

AGH

AKADEMIA GÓRNICZO-HUTNICZA IM. STANISŁAWA STASZICA W KRAKOWIE
WYDZIAŁ ZARZĄDZANIA

Samodzielna Pracownia Zastosowań Matematyki w Ekonomii

Praca dyplomowa licencjacka

*Analiza czynników wpływających na zmianę trenerów w
piłce nożnej na podstawie polskiej Ekstraklasy*

*Analysis of factors impacting the change of coaches in
football based on the Polish Ekstraklasa*

Autor:
Kierunek studiów:
Opiekun pracy:

Tomasz Bugajski
Informatyka i Ekonometria
dr Łukasz Lach

Kraków, 2020

Spis treści

Wstęp	2
1. Metodyka badań	4
1.1 Model regresji	4
1.1.1 Regresja liniowa.....	4
1.2 Regresja logistyczna	5
1.2.1 Iloraz szans	6
1.2.2 Metoda największej wiarygodności	6
1.3 Testowanie modelu regresji logistycznej.....	7
1.3.1 Badanie współliniowości.....	7
1.3.2 Test ilorazu wiarygodności.....	7
1.3.3 Test Hosmera-Lemeshowa	8
1.3.4 Współczynnik determinacji	9
1.3.5 Test Walda.....	9
1.4 Testowanie mocy predykcyjnej modelu	10
1.4.1 Analiza ROC.....	11
2. Wykorzystany zbiór danych.....	13
2.1 Charakterystyka zmiennych.....	13
2.2 Podstawowe statystyki opisowe	15
2.3 Trenerzy w Ekstraklasie	16
3. Badania Empiryczne.....	19
3.1 Przygotowanie danych.....	19
3.2 Wybór modelu.....	19
3.2.1 Model ze wszystkimi zmiennymi	19
3.2.2 Porównanie modeli.....	20
3.3 Charakterystyka modelu	21
3.4 Badanie współliniowości.....	22
3.5 Dopasowanie modelu	24
3.6 Interpretacja modelu	24
3.7 Moc predykcyjna modelu	26
Podsumowanie	29
Bibliografia	31
Spis tabel.....	32

Wstęp

Praca trenera to bardzo stresujący i niepewny zawód. Szczególnie widać to w piłce nożnej, gdzie są największe pieniądze, ale również największa presja społeczeństwa, trenerzy nigdy nie mogą być pewni swojego stanowiska. Niestety zdarza się, szczególnie w klubach będących własnością bogatych przedsiębiorców, że szkoleniowcy tracą pracę pod wpływem impulsu właściciela. Historia polskiego futbolu po roku 1989 zna wiele takich przypadków, jak Wisła Kraków zarządzana przez Bogusława Cupiała, czy Polonia Warszawa będąca własnością Józefa Wojciechowskiego, gdzie trenerzy byli zwalniani pomimo osiągnięcia dobrych wyników.

W ostatnim czasie w Ekstraklasie jest mniej takich klubów, jednak dalej liczba zwalnianych trenerów w Polsce jest większa niż w krajach lepiej radzących sobie na arenie międzynarodowej. Ta niestabilność wydaje się być jednym z czynników wpływających na kolejne niepowodzenia w europejskich pucharach. Warto wspomnieć o tym, że mistrzem Polski w ostatnim sezonie został Piast Gliwice, zespół w którym trener dostał trochę więcej czasu na ułożenie drużyny i nie został zwolniony przy pierwszym kryzysie, choć tutaj zdecydował raczej przypadek niż rozsądna polityka klubu.¹ Jednak ta sytuacja może być dobrym przykładem dla innych polskich klubów, że czasami warto dać trenerowi więcej czasu.

Celem niniejszej pracy jest znalezienie czynników istotnie wpływających na zwolnienia trenerów w piłce nożnej. Przede wszystkim autorowi pracy zależałoby na sprawdzeniu, czy na zwolnienia szkoleniowców w polskiej Ekstraklasie mają wpływ wyniki uzyskiwane w ostatnich meczach, jak również powiązana z nimi frekwencja na trybunach. Dodatkowo sprawdzone zostanie czy wpływ na badane zjawisko mają przyzwyczajenia zarówno klubów, jak i trenerów, zatem czy są kluby, w których nieustannie zmienia się menedżerów oraz, czy są trenerzy którzy w klubach zazwyczaj są zatrudniani tylko na krótki okres. Wykorzystany do tego zostanie model regresji logistycznej, ze względu na prostotę interpretacji czynników w modelu. Praca została podzielona na trzy części. Rozdział pierwszy zawierać będzie teoretyczną część badania. W drugim rozdziale zaprezentowany zostanie zbiór danych, na którym będzie przeprowadzana analiza.

¹ L. Błażyński, *Władze Piasta rozczarowane. Formalik może zostać zwolniony*, Przegląd sportowy, <https://www.przegladsportowy.pl/pilka-nozna/pko-ekstraklasa/piast-gliwice/waldemar-formalik-moze-zostac-zwolniony-z-piast-gliwice/4rv2x73> [dostęp: 29.04.2020]; <https://www.przegladsportowy.pl/pilka-nozna/pko-ekstraklasa/piast-gliwice/mateusz-borek-zdradza-kulisy-pracy-waldemara-formalika-w-piastcie-gliwice-misja-futbol/tnhxkmb> [dostęp: 29.04.2020]

Dodatkowo opisana zostanie sytuacja trenerów w polskiej Ekstraklasie. Trzeci rozdział to badanie empiryczne, które zostanie przeprowadzone w środowisku R.

1. Metodyka badań

Badanie empiryczne w pracy zostanie przeprowadzone z wykorzystaniem regresji logistycznej, jednak przed jej omówieniem należy zdefiniować pojęcie regresji.

1.1 Model regresji

Regresja jest to metoda statystyczna badająca zależności pomiędzy jedną lub kilkoma zmiennymi zależnymi (objaśnianymi), a zmiennymi niezależnymi (objaśniającymi). Głównym celem modelu regresji jest przewidywanie wartości zmiennej zależnej, na podstawie zmiennych niezależnych. Jako, że polega ona na stworzeniu funkcji opisującej wartość oczekiwaną zmiennej objaśnianej w zależności od zmiennych objaśniających, jej ogólna postać dana jest wzorem:

$$E(Y|X) = f(X, \beta) \quad (1.1)$$

$$Y = f(X, \beta) + \varepsilon \quad (1.2)$$

gdzie:

Y – zmienna zależna,

X – zmienne niezależne,

$E(Y|X)$ – warunkowa wartość oczekiwana Y przy danej wartości X ,

β – wektor współczynników regresji,

$f(X, \beta)$ – funkcja regresji,

ε – błąd losowy.²

1.1.1 Regresja liniowa

Model regresji liniowej dla wielu zmiennych niezależnych (regresja liniowa wielokrotna) przyjmuje postać:

$$Y = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \dots + \beta_k x_k + \varepsilon \quad (1.3)$$

gdzie:

Y – zmienna zależna,

β_0 – wyraz wolny (punkt przecięcia z osią Y),

$\beta_1, \beta_2, \dots, \beta_k$ – współczynniki kierunkowe nachylenia linii regresji,

$x_1, x_2, \dots, x_k$ – zmienne niezależne,

ε – błąd losowy.

² A. Kirpsza, *Zastosowanie regresji logistycznej w studiach nad Unią Europejską*, [w:] *Metody jakościowe i ilościowe w badaniu organizacji i działania Unii Europejskiej*, pod red. K. Ławniczaka, Warszawa 2013, s. 11-32.

W regresji liniowej celem jest znalezienie współczynników β poprzez wybranie hiperpłaszczyzny najbliższej wszystkim obserwacjom. Dokonuje się tego wykorzystując metodę najmniejszych kwadratów, która polega na znalezieniu najmniejszej sumy kwadratów odległości punktów obserwacji na wykresie od linii regresji.³

1.2 Regresja logistyczna

W przypadku gdy zmienna zależna jest binarna lub kategoryczna z kilku powodów nie można stosować regresji liniowej. Po pierwsze warunkowa wartość oczekiwana inaczej niż w regresji liniowej, nie może przyjąć wartości mniejszych od 0 i większych od 1. Jej krzywa natomiast nie jest liniowa tylko przybiera kształt litery S. Po drugie błąd losowy nie ma rozkładu normalnego i jego wariancja nie jest homoskedastyczna. W takim przypadku najczęściej stosowaną metodą jest regresja logistyczna. Model regresji logistycznej wygląda następująco:

$$P(Y) = \frac{e^{\beta_0 + \beta_1 x_1 + \dots + \beta_k x_k}}{1 + e^{\beta_0 + \beta_1 x_1 + \dots + \beta_k x_k}} = \frac{1}{1 + e^{-(\beta_0 + \beta_1 x_1 + \dots + \beta_k x_k)}} = \frac{1}{1 + e^{-\text{logit}(p)}} \quad (1.4)$$

gdzie:

e – stała Eulera,

$\text{logit}(p)$ – zostanie objaśniony w dalszej części pracy.

