: Gdańskie Wydawnictwo Psychologiczne.

Artykuły

1. Ahlbrecht, M., Weber, M. (1997). An Empirical Study on Intertemporal Decision Making Under Risk. *Management Science*, 43(6), 813-826.
2. Benartzi, S., Thaler, R. (1995). Myopic Loss Aversion and the Equity Premium Puzzle. *Quarterly Journal of Economic*, 110, 75-92.
3. Black, F. (1986). Noise. *Journal of Finance*, 41(3), 529-542.
4. Camerer, F.C. (2000). Prospect Theory in the Wild. Evidence from the Field. W: *Choices, Values and Frames*, 288-300, Kahneman, D., Tversky, A. (Red.). Cambridge: Cambridge University Press.
5. De Bondt, W.F.M. (1998). A Portrait of the Individual Investor. *European Economic Review*, 42, 831-844.
6. De Bondt, W.F.M., Thaler, R. (1985). Does the Stock Market Overreact? *Journal of Finance*, 49, 793-805.
7. De Bondt, W.F.M., Thaler, R. (1990). Do Security Analysts Overreact? *American Economic Review*, 80(2), 52-57.
8. De Long, J.B., Shleifer, A., Summers, L.H., Waldmann, R.J. (1990). Noise Trader Risk in Financial Markets. *Journal of Political Economy*, 98(4), 703-738.
9. Fama, E. (1998). Market Efficiency, Long-Term Returns, and Behavioral Finance. *Journal of Financial Economics*, 49(3), 283-306.
10. Frederick, S., Loewenstein, G. (1992). Hedonic Adaptation. W: *Well-Being: The Foundations of Hedonic Psychology*, 302-329, Kahneman, D., Diener, E., Schwarz, N. (Red.). New York: Russel Sage Foundation.

11. Frederick, S., Loewenstein, G., O'Donoghue, T. (2002). Time Discounting and Time Preference: A Critical Review. *Journal of Economic Literature*, 40, 351-401.
12. French, K.R., Poterba, J.M. (1991). Investor Diversification and International Equity Markets. *American Economic Review*, 81(2), 222-226.
13. Gilovich, T., Vallone, R., Tversky, A. (1985). The Hot Hand in Basketball: On the Misperception of Random Sequences. *Cognitive Psychology*, 17, 295-314.
14. Kahneman, D. (2000). Evaluation by Moments: Past and Future. W: *Choices, Values, and Frames*, 693-708, Kahneman, D., Tversky, A. (Red.). Cambridge: Cambridge University Press.
15. Kahneman, D., Knetsch, J.L., Thaler, R. (1991). Anomalies. The Endowment Effect, Loss Aversion and Status Quo Bias. *Journal of Economic Perspectives*, 5(1), 193-206.
16. Kahneman, D., Riepe, W. (1998). Aspects of Investor Psychology. *Journal of Portfolio Management*, 24(4), 52-65.
17. Kahneman, D., Tversky, A. (1972). Subjective Probability: A Judgment of Representativeness. *Cognitive Psychology*, 3, 430-454.
18. Kahneman, D., Tversky, A. (1973). On the Psychology of Prediction. *Psychological Review*, 80, 237-251.
19. Kahneman, D., Tversky, A. (1979). Prospect Theory: An Analysis of Decisions under Risk. *Econometrica*, 47(2), 263-292.
20. Kahneman, D., Tversky, A. (1984). Choices, Values and Frames. *American Psychologist*, 39(4), 341-350.
21. Kahneman, D., Tversky, A. (1986). Rational Choice and the Framing of Decisions. *Journal of Business*, 59(4), 251-278.
22. Kahneman, D., Tversky, A. (1992). Advances in Prospect Theory: Cumulative Representation of Uncertainty. *Journal of Risk and Uncertainty*, 5, 297-324.
23. Lee, C., Shleifer, A., Thaler, R. (1991). Investor Sentiment and the Closed-End Fund Puzzle. *Journal of Finance*, 46(1), 75-109.

