



8. PROGNOZOWANIE POPYTU



W rozdziale omówiono teorię prognozowania. Szczególną uwagę poświęcono prognozowaniu popytu. Jest to przewidywanie przyszłych zdarzeń (związanych z zapotrzebowaniem i popytem), którego celem jest zminimalizowanie ryzyka związanego z podejmowaniem decyzji biznesowych. Do najważniejszych poruszanych zagadnień tego rozdziału zaliczono:

- zasady i trendy prognozowania,
- prognozowanie szeregów czasowych,
- procedurę opracowania prognoz na podstawie szeregów czasowych,
- metody prognozowania i błędy prognoz,
- zagadnienie sztucznej inteligencji w prognozowaniu.

8.1. Wprowadzenie

Prognozowanie jest szeroko stosowaną, multidyscyplinarną nauką. Jest istotną czynnością, która jest wykorzystywana w celu podejmowania decyzji biznesowych w wielu obszarach planowania: ekonomicznego, przemysłowego i naukowego (Chatfield, 2001). Zbudowana prognoza wspomaga podjęcie decyzji mikro- i makroekonomicznych. Wspomaga także podejmowanie działań aktywizujących lub przeciwstawiających się jakiemuś zjawisku. Jest również źródłem cennych informacji. Prognozowanie można nazwać przewidywaniem; przewidywaniem przyszłego zapotrzebowania, przewidywaniem popytu, przewidywaniem sprzedaży lub nowego trendu. Zmiany warunków rynkowych, do których przedsiębiorstwo się musi dostosowywać można więc przewidzieć.

Jednak to przewidywanie nie może bazować wyłącznie na doświadczeniu i intuicji menadżerów, ponieważ niewłaściwe przewidywanie lub podejmowanie decyzji na podstawie niewłaściwie przygotowanych przewidywań może skutkować problemami finansowymi przedsiębiorstwa. Dlatego nie każde przewidywanie jest prognozowaniem,



ponieważ prognozowanie (zwane również predykcją) opiera się na racjonalnych, zwykle naukowych podstawach.

Prognozowanie to wnioskowanie o zdarzeniach nieznanym, na podstawie zdarzeń znanych (Cieślak, 2005). Przykładowo można przewidzieć, że: (1) zdarzenie nastąpi, ponieważ nastąpiło w przeszłości; (2) zdarzenie nastąpi, ponieważ wskazuje na to częstotliwość jego występowania; (3) zdarzenie nastąpi, ponieważ powiązane jest z innymi zdarzeniami, które wystąpiły (Dittmann, 2003).

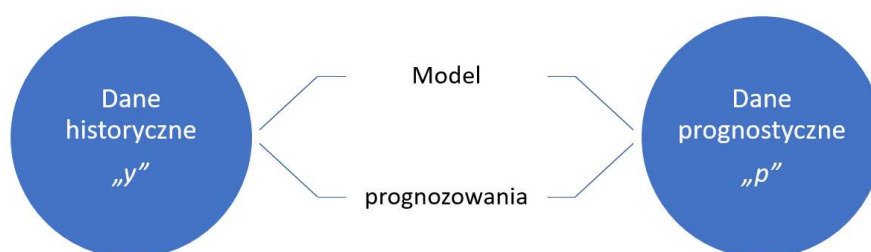
Prognozy są opracowywane (budowane) na podstawie przesłanek bardzo różnej natury. Biorąc jednak pod uwagę ich naukową naturę są konstruowane przede wszystkim na podstawie modeli statystycznych, ekonometrycznych oraz z wykorzystaniem badań operacyjnych. Prognozy opracowanie są z wykorzystaniem danych historycznych – zaistniałych w przeszłości. Zaś z punktu widzenia logistyki powiązane są z danymi nieodległej przeszłości. Szczególnie w złożonych branżach produktowych (np. przemysł motoryzacyjny), ale nie tylko, prognozy popytu mają kluczowe znaczenie dla obszaru sprzedaży oraz również dla wydajności systemu produkcyjnego.

Prognoza dotyczy zawsze konkretnego horyzontu prognostycznego. Horyzont prognozy to przedział (T, T_i) , gdzie: T – moment bieżący, T_i – moment końcowy.

W zależności od horyzontu czasowego problem prognozowania dzieli się ogólnie na trzy obszary: krótko-, średnio- i długoterminowe prognozowanie. Jak wspomniano wcześniej z punktu widzenia menadżera logistyki kluczowym jest prognozowanie **krótkoterminowe**. Obejmuje ono horyzonty predykcyjne od jednej godziny do tygodnia. Interesujące z punktu widzenia pracy menadżera logistyki jest także prognozowanie średnioterminowe, które odnosi się do prognoz od jednego miesiąca do maksymalnie roku. Wyróżnić wreszcie można prognozy długoterminowe, które charakteryzują się horyzontem predykcyjnym dłuższym niż rok. Mają one mniejsze znaczenie dla działalności operacyjnej związanej z logistyką. Teoria chaosu w dużej mierze pokazała, że długoterminowe prognozowanie jest daremnym wysiłkiem. Można więc przyjąć założenie, że dla szeroko rozumianej działalności logistycznej im horyzont prognozy dłuższy, tym prawdopodobieństwo zaistnienia zbudowanej prognozy maleje. Zmniejsza się jej pewność. Również prognozowanie dla wyrobu w dłuższej perspektywie niż cykl życia tego wyrobu nie ma sensu.



Wartość (znaczenie) modeli prognostycznych opiera się na ich zdolności do tworzenia dokładnych prognoz. Prognozy są więc tak dobre, jak założenia zastosowanego modelu. Ważne jest, aby mieć świadomość i wiedzieć, jakie są te założenia. W przypadku, gdy którekolwiek z tych założeń okaże się błędne, prognozy można ponownie ocenić, zmodyfikować i poprawić. Głównym problemem dokładności prognoz jest nieprzewidywalność trendów gospodarczych oraz zewnętrznych wydarzeń i kryzysów. Dlatego należy dosadnie stwierdzić, że konkretne wartości prognoz obarczone są **błędem i niepewnością**. Wynika to z faktu, że aby zbudować prognozę na przyszłość, korzystać należy z dostępnych danych historycznych; na podstawie dotychczasowych, przeszłych obserwacji. Tak więc, na podstawie posiadanej wiedzy o przeszłości następuje określenie przyszłości (rys.8.1).



Rysunek 8.1. Uogólniony model prognozowania

Źródło: (Dittmann, 2003)

Nie należy zapomnieć o potrzebie wymiany informacji w zarządzaniu łańcuchem dostaw, które ma kluczowe znaczenie dla sukcesu przewidywania popytu (Altendorfer i Felberbauer, 2023). Im bardziej dokładna informacja o zapotrzebowaniu, tym bardziej dokładna będzie opracowana prognoza. Również kluczowa jest bieżąca aktualizacja informacji o popycie (polega na zmianie wcześniejszych informacji, np. o wielkości zamówienia), dzięki której następują aktualizacje popytu w horyzoncie czasowym i eliminacja asymetrii informacji. Ali i inni wskazują, że dzielenie się pełnymi prognozami popytu, a nie ostateczną wielkością zamówienia, jest korzystne dla wydajności łańcucha dostaw (Ali i in., 2012).

Wyzwania związane z udanym prognozowaniem to coś więcej niż tylko trudności techniczne związane z opracowaniem dokładnego modelu prognostycznego. Modele prognostyczne muszą być opracowywane z jasnym zrozumieniem zarówno natury sytuacji, dla której prognoza jest pożądana, jak i zasobów dostępnych do sporządzenia prognozy. Ważne



jest, aby upewnić się, że wybrana zmienna odnosi się bezpośrednio do potrzebnych danych prognozowanych (Sheldon, 1993). Nie oznacza to, że prognozy są bezużyteczne, ale że ci, którzy z nich korzystają, powinni stale monitorować swoje środowisko operacyjne, aby wykryć wszelkie czynniki, które wskazują na niespójne lub nieregularne wzorce.

Tabela 8.1. Korzyści prognozowania popytu

Zidentyfikowana korzyść prognozowania	Uzasadnienie
Lepsza organizacja produkcji	znając prognozowaną wielkość sprzedaży wyrobów gotowych organizacja może z wyprzedzeniem zaplanować właściwą wielkość produkcji oraz właściwą wielkość zapotrzebowania w surowce i opakowania; eliminuje tym samym niedobory na linii produkcyjnej
Większa kontrola zapasów zabezpieczających	znając prognozowaną wielkość sprzedaży wyrobów gotowych można zaplanować zapas zabezpieczający, który zagwarantuje pokrycie zapotrzebowania rynku
Skuteczniejsze ograniczenie przestarzałego asortymentu	znając prognozowaną wielkość sprzedaży wyrobów gotowych można skupić się na obsłudze wyłącznie asortymentu niezbędnego do pokrycia zapotrzebowania; można eliminować zaleganie przestarzałych asortymentów, a w efekcie optymalizować koszty zamrożonego kapitału w zapas oraz koszty magazynowania
Lepsza satysfakcja klientów i poprawa wizerunku organizacji	znając prognozowaną wielkość sprzedaży wyrobów gotowych można zagwarantować utrzymywanie właściwego ilościowo poziomu zapasów w magazynie
Efektywniejsze wykorzystanie przestrzeni magazynowej	znając prognozowaną wielkość sprzedaży wyrobów gotowych można gromadzić tylko niezbędne zapasy asortymentu; można również istotnie ograniczyć wykorzystywaną przestrzeń magazynową
Skuteczniejsza kontrola i minimalizacja kosztów	znając prognozowaną wielkość sprzedaży wyrobów gotowych można dokładniej zaplanować budżet organizacji oraz podjąć kroki związane z dokładniejszą kontrolą wydatków

Źródło: (Wojciechowski i Wojciechowska, 2015; Wolny i Kmiecik, 2020)

Pomimo, że prognoza jest obarczona nieścisłością, stanowi ważną kierunkową przyszłego działania operacyjnego przedsiębiorstwa. Uzasadnieniem tworzenia prognoz w przedsiębiorstwie jest również cykliczność, która występuje w działalności przedsiębiorstw. Przewidujemy, że skoro jakieś zdarzenie wystąpiło w przeszłości to może również wystąpić w przyszłości. Jeśli natomiast zdarzenie występowało w przeszłości z określoną częstością to prawdopodobieństwo, że wystąpi ponownie wzrasta. Pomimo wielu niepewności zbudowana,



przy pomocy naukowych metod, prognoza jest przesłanką podjęcia racjonalnej decyzji w zakresie działalności przedsiębiorstwa.

Praktyka gospodarcza pokazuje również, że **prosta metoda prognozowania nie oznacza automatycznie metody gorszej** (Kucharski, 2013). Jak wskazuje Kucharski metody naiwne potrafią prognozować te same dane ze zbliżoną dokładnością. Są zaś o wiele łatwiejsze do zastosowania.

Prognozowanie wpływa m.in. na określenie mocy produkcyjnych, sposobów wytwarzania i świadczenia usług, a co za tym idzie wpływa także pośrednio na elementy takie jak liczba pracowników czy poziom kosztów. W wyniku działań związanych z prognozowaniem popytu możliwe jest uzyskanie wielu korzyści przez organizację (tab. 8.1).