Inaczej niż w przypadku regresji liniowej, wzór nie opisuje wartości oczekiwanej, a prawdopodobieństwo wystąpienia badanego zjawiska.⁴ Jednak w modelach regresji logistycznej nie używa się zwykle klasycznego prawdopodobieństwa, tylko pojęcie szansy (ang. *odds*), które wyraża się wzorem:

$$O = \frac{p}{q} = \frac{p}{1 - p} = e^{\beta_0 + \beta_1 x_1 + \dots + \beta_k x_k} \quad (1.5)$$

Szansa jest to zatem stosunek prawdopodobieństwa sukcesu p do prawdopodobieństwa porażki q . Szanse przyjmują wartości od 0 do nieskończoności, przy czym dla wartości większych od 1 szansa sukcesu jest większa niż porażki. Na przykład dla szansy równej 4, szansa sukcesu jest cztery razy większa od porażki. Analogicznie dla szansy 1/4, szansa porażki jest cztery razy większa od szansy sukcesu.⁵

Jednak z uwagi na fakt, że szanse nie przyjmują wartości ujemnych i nie mają rozkładu linearnego względem zmiennych niezależnych, nie jest możliwa łatwa interpretacja współczynników regresji β . Dlatego w modelach regresji logistycznej stosuje się

³ Xin Yan, *Linear Regression Analysis: Theory and Computing*, World Scientific, Singapur 2009, s. 9-12.

⁴ D. Hosmer, S. Lemeshow, *Applied Logistic Regression*, John Wiley & Sons, New York 2000, s. 5-6.

⁵ A. Agresti, *Categorical data analysis*, John Wiley & Sons, New Jersey 2013, s. 44-45.

transformację logitową, wyrażaną wzorem:

$$\text{logit}(p) = \ln\left(\frac{p}{1-p}\right) = \beta_0 + \beta_1 x_1 + \dots + \beta_k x_k \quad (1.6)$$

Transformacja ta jest to zatem logarytm naturalny szansy i pozwala ona interpretować współczynniki regresji podobnie jak w regresji liniowej.⁶

1.2.1 Iloraz szans

Iloraz szans (ang. *odds ratio*) informuje jak zmienia się zmienna zależna, w przypadku wzrostu predyktora o jedną jednostkę. Oblicza się go ze wzoru:

$$OR = \frac{O_1}{O_2} = \frac{p_1/(1-p_1)}{p_2/(1-p_2)} \quad (1.7)$$

gdzie:

p_i – prawdopodobieństwo zajścia zdarzenia (wartość 1 w zmiennej zależnej) w i -tej klasie obserwacji.

Podobnie jak w przypadku szans równanie przyjmuje wartości dodatnie. Iloraz szans interpretuje się następująco:

$OR > 1$ – predyktor ma stymulujący wpływ na zmienną zależną. Wartość ilorazu szans przekazuje informacje o ile zwiększa się szansa wystąpienia badanego zdarzenia;

$OR < 1$ – predyktor ma ograniczający wpływ na zmienną zależną. Wartość ilorazu szans przekazuje informacje o ile zmniejsza się szansa wystąpienia badanego zdarzenia;

$OR = 1$ – predyktor nie ma wpływu na zmienną zależną.⁷

1.2.2 Metoda największej wiarygodności

W regresji liniowej do estymacji współczynników β wykorzystuje się metodę najmniejszych kwadratów, jednak w regresji logistycznej, w której zmienna objaśniana jest kategorierna, nie jest możliwe wykorzystanie tej metody. W regresji logistycznej zatem wykorzystuje się inne narzędzie, zwane metodą największej wiarygodności. W metodzie tej należy skonstruować funkcję wiarygodności, składającą się z nieznanych parametrów β , która określa prawdopodobieństwo otrzymania zaobserwowanych danych. Następnie znajduje się współczynniki regresji, poprzez maksymalizację funkcji wiarygodności. Funkcję tę definiuje się jako:

$$L(\beta) = \prod_{i=1}^n p_i^{y_i} (1-p_i)^{1-y_i} \quad (1.8)$$

⁶ D. Hosmer, S. Lemeshow, *Applied Logistic...*, s. 6-7.

⁷ A. Agresti, *Categorical data...*, s. 44-46.

gdzie:

p_i – prawdopodobieństwo sukcesu dla i -tej jednostki obserwacji,

y_i – wartość zmiennej zależnej dla i -tej jednostki obserwacji.

Jednak, ponieważ łatwiej ze względów obliczeniowych pracuje się na jej logarytmie, zazwyczaj zastępuje się ją funkcją w następującej postaci:

$$\ln[L(\beta)] = \sum_{i=1}^n [y_i \ln(p_i) + (1 - y_i) \ln(1 - p_i)] \quad (1.9)$$

W celu znalezienia estymatorów największej wiarygodności β , należy zróżniczkować równanie (1.9) względem β_0 i β_k , a następnie przyrównać je do zera.⁸

1.3 Testowanie modelu regresji logistycznej

1.3.1 Badanie współliniowości

W pracy do zbadania współliniowości zostanie wykorzystany współczynnik VIF (Variance inflation factor). Współczynnik ten pokazuje, o ile wariancje parametrów są zawyżone z powodu zależności liniowych w badanym modelu regresji. Oblicza się go ze wzoru:

$$VIF_i = \frac{1}{1 - R_i^2} \quad (1.10)$$

gdzie:

R_i^2 – jest współczynnikiem determinacji wielokrotnej dla i -tej zmiennej w modelu regresji.

Przyjmuje się, że dla VIF większego od 10 można podejrzewać wystąpienie współliniowości.⁹

1.3.2 Test ilorazu wiarygodności

Funkcja wiarygodności służy do wyboru zmiennych do modelu. Zazwyczaj stosuje się jej logarytm pomnożony przez -2:

$$D = -2 \log L \quad (1.11)$$

Takie równanie nazywa się dewiancją. Dewiancja przyjmuje wartości dodatnie i im jej wartość jest niższa, tym model lepiej opisuje zmienną zależną.

⁸ D. Hosmer, S. Lemeshow, *Applied Logistic...*, s. 8-9.

⁹ Xin Yan, *Linear Regression Analysis...*, s. 85-86.

W teście ilorazu wiarygodności porównuje się dewiancję dla modelu tylko z wyrazem wolnym i modelu ze zmiennymi niezależnymi. Do tego celu używa się statystyki G, wyrażonej wzorem:

$$G = D_0 - D_m = -2 \log L(\text{model bez } x_k) + 2 \log L(\text{model z } x_k) = \quad (1.12)$$

$$= -2 \log \left[\frac{L(\text{model bez } x_k)}{L(\text{model z } x_k)} \right]$$

Statystyka G ma rozkład chi-kwadrat z k stopniami swobody. W teście stawia się hipotezę zerową o nieistotności statystycznej modelu ze zmiennymi niezależnymi, oraz hipotezę alternatywną o istotności modelu ze zmiennymi. Jeśli wartość p w teście jest mniejsza lub równa 0,05 odrzuca się H_0 i można uznać, że model opisuje zmienną niezależną lepiej niż w przypadku modelu tylko z wyrazem wolnym.¹⁰

1.3.3 Test Hosmera-Lemeshowa

Oceny jakości dopasowania modelu dokonuje się w teście Hosmera-Lemeshowa. Polega on na podziale obserwacji na dziesięć podgrup (grupy decylowe) i porównaniu liczby zaobserwowanych sukcesów oraz oczekiwanych sukcesów. Hipoteza zerowa w teście mówi, że zaobserwowane i oczekiwane liczebności są równe we wszystkich podgrupach, natomiast hipoteza alternatywna, że różnią się dla przynajmniej jednej podgrupy. W teście stosuje się statystykę wyrażoną wzorem:

$$HL = \sum_{k=1}^g \frac{(o_k - n_k \bar{\pi}_k)^2}{n_k \bar{\pi}_k (1 - \bar{\pi}_k)} \quad (1.13)$$

gdzie:

o_k - zdarzenia zaobserwowane,

n_k - liczba obserwacji,

$\bar{\pi}_k$ - prognozowane ryzyko,

g - liczba podgrup.

Statystyka ta ma rozkład chi-kwadrat z $g - 2$ stopniami swobody. Dla wartości p większych od 0,05 nie ma podstaw do odrzucenia hipotezy zerowej i uznaje się wtedy, że model jest dobrze dopasowany.¹¹ Problemem testu Hosmera-Lemeshowa jest, że na jego podstawie nie można porównywać modeli.¹²

¹⁰ S. Menard, *Applied Logistic Regression Analysis*, Sage University Paper Series on Quantitative Applications in the Social Sciences 07-106, Thousand Oaks 2002, s. 20-23.

¹¹ D. Hosmer, S. Lemeshow, *Applied Logistic...*, s. 147-151.

¹² M. Chlebus, *Pomiar ryzyka rynkowego za pomocą miary Value at Risk – podejście dwuetapowe*, Rozprawa doktorska, Uniwersytet Warszawski, 2014, s. 57-59.

1.3.4 Współczynnik determinacji

Miarą dopasowania modelu która pozwala porównywać modele jest współczynnik determinacji R^2 . W regresji logistycznej najczęściej stosuje się zmienioną wersję współczynnika determinacji nazywaną R^2 McFaddena (tzw. pseudo- R^2) i wyrażoną wzorem:

$$R_{McFaddena}^2 = 1 - \frac{\ln L_{FM}}{\ln L_0} \quad (1.14)$$

gdzie:

$\ln L_{FM}$ – logarytm funkcji wiarygodności w oszacowanym modelu,

$\ln L_0$ – logarytm funkcji wiarygodności w modelu bez zmiennych objaśniających.