24. Loewenstein, G., Prelec, D. (1992). Anomalies in Intertemporal Choice: Evidence and Interpretation. *Quarterly Journal of Economics*, 107(2), 573-597.
25. Loewenstein, G., Prelec, D. (1993). Preferences for Sequences of Outcomes. *Psychological Review*, 100(1), 91-108.
26. Loewenstein, G., Thaler, R. (1989). Anomalies: Intertemporal Choice. *Journal of Economic Perspectives*, 3(1), 181-193.
27. Odean, T. (1998a). Are Investors Reluctant to Realise Their Losses? *Journal of Finance*, 53, 1775-1798.
28. Odean, T. (1998b). Volume, Volatility, Price and Profit When All Traders Are Above Average. *Journal of Finance*, 53(6), 1887-1934.
29. Odean, T. (1999). Do Investors Trade Too Much? *American Economic Review*, 89, 1279-1298.
30. Prelec, D. (2000). Compound Invariant Weighting Function in Prospect Theory. W: *Choices, Values and Frames*, 67-92, Kahneman, D., Tversky, A. (Red.). Cambridge: Cambridge University Press.
31. Rabin, M. (2002). A Perspective on Psychology and Economics. *European Economic Review*, 46(4-5), 657-685.
32. Shefrin, H., Statman, M. (1984). Explaining Investor Preference for Cash Dividends. *Journal of Financial Economics*, 13, 253-282.
33. Shefrin, H., Statman, M. (1985). The Disposition to Sell Winners Too Early and Ride Losses Too Long. *Journal of Finance*, 40, 777-790.
34. Shefrin, H., Thaler, R. (1988). The Behavioral Life-Cycle Hypothesis. *Economic Inquiry*, 26, 609-643.
35. Shleifer, A., Summers, L.H. (1990). The Noise Trader Approach to Finance. *Journal of Economic Perspectives*, 4(2), 19-33.
36. Simon, H. (1955). A Behavioral Model of Rational Choice. *Quarterly Journal of Economics*, 69, 99-118.
37. Thaler, R. (1980). Toward a Positive Theory of Consumer Choice. *Journal of Economic Behavior and Organisation*, 81, 39-60.

38. Thaler, R. (1981). Some Empirical Evidence on Dynamic Inconsistency. *Economic Letters*, 8, 201-207.
39. Thaler, R. (1985). Mental Accounting and Consumer Choice. *Marketing Science*, 4(3), 199-214.
40. Thaler, R. (1999). Mental Accounting Matters. *Journal of Behavioral Decision Making*, 12, 183-206.
41. Thaler, R., Shefrin, H. (1981). An Economic Theory of Self-Control. *Journal of Political Economy*, 89(2), 391-406.
42. Tversky, A., Fox, C. (1995). Weighting Risk and Uncertainty. *Psychological Review*, 102(2), 269-283.
43. Tversky, A., Kahneman, D. (1971). Belief in the Law of Small Numbers. *Psychological Bulletin*, 2, 105-110.
44. Tversky, A., Kahneman, D. (1973). Availability: A Heuristic for Judging Frequency and Probability. *Cognitive Psychology*, 4, 207-232.
45. Tversky, A., Kahneman, D. (1974). Judgment under Uncertainty: Heuristics and Biases. *Science*, 185, 1124-1131.
46. Tversky, A., Kahneman, D. (1982). Causal Schemas in Judgments under Uncertainty. W: *Judgment under Uncertainty: Heuristics and Biases*, 117-128, Kahneman, D., Slovic, P., Tversky, A. (Red.). Cambridge: Cambridge University Press.
47. Tversky, A., Sattath, S., Slovic, P. (1988). Contingent Weighting in Judgment and Choice. *Psychological Review*, 80, 204-217.