Należy wskazać kilka własności prognoz:

1. Prognozy formułowane są z wykorzystaniem dorobku nauki (opracowane i zweryfikowane modele matematyczne);
2. Prognozy odnoszą się do określonej przyszłości;
3. Prognozy są weryfikowane empirycznie (po upływie określonego czasu);
4. Prognozy są akceptowalne przez osobę opracowującą prognozę.

Prognozy wspomagają w przedsiębiorstwie proces decyzyjny i jednocześnie spełniają różnorodne funkcje (Gajda, 2001):

- preparacyjną – prognoza jest impulsem do podjęcia konkretnego działania, ale nie ma wpływu na prognozowane zjawisko. Na jej podstawie podejmowane są wyłącznie decyzje gospodarcze,
- aktywizującą – prognoza jest impulsem do podjęcia konkretnego działania i jednocześnie wpływa na prognozowane zjawisko. Podejmowane są więc działania, które mają na celu urealnienie prognozy (prognozy samorealizujące lub sprzyjające, które wywołują działania sprzyjające realizacji prognoz) lub unicestwiające prognozy (prognozy ostrzegające, które wywołują działania przeciwdziałające ich realizacji).

Trzeba jednak pamiętać, że zbudowane prognozy mogą się łatwo załamać z powodu przypadkowych zmiennych, których nie można włączyć do modelu lub mogą być po prostu



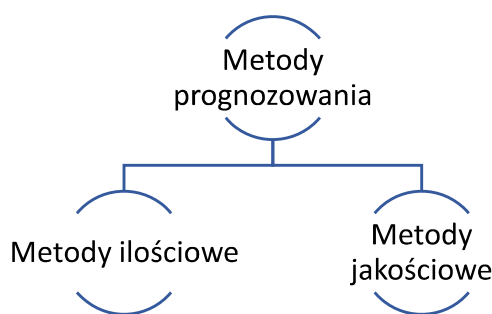
błędne od samego początku. Z tego powodu prognozowanie może być dla organizacji niebezpieczne. Wyróżnić można trzy problemy związane z prognozowaniem:

- dane, na podstawie których następuje prognozowanie zawsze będą stare, dotyczące okresów historycznych. Nigdy nie ma więc gwarancji, że warunki z przeszłości utrzymają się w przyszłości,
- nie można uwzględnić wyjątkowych lub nieoczekiwanych zdarzeń ani efektów zewnętrznych (przykład pandemii Covid-19; wpływ wojny i konfliktów zbrojnych; wpływ nieprzewidywanych kryzysów gospodarczych). Zdarzenia związane z czarnymi łabędziami stały się bardziej powszechne, ponieważ nasze zaufanie do prognoz wzrosło,
- prognozy nie mogą uwzględniać własnego wpływu. Kierownictwo staje się niewolnikiem danych historycznych i trendów prognostycznych.

Właściwie prowadzone prognozowanie pozwala przedsiębiorcom i menadżerom planować działalność z wyprzedzeniem, zwiększając jednocześnie szanse na utrzymanie konkurencyjności na rynkach.

8.2. Klasyfikacja metod prognozowania

Wyróżnić można dwie podstawowe grupy metod prognozowania – ilościowe i jakościowe (rys. 8.2). Prognoza zaliczana do tzw. ilościowych metod prognozowania przyjmuje postać wyrażoną konkretną liczbą (prognoza punktowa) lub ewentualnie przedziałem liczbowym (prognoza przedziałowa).



Rysunek 8.2. Metody prognozowania - rodzaje

Źródło: (Dittmann, 2000)



Prognozy jakościowe przyjmują postać nieliczbową. Odnoszą się do analizowanego zjawiska w przyszłości i oszacowania jego wzrostu, spadku lub braku zmiany. Prognozy jakościowe można traktować jako oparte na opiniach ekspertów rynkowych.

Z punktu widzenia specjalisty ds. logistyki kluczowe są jednak prognozy, które można określić przy pomocy liczb, a więc **prognozy ilościowe**. Prognozy ilościowe pomijają czynnik ekspercki i starają się usunąć czynnik ludzki z analizy. Podejścia te koncentrują się wyłącznie na danych.

Prognozy ilościowe	Modele szeregów czasowych
	Modele ekonometryczne
	Modele analogowe
	Modele zmiennych wiodących
	Modele analizy kohortowej
	Testy rynkowe

Rysunek 8.3. Metody ilościowe prognozowania

Źródło: (Dittmann, 2000)

Prognozy ilościowe można sklasyfikować uwzględniając stosowane modele (rys. 8.3). Na potrzeby książki skupiono się na **modelach szeregów czasowych**.

8.3. Prognozowanie szeregów czasowych

Jedną z najczęściej stosowanych metod prognostycznych do prognozowania popytu są metody oparte na modelach szeregów czasowych. Szeregi czasowe to metodologie eksploracji złożonych i sekwencyjnych typów danych. W modelach szeregów czasowych dane sekwencyjne, które składają się z ciągów danych liczbowych, są rejestrowane w równych odstępach czasu (np. na minutę, na godzinę lub na dzień). Popularność tych metody wynika z możliwości pozyskania poprzez prognozę informacji o przyszłym przebiegu zjawiska jakie jest obserwowane. W związku z tym nie ma potrzeby zbierania i analizowania kolejnych danych z innych źródeł. Prognozowanie przy pomocy szeregów czasowych jest również często stosowane ze względu na istnienie dużego prawdopodobieństwa jej wystąpienia. Praktyka gospodarcza pokazuje również, że opracowane przy pomocy modeli szeregów czasowych prognozy nie są gorsze od prognoz uzyskiwanych w oparciu o bardziej skomplikowane modele.



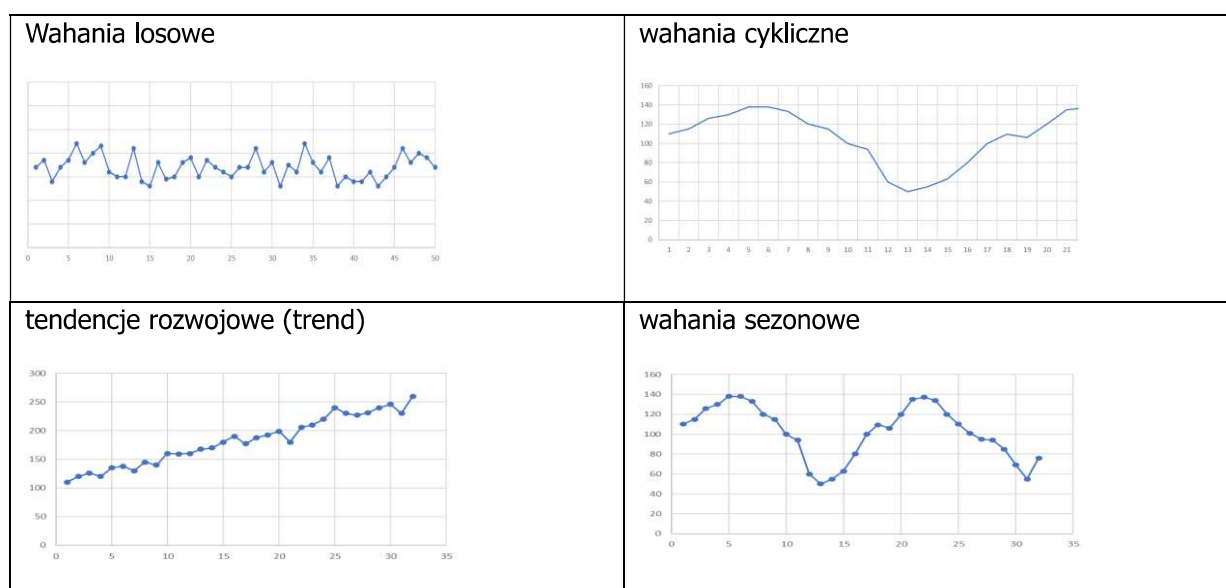
Doświadczenie pokazuje także, że modele szeregów czasowych mają potencjał rozwojowy. Każda kolejna modyfikacja metody lub kolejna metoda prognozowania według szeregów czasowych powinna z założenia poprawiać jakość jej wyników.

8.4. Dekompozycja szeregu czasowego

Prognozy są budowane z wykorzystaniem danych z szeregu czasowego. Dzieje się tak niezależnie od przyjętej metody prognozowania. Dane szeregu czasowego (zmiennie) są uporządkowane chronologicznie, od danych najstarszych, do danych najnowszych. Należy podkreślić, że ostatnia dana nie jest daną pokrywającą się z momentem budowania prognozy. W publikacjach i opracowaniach naukowych przyjęto, że y_t określa zawsze konkretną wartość szeregu y w okresie (momencie) t .

Składowymi szeregu czasowego są wahania przypadkowe, tendencja rozwojowa (trend), wahania cykliczne oraz wahania sezonowe (tab. 8.2).

Tabela 8.2. Składowe szeregu czasowego



Źródło: opracowanie własne

O każdym z nich można powiedzieć kilka słów (Cieslak, 1997):

- wahania przypadkowe – są to losowe, przypadkowe i nie dające się przewidzieć zmiany zmiennej szeregu o zróżnicowanej sile, które są obserwowane w czasie i nie wykazują żadnej wyraźnej tendencji. Powiązane są z błędami o charakterze statystycznym lub prognostycznym. Należy nadmienić, że wahania przypadkowe



nie wpływają w istotny sposób na charakter badanego zjawiska, są one wliczone (jako błąd prognozy) w schemat postępowania prognostycznego, który należy wyeliminować,

- tendencje rozwojowe (trend) – są to długookresowe skłonności danych szeregu do jednokierunkowych (monotonicznych) zmian prognozowanej zmiennej. Mają one charakter wzrostu lub spadku. Dotyczą najczęściej trwałego zjawiska, które ma wpływ na analizowane dane. W danych szeregu czasowego mogą występować zarówno tendencje rozwojowe oraz wahania przypadkowe. W celu wyodrębnienia tendencji rozwojowych potrzeba z reguły większej liczby danych historycznych. Stąd obserwujemy praktyczną zasadę: im dłuższy czas obserwacji danych historycznych, tym większa możliwość precyzyjnego określenia rodzaju trendu. Trend przedstawiany jest za pomocą funkcji matematycznej liniowej lub nieliniowej. W celu określania wielkości przyszłego popytu stosuje się cztery funkcje trendu: liniową, wykładniczą, logarytmiczną i potęgową. Najczęściej stosowaną funkcją tendencji jest jednak funkcja liniowa,
- wahania cykliczne – są to długookresowe, rytmiczne wahania wartości zmiennej układające się wokół trendu lub stałego poziomu, które utrzymują się dłuższy czas (dłużej niż rok). Są wynikiem cykliów koniunkturalnych. Mogą być obserwowane różne długości cykli oraz różna ich dynamika. Wahania cykliczne związane są więc ze zmianami w aktywności gospodarczej przedsiębiorstw, kryzysów lub ożywienia gospodarczego albo zamożności społeczeństwa. Do analizy wahań cyklicznych i budowy prognozy przyszłego zapotrzebowania potrzebne są dane historyczne miesięczne, kwartalne lub roczne z kilku ostatnich lat,
- wahania sezonowe – są to wahania wartości zmiennej szeregu czasowego wokół trendu lub stałego poziomu, które mają skłonność do powtarzania się w regularnych (sezonowych) odstępach, nie przekraczających roku. W takim wypadku dokładność budowanych prognoz będzie zależała od rodzaju i skali wahań sezonowych, liczby i rodzaju luk w dostępnych danych, horyzontu prognozy.