Współczynnik determinacji przyjmuje wartości między 0 i 1, jednak w przypadku R^2 McFaddena wartości te zazwyczaj są niskie i nie zbliżają się do 1.¹³ Z tego też powodu zostało stworzone kilka innych wersji współczynnika determinacji, w tym R^2 Nagelkerke'a które zostaną wykorzystane w tej pracy. Przyjmuje ono następującą postać:

$$R_{Nagelkerke}^2 = \frac{1 - \left(\frac{\ln L_0}{\ln L_{FM}} \right)^{2/n}}{1 - (\ln L_0)^{2/n}} \quad (1.15)$$

gdzie:

n – liczba obserwacji.

Ten współczynnik również może przyjmować wartości od 0 do 1, jednak został stworzony tak, aby wartości rozkładały się bardziej równomiernie pomiędzy 0 i 1.¹⁴

1.3.5 Test Walda

W celu analizy istotności statystycznej wpływu predyktorów na zmienną zależną wykorzystujemy się test Walda. Hipoteza zerowa w teście mówi, że wpływ predyktora na zmienną zależną jest nieistotny statystycznie. Statystyka testowa ma postać:

$$\text{Statystyka Walda} = \frac{\beta^2}{S.E. (\beta)^2} \quad (1.16)$$

gdzie:

β – współczynnik regresji,

S.E. – błąd standardowy.

¹³ P. D. Allison, *Measures of Fit for Logistic Regression*, Statistical Horizons LLC and the University of Pennsylvania, 2014.

¹⁴ J. Cohen, P. Cohen, S. G. West, L. S. Aiken, *Applied Multiple Regression/Correlation Analysis for the Behavioral Sciences*, Lawrence Erlbaum Associates, Publishers, Mahwey 2003, s. 502-504.

Statystyka ma rozkład chi-kwadrat z jednym stopniem swobody. Jeśli wartość p w teście jest mniejsza lub równa 0,05 odrzuca się H_0 i można uznać, że zmienna jest istotna statystycznie.¹⁵

1.4 Testowanie mocy predykcyjnej modelu

W modelach regresji logistycznej moc predykcyjną testuje się na podstawie tabel klasyfikacyjnych oraz na podstawie krzywej ROC (ang. *relative operating characteristics*) i pola pod jej wykresem.

Przykładowa tabela klasyfikacyjna została przedstawiona w tabeli 1. Tabele klasyfikacyjne są to tablice z dwoma wierszami i dwoma kolumnami. W wierszach zawarte są obserwacje rzeczywiste, natomiast w kolumnach obserwacje przewidywane. Tabele tworzy się na podstawie prognoz których dla punktu odcięcia równego 0,5 dokonuje się według poniższego wzoru:

$$\hat{y}_i = \begin{cases} 1 & \text{dla } \hat{p}_i > 0,5 \\ 0 & \text{dla } \hat{p}_i \leq 0,5 \end{cases} \quad (1.17)$$

gdzie:

$\hat{y}_i$ – przewidywana wartość zmiennej zależnej dla i -tej jednostki obserwacji,

$\hat{p}_i$ – prognozowane prawdopodobieństwo dla i -tej jednostki obserwacji.

Dla poprawy mocy predykcyjnej możliwa jest zmiana wartości punktu odcięcia i przyjmuje wtedy ona postać:

$$\hat{y}_i = \begin{cases} 1 & \text{dla } \hat{p}_i > p^* \\ 0 & \text{dla } \hat{p}_i \leq p^* \end{cases} \quad (1.18)$$

gdzie:

p^* – nowa wartość punktu odcięcia.¹⁶

Tabela 1. Przykładowa tabela klasyfikacyjna

Obserwacje empiryczne	Obserwacje przewidywane	
	Pozytywna	Negatywna
Pozytywna	TP (ang. <i>true positive</i>)	FN (ang. <i>fals negative</i>)
Negatywna	FP (ang. <i>false positive</i>)	TN (ang. <i>true negative</i>)

Źródło: opracowanie własne.

¹⁵ J. Cohen i inni, *Applied Multiple Regression...*, s. 507-508.

¹⁶ A. Sączewska-Piotrowska, *Zastosowanie krzywych ROC w analizie ubóstwa miejskich i wiejskich gospodarstw domowych*, Przegląd statystyczny r. LXIII, zeszyt 2, Katowice 2016.

Na podstawie tabel klasyfikacyjnych oblicza się trzy główne wskaźniki: trafność (ang. *accuracy*), czułość (ang. *sensitivity*) i specyficzność (ang. *specificity*). Pierwszy z nich przedstawia procent wszystkich poprawnie sklasyfikowanych przypadków. Oblicza się go ze wzoru:

$$ACC = \frac{TP + TN}{TP + TN + FP + FN} \quad (1.19)$$

Czułość przedstawia procent przypadków poprawnie sklasyfikowanych wśród przypadków pozytywnych i przyjmuje postać:

$$C = \frac{TP}{TP + FN} \quad (1.20)$$

Specyficzność natomiast przedstawia procent przypadków poprawnie sklasyfikowanych wśród przypadków negatywnych i wyraża się wzorem:

$$S = \frac{TN}{FP + TN} \quad (1.21)$$

Wskaźniki te pozwalają ocenić w jakim stopniu model potrafi poprawnie przewidzieć wartość zmiennej objaśnianej, a zatem jaka jest jego moc predykcyjna.¹⁷

1.4.1 Analiza ROC

Krzywe ROC tworzy się na podstawie czułości i specyficzności modelu. Dla wszystkich punktów odcięcia z przedziału 0-1 oblicza się te dwa wskaźniki i następnie nanosi je się na układ współrzędnych gdzie na osi rzędnych znajduje się czułość, natomiast na osi odciętych zazwyczaj stosuje się dopełnienie specyficzności do liczby 1. Uzyskane punkty łączy się ze sobą. Na wykresie 1 została zilustrowana przykładowa krzywa ROC. Optymalnym punktem odcięcia jest punkt znajdujący się najbliżej punktu o współrzędnych (0,1), a zatem taki dla którego czułość i specyficzność są jak największe.

Moc predykcyjną modelu na podstawie krzywej ROC uzyskuje się obliczając pole pod wykresem krzywej nazywane AUC (ang. *area under curve*). Model jest dobrze skonstruowany jeśli AUC jest większe niż 0,5, a zatem wykres krzywej znajduje się powyżej przekątnej $y = x$. Dla AUC przyjmuje się, że jeżeli:

$AUC = 0,5$ – brak klasyfikacji,

$0,6 \leq AUC < 0,7$ – akceptowalna klasyfikacja,

$0,7 \leq AUC < 0,8$ – dobra klasyfikacja,

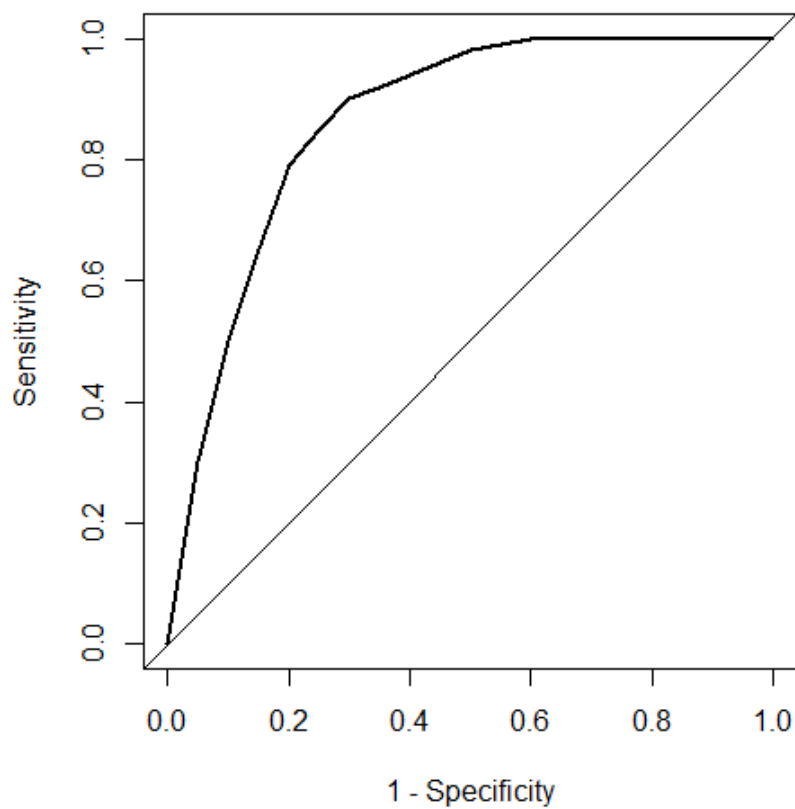
$0,8 \leq AUC < 0,9$ – bardzo dobra klasyfikacja,

$AUC \geq 0,9$ – wyśmienita klasyfikacja.

¹⁷ A. Kirpsza, *Zastosowanie regresji logistycznej...*

Jeśli model idealnie przewiduje obserwacje to jego krzywa ROC przechodzi przez punkt $(0,1)$, a wartość AUC jest równa 1.¹⁸

Wykres 1. Przykładowa krzywa ROC



Źródło: opracowanie własne.