Pozostałe publikacje

1. Barberis, N., Huang, M. (2001). *Mental Accounting, Loss Aversion, and Individual Stock Returns*. NBER Working Paper, nr 8190, <http://papers.nber.org/papers/w8190>, wersja z 27.04.2003.
2. Barberis, N., Thaler, R.H. (2001). *A Survey of Behavioral Finance*. NBER Working Paper, nr 9222, <http://www.nber.org/papers/w9222>, wersja z 03.09.2002.

3. Brown, P., Chappel, N., da Silva Rosa, R., Walter, T. (2002). *The Reach of the Disposition Effect: Large Sample Evidence Across Investor Classes*. Working Paper, University of New South Wales, Sydney, Australia.
4. Coval, J., Shumway, T. (2001). *Do Behavioral Biases Affect Prices?* Working Paper, University of Michigan, USA.
5. Fisher, K., Statman, M. (1999). *The Sentiment of Investors, Large and Small*. Working Paper, <http://212.67.202.199/~msewell/ta/sentiment/FiSt/99.pdf>, wersja z 27.04.03.
6. Harris, C., Laibson, D. (2001). *Hyperbolic Discounting and Consumption*. Working Paper, Harvard University, USA.
7. Lamont, O.A., Thaler, R.H. (2001). *Can the Market Add and Subtract? Mispricing in Stock Market Carve-Outs*. NBER Working Paper, nr 8302, <http://papers.nber.org/papers/w8302>, wersja z 27.04.2003.
8. Langer, T., Sarin, R., Weber, M. (2002). *The Retrospective Evaluation of Payment Sequences: Duration Neglect and Peak-and-End Effects*. Working Paper, University of Mannheim, 69131 Mannheim, Germany.
9. Palacios-Huerta, I. (2002). *Consumption and Portfolio Rules under Hyperbolic Discounting*. Working Paper, Brown University, Department of Economics, USA.
10. Shiller, R.J. (1998). *Human Behavior and the Efficiency of the Financial System*. NBER Working Paper, nr 6375, <http://papers.nber.org/papers/w6375>, wersja z 11.11.2002.
11. Shiller, R.J. (2002). *From Efficient Market Theory to Behavioral Finance*. Cowles Foundation Discussion Paper, nr 1385, http://papers.ssrn.com,abstract_id=349660, wersja z 08.11.2002.
12. Stephan, E., Kiel, G. (1998). *Urteilsprozesse bei Professionellen Akteuren im Finanzmarkt. (Abschluss einer Experimentellen Studie mit Trade- und Salespersonal einer Deutschen Geschäftsbank)*. Beitrag zum 41. Kongress der Deutschen Gesellschaft für Psychologie, Dresden.

13. Stephan, E., Kiell, G. (2000). *Decision Processes in Professional Investors: Does Expertise Moderate Judgmental Biases?* Paper presented at the IAREP/SABE 2000 Conference in Baden, Vienna, July 2000.

8

Literatura przedmiotu*

1. Ainslie, G. (1975). Specious Reward: A Behavioral Theory of Impulsiveness and Impulse Control. *Psychological Bulletin*, 82(4), 463-496.
2. Allais, M. (1953). Le Comportement de l'Homme Rationnel devant le Risque: Critique des Postulats et Axiomes de l'École Américaine. *Econometrica*, 21, 503-546.
3. Ariely, D. (1998). Combining Experience over Time: The Effects of Duration, Intensity Changes and On-Line Measurements on Retrospective Pain Evaluations. *Journal of Behavioural Decision Making*, 11, 19-45.
4. Arkes, H., Blumer, C. (1985). The Psychology of Sunk Cost. *Organizational Behavior and Human Decision Process*, 35, 124-140.
5. Benartzi, S., Thaler, R. (1999). Risk Aversion or Myopia? Choices in Repeated Gambles and Retirement Investment. *Management Science*, 45, 364-381.
6. Benzion, U., Rapoport, A., Yagil, J. (1989). Discount Rates Inferred from Decisions: An Experimental Study. *Management Science*, 35, 270-284.
7. Bernard, V., Thomas, J.K. (1989). Post Earnings Announcement Drift: Delayed Price Response or Risk Premium? *Journal of Accounting Research, Supplement*, 27, 1-36.
8. Bernoulli, D. ([1738] 1954). Exposition of a New Theory of the Measurement of Risk. *Econometrica*, 22, 23-26.
9. Black, F. (1976). Studies of Stock Price Volatility Changes. *Proceedings of the 1976 Meetings of the American Statistical Association, Business and Economics Section*, 177-181.