Identyfikacja i analiza wskazanych składowych szeregu czasowego nazywana jest dekompozycją szeregu czasowego.



8.5. Przygotowanie danych w szeregach czasowych

Zanim podejmie się kroki związane z budową prognoz warto wcześniej wykonać wstępną obróbkę danych, zwaną także **czyszczeniem danych**. Konieczna jest ich weryfikacja w celu wyeliminowania błędów lub wartości odstających. Pominięcie tego kroku może powodować zniekształcenie wyników prognoz i w efekcie błędy we wnioskowaniu. Trzeba pamiętać, że dane powiązane z przypadkami nietypowymi (odstającymi lub rzadko występującymi) to informacje prawdziwe, co do których prognosta nie ma wątpliwości. Są one sprawdzone i wiarygodne.

Tabela 8.3. Wybrane rodzaje wartości odstających

Addytywna wartość odstająca Występuje jako wartość zaskakująco duża lub zaskakująco mała dla pojedynczej obserwacji. Nie ma ona wpływu na kolejne obserwacje – wartość szeregu nie ulega dalszemu odstawianiu.	
Innowacyjna wartość odstająca Występuje jako odchylenie z dalszym wpływem na obserwacje. Zaobserwować można wpływ początkowy (pierwszy) z efektem opóźnienia i rozciągnięcia na kolejne obserwacje (malejące lub rosnące). Wpływ ten może zmniejszać się lub zwiększać w miarę upływu czasu.	
Wartość odstająca zmiany przemijającej Występuje w przypadku, kiedy wpływ wygasa się wykładniczo w miarę kolejnych obserwacji. Ostatecznie szereg powraca do normalnego poziomu.	
Wartość odstająca sezonowo addytywna Występuje jako zaskakująco duża lub zaskakująco mała wartość, która występuje okresowo (w regularnych odstępach).	

Źródło: opracowanie własne



Są różne strategie postępowania w przypadku zaobserwowania danych odstających. To między innymi:

- brak działania – polega na ignorowaniu danych nietypowych, ponieważ niektóre metody prognozowania są odporne na występowanie danych nietypowych,
- filtrowanie przypadków z danymi odstającymi – polega na usuwaniu tych danych; nie jest to jednak najlepsza strategia,
- zastępowanie danych nietypowych – jest to popularna strategia, w której wartości odstające są zastępowane: (1) wartością 0, (2) wartością średnią; (3) wartością filtra maksymalną/minimalną lub (4) inną wartością określoną na podstawie kryterium merytorycznego.

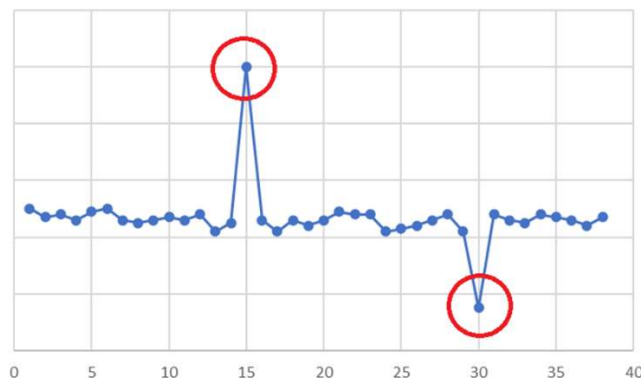
Pozostawienie w szeregu czasowym danych nietypowych zniekształca wynik ich analizy oraz utrudnia formułowanie wniosków, ponieważ dane nietypowe są wartościami ekstremalnie małymi lub ekstremalnie dużymi. Ze względu na swoją niespójność wartości te nazywamy wartościami odstającymi. W efekcie zwiększają rozstęp w szeregu czasowym (zakres minimum-maksimum). Dane nietypowe mają więc duży wpływ zniekształcający na budowaną wartość prognozy (tab. 8.3).

Zawsze decyzję o zmianie wielkości danej nietypowej lub jej usunięciu cechuje duży subiektywizm prognosty, dlatego wymaga ona rozważań. Ograniczenie subiektywizmu prognosty podczas usuwania przypadków nietypowych jest możliwe w odniesieniu do danych ilościowych. Przykładowo, można zastosować regułę odchylenia standardowego. Oznacza to, że jeśli dana historyczna jest nietypowa (np. wykracza poza zakres średniej grupowej ($\bar{x}$) powiększonej lub pomniejszonej o 2 lub 3 odchylenia standardowego), to ulega zmianie lub jest usuwana.

Warto każdy szereg danych poddać procedurze dekompozycyjnej. Wyszczególnić można kilka kroków:

1. Zidentyfikowanie formy funkcyjnej szeregu, co oznacza określenie rodzaju trendu.
2. Wyszukanie obserwacji odstających i zastąpienie ich wartościami uśrednionymi lub tak zwanymi filtrami górnym i/lub dolnym.
3. Weryfikacja czy ostatnia obserwowana dana z szeregu zachowuje się typowo; jeśli nie oczyszczenie jej.
4. Identyfikacja współczynnika kierunkowego trendu, w celu określenia stabilności głównego trendu obserwowanego w szeregu.

5. Zbadanie stabilności bieżącego, krótkookresowego trendu.

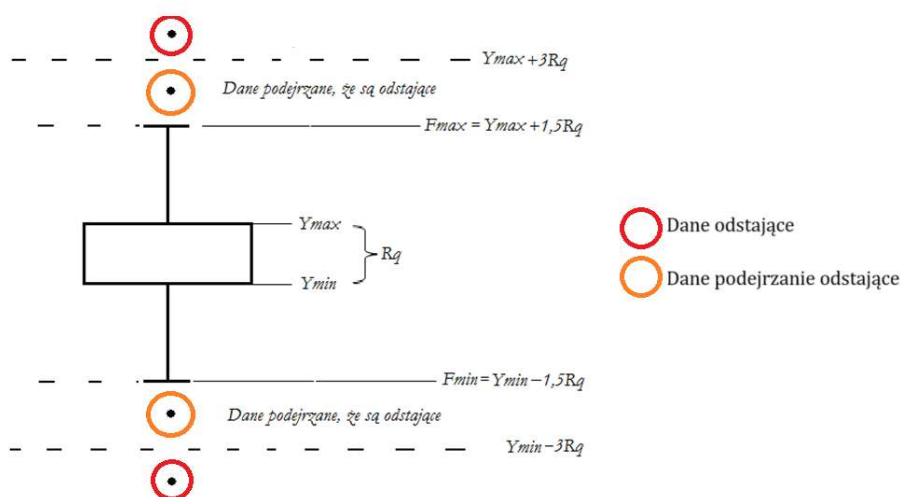


Rysunek 8.4. Wyraźne niezgodne wielkości z ogólną prawidłowością szeregu czasowego

Źródło: opracowanie własne

Filtry (minimalny lub maksymalny) są dobrym rozwiązaniem w sytuacji, gdy pojawiają się obserwacje odstające (rys. 8.4).

Filtr służy do tego, aby skorygować dane przed przystąpieniem do budowy kolejnej prognozy. Wartość odbiegająca od prawidłowego szeregu jest zastępowana filtrem minimalnym lub maksymalnym (rys. 8.5). Wartości ekstremalnie odstające zawarte są powyżej $y_{max} + 3R_q$ lub poniżej $y_{min} - 3R_q$. Wartości podejrzane, że są odstające zawarte są w przedziałach $(y_{min} - 1,5R_q; y_{min} - 3R_q)$ oraz $(y_{max} + 1,5R_q; y_{max} + 3R_q)$.



Rysunek 8.5. Wartości filtrów i dane odstające

Źródło: (Grzybowska, 2009)



Przedstawia to model:

$$F_{min} = y_{min} - 1,5R_q$$

$$F_{max} = y_{max} + 1,5R_q$$

$$R_q = y_{max} - y_{min}$$

gdzie:

F_{min} – wartość filtra minimalnego,

F_{max} – wartość filtra maksymalnego,

y_{min} – wartość minimalna określona z szeregu czasowego,

y_{max} – wartość maksymalna określona z szeregu czasowego,

R_q – rozstęp międzykwartylowy.

8.6. Metody prognozowania szeregów czasowych

Metody prognozowania szeregów czasowych podzielone zostały ze względu na występujący trend danych. Wyszczególnić można prognozy dla popytu stałego, o charakterze trendu (rosnącego lub malejącego) oraz dla popytu sezonowego (rys. 8.6).

Popyt o charakterze stałym	Metody naiwne
	Metody średniej
	Metoda wygładzania wykładniczego (Browna)
	Model ARMA
Popyt o charakterze trendu	Metoda wygładzania wykładniczego (Holta)
	Metoda prostej regresji liniowej
	Model ARIMA
	Model RW
Popyt o charakterze sezonowym	Metoda wygładzania wykładniczego (Holta-Wintersa)

Rysunek 8.6. Metody ilościowe prognozowania szeregów czasowych

Źródło: opracowanie własne

W prezentowanych modelach parametr y odnosi się zawsze do wartości rzeczywistych popytu, zaś paramet p odnosi się zawsze do zbudowanej prognozy.



Metody naiwne

Metody naiwne prognozowania są charakteryzowane jako proste, szybkie i tanie. Pozwalają na opracowanie prognoz z niewielkiej liczby danych historycznych. Metody naiwne służą również jako punkt odniesienia dla innych metod prognozowania (Kucharski, 2013).

Metody naiwne są najprostszymi metodami mechanicznymi. Zostały opracowane na założeniu, że w przyszłości nie nastąpią wyraźne zmiany w popycie. Bardzo dobrze nadają się wszędzie tam, gdzie nie ma dużych wahań prognozowanej zmiennej. Opierają się wyłącznie na historycznych obserwacjach. Modele naiwne mają pamięć tylko jednej (ostatniej) obserwacji, dlatego nie odfiltrują szumu w danych, ale raczej skopiują go w przyszłość.

Modele naiwne składają się z prostych modeli projekcyjnych. Oznacza to, że wymagają danych wejściowych z ostatnich obserwacji i nie jest wykonywana żadna analiza statystyczna. Są niezwykle proste i jednocześnie zaskakująco skuteczne. Zaletą tych metod jest szybka decyzja o wielkości prognozowanej wartości. Wadą jest jednak brak możliwości analizy związków przyczynowo-skutkowych, które leżą u podstaw prognozowanej zmiennej.

Na potrzeby opracowania zostaną przedstawione 3 metody naiwne: (1) metoda naiwna, (2) sezonowa metoda naiwna; (3) metoda dryfu.