¹⁸ A. Sączewska-Piotrowska, *Zastosowanie krzywych ROC...*

2. Wykorzystany zbiór danych

2.1 Charakterystyka zmiennych

W celu przeprowadzenia badania została stworzona baza danych zawierająca wszystkie kolejki w których następowało więcej niż jedno zwolnienie trenera w polskiej Ekstraklasie w sezonach 2015/16-2019/20. Baza pokazuje jak wyglądała sytuacja w danych meczach we wszystkich klubach Ekstraklasy. Dane pochodzą ze stron Transfermarkt¹⁹ oraz FBref²⁰. W tym okresie wystąpiło 9 takich kolejek, co przełożyło się na 144 obserwacje. Zmienna objaśniana przyjęła wartości:

- $y = 1$ dla klubów, które w danej kolejce zwolniły trenera,
- $y = 0$ w pozostałych przypadkach.

Ze względu na dostępność danych oraz mając na celu jak najbardziej kompleksowe badanie, przyjęto następujący zestaw zmiennych objaśniających:

pkt5 – liczba punktów zdobyta przez drużynę w 5 ostatnich meczach, w tym w podanej kolejce;

pkt3 – liczba punktów zdobyta przez drużynę w 3 ostatnich meczach, w tym w podanej kolejce;

gole5 – liczba goli zdobytych przez drużynę w 5 ostatnich meczach, w tym w podanej kolejce;

gole3 – liczba goli zdobytych przez drużynę w 3 ostatnich meczach, w tym w podanej kolejce;

stracone5 – liczba goli straconych przez drużynę w 5 ostatnich meczach, w tym w podanej kolejce;

stracone3 – liczba goli straconych przez drużynę w 3 ostatnich meczach, w tym w podanej kolejce;

stracone – liczba goli straconych przez drużynę w podanej kolejce;

bilans – różnica między bramkami zdobytymi a straconymi w podanej kolejce;

frekwencja3 – różnica między średnią liczbą kibiców w 3 ostatnich meczach na własnym stadionie, w tym w podanej kolejce jeśli drużyna grała u siebie, a średnią kibiców na stadionie w całym sezonie;

¹⁹ <https://www.transfermarkt.pl/> [dostęp: 30.03.2020]

²⁰ <https://fbref.com/en/comps/36/history/Ekstraklasa-Seasons> [dostęp: 30.03.2020]

frekwencja5 – różnica między średnią liczbą kibiców w 5 ostatnich meczach na własnym stadionie (w tym w podanej kolejce jeśli drużyna grała u siebie) a średnią kibiców na stadionie w całym sezonie;

trenerzy10 – liczba trenerów którzy pracowali w klubie w 10 sezonach poprzedzających podany (bez uwzględnienia trenerów tymczasowych);

trenerzy5 – liczba trenerów którzy pracowali w klubie w 5 sezonach poprzedzających podany (bez uwzględnienia trenerów tymczasowych);

różnica3 – różnica między liczbą zgromadzonych punktów do podanej kolejki, a średnią liczbą punktów zgromadzonych w poprzednich trzech sezonach w analogicznym okresie (dla drużyn które w poprzednich sezonach nie grały w najwyższej klasie rozgrywkowej liczba punktów zgromadzona w niższych ligach dzielona jest na pół);

różnica2 – różnica między liczbą zgromadzonych punktów do podanej kolejki, a średnią liczbą punktów zgromadzonych w poprzednich dwóch sezonach w analogicznym okresie (dla drużyn które w poprzednich sezonach nie grały w najwyższej klasie rozgrywkowej liczba punktów zgromadzona w niższych ligach dzielona jest na pół);

staż – średnia liczba dni przepracowanych jako trener seniorskich drużyn w poprzednich klubach w których pracował.

Zmienne *pkt5*, *pkt3*, *stracone5*, *stracone3* charakteryzują efektywność osiąganych wyników przez drużynę w ostatnich kolejkach. Jeśli zarząd nie wymaga od trenera efektywnej gry, tylko stawia na pragmatyczność, najważniejszym czynnikiem jest zdobywanie jak największej liczby punktów i strata małej liczby bramek. Zmienne *stracone*, *bilans* charakteryzują postawę drużyny w ostatnim meczu. Często zdarza się, że trenerzy tracą pracę z powodu blamażu drużyny w jednym meczu. Zmienne *gole5*, *gole3*, *frekwencja3*, *frekwencja5* odzwierciedlają atrakcyjność futbolu prezentowanego przez drużynę w ostatnim czasie. Niewielka liczba bramek i spadek liczby widzów na trybunach mogą świadczyć o nieefektywnym stylu gry i przez to obniżeniu zainteresowania klubem wśród fanów. Zmienne *trenerzy10*, *trenerzy5* charakteryzują politykę klubu w ostatnich latach w kwestii trenerów. Można przypuszczać, że kluby często zmieniające trenerów z większym prawdopodobieństwem zwolnią kolejnego trenera. Zmienne *różnica3*, *różnica2* znamionują jak drużyna radzi sobie z oczekiwaniami kibiców i wszystkich ludzi związanych z klubem. Jeśli drużyna zdobywa więcej punktów niż w poprzednich latach, fani prawdopodobnie będą zadowoleni z rozwoju drużyny. Zmienna *staż* charakteryzuje stabilność i cierpliwość trenera. Krótki średni staż szkoleniowca w poprzednich klubach

może świadczyć o jego konfliktowości i zwiastować kolejne szybkie zwolnienie ze stanowiska.

2.2 Podstawowe statystyki opisowe

Tabela 2 przedstawia statystyki opisowe zmiennych niezależnych. Można w niej zauważyć, że bywały drużyny, które miały serię 5 porażek z rzędu. Zazwyczaj takie wyniki sprawiają, że trener może obawiać się zwolnienia. W przypadku tej samej zmiennej, czyli *pkt5* widać, że żadnej z drużyn nie udało się wygrać 5 meczów z rzędu, co może być dość zaskakujące. Może tu jednak mieć znaczenie charakterystyka polskiej ligi, w której zazwyczaj nie ma jednej, lub dwóch drużyn, które zdecydowanie zdominowałyby ligę i wygrały znaczącą większość meczów. Średnie goli dla pięciu i trzech meczów wynoszą 6,74 oraz 4,04 i przedstawiają się bardzo podobnie do średniej goli straconych, które odpowiednio wynoszą 6,76 i 4,07, a zatem drużyny strzelają i tracą średnio więcej niż gola na mecz, co wskazywałoby na sporą liczbę bramek w meczach i zaciętą walkę o zwycięstwo. O zaciętości pojedynków może również mówić średnia i mediana zmiennej *bilans*, która w obu przypadkach wynosi 0. *frekwencja5* i *frekwencja5* charakteryzują się bardzo dużym odchyleniem standardowym co odzwierciedla usposobienie kibiców, którzy dużo chętniej chodzą na mecze, kiedy ich drużyna jest w formie i gra z mocnym przeciwnikiem, natomiast w innych przypadkach kibiców na meczach jest zwykle zdecydowanie mniej. Średnie dla zmiennych *trenerzy10* i *trenerzy5* to odpowiednio 10,94 i 5,92. Ilustrują one złą tendencję polskich klubów do bardzo częstych zmian trenerów, a przypadki takie jak dwudziestu trenerów w ciągu 10 lat, lub dwunastu w ciągu 5 lat (warte podkreślenia, że chodzi o sytuacje w dwóch innych klubach) stanowią bardzo ciekawe precedensy. Pewnym zaprzeczeniem takiego obrazu polskich klubów, może być zmienna *staż*, w której kwartył dolny wyniósł 367,40, a zatem 75% trenerów pracujących w Ekstraklasie pracuje w klubie dłużej niż rok. Zmienna ta charakteryzuje się jeszcze jednym ważnym aspektem, a mianowicie ma 12 braków w danych i jest to spowodowane tym, że dla części trenerów praca w Ekstraklasie była ich pierwszą pracą jako menadżerowie seniorskiej drużyny. Problem ten zostanie przeanalizowany w dalszej części pracy.

Tabela 2. Statystyki opisowe

Zmienna	Min	Kwartył pierwszy	Mediana	Średnia	Kwartył trzeci	Max	Odchylenie Standardowe	NA
Pkt5	0,00	5,00	7,00	6,80	9,00	13,00	3,01	0
Pkt3	0,00	2,00	4,00	4,09	6,00	9,00	2,35	0
Gole5	1,00	5,00	6,00	6,74	8,00	18,00	2,79	0
Gole3	0,00	2,00	4,00	4,04	5,00	13,00	2,25	0
Stracone5	1,00	5,00	7,00	6,76	9,00	13,00	2,68	0
Stracone3	0,00	2,00	4,00	4,07	5,00	10,00	2,18	0
Stracone	0,00	0,00	1,00	1,35	2,00	5,00	1,25	0
Bilans	-4,00	-1,00	0,00	0,00	1,00	4,00	1,77	0
Frekwencja5	-7449,00	-1082,25	-82,00	-75,57	813,50	7463,00	2131,40	0
Frekwencja3	-8840,00	-1486,00	-125,50	-117,70	1237,20	9242,00	2571,61	0
Trenerzy10	6,00	9,00	11,00	10,94	12,25	20,00	2,91	0
Trenerzy5	2,00	5,00	6,00	5,92	7,00	12,00	2,02	0
Różnica3	-16,67	-3,50	1,33	0,92	5,33	19,33	6,66	0
Różnica2	-17,00	-4,50	0,75	0,68	5,06	20,25	7,64	0
Staż	84,00	367,40	452,20	515,00	564,70	1542,00	284,82	12

Źródło opracowanie własne.

