

Marek Biernacki*, Wiktor Ejsmont**

EFEKTYWNOŚĆ KSZTAŁCENIA WE WROCŁAWSKICH LICEACH

Streszczenie: Celem artykułu jest porównanie efektywności nauczania we wrocławskich liceach. Jedną z metod oceny oparta jest na pionierskim artykule M. Aitkina i N. Longforda „Statistical Modelling Issues in School Effectiveness Studies”. Modele wykorzystane przez powyższych autorów służą do analizy danych panelowych. Dokonamy też pomiaru efektywności kształcenia tych szkół za pomocą miernika, w którym wykorzystuje się współczynnik Giniego. W polskiej literaturze pojęcie efektywności kształcenia jest również nazywane dodaną wartością edukacji.

Nasze badania oparte zostały na wynikach gimnazjalnych i maturalnych uczniów wrocławskich liceów z rozbiciem na część humanistyczną i ścisłą. Spróbowaliśmy też znaleźć zależności między charakterystykami związanymi ze szkołą (staż pracy nauczycieli) i efektywnością kształcenia.

1. WPROWADZENIE

Celem artykułu jest porównanie efektywności nauczania we wrocławskich liceach. Jedną z metod oceny oparta jest na pionierskim artykule Aitkina i Longforda [1981] *Statistical Modeling Issues in School Effectiveness Studies*. Modele wykorzystane przez powyższych autorów służą do analizy danych panelowych. Dokonamy też pomiaru efektywności kształcenia tych szkół za pomocą miernika, w którym wykorzystuje się współczynnik Giniego. W polskiej literaturze pojęcie efektywności kształcenia jest również nazywane dodaną wartością edukacji.

Nasze badania oparte zostały na wynikach gimnazjalnych i maturalnych uczniów wrocławskich liceów z rozbiciem na część humanistyczną i ścisłą. Spróbowaliśmy też znaleźć zależności między charakterystykami związanymi ze szkołą (staż pracy nauczycieli) i efektywnością kształcenia.

2. OPIS DANYCH

Zebrane dane opisują wyniki gimnazjalne oraz wyniki maturalne absolwentów wrocławskich liceów z lat 2007-2009. Dane zostały podzielone na dwie części: pierwsza część dotyczy wyników z humanistycznego egzaminu gimnazjalnego oraz z matury podstawowej z języka polskiego, natomiast druga część to punkty z egzaminu gimnazjalnego z części matematyczno-przyrodniczej oraz z matury rozszerzonej z matematyki. Matura podstawowa z języka polskiego jest obowiązkowa, zaś matura z matematyki na poziomie rozszerzonym w roku 2009 miała charakter fakultatywny. W konsekwencji

* Dr, Katedra Matematyki i Cybernetyki, Uniwersytet Ekonomiczny we Wrocławiu.

** Mgr, Katedra Matematyki i Cybernetyki, Uniwersytet Ekonomiczny we Wrocławiu.

liczba analizowanych liceów z danymi humanistycznymi jest dwukrotnie większa od liczby liceów, w których można wybrać liczną próbę uczniów zdających na maturze

matematykę rozszerzoną. Każdy uczeń, który pisze maturę rozszerzoną przystępuje wcześniej do egzaminu podstawowego z tego przedmiotu. Liczba uczniów przystępujących tylko do matury podstawowej z matematyki w latach 2007-2009 we wrocławskich liceach jest prawie trzykrotnie mniejsza od tych uczniów, którzy zdawali maturę zarówno na poziomie rozszerzonym jak i podstawowym.

Oznaczenia:

- LO + liczba – Liceum Ogólnokształcące, którego numer jest wyznaczany przez tę liczbę;
- DZDZ – Dolnośląski Zakład Doskonalenia Zawodowego (Zespół Szkół);
- LO S – Prywatne Salezjańskie Liceum Ogólnokształcące;
- LO SU – Liceum Ogólnokształcące Sióstr Urszulanek.

Dwie ostatnie kolumny tabel 1 i 2 dotyczą danych przeskalowanych, których istota zostanie wyjaśniona w drugim podrozdziale.