* „Literatura przedmiotu” obejmuje opracowania, których autorka nie przestudiowała w sposób bezpośredni, a jedynie poprzez odniesienia i komentarze zawarte w innych publikacjach. Pozycje tu zawarte wymienione zostały w niektórych miejscach pracy. Dlatego też za zasadne uznano umieszczenie odrębnej listy tych tytułów.

10. Bodurtha, J., Kim, D., Lee, C.M. (1993). Closed-End Country Funds and U.S. Market Sentiment. *Review of Financial Studies*, 8, 879-918.
11. Borges, B., Goldstein, D.G., Ortmann, A., Gigerenzer, G., Todd, P.M. (1999). Can Ignorance Beat the Stock Market? W: *Simple Heuristics that Make Us Smart*, 59-72, Gigerenzer, G., Todd, P.M., ABC Group (Red.). Oxford: Oxford University Press.
12. Brickman, P., Coates, D., Janoff-Bulman, R. (1978). Lottery Winners and Accident Victims: Is Happiness Relative? *Journal of Personality and Social Psychology*, 36(8), 917-927.
13. Campbell, J.Y., Shiller, R.J. (1988). The Dividend-Price Ratio and Expectations of Future Dividends and Discount Factors. *Review of Financial Studies*, 1, 195-228.
14. Chan, L., Jegadeesh, N., Lakonishok, J. (1996). Momentum Strategies. *Journal of Finance*, 51, 1681-1713.
15. Chapman, G.B., Elstein, A.S. (1995). Valuing the Future: Temporal Discounting of Health and Money. *Medical Decision Making*, 15(4), 373-386.
16. Chapman, L.J., Chapman, J.P. (1969). Illusory Correlation as an Obstacle to the Use of Valid Psychodiagnostic Signs. *Journal of Abnormal Psychology*, 74, 271-280.
17. Chopra, N., Lakonishok, J., Ritter, J.R. (1992). Measuring Abnormal Performance. Do Stocks Overreact? *Journal of Financial Economics*, 31, 235-268.
18. Constantinides, G. (1984). Optimal Stock Trading with Personal Taxes: Implications for Prices and the Abnormal January Returns. *Journal of Financial Economics*, 13, 65-69.
19. Constantinides, G. (1990). Habit Formation: A Resolution of the Equity Premium Puzzle. *Journal of Political Economy*, 98, 519-543.
20. De Bondt, W.F.M., Thaler, R.H. (1987). Further Evidence on Investor Overreaction and Stock Market Seasonality. *Journal of Finance*, 42, 557-581.