2.3 Trenerzy w Ekstraklasie

Jak już zostało wspomniane we wstępie trenerów w polskiej Ekstraklasie zmienia się bardzo często. Tabela nr 3 obrazuje liczbę zmian trenerów w wybranych ligach europejskich oraz jaka ich część następowała w trakcie sezonu. Dodatkowo przedstawiono miejsce każdej z lig w rankingu UEFA, obliczanym na podstawie wyników drużyn w europejskich pucharach. Więcej razy niż w polskiej lidze, w ostatnich czterech sezonach, zmieniano trenerów tylko w ligach hiszpańskiej oraz włoskiej, jednak trzeba wziąć pod uwagę, że w tych dwóch ligach podobnie jak w ligach angielskiej i francuskiej występuje po 20 drużyn, natomiast w lidze polskiej tylko 16. Zatem jeśli podzielimy średnią liczbę zmian trenerów przez liczbę drużyn w nich występujących okaże się, że w żadnej z pięciu najmocniejszych lig w Europie nie zmienia się trenerów częściej niż w Ekstraklasie. Jako że polskie kluby nie mogą się równać poziomem finansów, czy popularnością z klubami krajów Europy Zachodniej, dlatego do tabeli wprowadzone zostały również ligi z trzech krajów sąsiadujących z Polską, z którymi polska liga mogłaby skutecznie rywalizować. Z tych trzech lig, tylko w lidze słowackiej w ostatnich czterech sezonach zmieniano

trenerów częściej niż w lidze polskiej. Zatem z ujętych w tabeli dziewięciu lig tylko w jednej zmienia się trenerów częściej niż w Ekstraklasie, jednak w kilku z przedstawionych krajów liczba zmian szkoleniowców na drużynę była tylko niewiele mniejsza niż w polskiej lidze. Tym, co zdecydowanie wyróżnia Ekstraklasę na tle przedstawionych lig jest liczba zmian trenerów w trakcie sezonu. W polskiej najwyższej klasie rozgrywkowej w sezonach 2015/16-2018/19 w trakcie sezonu dokonano 51 zmian trenerów. Tylko w lidze włoskiej przeprowadzono taką samą liczbę roszad na stanowisku trenera, a trzeba pamiętać, że występuje w niej 20 zespołów. Jeśli weźmiemy pod uwagę średnią na sezon podzieloną przez liczbę drużyn, to polska liga nie ma sobie równych. Jest to bardzo niepokojąca tendencja. Zmianę trenera pomiędzy sezonami można uznać za pewien naturalny proces, trener mógł na przykład nie wypełnić postawionych przed nim celów, relacja z piłkarzami mogła się wypalić lub klub znalazł kandydata, który według nich lepiej sprawdzi się w tej roli. Jednak częste zmiany trenerów w trakcie sezonu świadczą raczej o nieudolności prezesów i ciągłym chaosie. W Polsce pokutuje mit tak zwanego „efektu nowej miotły”, to znaczy pewne przekonanie, że drużyna zawsze gra lepiej w pierwszych meczach po zmianie trenera. Jest to jednak tylko pozorne, ponieważ po kilku lepszych meczach, drużyny zazwyczaj wracają do poziomu prezentowanego przed zmianą. Krótkowzroczne podejście ludzi zarządzających polską piłką jest prawdopodobnie jednym z głównych powodów, dla których polskie drużyny nie osiągają wyników na miarę dużego europejskiego kraju, w którym piłka nożna jest jednym z najpopularniejszych sportów.

Tabela 3. Liczba zmian trenerów w wybranych ligach europejskich w sezonach 2015/16-2018/19

Liga	Zmiany trenerów	Średnia na sezon	Średnia / lb. drużyn	Zmiany trenerów w trakcie sezonu	Średnia na sezon	Średnia / lb. drużyn	Miejsce w rankingu UEFA
polska	65	16,25	1,02	51	12,75	0,80	29
hiszpańska	79	19,75	0,99	47	11,75	0,59	1
angielska	50	12,50	0,63	33	8,25	0,41	2
niemiecka	51	12,75	0,71	31	7,75	0,43	3
włoska	80	20,00	1,00	51	12,75	0,64	4
francuska	59	14,75	0,74	35	8,75	0,44	5
czeska	59	14,75	0,92	42	10,50	0,66	15
białoruska ²¹	36	9,00	0,56	13	3,25	0,20	25
słowacka	54	13,50	1,13	28	7,00	0,58	30

Źródło: Opracowanie własne na podstawie Transfermarkt.

²¹ W lidze białoruskiej gra się systemem wiosna-jesień, dlatego do tabeli wzięto dane z sezonów 2016-2019.

Obecnie najdłużej pracującym trenerem w polskiej Ekstraklasie jest Marek Papszun, szkoleniowiec Rakowa Częstochowa, który rozpoczął pracę z tą drużyną 18 kwietnia 2016 roku. Raków występował wtedy jeszcze w 2. lidze, czyli na trzecim poziomie rozgrywkowym. Drugim najdłużej pracującym trenerem jest szkoleniowiec Górnika Zabrze, Marcin Brosz, który pracuje tam od 1 lipca 2016 roku. Następny jest Michał Probiez, szkoleniowiec Cracovii, pracujący w niej od 21 czerwca 2017 roku. Na czwartym miejscu znajduje się obecny mistrz Polski z Piastem Gliwice, czyli Waldemar Fornalik, który pracuje z tą drużyną od 19 września 2017 roku.

3. Badania Empiryczne

Badanie zostanie przeprowadzonych z wykorzystaniem regresji logistycznej oraz testów przedstawionych w rozdziale pierwszym. Zbiór danych potrzebnych do badania został omówiony w rozdziale drugim.

3.1 Przygotowanie danych

Jak już wcześniej pokazano w tabeli 2 w zmiennej *Staż* występuje 12 braków w danych. Jako, że modele regresji logistycznej słabo sobie radzą z obserwacjami brakującymi, podjęta została decyzja o usunięciu tych obserwacji. W wyniku tej operacji liczba obserwacji zmniejszyła się do 132, z czego w 17 przypadkach Y przyjął wartość 1, a zatem doszło do zwolnienia trenera.

3.2 Wybór modelu

Na potrzeby badania zostały wykonane trzy modele. Zmienne do modeli zostały wybrane przy użyciu dwóch metod: metody krokowej wstecznej oraz poprzez znalezienie modelu z najmniejszym kryterium informacyjnym Akaikego. W metodzie krokowej wstecznej wartościami progowymi eliminacji zmiennych było 5% oraz 10% poziomu istotności w teście Walda.

3.2.1 Model ze wszystkimi zmiennymi

Do obu metod na początek potrzebny był model zawierający wszystkie zmienne objaśniające. Zatem stworzono model zawierający piętnaście zmiennych niezależnych. Tabela 4 przedstawia współczynniki regresji, błędy standardowe, wyniki testów Walda oraz wartości p dla tego testu. Istotne w teście Walda z poziomem istotności 5% zostały uznane zmienne *stracone*, *bilans*, *trenerzy10* oraz *staż*. Prawdopodobnie te zmienne znajdują się w ostatecznym modelu, choć nie jest to przesądzone. Kryterium informacyjne Akaikego wyniosło w modelu 78,15. Ta wartość będzie punktem odniesienia dla kolejnych modeli, które będą dążyć do zmniejszenia tej wartości, a co za tym idzie do lepszego dopasowania modelu. Współczynnik determinacji R^2 Negelkerke'a wyniósł 0,6379, zatem 63,79% zmienności zmiennej zależnej jest wyjaśniane przez model, jednak prawdopodobnie jest to wynik zawyżony przez zbyt dużą liczbę zmiennych. Jeśli w modelu finalnym wartość współczynnika determinacji utrzymałaby się na zbliżonym poziomie, można by to uznać za wynik satysfakcjonujący,

Tabela 4. Wyniki regresji logistycznej

Zmienna zależna: zwolnienia trenerów (1 – trener zwolniony, 0 – trener nie zwolniony)				
Zmienne niezależne	β	S. E.	Statystyka Walda	Wartość p
Wyraz wolny	10.1143	5.1472	1.9650	0.0494**
Pkt5	-0.6395	0.7439	-0.8597	0.3900
Pkt3	0.6443	0.8669	0.7433	0.4573
Pole5	0.2708	0.7106	0.3810	0.7032
Gole3	-0.7921	0.8441	-0.9385	0.3480
Stracone5	-0.5419	0.5977	-0.9067	0.3645
Stracone3	0.9245	0.7679	1.2039	0.2286
Stracone	-2.0686	0.9537	-2.1690	0.0301**
Bilans	-1.7227	0.7509	-2.2943	0.0218**
Frekwencja5	-0,0001	0,0004	-0.1331	0.8941
Frekwencja3	-0,0004	0,0004	-0.9813	0.3265
Trenerzy10	-0.7416	0.314	-2.3614	0.0182**
Trenerzy5	0.2409	0.3484	0.6912	0.4894
Różnica3	-0.4995	0.3001	-1.6642	0.0961*
Różnica2	0.2673	0.2298	1.1632	0.2447
Staż	-0.0054	0.0026	-2.0851	0.0371**

Źródło: opracowanie własne. Poziomy istotności: *** - $p < 0,01$; ** - $p < 0,05$; * - $p < 0,1$. Pogrubione zostały współczynniki istotne statystycznie.