Tab. 1. Zestawienie uczniów zdających maturę z języka polskiego w latach 2007-2009

Szkoła	Liczba uczniów	Wejście - średni wynik gimnazjalny	Odchylenie standardowe gimnazjum	Wyjście - średni wynik maturalny	Odchylenie standardowe matura	Średni wynik matury skalowanej	Odchylenie standardowe matury skalowanej
LO DZDZ	103	62,72	13,54	43,71	13,63	57,07	12,46
LO I	423	72,74	10,24	58,89	11,83	69,76	10,56
LO II	671	75,90	9,45	59,57	11,74	72,78	9,59
LO III	258	86,79	6,88	69,31	10,22	84,88	7,47
LO IV	489	74,35	11,06	60,73	12,75	71,88	11,33
LO IX	571	79,53	9,39	63,80	11,02	77,08	9,61
LO S	70	72,57	10,64	53,97	13,60	67,77	11,65
LO SU	111	77,50	10,37	63,85	11,28	75,04	11,02
LO V	264	83,29	7,76	70,06	12,05	81,94	8,47
LO VI	446	71,79	10,84	56,98	11,12	68,56	10,34
LO VII	790	82,91	7,97	65,77	11,11	80,49	8,58
LO VIII	482	79,32	9,37	61,54	11,10	76,58	9,15
LO X	592	75,69	10,37	60,33	11,16	72,95	10,24
LO XI	341	72,30	10,49	62,20	12,42	70,50	10,46
LO XII	554	83,45	9,09	64,93	10,83	81,09	9,39
LO XIII	542	78,38	10,02	61,99	11,18	75,66	10,12
LO XIV	363	84,64	8,48	68,46	10,64	82,96	8,69
LO XV	774	73,51	9,82	58,80	10,99	70,56	9,57
LO XVI	199	64,81	11,75	48,38	12,35	59,67	12,18
LO XVII	315	73,18	9,82	57,80	12,53	70,10	9,76
LO XXI	119	67,18	11,67	58,34	11,42	65,11	11,39
LO XXIV	238	71,31	11,82	56,76	13,15	68,05	11,96
LO XXIX	118	67,68	10,60	48,53	10,70	62,38	10,29
LO XXX	125	66,38	14,00	52,08	13,30	62,85	13,91
SUMA	8958	76,66	11,21	61,08	12,49	73,91	11,58

Źródło: Okręgowa Komisja Egzaminacyjna we Wrocławiu [2009].

W tabeli 1 przedstawiono ogólne ujęcie wyników gimnazjalnych oraz maturalnych z części humanistycznej. Najmniejszym odchyleniem standardowym zarówno na poziomie gimnazjalnym (przyjęcia) jak i maturalnym charakteryzuje się LO III. Świadczy to o bardzo równym poziomie kształcenia w tym liceum. Na poziomie gimnazjum (przyję-

cia) drugie miejsce pod względem „rozproszenia” (odchylenie standardowe) zajmuje LO V, ale pod względem wyników maturalnych plasuje się dopiero na 18 pozycji. Należy wziąć pod uwagę fakt, iż odchylenia standardowe dla poszczególnych szkół nie różnią się istotnie na poziomie wyników maturalnych.

Tab. 2. Zestawienie wyników uczniów zdających maturę rozszerzoną z matematyki w latach 2007-2009

Szkoła	Liczba uczniów	Wejście średni wynik gimnazjalny	Odchylenie standardowe gimnazjum	Wyjście - średni wynik maturalny	Odchylenie standardowe matura	Średni wynik matury skalowanej	Odchylenie standardowe matury skalowanej
LOI	119	72,30	11,71	38,99	19,18	49,51	20,54
LOII	233	80,13	8,79	44,95	22,12	56,84	22,30
LOIII	398	91,75	7,06	73,97	18,73	84,78	14,15
LOIV	105	77,10	10,68	48,91	21,87	59,66	20,50
LOIX	271	80,30	10,64	54,82	21,71	65,36	19,77
LOV	141	82,87	8,70	54,34	22,03	66,13	20,59
LOVI	81	75,65	11,88	38,22	20,45	49,47	22,91
LOVII	384	85,40	8,01	59,95	21,21	71,78	17,96
LOVIII	229	80,12	9,18	50,76	17,23	63,32	15,78
LOX	193	79,55	10,73	51,58	23,18	62,09	21,89
LOXII	321	83,93	8,13	60,26	19,61	71,80	16,18
LOXIII	258	79,37	9,88	51,56	20,81	62,69	18,79
LOXIV	284	88,06	9,05	71,13	20,93	80,75	16,69
LOXV	183	75,51	11,41	35,36	20,05	46,52	22,48
LOXVII	69	74,78	11,26	48,93	19,23	58,31	18,10
SUMA /ŚREDNIA	3269	82,40	10,68	55,86	23,28	66,97	21,55

Źródło: Okręgowa Komisja Egzaminacyjna we Wrocławiu [2009].