Metoda naiwna prosta (Naive Forecast)

Prognoza naiwna to taka, w której wartość prognozowana na dany okres jest po prostu równa wartości zaobserwowanej w poprzednim okresie. Przedstawia to formalny model:

$$p_t = y_{t-1}$$

gdzie:

p_t – prognoza na przyszły okres

y_{t-1} – wartość rzeczywista popytu z okresu poprzedniego.



Formuła stosowana w Excelu:

prognoza_(t) = popyt_(t-1)



Sezonowa metoda naiwna (Seasonal Naive method)

Również metoda naiwna jest przydatna w przypadku danych o małych wahaniami sezonowych. W takiej sytuacji każda prognoza jest równa ostatniej obserwowanej wartości z tego samego sezonu (np. z tego samego miesiąca poprzedniego roku). Prognozy przyjmują wartość zaobserwowaną w sezonie wcześniejszym. Model jest przydatny, kiedy występują małe wahania przypadkowe oraz sezonowość addytywna. Przedstawia to formalny model:

$$p_{t+h|T} = y_{t+h-m(k+1)}$$

gdzie:

$p_{t+h|T}$ – prognoza na przyszły okres,

m – okres sezonowy,

k – część całkowita $(h-1)/m$ (tj. liczba pełnych lat w okresie prognozy poprzedzającym $T+h$).

Model wygląda na bardziej skomplikowany niż jest w rzeczywistości. Przykładowo, jeśli budowana jest prognoza na podstawie danych miesięcznych, odnosi się ona do wszystkich przyszłych wartości miesięcznych i jest równa ostatniej zaobserwowanej wartości z tego miesiąca poprzedniego roku. W przypadku danych kwartalnych prognoza wszystkich przyszłych wartości $Q2$ jest równa ostatniej zaobserwowanej wartości $Q2$ (gdzie $Q2$ oznacza drugi kwartał). Podobne zasady obowiązują w innych miesiącach i kwartałach oraz w innych okresach sezonowych.



Formuła stosowana w Excelu:

prognoza(t) = popyt(t, rok poprzedni)

Prognoza na okres t jest równa popytowi z adekwatnego okresu roku poprzedniego. Sezonowa metoda naiwna będzie wymagała rocznego opóźnienia.

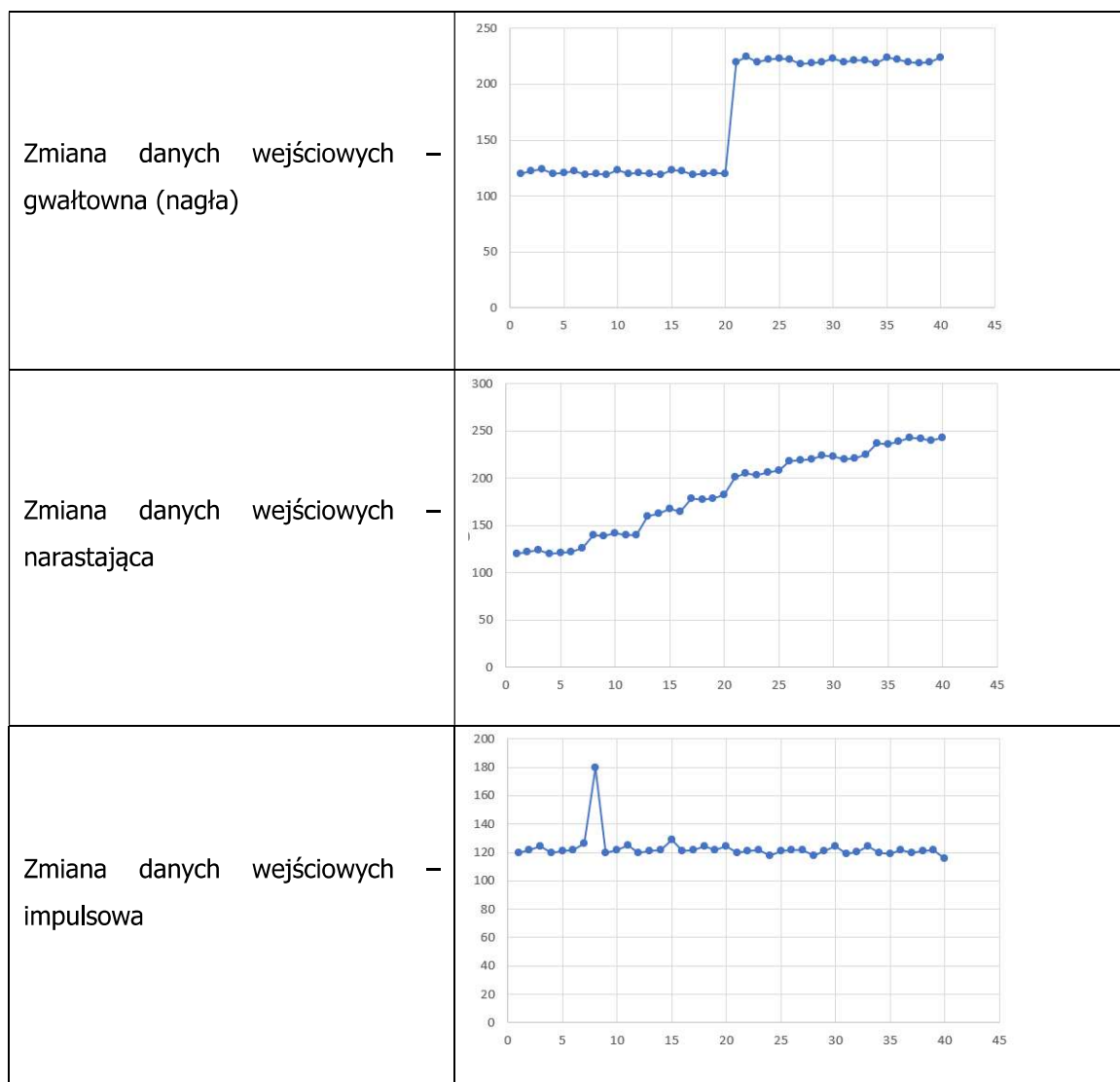
Metoda dryfu (Drift method)

Metoda dryfu jest odmianą sezonowej metody naiwnej polega na umożliwieniu wzrostu lub spadku prognoz w czasie, gdzie wielkość zmiany w czasie (zwana dryfem) jest ustawiona



jako średnia zmiana widoczna w danych historycznych. Metoda wykorzystuje dodatkowy komponent o nazwie dryf.

Tabela 8.4. Rodzaje dryfu



Źródło: opracowanie własne

Dryf odnosi się do spadku wydajności modelu z powodu zmian danych i relacji między zmiennymi wejściowymi i wyjściowymi. To natomiast może spowodować pogorszenie jakości modelu prognostycznego, a w efekcie niedokładnymi prognoząmi. Dryf zmiennych odnosi się do zmian wartości wejściowych, np. wynikający z gwałtownych zmian trendów sprzedaży. Zmiana danych wejściowych może być (tab. 8.4):

- gwałtowna (nagła) np. z powodu lockdownów w okresie pandemii COVID-19,
- narastająca (zmieniająca się powoli),



- impulsowa (jednorazowa) np. w przypadku błędnie zasilonych danych.

Przedstawia to formalny model:

$$p_{t+h|T} = y_t + h\left(\frac{y_t - y_1}{t - 1}\right)$$

gdzie:

$p_{t+h|T}$ – prognoza na przyszły okres,

$h\left(\frac{y_t - y_1}{t - 1}\right)$ – komponent dryfu.

Jest to równoznaczne z nakreśleniem granicy między pierwszą a ostatnią obserwacją i ekstrapolacją jej na przyszłość.

Formuła stosowana w Excelu:

prognoza(t) = popyt(z ostatniego okresu roku poprzedniego) + [h *(popyt(z ostatniego okresu roku poprzedniego) – popyt(z pierwszego okresu roku poprzedniego)]/n – 1



Prognoza na okres t jest równa popytowi z ostatniego okresu roku poprzedniego + komponent dryfu.

W komponencie dryfu występuje liczba h , która odnosi się do kolejnej liczby zbudowanych prognoz, zaś t określa liczbę badanych okresów w roku.

Metody średniej

Metody średniej ruchomej działają jak filtr, ponieważ eliminują z szeregu danych wahania krótkookresowe. Do budowy prognozy, metody średniej ruchomej korzystają z określonej liczby sąsiadujących danych o popycie.

Ze wzrostem liczby danych historycznych, na podstawie których budowana jest prognoza rośnie efekt wygładzenia. Oznacza to, że korzystanie z większej liczby danych w modelu wygładza silniej szereg. Sprawia także, że następuje wolniejsza reakcja na zmiany poziomu prognozowanej zmiennej. I na odwrót. Korzystanie z mniejszej liczby danych historycznych wpływa na szybsze odzwierciedlanie zmian popytu z ostatnich okresów. Prognoza staje się wtedy bardziej wrażliwsza na wahania przypadkowe.



Na potrzeby opracowania zostaną przedstawione 4 metody średniej (1) metoda średniej globalnej; (2) metoda prostej średniej ruchomej; (3) metoda wykładniczej średniej ruchomej; (4) metoda średniej ruchomej ważonej.

Średnia globalna (Global Mean or Global Average)

Prognoza, przy pomocy metody średniej globalnej, budowana jest na wszystkich dostępnych obserwacjach historycznych zawartych w szeregu. Metoda ta określa tendencję centralną, którą jest lokalizacja środka danych szeregu w rozkładzie statystycznym.

Powyższe przedstawia formalny model:

$$p_t = \frac{\sum_{t=1}^n y_t}{n}$$

gdzie:

n – liczba analizowanych okresów.

Wartość prognozy przy pomocy metody średniej globalnej należy obliczać dodając do siebie wartości wszystkich danych historycznych i dzieląc obliczoną wartość przez liczbę analizowanych okresów.



Formuła stosowana w Excelu:

$$\text{prognoza}_{(t)} = \Sigma \text{popytu} / n$$

Stosowanie średniej arytmetycznej z całego szeregu danych powoduje zakłamanie wyniku budowanej prognozy. Wykorzystywane w metodzie dane są zbyt przestarzałe, co zniekształca właściwy obraz przyszłego popytu.

Metoda prostej średniej ruchomej (Simple Moving Average, SMA)

Metoda prostej średniej ruchomej (kroczącej) wykorzystuje typową średnią arytmetyczną, ale tylko z pewnej określonej liczby danych historycznych. Wyselekcjonowana liczba danych pozwala na wygładzenie danych dotyczących popytu, zmniejszając wpływ przypadkowych fluktuacji i przestarzałych danych. Nazwa średnia ruchoma oznacza, że każda



prognoza jest obliczana na podstawie danych z poprzednich x okresów. Prosta średnia ruchoma nie wyróżnia danych historycznych i nie nadaje wag tym danym.

Prosta średnia ruchoma to arytmetyczna średnia krocząca obliczona przez dodanie ostatnich danych popytu, a następnie podzielenie uzyskanej wartości przez liczbę okresów w średniej obliczeniowej. Powyższe przedstawia formalny model:

$$p_t = \frac{\sum_{t+1-m}^t y}{m}$$

w którym

m – liczba analizowanych okresów.