21. Dimson, E., Marsh, P.R., Staunton, M. (2000). *The Millennium Book: A Century of Investment Returns*. London: ABN-Amro/London School of Business.
22. Easterlin, R.A. (1995). Will Raising the Incomes of All Increase the Happiness of All? *Journal of Economic Behavior and Organisation*, 27, 35-47.
23. Edwards, W. (1968). Conservatism in Human Information Processing. W: *Formal Representation of Human Judgment*. Kleinmuntz, B. (Red.). New York: Wiley and Sons.
24. Einhorn, H.J, Hogarth, R.M. (1978). Confidence in Judgment: Persistence of Illusion of Validity. *Psychological Review*, 85, 395-416.
25. Ellsberg, D. (1961). Risk, Ambiguity, and the Savage Axioms. *Quarterly Journal of Economics*, 75, 643-669.
26. Epstein, L.G., Zin, S.E. (1990). "First-Order" Risk Aversion and the Equity Premium Puzzle. *Journal of Monetary Economics*, 26(3), 387-407.
27. Fama, E.F., French, K.R. (2000). *The Equity Premium*. Working Paper, Center for Research in Security Prices, University of Chicago, Graduate School of Business, USA.
28. Figlewski, S. (1979). Subjective Information and Market Efficiency in a Betting Market. *Journal of Political Economics*, 87, 75-88.
29. Fischhoff, B. (1975). Hindsight: Thinking backward. *Psychology Today*, 8, 71-76.
30. Fischhoff, B. (1980). For These Condemned to Study the Past: Reflections on Historical Judgment. W: *New Directions for Methodology of Behavioral Science: Fallible Judgment in Behavioral Research*, 79-93, Shweder, R.A., Fiske, D.W. (Red.). San Francisco: Jossey-Bass.
31. Fischhoff, B., Slovic, P., Lichtenstein, S. (1977). Knowing with Certainty. The Appropriateness of Extreme Confidence. *Journal of Experimental Psychology: Human Perception and Performance*, 3, 552-564.
32. Fisher, I. (1930). *The Theory of Interest*. New York: Macmillan.

33. Fisher, K., Statman, M. (1997). Investment Advice from Mutual Fund Companies. *Journal of Portfolio Management*, 24(1), 9-25.
34. Fredrickson, B.L., Kahneman, D. (1993). Duration Neglect in Retrospective Evaluations of Affective Episodes. *Journal of Personality and Social Psychology*, 65(1), 45-55.
35. French, K., Schwert, W., Stambaugh, R. (1987). Expected Stock Returns and Volatility. *Journal of Financial Economics*, 19, 3-30.
36. Friedman, M. (1953). The Case for Flexible Exchange Rates. W: *Essays in Positive Economics*. Chicago: University of Chicago Press.
37. Froot, K.A., Dabora, E. (1999). How Are Stock Prices Affected by the Location of Trade. *Journal of Financial Economics*, 53, 189-216.
38. Gerke, W., Bienert, H. (1993). Überprüfung des Dispositionseffektes und seiner Auswirkungen in computerisierten Börsenexperimenten. *Zeitschrift für betriebswirtschaftliche Forschung – Sonderheft*, 31, 169-194.
39. Gigerenzer, G., Todd, P.M., ABC Research Group. (1999). Probabilistic Mental Models: A Brunswikian Theory of Confidence. *Psychological Review*, 98, 506-28.
40. Gigerenzer, G. (1996). On Narrow Norms and Vague Heuristics: A Reply to Kahneman and Tversky (1996). *Psychological Review*, 103(3), 592-596.
41. Goldstein, D.G., Gigerenzer, G. (1999). The Recognition Heuristic: How Ignorance Makes Us Smart. W: *Simple Heuristics that Make Us Smart*, 37-58, Gigerenzer, G., Todd, P.M., ABC Group (Red.). Oxford: Oxford University Press.
42. Gourville, J.T., Soman, D. (1998). Payment Depreciation: The Effects of Temporally Separating Payments from Consumption. *Journal of Consumer Research*, 25(2), 160-174.
43. Griffin, D., Tversky, A. (1992). The Weighting of Evidence and the Determinants of Confidence. *Cognitive Psychology*, 24(3), 411-435.