W Tabeli 2 pokazano wyniki gimnazjalne oraz maturalne z przedmiotów ścisłych. Zauważmy, że średni wynik gimnazjalny z części matematyczno-przyrodniczej wypada trochę lepiej niż średni wynik z części humanistycznej. Na poziomie wszystkich uczniów odchylenie standardowe wyników maturalnych wynosi 23,28, które świadczy o istotnym zróżnicowaniu wyników, zarówno wewnątrz szkół jak i na poziomie wszystkich uczniów.

3. METODY BADAŃ

3.1. Model VC

Oznaczenia:

- x_{ij} liczba punktów gimnazjalnych i -tego ucznia w j -tej szkole (wejście),
- y_{ij} liczba punktów maturalnych uzyskanych przez i - tego ucznia w j -tej szkole (wyjście),
- n_j liczba uczniów w szkole j -tej, $j \in \{1, \dots, k\}$,
- n liczba wszystkich uczniów tzn. $n = n_1 + \dots + n_k$,
- k liczba szkół,
- $\bar{x}$ średni wynik gimnazjalny badanych uczniów,

- $\bar{y}$ średni wynik maturalny badanych uczniów,
- $\bar{x}_j, \bar{y}_j$ średni wynik uczniów odpowiednio z gimnazjum oraz z matury j - tej szkoły.

Zastosowany model mierzy przyrost (spadek) wiedzy opierając się przede wszystkim na różnicy pomiędzy wynikiem z matury a wynikiem z gimnazjum wzór (16). Pojawia się jednak problem, który ilustruje poniższy przykład. Jeżeli dany uczeń uzyskał na teście gimnazjalnym 30, a na maturalnym 40 punktów, to jego przyrost „absolutny” wiedzy jest taki sam jak u ucznia, który w gimnazjum uzyskał wynik 90 zaś na maturze 100 punktów. Przyrost bezwzględny punktów w obu przypadkach jest równy 10; nie można jednak mówić o identycznym przyroście wiedzy uczniów. Z tego powodu dane zostały przeskalowane tak, by uwzględniły początkową wiedzę danego ucznia (wynik egzaminu gimnazjalnego), oraz dowartościowały tych uczniów, którzy mieli lepsze wyniki z matur. To samo rozumowanie dotyczy spadku punktów w stosunku do wyniku gimnazjalnego. Zaprezentowane rozumowanie jest powodem przeskalowania danych, które przedstawiono poniżej:

$$y'_{ij} = \begin{cases} x_{ij} + (y_{ij} - x_{ij})y_{ij}/100 & \text{o ile } (y_{ij} - x_{ij}) \geq 0 \\ x_{ij} + (y_{ij} - x_{ij})(100 - y_{ij})/100 & \text{o ile } (y_{ij} - x_{ij}) < 0 \end{cases} \quad (1)$$

Dwie ostatnie kolumny tabel 1 oraz 2 przedstawiają średnie wyniki maturalne po przeskalowaniu. Zauważmy, że średnia maturalna wyników przeskalowanych jest wyższa od średniej dla wyników maturalnych rzeczywistych. Powodem takiego zachowania się przeskalowanych wyników maturalnych jest to, że zdecydowana większość rozważanych uczniów uzyskała mniej punktów na egzaminie maturalnym niż na gimnazjalnym.

Dalsza część tego podrozdziału opiera się na dopasowaniu modelu do punktów przekształconych tzn. postaci. (x_{ij}, y'_{ij}) . Oznaczenia średnich dotyczące y'_{ij} są analogiczne do oznaczeń średnich y_{ij} (na poziomie całej populacji jak i szkoły).

Model, który zastosowano jest modelem z czynnikami losowymi, który znany jest też pod nazwą modelu komponentów wariancyjnych (*variance components model* – VC lub również *error component model*), ponieważ dane na podstawie których chcemy zbudować model są próbką reprezentatywną, oraz niejednorodną¹. Model ten jest postaci:

$$y'_{ij} = \alpha + \beta x_{ij} + \xi_j + e_{ij}, \quad (2)$$

gdzie:

- e_{ij} zmienna losowa z rozkładu $N(0, \sigma^2)$,
- ξ_j zmienna losowa z rozkładu $N(0, \sigma_j^2)$ - interpretacja tego składnika modelu jest taka, że każda szkołę możemy traktować jak zmienną losową z wariancją σ_j^2 oraz średnią α ,
- zakładamy, że składniki losowe pochodzące z różnych szkół i dla różnych uczniów są nieskorelowane,

¹ Por. P. Balestra, M. Nerlove M., [1966], *Pooling Cross Section and Time Series Data in the Estimation of a Dynamic Model: The Demand for Natural Gas*, *Econometrica*, Vol. 34, No. 3, s. 585-612.