3.2.2 Porównanie modeli

W modelu z poziomem istotności 10% w teście Walda za istotne zostały uznane zmienne *gole3*, *stracone*, *bilans*, *frekwencja3*, *trenerzy10*, *różnica3* oraz *staż*. W modelu z poziomem istotności w teście Walda równym 5% zostały uznane za istotne zmienne *gole3*, *bilans*, *trenerzy10*, *różnica3* i *staż*. W ostatnim modelu, który minimalizował kryterium informacyjne Akaikego, znalazły się zmienne *gole3*, *stracone3*, *stracone*, *bilans*, *frekwencja3*, *trenerzy10*, *różnica3*, *różnica2* oraz *staż*. Modele zostały porównane pod względem dopasowania, a także pod względem mocy predykcyjnej. Tabela 5 pokazuje kryterium informacyjne Akaikego, dewiancję, R^2 Nagelkerke'a oraz celność, czułość i specyficzność. Modele charakteryzują się zbliżonymi właściwościami. Wskazuje to na pożądaną w modelach ekonometrycznych małą wrażliwość na reparametryzację. Trzeci model charakteryzuje się jednak najlepszym dopasowaniem. Ma

najniższe AIC oraz dewiancję, a zarazem najwyższy współczynnik determinacji. Jeśli natomiast chodzi o moc predykcyjną, to ma on tylko nieznacznie gorszą trafność, czyli procent trafnie sklasyfikowanych wszystkich obserwacji i specyficzność, zatem procent poprawnie sklasyfikowanych obserwacji negatywnych, od modelu pierwszego i dokładnie taką samą czułość, czyli procent poprawnie sklasyfikowanych przypadków pozytywnych. Mając to wszystko na względzie oraz biorąc pod uwagę, że metoda minimalizacji kryterium informacyjnego Akaikego jest mniej kontrowersyjna, ponieważ nie porównuje modeli budowanych dla różnych długości prób, w dalszej części szczegółowo zostanie przeanalizowany trzeci model.

Tabela 5. Porównanie modeli

	Model z poziomem istotności w teście Walda 10%	Model z poziomem istotności w teście Walda 5%	Model z najmniejszym AIC
AIC	69,364	71,521	68,702
Dewiancja	53,364	59,521	48,702
R ² Nagelkerke'a	0,569	0,507	0,614
ACC	93,18%	91,67%	90,91%
C	58,82%	52,94%	58,82%
S	98,26%	97,39%	95,65%

Źródło: opracowanie własne.

3.3 Charakterystyka modelu

W modelu znalazły się zmienne *gole3*, *stracone3*, *stracone*, *bilans*, *frekwencja3*, *trenerzy10*, *różnica3*, *różnica2* oraz *staż*. Tabela 6 przedstawia wyniki regresji logistycznej. Za istotne na poziomie 5% zostały uznane zmienne *stracone*, *bilans*, *trenerzy10*, *różnica3* oraz *staż*. Na poziomie 10% istotne okazały się zmienne *gole3* i *frekwencja3*, natomiast zmienne *stracone3* oraz *różnica2* w teście Walda zostały uznane za nieistotne. Model przyjął następującą postać:

$$P(Y) = \frac{1}{1 + e^{-\text{logit}(Z)}}$$

gdzie:

$$Z = 6,83 + 3,4 \cdot 0,55 \cdot \text{Gole3} + 3,5 \cdot 1,76 \cdot \text{Stracone} + 3,6 \cdot 1,52 \cdot \text{Bilans} + 3,7 \cdot 0,0003 \cdot \text{Frekwencja3} + 0,35 \cdot \text{Stracone3} + 3,8 \cdot 0,55 \cdot \text{Trenerzy10} + 3,9 \cdot 0,52 \cdot \text{Różnica3} + 0,30 \cdot \text{Różnica2} + 3,10 \cdot 0,005 \cdot \text{Staż}$$

Tabela 6. Wyniki regresji logistycznej

Zmienna zależna: zwolnienia trenerów (1 – trener zwolniony, 0 – trener nie zwolniony)				
Zmienne niezależne	β	S. E.	Statystyka Walda	Wartość p
Wyraz wolny	6,8262	3,0796	2,217	0,0267**
Gole3	-0,5522	0,2823	-1,956	0,0505*
Stracone3	0,3496	0,2171	1,611	0,1072
Stracone	-1,7578	0,7942	-2,213	0,0269**
Bilans	-1,5224	0,6793	-2,241	0,0250**
Frekwencja3	-0,0003	0,0002	-1,946	0,0517*
Trenerzy10	-0,5494	0,2322	-2,366	0,0180**
Różnica3	-0,5153	0,2497	-2,063	0,0391**
Różnica2	0,3033	0,1920	1,579	0,1143
Staż	-0,0046	0,0023	-2,018	0,0436**

Źródło: opracowanie własne. Poziomy istotności: *** - $p < 0,01$; ** - $p < 0,05$; * - $p < 0,1$. Pogrubione zostały współczynniki istotne statystycznie.

3.4 Badanie współliniowości

Tabela 7 ukazuje wartości współczynnika VIF. Wykryta została (dosyć przewidywalna) współliniowość pomiędzy zmiennymi *różnica3* i *różnica2*. W przypadku obu tych zmiennych współczynnik VIF wyniósł więcej niż 10. Podjęta została decyzja o usunięciu z modelu zmiennej *różnica2*, ponieważ test Walda wykazał, że jest ona mniej istotna statystycznie. Również dla zmiennych *stracone* oraz *bilans* współczynnik przyjął wartości na pograniczu, jednak na razie te dwie zmienne zostaną zachowane w modelu i decyzja co do nich zostanie podjęta po ponownym obliczeniu współczynnika VIF.

Tabela 7. Współczynnik VIF

Zmienne	gole3	stracone3	stracone	bilans	frekwencja3
VIF	1.4810	1.7015	9.4302	9.3498	1.8308
Zmienne	trenerzy10	różnica3	różnica2	staż	-
VIF	1.6839	12.6489	12.2242	1.2976	-

Źródło: opracowanie własne.

W tabeli 8 po raz kolejny zostały przedstawione wyniki regresji logistycznej. Po usunięciu zmiennej *różnica2* w modelu pozostała jedna zmienna nieistotna, czyli

stracone3. Natomiast zmienna *staż* jest teraz istotna na najwyższym poziomie, czyli na poziomie 1%.

Tabela 8. Wyniki regresji logistycznej

Zmienna zależna: zwolnienia trenerów (1 – trener zwolniony, 0 – trener nie zwolniony)				
Zmienna niezależna	β	S. E.	Statystyka Walda	Wartość p
Wyraz wolny	6,9888	2,9317	2,384	0,0171**
Gole3	-0,5099	0,2721	-1,874	0,0609*
Stracone3	0,3049	0,2145	1,421	0,1552
Stracone	-1,4064	0,6724	-2,091	0,0365**
Bilans	-1,2993	0,5853	-2,220	0,0264**
Frekwencja3	-0,0003	0,0002	-2,079	0,0377**
Trenerzy10	-0,5427	0,2212	-2,454	0,0141**
Różnica3	-0,1656	0,0783	-2,115	0,0344**
Staż	-0,0054	0,0020	-2,674	0,0075***

Źródło: opracowanie własne. Poziomy istotności: *** - $p < 0,01$; ** - $p < 0,05$; * - $p < 0,1$. Pogrubione zostały współczynniki istotne statystycznie.

Po ponownym obliczeniu współczynnika VIF (tabela 9) wyniki dla zmiennych *stracone* oraz *bilans* obniżyły się, dlatego pozostaną one w modelu. Zatem w modelu wykoprzystane zostaną zmienne *gole3*, *stracone3*, *stracone*, *bilans*, *frekwencja3*, *trenerzy10*, *różnica3* oraz *staż*. Ostateczna postać modelu będzie zatem wyglądać następująco:

$$P(y) = \frac{1}{1 + e^{-\logit(Z)}}$$

gdzie:

$$Z = 6,83 + 3,11 \cdot Gole3 + 0,55 \cdot Stracone3 + 1,76 \cdot Stracone + 3,13 \cdot Bilans + 0,0003 \cdot Frekwencja3 + 0,35 \cdot Trenerzy10 + 0,52 \cdot Różnica3 + 0,005 \cdot Staż$$

Tabela 9. Współczynnik VIF

Zmienna	Gole3	Stracone3	Stracone	Bilans
VIF	1,3920	1,7285	7,3622	7,3761
Zmienna	Frekwencja3	Trenerzy10	Różnica3	Staż
VIF	1,7906	1,6261	1,2362	1,3475

Źródło: opracowanie własne.

3.5 Dopasowanie modelu

W tabeli 10 przedstawiono cztery miary dopasowania modelu: test ilorazu wiarygodności, test Hosmera-Lemeshowa, współczynnik determinacji R^2 Nagelkerke'a oraz kryterium informacyjne Akaikego. Wartość statystyki G dla dewiancji wyniosła 50,163. Wartość krytyczna rozkładu χ^2 dla $\alpha = 0,05$ i ośmiu stopni swobody wynosi 15,507, zatem można odrzucić hipotezę H_0 i uznać model za lepszy od modelu naiwnego. W teście Hosmera-Lemeshowa wartość współczynnika wyniosła 5,349, a zatem mniej od wartości krytycznej, która dokładnie tak samo, jak w przypadku dewiancji wynosi 15,507 dla $\alpha = 0,05$ i ośmiu stopni swobody. Zatem nie ma podstaw do odrzucenia hipotezy zerowej, z czego wynika, że model jest dobrze dopasowany. Współczynnik determinacji R^2 Nagelkerke'a wyniósł 0,590, a zatem 59% zmienności zmiennej zależnej jest wyjaśniane przez zmienne objaśniające. Jest to dosyć wysoki wynik, zważywszy na nieprzewidywalność sportu oraz złożoność relacji międzyludzkich, które również mają duży wpływ na to, czy trener zachowa swoją pracę. Kryterium informacyjne Akaikego po usunięciu z modelu zmiennej *różnica2* wynosi 69,233, jednak nadal jest ono niższe niż w pozostałych modelach, które były opisane wcześniej.