Powyższe wyjaśnia formuła stosowana w Excelu:

dla średniej 3-elementowej:



prognoza_(t) = ŚREDNIA(popyt_(t-1); popyt_(t-2); popyt_(t-3);)

Aby obliczyć średnią ruchomą można użyć prostej formuły opartej na funkcji ŚREDNIA z odniesieniami względnymi.

Gdy formuły są kopiowane w dół kolumny, zakres zmienia się w każdym wierszu, aby uwzględnić wartości potrzebne dla każdej średniej. Im dłuższa średnia krocząca, tym większe opóźnienie. Można to wytłumaczyć w następujący sposób:

- Prognoza na bazie 3 danych historycznych to krótkoterminowa średnia ruchoma; jest jak łódź motorowa – zwinna i szybko się zmieniająca.
- Prognoza na bazie 50 danych historycznych to długoterminowa średnia ruchoma; jest jak tankowiec oceaniczny – ospała i wolno się zmieniająca.
- Należy więc pamiętać o czynniku opóźnienia przy wyborze odpowiedniej liczby okresów historycznych (nie korzystać ze zbyt dużej liczby danych).

Metoda wykładniczej średniej ruchomej (Exponential Moving Average, EMA)

Metoda wykładniczej średniej ruchomej pozwala zmniejszyć opóźnienie przykładając większą wagę do ostatnich wartości danych historycznych. Dzięki temu metoda lepiej reaguje



na ostatnie wartości danych. Wykładnicza średnia ruchoma jest zwykle bardziej wrażliwa na ostatnie zmiany popytu w porównaniu z prostą średnią ruchomą. Powyższe przedstawia formalny model:

$$p_t = p_{t-1} + \alpha(y_{t-1} - p_{t-1})$$

$$\alpha = \frac{2}{n+1}$$

w którym:

α – mnożnik,

n – wybrany przedział czasu.

Obliczanie wykładniczej średniej ruchomej składa się z trzech kroków:

1. Obliczenie prostej średniej ruchomej dla okresu (prognoza wstępna). SMA jest wymagana jedynie w celu dostarczenia wartości początkowej do dalszych obliczeń.
2. Obliczenie mnożnika do ważenia wykładniczej średniej ruchomej

Przykład: jeśli prognoza jest budowana z 3 okresów, mnożnik zostanie obliczony w następujący sposób: $\text{mnożnik} = \alpha = \frac{2}{n+1} = \frac{2}{3+1} = 0,5$

3. Obliczenie aktualnej prognozy według wykładniczej średniej ruchomej.

Przypomnienie: Pierwsza obliczona prognoza nazwana jest prognozą wstępną.



Formuła stosowana w Excelu:

$$\text{prognoza}_{(t)} = \text{prognoza}_{(t-1)} + \text{mnożnik} * (\text{popyt}_{(t-1)} - \text{prognoza}_{(t-1)})$$

Metoda EMA służy do uchwycenia krótszych ruchów trendu, ze względu na skupienie się na najnowszych danych i ostatnich prognozach.

Metoda średniej ruchomej ważonej (Weighted Moving Average, WMA)

Jest to wariant prostej średniej ruchomej (SMA). Metoda średniej ruchomej ważonej to średnia krocząca, w której przypisuje się wagę najnowszym wartościom popytu. To sprawia, że najnowsze dane mają najsilniejszy wpływ na wartość prognozy niż dane starsze. Jest to możliwe poprzez stosowanie współczynnika ważonego. Zastosowanie współczynników



ważonych pozwala na uzyskanie dokładniejszych prognoz. Metoda jest uważana za bardziej wrażliwą na zmiany popytu.

Prognozy według metody średniej ruchomej ważonej uzyskuje się mnożąc każdą wartość popytu przez z góry określony współczynnik ważony i sumując otrzymane wartości. Powyższe przedstawia formalny model:

$$p_t = \sum_{i=t+1-m}^t (y_i \cdot \varpi_i)$$

, w którym

ϖ – współczynnik ważony.

Aby określić ich wartość należy pamiętać o kilku zasadach:

- Wartości współczynników ważonych mieszczą się w przedziale $< 0,1 >$,
- Każdy kolejny zastosowany współczynnik ważony jest większy od swojego poprzednika $\varpi_i < \varpi_{i+1} < \varpi_{i+2}$. Jest to bardzo istotna zasada, gdyż różnicuje ona znaczenie wykorzystywanych danych historycznych. Te starsze mają niższy współczynnik ważony, dane świeższe są ważniejsze i mają wyższy współczynnik ważony,
- Suma wszystkich współczynników ważonych musi być równa 1: $\sum_1^n \varpi_i = 1$,
- Liczba współczynników ważonych zależy od liczby analizowanych okresów historycznych z szeregu czasowego.

Formuła stosowana w Excelu:

dla średniej ważonej 3-elementowej:



prognoza_(t) = (popyt_(t-3) * $\omega_{(1)}$) + (popyt_(t-2) * $\omega_{(2)}$) + (popyt_(t-1) * $\omega_{(3)}$)

dla średniej ważonej 5-elementowej:

prognoza_(t) = (popyt_(t-5) * $\omega_{(1)}$) + (popyt_(t-4) * $\omega_{(2)}$) + (popyt_(t-3) * $\omega_{(3)}$) + (popyt_(t-2) * $\omega_{(4)}$) + (popyt_(t-1) * $\omega_{(5)}$)



W metodzie średniej ruchomej ważonej należy określić wartość stałej wygładzania (z ilu danych okresów historycznych należy korzystać) oraz określić poziomy poszczególnych wag współczynników ważonych.

Metody wygładzania wykładniczego

Metody wygładzania wykładniczego (Exponential Smoothing, ES) są szeroko stosowane (Chatfield i in., 2001). Istnieje 15 różnych metod. Każdy wariant jest określony dla innego scenariusza prognozowania. Najbardziej znane warianty metody wygładzania wykładniczego to SES (brak trendu, brak sezonowości), metoda liniowa Holta (trend addytywny, brak sezonowości), metoda addytywna Holta-Wintersa (trend addytywny, sezonowość addytywna) oraz metoda multiplikatywna Holta-Wintersa (trend addytywny, sezonowość multiplikatywna) (De Gooijer i Hyndman, 2006). Badacze zaproponowali liczne warianty oryginalnych metod wygładzania wykładniczego, np. Carreno i Madinaveitia (1990) zaproponowali modyfikacje w celu radzenia sobie z nieciągłościami, a Rosas i Guerrero (1994) przyjrzyli się prognozom wygładzania wykładniczego podlegającym jednemu lub kilku ograniczeniom.

Na potrzeby opracowania zostaną przedstawione 3 metody wygładzania wykładniczego: (1) metoda proste wygładzanie wykładnicze (model Browna); (2) metoda liniowe wygładzanie wykładnicze (model Holta); (3) metoda sezonowe wygładzanie wykładnicze (model Holt-Winters).

Metoda prostego wygładzania wykładniczego (Simple Exponential Smoothing, SES), model Browna

Proste wygładzanie wykładnicze jest podstawową formą wygładzania wykładniczego. Metoda wygładzania wykładniczego (model Browna) jest relatywnie dokładną metodą prognozowania popytu. Uwzględnia ona współczynnik wygładzania wykładniczego (α). Współczynnik ten kontroluje tempo wpływu danych na budowane prognozy. Jednocześnie metoda przypisuje większą wagę nowszym danym. Przypisuje zaś wykładniczo malejące wagi kiedy dane stają się bardziej odległe.

Aby ustalić wartość współczynnika wygładzania wykładniczego należy pamiętać o następujących zasadach:

- wartość współczynnika wygładzania wykładniczego mieści się w przedziale $(0,1)$,



- wartość współczynnika wygładzania wykładniczego dobierana jest eksperymentalnie. Należy korzystać z założenia: im współczynnik bliższy 0 tym dane ulegają silniejszemu wygładzeniu (prognoza jest mniej wrażliwa na zmiany popytu), a prognosta bardziej ufa zbudowanej w poprzednim okresie prognozie; wskaźnik bliższy wartości 1 oznacza, że prognoza jest bardziej wrażliwa na zmiany w popycie, a prognosta bazuje na sytuacji rzeczywistej jaka zaistniała w poprzednim okresie.

Przedstawia to formalny model:

$$p_t = a \cdot y_{t-1} + (1 - a) \cdot p_{t-1}$$

gdzie:

a – współczynnik wygładzania wykładniczego.

Obliczanie prognozy według metody prostego wygładzania wykładniczego polega na trzech krokach:

1. Obliczenie prognozy naiwnej dla pierwszego okresu (prognoza wstępna). Prognoza wstępna jest wymagana jedynie w celu dostarczenia wartości początkowej do dalszych obliczeń.
2. Określenie współczynnika wygładzania wykładniczego.

Współczynnik wygładzania wykładniczego $a = \langle 0,1 \rangle$

3. Obliczenie prognozy według metody prostego wygładzania wykładniczego.



Formuła stosowana w Excelu:

$$\text{prognoza}_{(t)} = \text{alfa} * \text{popyt}_{(t-1)} + [(1 - \text{alfa}) * \text{prognoza}_{(t-1)}]$$

Metoda prostego wygładzania wykładniczego jest stosowana do prognozowania na podstawie danych nieobarczonych znaczącym trendem ani sezonowością.



Metoda podwójnego wygładzania wykładniczego (Linear Exponential Smoothing, LES), model Holta

Metoda podwójnego wygładzania wykładniczego (tak zwany model Holta) wychwytuje trendy liniowe w danych. Jest właściwym modelem dla popytu, w którym można zaobserwować stały trend wzrostowy lub spadkowy. Nie występuje jednak sezonowość.

Model podwójnego wygładzania wykładniczego bazuje na dwóch współczynnikach wygładzania. Jeden z nich odnosi się do wygładzania poziomu zmiennej (wahania losowe), a drugi do jej przyrostu (wahania trendu). Oba współczynniki powinny być zawarte w przedziale: $\langle 0,1 \rangle$. Przedstawia to formalny model:

$$p_t = L_{t-1} + T_{t-1}$$

$$L_t = \alpha y_{t-1} + (1 - \alpha)(L_{t-1} + T_{t-1})$$

$$T_t = \beta(L_{t-1} - L_{t-2}) + (1 - \beta)T_{t-1}$$

, w którym

α – współczynnik wygładzania poziomu zmiennej,

β – współczynnik wygładzania przyrostu.

Zgodnie z modelem potrzebne są dwie wartości startowe L_t i T_t :

$L_t = y_{t-1}$ – zgodnie z metodą naiwną

$T_t = y_{t-1} - y_{t-2}$

Formuła stosowana w Excelu:



prognoza_(t) = [alfa * popyt_(t-1) + (1 – alfa)(wahanie losowe_(t-1) + wahanie trendu_(t-1))] + [beta * (wahanie losowe_(t-1) – (wahanie losowe_(t-2)) + (1 – beta) * wahanie trendu_(t-1)]

Metoda sezonowego wygładzania wykładniczego (model Holt-Winters)

Metoda sezonowego wygładzania wykładniczego, zwana również modelem Holta-Wintersa, bardzo dobrze nadaje się do prognozowania popytu dla danych charakteryzujących się zarówno trendem i sezonowością (www_8.1).