– ponadto zakładamy, że indywidualny składnik losowy ξ_j jest nieskorelowany z składnikiem losowym e_{ij} (tzn. $E(\xi_j, e_{is}) = 0$).

Z powyższych założeń otrzymujemy:

$$\begin{aligned} \text{var}(y'_{ij}) &= \text{var}(\xi_j + e_{ij}) = E(\xi_j + e_{ij})^2 - E^2(\xi_j + e_{ij}) \\ &= E(\xi_j^2 + 2\xi_j e_{ij} + e_{ij}^2) = \sigma_I^2 + \sigma^2, \end{aligned} \quad (3)$$

$$\begin{aligned} \text{cov}(y'_{ij}, y'_{pj}) &= \text{cov}((\xi_j + e_{ij}), (\xi_j + e_{pj})) \\ &= E(\xi_j^2 + \xi_j e_{ij} + \xi_j e_{pj} + e_{ij} e_{pj}) = \sigma_I^2, \end{aligned} \quad (4)$$

$$\rho = \text{cor}(y'_{ij}, y'_{pj}) = \frac{\sigma_I^2}{\sigma_I^2 + \sigma^2}. \quad (5)$$

Współczynniki tak określonego modelu szacujemy za pomocą największej wiarygodności (np. Atkins i Longford [1986]) lub uogólnioną metodą najmniejszych kwadratów (np. Baltagi [1995]). Estymator parametrów α oraz β jest postaci:

$$\begin{bmatrix} \hat{\alpha} \\ \hat{\beta} \end{bmatrix} = (\mathbf{X}^T \mathbf{V}^{-1} \mathbf{X})^{-1} \mathbf{X}^T \mathbf{V}^{-1} \mathbf{y}, \quad (6)$$

gdzie:

$$\begin{aligned} \mathbf{X}^T &= \begin{bmatrix} 1 & , \dots, & 1 & , \dots, & 1 & , \dots, & 1 \\ x_{1,1} & , \dots, & 1 & , \dots, & x_{1,k} & , \dots, & x_{k,n_k} \end{bmatrix}, \\ \mathbf{y}^T &= [y'_{1,1} \quad , \dots, \quad y'_{n_1} \quad , \dots, \quad y'_{k,1} \quad y'_{k,n_k}], \end{aligned}$$

oraz $\mathbf{V}$ jest macierzą kowariancji wektora $\mathbf{y}$. Zgodnie z założeniami do modelu macierz ta jest postaci:

$$\mathbf{V} = \begin{bmatrix} \Omega_1 & 0 & \dots & 0 \\ 0 & \Omega_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \Omega_k \end{bmatrix},$$

gdzie:

$$\Omega_j = \begin{bmatrix} \sigma^2 + \sigma_I^2 & \sigma_I^2 & \dots & \sigma_I^2 \\ \sigma_I^2 & \sigma^2 + \sigma_I^2 & \dots & \sigma_I^2 \\ \vdots & \vdots & \ddots & \vdots \\ \sigma_I^2 & \sigma_I^2 & \dots & \sigma^2 + \sigma_I^2 \end{bmatrix}_{n_j \times n_j}.$$

Po uproszczeniach wzoru (6) otrzymujemy (np. zastosowanych przez Atkinsa oraz Longforda [1986]):

$$\begin{bmatrix} \hat{\alpha} \\ \hat{\beta} \end{bmatrix} = \begin{bmatrix} \sum_{j=1}^k w_j & \sum_{j=1}^k w_j \bar{x}_j \\ \sum_{j=1}^k w_j \bar{x}_j & \sum_{j=1}^k \sum_{i=1}^{n_j} (x_{ij} - \bar{x}_j)^2 + \sum_{j=1}^k w_j \bar{x}_j^2 \end{bmatrix}^{-1} \begin{bmatrix} \sum_{j=1}^k w_j \bar{y}'_j \\ \sum_{j=1}^k \sum_{i=1}^{n_j} (y'_{ij} - \bar{y}'_j)(x_{ij} - \bar{x}_j) + \sum_{j=1}^k w_j \bar{x}_j \bar{y}'_j \end{bmatrix},$$