44. Hardie, B.G.S., Johnson, E.J., Fader, P.S. (1993). Modeling Loss Aversion and Reference Dependence Effects on Brand Choice. *Marketing Science*, 12, 378-394.
45. Hardouvelis, G., La Porta, R., Wizman, T. (1994). What Moves the Discount on Country Equity Funds? W: *The Internationalisation of Equity Markets*. Frankel, J. (Red.). Chicago: The University of Chicago Press.
46. Hartman, R.S., Doane, M.J., Woo, C.K. (1991). Consumer Rationality and the Status Quo. *Quarterly Journal of Economics*, 106(1), 141-162.
47. Harvey, C.M. (1986). Value Functions for Infinite Period Planning. *Management Science*, 32, 1123-1139.
48. Haugen, R., Talmor, E., Torous, W. (1991). The Effect of Volatility Changes on the Level of Stock Prices and Subsequent Expected Returns. *Journal of Finance*, 46(3), 985-1007.
49. Helson, H. (1964). *Adaptation Level Theory: An Experimental and Systematic Approach to Behaviour*. New York: Harper & Row.
50. Hershey, J., Johnson, E., Meszaros, J., Robinson, M. (1990, June). *What is the Right to Sue Worth?* Wharton School, University of Pennsylvania.
51. Hoffrage, U., Hertwig, R. (1999). Hindsight Bias: A Price Worth Paying for Fast and Frugal Memory. W: *Simple Heuristics that Make Us Smart*, 191-208, Gigerenzer, G., Todd, P.M., ABC Group (Red.). Oxford: Oxford University Press.
52. Holcomb, J.L., Nelson, P.S. (1992). Another Experimental Look at Individual Time Preference. *Rationality Society*, 4(2), 199-220.
53. Hsee, C.K., Abelson, R.P., Salovey, P. (1991). The Relative Weighting of Position and Velocity of Satisfaction. *Psychological Science*, 2(4), 263-266.
54. Kahneman, D., Lovallo, D. (1993). Timid Choices and Bold Forecasts: A Cognitive Perspective on Risk Taking. *Management Science*, 39, 17-31.
55. Kahneman, D., Tversky, A. (1996). On the Reality of Cognitive Illusions. *Psychological Review*, 103(3), 582-591.

56. Knetsch, J.L., Sinden, J.A. (1984). Willingness to Pay and Compensation Demanded: Experimental Evidence of an Unexpected Disparity in Measures of Value. *Quarterly Journal of Economics*, 99, 507-521.
57. Knight, F.H. (1921). *Risk, Uncertainty and Profit*. Boston: Houghton and Mifflin.
58. Lakonishok, J., Smidt, S. (1986). Volume for Winners and Losers: Taxation and Other Motives for Stock Trading. *Journal of Finance*, 41, 951-976.
59. Le Roy, S., Porter, R. (1981). Stock Price Volatility: A Test Based on Implied Variance Bounds. *Econometrica*, 49, 97-113.
60. Lichtenstein, S., Slovic, P. (1973). Response-Induced Reversals of Preferences in Gambling: An Extended Replication in Las Vegas. *Journal of Experimental Psychology*, 101, 16-20.
61. Loewenstein, G. (1988). *The Weighting of Waiting: Response Mode Effect in Intertemporal Choice*. Working Paper, Center for Decision Research, University of Chicago, USA.
62. Mankiw, N.G., Zeldes, S.P. (1991). The Consumption of Stockholders and Nonstockholders. *Journal of Financial Economics*, 29, 97-112.
63. Markowitz, H. (1952). Portfolio Selection. *Journal of Finance*, 7(1), 77-91.
64. Mazur, J.E. (1987). An Adjustment Procedure for Studying Delayed Reinforcement. W: *The Effect of the Delay and Intervening Events on Reinforcement Value*. Commons, M., Mazur, J.E., Nevin, A., Rachlin, H. (Red.). Hillsdale, NJ: Erlbaum.
65. Mehra, R., Prescott, E. (1985). The Equity Premium: A Puzzle. *Journal of Monetary Economics*, 2, 145-161.
66. Nofsinger, J.R. (2001). *Investment Madness: How Psychology Affects Your Investing*. London: Prentice Hall.
67. Northcraft, G.B., Neale, M.A. (1987). Experts, Amateurs and Real Estate: An Anchoring-and-Adjustment Perspective on Property Pricing De-