Konieczne jest jednak pozyskanie długich szeregów danych dotyczących popytu, ponieważ konieczne jest zweryfikowanie powtarzających się wahań cyklicznych (potwierdzenie, że w szeregu występuje popyt o charakterze sezonowym). Wahania sezonowe występują w wersji addytywnej lub multiplikatywnej (Kucharski, 2013).

Wahania addytywne występują, gdy w poszczególnych podokresach cyklu sezonowości zaobserwować można stałe pod względem wartości bezwzględnej odchylenie poziomu analizowanego zjawiska od średniego poziomu lub trendu.

Przedstawia to model addytywny:

$$p_t = L_{t-1} + T_{t-1} + S_{t-p}$$

$$L_t = \alpha(y_t - S_{t-p}) + (1 - \alpha)(L_{t-1} + T_{t-1})$$

$$T_t = \beta(L_t - L_{t-1}) + (1 - \beta)T_{t-1}$$

$$S_t = \delta(y_t - L_t) + (1 - \delta)S_{t-p}$$

, w którym

α – współczynnik wygładzania poziomu zmiennej,

β – współczynnik wygładzania trendu,

δ – współczynnik składnika sezonowego,

L_t – składnik poziomu w czasie t ,

T_t – składnik trendu w czasie t ,

S_t – składnik sezonowy w czasie t ,

p – okres sezonowy.

Formuła stosowana w Excelu:

**prognoza_(t) = składnik poziomu_(t-1) + składnik trendu_(t-1) +
składnik sezonowości_(t-p)**



**składnik poziomu_(t) = alpha * [popyt_(t) - składnik sezonowości_(t)] +
(1 – alpha) * (składnik poziomu_(t-1) + składnik trendu_(t-1))**

**składnik trendu_(t) = beta * (składnik poziomu_(t) - składnik
poziomu_(t-1)) + [(1 – beta) * składnik trendu_(t-1)]**



$$\text{składnik sezonowości}_{(t)} = \text{gamma} * (\text{popyt}_{(t)} - \text{składnik poziomu}_{(t)}) + (1 - \text{gamma}) * \text{składnik sezonowości}_{(t-p)}$$

Składnik sezonowości_(t-p) dla czterech kwartałów

$$\text{składnik sezonowości}_{(1)} = \text{popyt}_{(1)} / \text{średni popyt}_{(1-4)}$$

$$\text{składnik sezonowości}_{(2)} = \text{popyt}_{(2)} / \text{średni popyt}_{(1-4)}$$

$$\text{składnik sezonowości}_{(3)} = \text{popyt}_{(3)} / \text{średni popyt}_{(1-4)}$$

$$\text{składnik sezonowości}_{(4)} = \text{popyt}_{(4)} / \text{średni popyt}_{(1-4)}$$

Wahania multiplikatywne występują, gdy w poszczególnych podokresach cyklu zaobserwować można zjawisko odchylania się od średniego poziomu bądź trendu o pewną stałą wielkość względną (Sobczyk, 2006). Przedstawia to model multiplikatywny:

$$p_t = (L_{t-1} + T_{t-1})S_{t-p}$$

$$L_t = \alpha \left(\frac{y_t}{S_{t-p}} \right) + (1 - \alpha)(L_{t-1} + T_{t-1})$$

$$T_t = \beta(L_t - L_{t-1}) + (1 - \beta)T_{t-1}$$

$$S_t = \delta \left(\frac{y_t}{L_t} \right) + (1 - \delta)S_{t-p}$$

, w którym

α – współczynnik wygładzania poziomu zmiennej,

β – współczynnik wygładzania trendu,

δ – współczynnik składnika sezonowego,

L_t – składnik poziomu w czasie t ,

T_t – składnik trendu w czasie t ,

S_t – składnik sezonowy w czasie t ,

p – okres sezonowy.

Warto pamiętać, że: trend addytywny dotyczy podwójnego wygładzania wykładniczego z trendem liniowym, zaś trend multiplikatywny związany jest z podwójnym wygładzaniem wykładniczym z trendem wykładniczym (www_8.2).



Formuła stosowana w Excelu:

$$\text{prognoza}_{(t)} = (\text{składnik poziomu}_{(t-1)} + \text{składnik trendu}_{(t-1)}) * \text{składnik sezonowości}_{(t-p)}$$

$$\text{składnik poziomu}_{(t)} = \alpha * [\text{popyt}_{(t)} / \text{składnik sezonowości}_{(t)}] + (1 - \alpha) * (\text{składnik poziomu}_{(t-1)} + \text{składnik trendu}_{(t-1)})$$

$$\text{składnik trendu}_{(t)} = \beta * (\text{składnik poziomu}_{(t)} - \text{składnik poziomu}_{(t-1)}) + [(1 - \beta) * \text{składnik trendu}_{(t-1)}]$$



$$\text{składnik sezonowości}_{(t)} = \gamma * (\text{popyt}_{(t)} / \text{składnik poziomu}_{(t)}) + (1 - \gamma) * \text{składnik sezonowości}_{(t-p)}$$

składnik sezonowości_(t-p) dla czterech kwartałów

$$\text{składnik sezonowości}_{(1)} = \text{popyt}_{(1)} / \text{średni popyt}_{(1,2,3,4)}$$

$$\text{składnik sezonowości}_{(2)} = \text{popyt}_{(2)} / \text{średni popyt}_{(1,2,3,4)}$$

$$\text{składnik sezonowości}_{(3)} = \text{popyt}_{(3)} / \text{średni popyt}_{(1,2,3,4)}$$

$$\text{składnik sezonowości}_{(4)} = \text{popyt}_{(4)} / \text{średni popyt}_{(1,2,3,4)}$$

Metody autoregresyjne

Wśród modeli prognozowania szczególne miejsce zajmują modele stacjonarnych szeregów ARMA (AutoRegressive Moving Average) oraz niestacjonarnych szeregów ARIMA (AutoRegressive Integrated Moving Average model). Są to modele oparte na zjawisku autokorelacji, powstałe przez integrację modelu autoregresyjnego AR (AutoRegressive model) oraz modelu średniej ruchomej MA (Moving Average) (Grzelak, 2019).

Na potrzeby opracowania zostaną przedstawione 3 metody autoregresyjne: (1) metoda autoregresyjna średnia ruchoma (ARMA); (2) metoda autoregresyjna zintegrowana średnia ruchoma (ARIMA); (3) metoda błędzenia losowego (RW).

Modele ARMA i ARIMA mają wiele podobieństw. Komponenty AR(p) – ogólny model autoregresyjny i MA – ogólny model średniej ruchomej MA(q) są takie same. Tym, co odróżnia oba modele ARMA i ARIMA jest różnica. Jeśli w modelu ARMA nie ma różnic, staje się on po prostu modelem ARIMA.



Autoregresywna średnia ruchoma (Autoregressive Moving Average, ARMA)

Wymogiem prognozowania według metody ARMA jest szereg danych, który charakteryzuje się stacjonarnością. Oznacza, że w szeregu tym można wyróżnić stałą średnią, stałą wariancję i stałą kowariancję, która zależy tylko od odstępu czasu między wartościami (Schaffer i in., 2021).

Według modelu ARMA wartość prognozowana w czasie t zależy od przeszłych jej wielkości oraz od różnic między przeszłymi wartościami rzeczywistymi zmiennej prognozowanej, a jej wartościami uzyskanymi z modelu – błędów prognoz. Model ten określany jest jako **model ARMA(p, q)**, w którym p odnosi się do rzędu wielomianu autoregresyjnego, zaś q jest rzędem wielomianu średniej ruchomej. Powyższe przedstawia formalny ogólny model:

$$p_t = \varepsilon_t + (\alpha y_{t-1} + \varepsilon_t) + (\beta y_{t-2} - \alpha y_{t-1} + \varepsilon_t) + (\varepsilon_t + \alpha \varepsilon_{t-1})$$

gdzie:

α – parametr modelu autoregresyjnego,

β – parametr modelu średniej ruchomej,

ε – błąd modelu (biały szum).

Formuła stosowana w Excelu:

**prognoza_(t) = błąd modelu_(t) + komponent 1 + komponent 2 +
komponent 3**



błąd modelu_(t) = NORM.S.INV(rand())

komponent 1 = alfa * komponent 1_(t-1) + błąd modelu_(t)

**komponent 2 = beta * komponent 2_(t-2) - alfa * komponent 2_(t-1) +
błąd modelu_(t)**

komponent 3 = błąd modelu_(t) + alfa * błąd modelu_(t-1)



Autoregresyjna zintegrowana średnia ruchoma (Autoregressive Integrated Moving Average, ARIMA)

W przypadku modeli ARIMA uwagę zwraca się na niestacjonarność szeregu. W modelu wyróżnia się trzy parametry: parametr autoregresyjny (p), parametr średniej ruchomej (q) oraz rząd różnicowania (d). Model **ARIMA**(p, q, d) opisywany jest także za pomocą cyfr, przykładowo: (1,1,0), co oznacza, że w szeregu $p=1$ występuje jednokrotny parametr autoregresyjny, $q=1$ występuje jednokrotny parametr średniej ruchomej oraz $d=0$ nie występuje różnicowanie (Malska i Wachta, 2015). Dla podanego przykładu formalny ogólny model jest następujący:

$$p_t = \alpha + y_{t-1} + \beta_1(y_{t-1} - y_{t-2})$$

gdzie:

α – parametr modelu autoregresyjnego,

β – parametr średniej ruchomej.



Formuła stosowana w Excelu:

$$\text{prognoza}_{(t)} = \text{alfa} + \text{popyt}_{(t-1)} + \text{beta} * [\text{popyt}_{(t-1)} - \text{popyt}_{(t-2)}]$$

Model błędzenia losowego (Random Walk, RW)

Model błędzenia losowego jest podkategorią modelu ARIMA. W prostym modelu RW zakłada się, że każda prognoza jest sumą ostatniej obserwacji i składnika błędu losowego. Zakłada więc, że najnowsza obserwacja jest najlepszą wskazówką do budowy najbliższej następnej prognozy. Ten model jest dość prosty do zrozumienia i wdrożenia. Jest stosowana w sytuacji, kiedy w ciągu danych zostanie zaobserwowana tendencja rozwojowa. Przedstawia to model:

$$p_t = y_{t-1} + \varepsilon_t$$

$$\varepsilon_t = y_{t-1} - y_{t-2}$$

gdzie:

ε – błąd modelu (biały szum).



Formuła stosowana w Excelu:

$$\text{prognoza}_{(t)} = \text{popyt}_{(t-1)} + [\text{popyt}_{(t-1)} - \text{popyt}_{(t-2)}]$$

Zaobserwowano, że wiele złożonych metod prognostycznych opartych na strukturze liniowej nie jest w stanie pokonać naiwnego modelu RW (Adhikari i Agrawal, 2014).