$$\hat{\beta} = \frac{\sum_{j=1}^k \sum_{i=1}^{n_j} (y'_{ij} - \bar{y}'_j)(x_{ij} - \bar{x}_j) + A_{xy}^*}{\sum_{j=1}^k \sum_{i=1}^{n_j} (x_{ij} - \bar{x}_j)^2 + A_{xx}^*}, \quad (7)$$

gdzie:

$$A_{xx}^* = \sum_{j=1}^k w_j (\bar{x}_j - \bar{x}^*)^2; \quad A_{xy}^* = \sum_{j=1}^k w_j (\bar{x}_j - \bar{x}^*)(\bar{y}'_j - \bar{y}^{*'}), \quad (8)$$

$$\bar{x}^* = \sum_{j=1}^k w_j \bar{x}_j / \sum_{j=1}^k w_j; \quad \bar{y}^{*'} = \sum_{j=1}^k w_j \bar{y}'_j / \sum_{j=1}^k w_j, \quad (9)$$

oraz

$$w_j = n_j \sigma^2 / (\sigma^2 + n_j \sigma_l^2), \quad (10)$$

$$\text{var}(\hat{\beta}) = \frac{\sigma^2 + \sigma_l^2}{\sum_{j=1}^k \sum_{i=1}^{n_j} (x_{ij} - \bar{x}_j)^2 + A_{xx}^*}. \quad (11)$$

Ze wzoru (7) oraz po przekształceniach otrzymamy:

1. Dla $\sigma_l^2 = 0$ ($w_j = n_j$):

$$\hat{\beta}_1 = \frac{\sum_{j=1}^k \sum_{i=1}^{n_j} (y'_{ij} - \bar{y}'_j)(x_{ij} - \bar{x}_j) + \sum_{j=1}^k \sum_{i=1}^{n_j} (\bar{x}_j - \bar{x})(\bar{y}'_j - \bar{y}')}{\sum_{j=1}^k \sum_{i=1}^{n_j} (x_{ij} - \bar{x}_j)^2 + \sum_{j=1}^k \sum_{i=1}^{n_j} (\bar{x}_j - \bar{x})^2} = \frac{\sum_{j=1}^k \sum_{i=1}^{n_j} (x_{ij} - \bar{x})(y'_{ij} - \bar{y}')}{\sum_{j=1}^k \sum_{i=1}^{n_j} (x_{ij} - \bar{x})^2}. \quad (12)$$

Otrzymamy współczynnik nachylenia taki sam jak za pomocą prostej regresji liniowej dopasowany do wszystkich danych (bez rozróżniania ze względu na szkołę). Jeżeli y'_{ij} nie są skorelowane na poziomie danej szkoły, wówczas model ten spełnia założenia klasycznej regresji liniowej.

2. Jeżeli σ_l^2 jest duże w stosunku do σ^2 tzn. $\sigma^2 / \sigma_l^2 \rightarrow 0$ wówczas w_j zbiegają do zera istotnie $w_j = n_j (\sigma / \sigma_l)^2 / ((\sigma / \sigma_l)^2 + n_j) \xrightarrow{(\sigma / \sigma_l)^2 \rightarrow 0} 0$. To zaś powoduje, że wielkości A_{xx}^* oraz A_{xy}^* zbiegają do zera. Ze wzoru (7) otrzymujemy:

$$\hat{\beta}_2 = \frac{\sum_{j=1}^k \sum_{i=1}^{n_j} (y'_{ij} - \bar{y}'_j)(x_{ij} - \bar{x}_j)}{\sum_{j=1}^k \sum_{i=1}^{n_j} (x_{ij} - \bar{x}_j)^2}. \quad (13)$$

Przejdziemy teraz do zdefiniowania efektywności uczenia (zastosowanych przez Aitkina i Longforda). Zestawienie szkół odbywa się za pomocą porównania wartości oczekiwanej składnika losowego ξ_j (wzór 2). Ponieważ składniki σ^2 oraz σ_l^2 są znane przed oszacowaniem modelu (w dalszej części artykułu zostanie przedstawiona procedura ich estymacji) więc możemy tą informację wykorzystać jako informację a priori. Wy-

znaczymy rozkład warunkowej zmiennej losowej ξ_j pod warunkiem $\bar{y}'_j$. Ze wzoru (2) średnia na poziomie j -tej szkoły wyraża się wzorem:

$$\bar{y}'_j = \alpha + \beta \bar{x}_j + \xi_j + \bar{e}_j. \quad (14)$$

Przy poczynionych założeniach $\bar{y}'_j$ ma rozkład normalny $N(\alpha + \beta \bar{x}_j, \sigma_I^2 + \sigma^2 / n_j)$. Ten rozkład przyjmujemy jako rozkład a priori. Ponieważ ξ_j jest zmienną losową z rozkładu $N(0, \sigma_I^2)$, więc rozkład warunkowy $f(\xi_j / \bar{y}'_j)$ też będzie rozkładem normalnym.