- cisions. *Organisational Behavior and Human Decision Processes*, 39, 84-97.
68. Read, D. (2001). Is Time-Discounting Hyperbolic or Subadditive? *Journal of Risk and Uncertainty*, 23, 5-32.
69. Redelmeier, D., Heller, D. (1993). Time Preference in Medical Decision Making and Cost-Effectiveness Analysis. *Medical Decision Making*, 13(3), 212-217.
70. Redelmeier, D., Tversky, A. (1992). On the Framing of Multiple Prospects. *Psychological Science*, 3(3), 191-193.
71. Samuelson, P. (1937). A Note on Measurement of Utility. *Review of Economic Studies*, 4, 155-161.
72. Samuelson, W., Zeckhauser, R. (1988). Status Quo Bias in Decision Making. *Journal of Risk and Uncertainty*, 1, 7-59.
73. Samuelson, P. (1963). Risk and Uncertainty: A Fallacy of Large Numbers. *Scientia*, 98, 108-113.
74. Savage, L.J. (1954). *The Foundations of Statistics*. New York: Wiley.
75. Shiller, R.J. (1981). Do Stock Prices Move Too Much to Be Justified by Subsequent Movements in Dividends? *American Economic Review*, 71(3), 421-436.
76. Shiller, R.J. (1988). Portfolio Insurance and Other Investor Fashions as Factors in the 1987 Stock Market Crash. W: *NBER Macroeconomic Annual*, 287-295. Cambridge (USA, MA): National Bureau of Economic Research.
77. Shiller, R.J. (1990a). Market Volatility and Investor Behavior. *American Economic Review*, 80, 58-62.
78. Shiller, R.J. (1990b). Speculative Prices and Popular Models. *Journal of Economic Perspectives*, 4, 55-65.
79. Shleifer, A., Vishny, R. (1997). The Limits of Arbitrage. *Journal of Finance*, 52, 35-55.
80. Siegel, J.J. (1994). *Stocks for the Long Run*. New York: Irvin.
81. Simon, H. (1957). *Models of Man*. New York: Wiley.

82. Slovic, P., Lichtenstein, S. (1971). Comparison of Bayesian and Regression Approaches in the Study of Information Processing and Judgment. *Organisational Behavior and Human Performance*, 6, 649-744.
83. Stephan, E. (1999). Die Rolle von Urteilsheuristiken bei Finanzentscheidungen: Ankereffekte und kognitive Verfügbarkeit. W: *Finanzpsychologie*, 101-134, Fischer, L., Kutsch, T., Stephan, E. (Red.). München, Wien: R. Oldenbourg Verlag.
84. Stevenson, M.K. (1992). The Impact of Temporal Context and Risk on the Judged Value of Future Outcomes. *Organisational Behavior and Human Decision Process*, 52, 455-491.
85. Thaler, R. (1992). *The Winner's Curse: Paradoxes and Anomalies of Economic Life*. New York: The Free Press.
86. Tversky, A., Heath, C. (1991). Preferences and Beliefs: Ambiguity and Competence in Choice under Uncertainty. *Journal of Risk and Uncertainty*, 4, 5-28.
87. Tversky, A., Kahneman, D. (1972). Belief in the Law of Small Numbers. *Psychological Bulletin*, 2, 105-110.
88. Tversky, A., Kahneman, D. (1981). The Framing of Decisions and the Rationality of Choice. *Science*, 211, 453-458.
89. Tversky, A., Kahneman, D. (1983). Extensional Versus Intuitive Reasoning: The Conjunction Fallacy in Probability Judgment. *Psychological Review*, 90, 239-315.
90. Varey, C.A., Kahneman, D. (1992). Experiences Extended Across Time: Evaluation of Moments and Episodes. *Journal of Behavioral Decision Making*, 5(3), 169-185.
91. von Neumann J., Morgenstern, O. (1944). *Theory of Games and Economic Behavior*. Princeton: Princeton University Press.
92. Weber, M., Camerer, C. (1998). The Disposition Effect in Securities Trading: An Experimental Analysis. *Journal of Economic Behaviour and Organisation*, 33, 167-184.

93. Wood, G. (1978). The Knew-It-All-Effect. *Journal of Experimental Psychology: Human Perception and Performance*, 4, 345-353.