Metody regresji

W modelach regresji nie można mówić o wpływie jednej zmiennej na drugą. Za pomocą zmiennej bądź zestawu zmiennych wyjaśniana jest inna zmienna. Do zastosowania metod regresji potrzebna jest większa liczba danych historycznych – im dłuższy czas obserwacji danych historycznych, tym większa możliwość precyzyjnego określenia prognoz.

Znanych jest wiele wariantów modeli regresji: regresja liniowa, regresja nieliniowa, regresja logistyczna, regresja krokowa, regresja porządkowa.

Wzór na ogólną postać regresji ma postać:

$$p_t = f(X, \beta) + \varepsilon_t$$

gdzie:

X – zmienna wyjaśniająca, przewidywana,

β – współczynnik regresji,

$f(X, \beta)$ – równanie regresji,

ε_t – błąd losowy.

Na potrzeby opracowania zostaną przedstawione 3 metody wygładzania wykładniczego: (1) metoda projekcji trendu; (2) prosta metoda regresji liniowej; (3) metoda wielokrotnej regresji liniowej.

Metoda projekcji trendu (Trend Projection Method)

Metoda Projekcji Trendu to odmiana metody linii prostej. Jest to najbardziej klasyczna metoda prognozowania biznesowego, która zajmuje się ruchem zmiennych w czasie. Można rozróżnić **metodę graficzną** – w której dane są wyświetlane na wykresie, a przez nie ręcznie



wykreślana jest linia. Linia jest rysowana z zachowaniem najmniejszej odległości między wyznaczonymi punktami a linią; **metodą dopasowania równania trendu** przy użyciu danych i równania linii prostej lub wykładniczej.

Prosta metoda regresji liniowej (Simple linear regression method)

Metoda regresji liniowej jest najprostszym wariantem regresji. Zadaniem metody regresji liniowej jest dopasowanie prostej linii do danych. Zatem należy znaleźć rozwiązanie, które pozwoli na znalezienie optymalnej prostej, która najlepiej pokaże zależność pomiędzy danymi. W metodzie prostej regresji liniowej model prognostyczny bazuje na tendencji liniowej. Aby wyznaczyć linię regresji, a tym samym wartości w modelu regresji liniowej, należy obliczyć współczynniki linii prostej a oraz b . Powyższe przedstawia formalny model:

$$p_t = at + b$$

w którym

a – wartość zmiennej w analizowanym okresie

b – wartość przyrostu lub spadku zmiennej zależnej

n – numer porządkowy analizowanego i prognozowanego okresu.

Aby określić parametry a oraz b należy obliczyć układ dwóch równań. Ma on następującą postać:

$$\begin{cases} a \sum_{i=1}^n t_i^2 + b \sum_{i=1}^n t_i = \sum_{i=1}^n t_i \cdot y_i \\ a \sum_{i=1}^n t_i + b \cdot n = \sum_{i=1}^n y_i \end{cases}$$

w którym

t_i – numer porządkowy okresu ($t = 1, 2, 3, \dots$), który jest wartością zmiennej czasowej niezależnej,

y_i – zmienna zależna (np. zapotrzebowanie na dane dobro w ustalonym przedziale czasowym),

a – wartość zmiennej w analizowanym okresie,

b – wartość przyrostu lub spadku zmiennej zależnej,

n – liczba wszystkich analizowanych okresów.



Jednak najszybszym sposobem na obliczenie powyższych współczynników regresji jest skorzystanie z funkcji REGLINP.

Składnia: **REGLINP(znane_y;[znane_x];[stała];[statystyka])**

Metoda wielokrotnej regresji liniowej (Multiple linear regression method)

Wielokrotna regresja liniowa pozwala budować modele zależności liniowej między wieloma zmiennymi. Metoda wielokrotnej regresji liniowej jest stosowana w analizie danych w celu zbadania złożonych relacji między wieloma zmiennymi. Powyższe przedstawia formalny model:

$$p_t = a + b_1y_1 + b_2y_2 + \dots + b_iy_i + \varepsilon$$

w którym:

a – wartość zmiennej w analizowanym okresie,

b – wartość przyrostu lub spadku zmiennej zależnej,

ε – błąd modelu (biały szum).

8.7. Błędy prognozy

Aby ocenić trafność prognozy należy mierzyć błędy prognozy. Biorąc pod uwagę, że nie ma gwarancji doskonałego przewidywania przyszłego popytu (Hopp i Spearman, 1999) każda prognoza obarczona jest błędem.

Badacze rozróżniają systematyczne i niesystematyczne skutki błędów prognoz (Zeiml, i in., 2019). Wydajność systemu prognozowania jest zazwyczaj mierzona za pomocą różnych miar błędu prognozy (tab. 8.5).



Tabela 8.5. Wybrane błędy prognoz

Typ błędu prognozy	Model	Interpretacja
Błąd średniokwadratowy Mean Squared Error (MSE)	$MSE = \frac{\sum (y_t - p_t)^2}{n}$	Błąd średniokwadratowy może być tylko dodatni, a jego wartość powinna być jak najmniejsza. Wartość tego błędu, która wynosi 0 wskazuje na doskonałą dokładność prognozy.
Pierwiastek błędu średniokwadratowego Root Mean Squared Error (RMSE)	$RMSE = \sqrt{\frac{\sum (p_t - y_t)^2}{n}}$	Wartość błędu powinna być jak najbardziej zbliżona do 0. Im niższa wartość pierwiastka błędu średniokwadratowego, tym lepszy jest model. A model doskonały ma wartość równą 0.
Średni bezwzględny błąd procentowy Mean Absolute Percentage Error (MAPE)	$MAPE = \frac{\sum E_t - A_t }{A_t n}$	Jest to jedna z bardziej popularnych miar błędu. Wartość 0 wskazuje, że model nie zawiera błędu średniego, co oznacza, że wartość powinna być jak najbliżej 0. W przypadku tych samych błędów prognoz mniejsze wartości rzeczywiste powodują, że błąd względny staje się większy.
Średni błąd procentowy Mean Percentage Error (MPE)	$MPE = \frac{100\%}{n} \sum \frac{a_t - f_t}{a_t}$	Jest to średni błąd procentowy (lub odchylenie). Informuje o ile średnio w okresie prognoz będzie wynosić odchylenie od wartości rzeczywistej. MPE jest przydatny, ponieważ pozwala sprawdzić czy model prognozy systematycznie zaniża (więcej ujemnego błędu) lub przeszacowuje (dodatni błąd).
Średnie odchylenie bezwzględne Mean Absolute Deviation (MAD)	$MAD = \sum \frac{ x_t - \bar{x}_t }{n}$	Jest prostym rozszerzeniem wariancji bezwzględnej. Średnie odchylenie bezwzględne jest wykorzystywane jako miara zróżnicowania danych.

Źródło: (Hyndman i Koehler, 2006; Zeiml, i in., 2019)

8.8. Zalety prognozowania w Excelu

Microsoft Excel może być używany do prognozowania. Przy użyciu opracowanych algorytmów i na podstawie zebranych danych z przeszłości można przygotować formularze w celu budowy prognoz, a w efekcie podejmowania właściwych decyzji w biznesie. Excel jest podstawowym narzędziem do prognozowania.

Excel jest używany w prognozowaniu przez wiele przedsiębiorstw, zarówno małych, średnich, jak i dużych (w tym koncernów), ponieważ ma do dyspozycji szereg odpowiednich narzędzi. Dane można łatwo przechowywać i obliczać w skoroszybie programu Excel. Excel



może wizualizować pozyskane i opracowane dane na różne sposoby, co jest użyteczne i przydaje się do łatwiejszego prognozowania (www_8.3).

Tabela 8.6. Wybrane funkcje programu Excel

Funkcja	Wyjaśnienie
=ŚREDNIA	Funkcja pozwala obliczyć średnią na podstawie istniejących wartości.
=SUMA	Funkcja pozwala obliczyć sumę na podstawie istniejących wartości.
=REGLINX	Funkcja pozwala przewidzieć przyszłą wartość na podstawie istniejących wartości przy użyciu regresji liniowej.
=REGLINX.ETS	Oblicza lub przewiduje przyszłą wartość na podstawie istniejących (historycznych) wartości przy użyciu wersji algorytmu wygładzania wykładniczego (ETS).
=REGLINP	Oblicza statystykę dla linii, korzystając z metody najmniejszych kwadratów.
=REGLINX.LINIOWA	Oblicza lub przewiduje przyszłą wartość na podstawie istniejących wartości przy użyciu regresji liniowej.
=REGLINW	Posłuży do wyznaczenia trendu liniowego.
=REGLINX.ETS.SEZONOWOŚĆ	Oblicza długość wzorca sezonowego na podstawie istniejących wartości i osi czasu.

Źródło: opracowanie własne

W końcu Excel wykorzystuje wiele formuł, których można użyć w skoroszycie programu, aby ułatwić obliczanie przewidywanych wartości. Excel obsługuje kilka różnych funkcji, które pozwalają na praktyczne wykorzystanie oprogramowania (tab. 8.6). Zrozumienie ich jest kluczem do maksymalnego wykorzystania programu Excel.

Do wad Excela zaliczyć należy konieczność ręcznego synchronizowania danych i ręcznego aktualizowania danych. Częstym problemem jest również możliwość popełnienia błędów w wyniku niepoprawnie przeprowadzonego importu danych lub w wyniku przerywania formuły. Ponieważ w programie Excel wprowadzanie danych odbywa się ręcznie, dane używane do prognozowania nie są danymi w czasie rzeczywistym.



8.9. Sztuczna inteligencja w prognozowaniu

Dane dotyczące łańcucha dostaw są wielowymiarowe i generowane w wielu punktach łańcucha, w różnych celach, w dużych ilościach (ze względu na mnogość dostawców, produktów i klientów) oraz generowane z dużą szybkością (co odzwierciedla wiele transakcji stale przetwarzanych w sieciach łańcucha dostaw). Ta złożoność i wielowymiarowość łańcuchów dostaw sprawia, że następuje odejście od konwencjonalnych (statystycznych) podejść do prognozowania popytu, które opierają się na identyfikacji statystycznie średnich trendów (charakteryzujących się atrybutami średniej i wariancji) (Michna, i in., 2020), w kierunku inteligentnych prognoz, które mogą uczyć się na podstawie danych historycznych i inteligentnie ewoluować, aby dostosować się do przewidywania stale zmieniającego się popytu w łańcuchach dostaw.

Odpowiedzią na nowe potrzeby i wyzwania staje się logistyka wyprzedzająca, która wspiera takie procesy jak prognozowanie popytu. U jej podstaw leży możliwość wykorzystania sztucznej inteligencji (Artificial Intelligence, AI). Jest to połączenie współczesnych technologii takich jak Big Data (BD), uczenie maszynowe (Machine Learning, ML) oraz sztucznej inteligencji (Sczaniecka i Smarzyńska, 2018). Postęp technologiczny doprowadził w ostatnich latach do coraz częstszego generowania i przechowywania ogromnych ilości danych. Dane te są przechwytywane w czasie w różnych punktach (np. w różnych ogniwach łańcucha dostaw) oraz przechowywane w różnych miejscach. Powinny więc być sprawnie przetwarzane, aby wydobyć z nich przydatną i wartościową wiedzę (Galicia, i in., 2019).