Tabela 10. Miary dopasowania modelu

	-2log(L)	Chi-kwadrat	Test Hosmera-Lemeshowa	R^2 Nagelkerke'a	AIC
Wartość współczynnika	51,233	50,163	5,349	0,590	69,233
Poziom istotności - p	-	0,000	0,720	-	-
Stopnie swobody	-	8	8	-	-

Źródło: opracowanie własne.

3.6 Interpretacja modelu

Interpretacja modelu zostanie dokonane na podstawie ilorazu szans współczynników finalnego modelu. W tabeli 11 przedstawiono współczynniki regresji, błędy standardowe, statystyki testów Walda, wartości p dla testów Walda, ilorazy szans oraz ich przedziały ufności. Dla zmiennej *gole3*, która okazała się istotna w teście Walda na poziomie 10%, iloraz szansy wyniósł 0,6005. Ponieważ jest to mniej niż 1, zmienna ta działa ograniczająco na zmienną zależną. Przy każdym голу więcej w ciągu trzech ostatnich meczów, szansa zwolnienia trenera spada o około 40%. Nie jest to wynik zaskakujący, ponieważ zazwyczaj, jeśli drużyna zdobywa dużo bramek, to nie zwalnia się trenera.

Przedziały ufności tej zmiennej wyniosły od 0,3303 do 0,9742, zatem z 95% prawdopodobieństwem iloraz szans tej zmiennej będzie się mieścić w tym przedziale. Zmienna *stracone3* jako jedyna została uznana za nieistotną w teście Walda, zatem jej wyników nie można brać pod uwagę, jeśli jednak przyjrzeć się jej ilorazowi szansy, to można zauważyć, że również jako jedyna działa ona stymulująco na zmienną objaśnianą. Z każdym straconym golem w trzech ostatnich meczach, szansa zwolnienia trenera rośnie o około 35%. Podobnie jak w przypadku zmiennej *gole3* jest to wynik spodziewany, ponieważ drużyny, które tracą wiele bramek dosyć często zmieniają trenerów. Widać, że ta zmienna nie daje wiarygodnych wyników, ponieważ według przedziałów ufności jej iloraz szans może spadać poniżej 1, zatem w niektórych przypadkach działać ograniczająco. Zmienna *stracone* została uznana za istotną w teście Walda na poziomie istotności 5%. Wynik jej ilorazu szans jest jednak bardzo zaskakujący i trudny do wyjaśnienia. Wyniósł on 0,2450, zatem wynika z niego, że po każdym голу straconym w meczu, szansa na zwolnienie trenera po tym meczu spada o około 75%. Potwierdzają to przedziały ufności, które nie wychodzą powyżej 1, zatem z 95% prawdopodobieństwem ta zmienna działa destymulująco. Wydaje się, że duża liczba bramek straconych w meczu przybliży szkoleniowca do straty posady, jednak wyniki modelu pokazują inaczej. Zmienne *bilans*, *frekwencja3*, *trenerzy10* oraz *różnica3* oceniono jako istotne w teście Walda z poziomem istotności 5%. Iloraz szans *bilansu* wyniósł 0,2727, zatem przy każdej poprawie różnicy bramek w ostatnim meczu o jednego gola, szanse na zwolnienie trenera spadają o około 73%. Iloraz szans dla zmiennej *frekwencja3* jest bardzo bliski 1, jednak jak wynika z przedziałów ufności, z 95% prawdopodobieństwem nigdy jej nie przekroczy. Wynik 0,9997 wykazuje, że jeśli średnia kibiców w ostatnich trzech meczach na własnym stadionie wzrośnie o jednego kibica w porównaniu ze średnią z całego sezonu, to szansa na zwolnienie trenera spadnie o 0,03%. Efekt tego predyktora wydaje się minimalny, jednak liczba kibiców na meczach drużyn Ekstraklasowych, może zmieniać się o kilka tysięcy, stąd ważna jest efektywna gra i przyciąganie jak największej liczby kibiców. Wynik zmiennej *trenerzy10* jest dość niespodziewany, ponieważ iloraz szans równy 0,5812 świadczy, że dla każdego trenera więcej w ostatnich dziesięciu sezonach szansa na zmianę trenera maleje o około 42%. Wydawać by się mogło, że kluby często zmieniające trenerów z większym prawdopodobieństwem kolejny raz zmienią szkoleniowca, jednak dlatego, że w polskiej lidze większość drużyn ustawicznie zwalnia kolejnych trenerów to po takiej serii zwolnień, następny trener dostaje więcej czasu na popracowanie. Iloraz szans zmiennej *różnica3* wyniósł 0,8474, zatem jeśli liczba punktów jest większa o jeden od średniej punktów

zdobytych w ostatnich trzech sezonach w analogicznej liczbie kolejek to szanse na zwolnienie trenera maleją o około 15%. Jest to logiczny wynik, ponieważ jeśli drużyna punktuje lepiej niż w poprzednich sezonach, a zatem gra powyżej oczekiwań, to stabilność posady trenera wzrasta. Ostatnia zmienna, czyli *staż* jako jedyna została uznana za istotną na poziomie 1% w teście Walda. Iloraz szans tej zmiennej wyniósł 0,9946, zatem każdy dzień więcej średniej długości pracy w poprzednich klubach zmniejsza szansę zwolnienia trenera o około 0,5%. Ciekawe, że ta zmienna została uznana za najbardziej istotną. Wynika z niej, że trenerzy w Polsce dzielą się na takich którzy pracują w klubach dłużej i takich którzy są zatrudniani na chwilę, aby poprawić wyniki, a później dosyć szybko zostają zwolnieni. Wpływ na wyniki mogło mieć usunięcie części obserwacji, w których zatrudnieni byli nowicjusze.

Tabela 11. Wyniki regresji logistycznej wraz z ilorazami szans

Zmienna zależna: zwolnienia trenerów (1 – trener zwolniony, 0 – trener nie zwolniony)							
Zmienne niezależne	β	S. E.	Statystyka Walda	Wartość p	Iloraz szans	Przedziały ufności (95% CI)	
						Dolny	Górny
Wyraz wolny	6,9888	2,9317	2,384	0,0171**	1084,4711	5,4592	638394,6737
Gole3	-0,5099	0,2721	-1,874	0,0609*	0,6005	0,3303	0,9742
Stracone3	0,3049	0,2145	1,421	0,1552	1,3565	0,9017	2,1352
Stracone	-1,4064	0,6724	-2,091	0,0365**	0,2450	0,0557	0,8258
Bilans	-1,2993	0,5853	-2,220	0,0264**	0,2727	0,0727	0,7554
Frekwencja3	-0,0003	0,0002	-2,079	0,0377**	0,9997	0,9993	1,0000
Trenerzy10	-0,5427	0,2212	-2,454	0,0141**	0,5812	0,3551	0,8547
Różnica3	-0,1656	0,0783	-2,115	0,0344**	0,8474	0,7131	0,9776
Staż	-0,0054	0,0020	-2,674	0,0075***	0,9946	0,9893	0,9979

Źródło: opracowanie własne. Poziomy istotności: *** - $p < 0,01$; ** - $p < 0,05$; * - $p < 0,1$. Pogrubione zostały współczynniki istotne statystycznie.

3.7 Moc predykcyjna modelu

W tabeli 12 zostały przedstawione zdolności klasyfikacyjne modelu. Wykorzystując ją można obliczyć, że celność modelu (ACC), czyli procent poprawnie sklasyfikowanych wszystkich obserwacji wynosi 90,91% $((111 + 9)/132 * 100\%)$ przy punkcie odcięcia 0,5. Czulość, czyli procent poprawnie sklasyfikowanych obserwacji pozytywnych, w modelu wynosi 52,94% $(9/17 * 100\%)$, natomiast specyficzność, czyli procent poprawnie sklasyfikowanych obserwacji negatywnych

wynosi 96,52% ($111/115 * 100\%$). Model zdecydowanie lepiej radzi sobie z sytuacjami kiedy do zwolnienia nie dochodziło, jednak wpływ na to może mieć nierównomierne rozłożenie obserwacji, ponieważ sytuacji w których do zwolnienia nie dochodziło jest o wiele więcej.