Big Data odnosi się do zbiorów danych dynamicznych o dużej objętości, dużej szybkości i dużej różnorodności, które przekraczają możliwości przetwarzania tradycyjnych podejść do zarządzania danymi (Chen, i in., 2014). Big Data może dostarczyć wiele unikalnych informacji na temat m.in. trendów rynkowych, wzorców zakupowych klientów i cykli konserwacji, a także sposobów obniżania kosztów i podejmowania bardziej ukierunkowanych decyzji biznesowych (Wang, i in., 2016). Badania wskazują, że analiza dużych zbiorów danych (Big Data Analytics, BDA) może być sposobem na uzyskanie bardziej precyzyjnych prognoz, które lepiej odzwierciedlają potrzeby klientów, ułatwiają ocenę wydajności łańcuchów dostaw, poprawiają wydajność łańcuchów dostaw, skracają czas reakcji oraz wspierają ocenę ryzyka łańcucha dostaw (Awwad, i in., 2018).



Analityka dużych zbiorów danych (BDA) w zarządzaniu łańcuchem dostaw (SCM) cieszy się coraz większym zainteresowaniem (Seyedan i Mafakheri, 2020). Analiza danych dotyczących łańcucha dostaw stała się złożonym zadaniem ze względu na (Awwad, i in., 2018):

- rosnącą liczbę jednostek SC,
- rosnącą różnorodność konfiguracji SC w zależności od jednorodności lub heterogeniczności produktów,
- współzależności między tymi podmiotami,
- niepewność co do dynamicznego zachowania tych komponentów,
- brak informacji dotyczących podmiotów łańcucha dostaw,
- usieciowione podmioty produkcyjne ze względu na ich rosnącą koordynację i współpracę w celu osiągnięcia wysokiego poziomu dostosowania i adaptacji do różnych potrzeb klientów,
- coraz częstsze stosowanie praktyk cyfryzacji łańcucha dostaw (i wykorzystanie technologii Blockchain) do śledzenia aktywności w łańcuchach dostaw.

Metody sztucznej inteligencji są rozwijane na świecie niezwykle intensywnie i to zarówno od strony teoretycznej, jak i aplikacyjnej (Trojanowska i Malopolski, 2004). Spośród metod sztucznej inteligencji do prognozowania wykorzystuje się sztuczne sieci neuronowe. Sieci neuronowe mają szereg cech przydatnych do analizy i prognozowania szeregów czasowych. Ich akutechność dotyczy przede wszystkim dostrajania przyjętej struktury na podstawie danych. Proces budowy modelu neuronowego polega na eksploracji dostępnych zbiorów danych i umożliwia automatyczne szacowanie modelu. Dodatkową zaletą sieci neuronowych jest łatwość ich adaptacji do zmieniających się warunków rynkowych (Cieźak i in., (2006).

W przeciwieństwie do tradycyjnych metod opartych na modelach, **sztuczne sieci neuronowe** (Artificial Neural Networks, ANNs) są samodostosowującymi się metodami opartymi na danych. Uczą się na przykładach i wychwytyują subtelne relacje funkcjonalne między danymi, nawet jeśli podstawowe relacje są nieznane lub trudne do opisanie. Sztuczne sieci neuronowe mogą uogólniać (Hornik i in., 1989). Po zapoznaniu się z przedstawionymi danymi mogą wywnioskować niewidoczną część populacji, nawet jeśli dane próbki zawierają zaszumione informacje. Mają również bardziej ogólne i elastyczne formy funkcjonalne niż tradycyjne metody statystyczne (Zhang i in., 1998).



Pytania do rozdziału

1. Jakie są ograniczenia stosowania metod szeregów czasowych w prognozowaniu popytu?
2. Jakie są główne składowe szeregu czasowego i jakie mają one znaczenie w procesie prognozowania?
3. Jakie są różne strategie postępowania w przypadku zaobserwowania danych odstających?

LITERATURA

Adhikari R., Agrawal R.K. (2014) A combination of artificial neural network and random walk models for financial time series forecasting. *Neural Comput i Applic* 24, 1441-1449.

Ali M.M., Boylan J.E. i Syntetos A.A. (2012). Forecast errors and inventory performance under forecast information sharing, *International Journal of Forecasting*, 28(4), 830-841.

Altendorfer K. i Felberbauer T. (2023) Forecast and production order accuracy for stochastic forecast updates with demand shifting and forecast bias correction. *Simulation Modelling Practice and Theory*, 125, 102740.

Awwad, M., Kulkarni, P., Bapna, R. i Marathe, A. (2018). Big data analytics in supply chain: a literature review. *Proceedings of the International Conference on Industrial Engineering and Operations Management Washington DC, USA, Vol. 2018*, 418-25.

Carreno J.J. i Madinaveitia J. (1990) A modification of time series forecasting methods for handling announced price increases. *International Journal of Forecasting*, 6(4), 479-484.

Chatfield Ch., Koehler A.B., Ord J.K. i Snyder R.D. (2001) A New Look at Models For Exponential Smoothing, *Journal of the Royal Statistical Society: Series D (The Statistician)*, 50(2), 147-159.

Chatfield, C. (2001). *Time Series Forecasting*, Chapman i Hall/CRC, Floryda: Boca Raton

Chen, M., Mao, S., Zhang, Y. i Leung, V.C. (2014). Big data: related technologies, challenges and future prospects (Vol. 100). Heidelberg: Springer

Cieślak, M. (Red.) (2005): *Prognozowanie gospodarcze. Metody i zastosowania*. Wyd. Naukowe PWN. Warszawa



Cieżak W., Siwoń Z. i Cieżak J. (2006) Zastosowanie sztucznych sieci neuronowych do prognozowania szeregów czasowych krótkotrwałego poboru wody w wybranych systemach wodociągowych. *Ochrona Środowiska*, 28(1), 39-44.

De Gooijer J.G. i Hyndman R.J. (2006) 25 years of time series forecasting, *International Journal of Forecasting*, 22(3), 443-473.

Dittmann, P. (2003). *Prognozowanie w przedsiębiorstwie*. Kraków: Oficyna Ekonomiczna

Gajda, J.B. (2001). *Prognozowanie i symulacja a decyzje gospodarcze*, Warszawa: Wydawnictwo C.H. Beck

Galicia A., Talavera-Llames R., Troncoso A., Koprinska I. i Martínez-Álvarez F. (2019). Multi-step forecasting for big data time series based on ensemble learning, *Knowledge-Based Systems*, 163, pp. 830-841.

Grzelak M. (2019) Zastosowanie modelu ARIMA do prognozowania wielkości produkcji w przedsiębiorstwie, *Systemy Logistyczne Wojsk*, 50.

Grzybowska, K. (2009). *Gospodarka Zapasami i Magazynem*, cz. 1, Difin.

Güllü R. (1996) On the value of information in dynamic production/inventory problems under forecast evolution. *Naval Research Logistics (NRL)*, 43(2), 289-303.

Hopp W.J. i Spearman M.L. (1999). *Factory physics*. 2nd. edn. McGraw-Hill / Irwin: Boston

Hornik K., Stinchcombe M. i White H. (1989) Multilayer feedforward networks are universal approximators. *Neural networks*, 2(5), 359-366.

Hyndman R.J. i Koehler A.B. (2006) Another look at measures of forecast accuracy. *International journal of forecasting*, 22(4), 679-688.

Kucharski, A. (2013). *Prognozowanie szeregów czasowych metodami ewolucyjnymi*. Łódź: Wydawnictwo Uniwersytetu Łódzkiego

Malska W. i Wachta H. (2015) Wykorzystanie modelu ARIMA do analizy szeregu czasowego. *Advances in IT and Electrical Engineering*, 23(34-3), 23-30.

Michna Z., Disney S.M., Nielsen P. (2020). The impact of stochastic lead times on the bullwhip effect under correlated demand and moving average forecasts, *Omega*, 93, 102033.



- Rosas A.L. i Guerrero V.M. (1994) Restricted forecasts using exponential smoothing techniques. *International Journal of Forecasting*, 10(4), 515-527.
- Schaffer A.L., Dobbins T.A. i Pearson S.A. (2021) Interrupted time series analysis using autoregressive integrated moving average (ARIMA) models: a guide for evaluating large-scale health interventions. *BMC Med Res Methodol* 21(58).
- Sczaniecka, E. i Smarzyńska, N. (2018). Logistyka wyprzedzająca, czyli innowacyjne podejście do branży e-commerce. *Journal of TransLogistics*, 4(1), 119-128
- Seyedan M. i Mafakheri F. (2020) Predictive big data analytics for supply chain demand forecasting: methods, applications, and research opportunities. *J Big Data* 7(53).
- Sheldon, P.J. (1993). Forecasting tourism: expenditure versus arrivals, *Journal of Travel Research*, 32, 13-20
- Sobczyk M. (2006) Statystyka, aspekty praktyczne i teoretyczne. Lublin: Wydaw. UMCS.
- Tan T., Güllü R. i Erkip N. (2009) Using imperfect advance demand information in ordering and rationing decisions, *International Journal of Production Economics*, 121(2), 665-677.
- Trojanowska M. i Malopolski J. (2004) Zastosowanie metod sztucznej inteligencji do prognozowania miesięcznej sprzedaży energii elektrycznej na wsi. *Acta Scientiarum Polonorum. Technica Agraria*, 3(1-2).
- Wang G., Gunasekaran A., Ngai E.W.T., Papadopoulos T. (2016). Big data analytics in logistics and supply chain management: Certain investigations for research and applications. *International Journal of Production Economics*, 176, 98-110.
- Wojciechowski, A. i Wojciechowska, N. (2015). Zastosowanie klasycznych metod prognozowania popytu w logistyce dużych sieci handlowych. *Zeszyty Naukowe Uniwersytetu Szczecińskiego*, 41, 545-554
- Wolny, M. i Kmiecik, M. (2020). Forecasting demand for products in distribution networks using R software. *Zeszyty Naukowe. Organizacja i Zarządzanie/Politechnika Śląska*, 107-116
- Zeiml S., Altendorfer K., Felberbauer T. i Nurgazina J. (2019) Simulation based forecast data generation and evaluation of forecast error measures. In *2019 Winter Simulation Conference (WSC)* (pp. 2119-2130). IEEE.



Zhang G., Patuwo B.E. i Hu M.H. (1998) Forecasting with artificial neural networks:: The state of the art, International Journal of Forecasting, 14(1), 35-62.

(www_8.1) <https://online.stat.psu.edu/stat501/> (dostęp: 08.02.2024)

(www_8.2) <https://machinelearningmastery.com/exponential-smoothing-for-time-series-forecasting-in-python/> (dostęp: 08.02.2024)

(www_8.3) <https://www.inventory-planner.com/forecasting-in-excel/> (dostęp: 08.02.2024)

(www_8.4) [IPM Insights Metrics \(oracle.com\)](https://www.oracle.com/inventory-planner/) (dostęp: 05.02.2024)