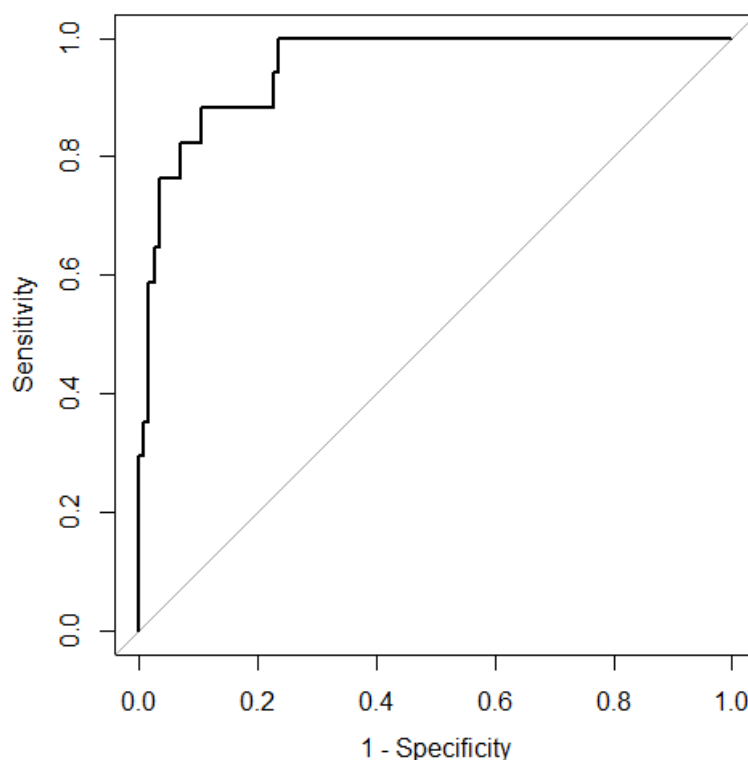
Tabela 12. Tabela klasyfikacyjna

Obserwacje empiryczne	Obserwacje przewidywane		Suma
	0 (nie zwolniony)	1 (zwolniony)	
0 (nie zwolniony)	111	4	115
1 (zwolniony)	8	9	17
Suma	119	13	132

Źródło: opracowanie własne.

Na wykresie 2 zaprezentowano krzywą ROC. Na osi x znajduje się 1 – specyficzność, natomiast na osi y znajduje się czułość. Pole pod wykresem (AUC) wyniosło 0,939. Jest to wynik bliski jeden, zatem można stwierdzić, że model wysmienie klasyfikuje przypadki.

Wykres 2. Krzywa ROC uzyskana w pakiecie R



Źródło: opracowanie własne.

Tabela 13 przedstawia czułość (C) oraz 1 – specyficzność (1 – S) dla wybranych punktów odcięcia. Na jej podstawie oraz na podstawie wykresu ROC, można określić punkt odcięcia dla którego poprawi się pozytywna moc predykcyjna modelu, która była zdecydowanie gorsza od specyficzności. Należy znaleźć taki punkt dla którego C będzie jak najwyższe, natomiast 1 – S jak najniższe, czyli punkt jak najbliższy na wykresie miejscu (0, 1). Najbliżej tego jest dla punktu odcięcia 0,241, dla którego czułość wyniosłaby 0,882, natomiast 1 – S wyniosłoby 0,078. Zatem przyjmując ten punkt odcięcia zdecydowanie podniosłaby się moc predykcyjna modelu.

Tabela 13. Punkty odcięcia prawdopodobieństwa na krzywej ROC

Punkty odcięcia	0	0.034	0,038	0.043	0,044	0.241	0.279	0.347	0.489	0.689	1
C	1	1.000	0,941	0.941	0.882	0.882	0.824	0.824	0.588	0.294	0
1 - S	1	0.304	0,304	0.287	0.287	0.078	0.078	0.061	0.035	0.009	0

Źródło: opracowanie własne.

Podsumowanie

W pracy podzielonej na trzy części, w pierwszym rozdziale zostały omówione podstawowe założenia modeli regresji logistycznej oraz testy przydatne przy ich weryfikacji. W drugim rozdziale został przedstawiony przygotowany zbiór danych oraz została nakreślona sytuacja trenerów w polskiej Ekstraklasie. W ostatnim rozdziale na podstawie przygotowanej bazy danych oraz z wykorzystaniem regresji logistycznej zostało przeprowadzone badanie empiryczne.

Celem niniejszej pracy było znalezienie czynników istotnie wpływających na zwolnienia trenerów w piłce nożnej. Z analizy przeprowadzonej na podstawie rozgrywek polskiej Ekstraklasy wynika, że za istotne można uznać czynniki takie jak liczba goli strzelonych w trzech meczach poprzedzających zwolnienie, liczba bramek straconych oraz bilans bramek w ostatnim meczu, liczba kibiców w trzech meczach przed decyzją o zmianie trenera, liczba trenerów jaka pracowała w klubie w poprzednich 10 sezonach, liczba zdobytych punktów w badanym sezonie w porównaniu z trzema poprzednimi sezonami oraz średni staż trenera w jego poprzednich klubach. Zmienna charakteryzująca liczbę bramek straconych w badanej kolejce oraz dwóch wcześniejszych została uwzględniona w modelu, pomimo że w teście Walda została uznana za nieistotną.

Zatem z uzyskanych wyników można wnioskować, że postawa drużyny w ostatnich meczach ma istotny wpływ na decyzję o zwolnieniu szkoleniowca, jednak tylko w bardzo krótkim okresie, czyli w maksymalnie trzech meczach poprzedzających zwolnienie. Potwierdza to opinie o polskich klubach, które decyzję o zmianie trenera podejmują pod wpływem impulsu. Widać, że wpływ na to mają również kibice, których mniejsza liczba na stadionie może zwiastować zmianę trenera. W pracy badane również było czy istotny wpływ na rozszady na stanowisku trenera mają przyzwyczajenia zarówno klubów, jak i menedżerów. Wyniki co do nawyku klubów do częstych zmian trenerów są dość nieoczywiste, ponieważ, pomimo że zmienna mówiąca o liczbie trenerów zatrudnionych w klubie w ostatnich 10 sezonach, w teście Walda okazała się istotna, to wartość jej ilorazu szans jest zaskakująca, ponieważ wynika z niej, że im częściej klub zmieniał trenera w ostatnich latach, tym mniejsza szansa, że obecny szkoleniowiec zostanie zwolniony. W przypadku trenerów zarysował się podział na takich, którzy jeśli już podejmują się pracy, to zazwyczaj pracują w danym klubie przez dłuższy okres oraz na takich, którzy nie zagrzewają zbyt długo miejsca w jednym klubie i jeśli dostaną pracę to można się spodziewać, że niedługo z niej zrezygnują lub zostaną zwolnieni. Zmienna dotycząca

stażu jako jedyna w teście Walda została uznana za istotną na poziomie 1%, zatem okazuje się, że w polskiej lidze zatrudniając trenera, należy dużą wagę przywiązywać do czasu, jaki ten szkoleniowiec średnio spędzał w poprzednich klubach.

Bibliografia

1. Agresti A., *Categorical data analysis*, John Wiley & Sons, New Jersey 2013
2. Allison P. D., *Measures of Fit for Logistic Regression*, Statistical Horizons LLC and the University of Pennsylvania, 2014.
3. Błażyński L., *Władze Piasta rozczarowane. Fornalik może zostać zwolniony*, Przegląd sportowy, <https://www.przegladsportowy.pl/pilka-nozna/pko-ekstraklasa/piast-gliwice/waldemar-fornalik-moze-zostac-zwolniony-z-piast-gliwice/4rv2x73> [dostęp: 29.04.2020].
4. Chlebus M., *Pomiar ryzyka rynkowego za pomocą miary Value at Risk – podejście dwuetapowe*, Rozprawa doktorska, Uniwersytet Warszawski, 2014.
5. Cohen J., Cohen P., West S. G., Aiken L. S., *Applied Multiple Regression/Correlation Analysis for the Behavioral Sciences*, Lawrence Erlbaum Associates, Publishers, Mahway 2003.
6. Hosmer D., Lemeshow S., *Applied Logistic Regression*, John Wiley & Sons, New York 2000.
7. Kirpsza A., *Zastosowanie regresji logistycznej w studiach nad Unią Europejską*, [w:] *Metody jakościowe i ilościowe w badaniu organizacji i działania Unii Europejskiej*, pod red. K. Ławniczaka, Warszawa 2013.
8. Menard S., *Applied Logistic Regression Analysis*, Sage University Paper Series on Quantitative Applications in the Social Sciences 07-106, Thousand Oaks 2002.
9. Sączewska-Piotrowska A., *Zastosowanie krzywych ROC w analizie ubóstwa miejskich i wiejskich gospodarstw domowych*, Przegląd statystyczny r. LXIII, zeszyt 2, Katowice 2016.
10. Yan Xin, *Linear Regression Analysis: Theory and Computing*, World Scientific, Singapur 2009.
11. <https://www.przegladsportowy.pl/pilka-nozna/pko-ekstraklasa/piast-gliwice/mateusz-borek-zdradza-kulisy-pracy-waldemara-fornalika-w-piast-gliwice-misja-futbol/tnhxkmb> [dostęp: 29.04.2020].

Spis tabel

Tabela 1. Przykładowa tabela klasyfikacyjna.....	10
Tabela 2. Statystyki opisowe	16
Tabela 3. Liczba zmian trenerów w wybranych ligach europejskich w sezonach 2015/16-2018/19	17
Tabela 4. Wyniki regresji logistycznej	20
Tabela 5. Porównanie modeli	21
Tabela 6. Wyniki regresji logistycznej	22
Tabela 7. Współczynnik VIF	22
Tabela 8. Wyniki regresji logistycznej	23
Tabela 9. Współczynnik VIF	23
Tabela 10. Miary dopasowania modelu.....	24
Tabela 11. Wyniki regresji logistycznej wraz z ilorazami szans	26
Tabela 12. Tabela klasyfikacyjna.....	27
Tabela 13. Punkty odcięcia prawdopodobieństwa na krzywej ROC	28