Uwaga: Znany jest następujący fakt z rachunku prawdopodobieństwa. Jeżeli zmienne losowe $X_1 \sim N(\mu_1, \sigma_1^2)$ i $X_2 \sim N(\mu_2, \sigma_2^2)$ oraz $\rho_{1,2} = \text{cor}(X_1, X_2)$, to rozkład warunkowy X_1 / X_2 jest postaci:

$$N\left(\mu_1 + \rho_{1,2} \frac{\sigma_1}{\sigma_2} (X_2 - \mu_2), \sigma_1^2 (1 - \rho_{1,2}^2)\right).$$

Stąd uwzględniając fakt $\rho' = \text{cor}(\xi_j, \bar{y}_j) = \sigma_I^2 / (\sigma_I^2 + \sigma^2 / n_j)$ mamy: $f(\xi_j / \bar{y}'_j)$ ma rozkład normalny w postaci:

$$N\left(\rho' \frac{\sigma_I}{\sqrt{\sigma_I^2 + \sigma^2 / n_j}} (\bar{y}'_j - \alpha - \beta \bar{x}_j), \sigma_I^2 (1 - \rho'^2)\right),$$

czyli:

$$N(\rho n_j^* (\bar{y}'_j - \alpha - \beta \bar{x}_j), n_j^* (1 - \rho) \sigma_I^2 / n_j), \quad (15)$$

gdzie

$$n_j^* = w_j / (1 - \rho).$$

Porównanie szkół będzie się opierało na porównaniu wartości średnich z rozkładu warunkowego zadanego wzorem (15). Stąd efekt kształcenia lub edukacją wartość dodatkową (EWD) definiujemy w postaci

$$e_j = \hat{\rho} n_j^* (\bar{y}'_j - \hat{\alpha} - \hat{\beta} \bar{x}_j). \quad (16)$$

Estymacja punktowa EWD może być obarczona błędami (na poziomie egzaminu gimnazjalnego i matury): po pierwsze, pytania i zadania nie badają całkowitego przyrostu wiedzy i po drugie oceny dokonywane są przez różne osoby, o być może zróżnicowanych kryteriach; co jest zwłaszcza widoczne w ocenach matury z języka polskiego (duże odchylenie standardowe). To powoduje konieczność wyznaczenia przedziału ufności oszacowanej wartości EWD dla analizowanej grupy uczniów. Przedział ufności, to przedział, który z określonym współczynnikiem ufności zawiera prawdziwą wartość interesującego nas parametru. Wyznaczone dla EWD przedziały ufności możemy traktować jako regułę decyzyjną. Jeżeli chcemy w sposób odpowiedzialny formułować na podstawie EWD oceny typu *szkoła A lepiej uczy w zakresie sprawdzanym przez egzamin gimnazjalny od szkoły B*, to warto wiedzieć, jakie jest ryzyko popełnienia błędu. Przedziały ufności pozwalają nam to ryzyko oszacować. Jeżeli wyznaczymy 95% prze-

działy ufności EWD dla porównywanych szkół i przedziały te są rozłączne, to ryzyko sformułowania nietrafnej oceny jest niewielkie. Gdy przedziały częściowo pokrywają się, formułowanie oceny staje się bardziej ryzykowne - należy się od niej powstrzymać. W przypadku EWD szkoły przyjmujemy 95% współczynnik ufności. Przyjęcie niższego współczynnika ufności dla tych grup związane jest z mniej dotkliwymi konsekwencjami sformułowania nietrafnej oceny.

Ze wzoru (15) wynika, że statystyka:

$$\frac{f(\xi_j / \bar{y}'_j) - \rho n_j^* (\bar{y}'_j - \alpha - \beta \bar{x}_j)}{\sqrt{n_j^* (1 - \rho) \sigma_I^2 / n_j}},$$

ma rozkład normalny $N(0,1)$. Stąd przedział ufności dla efektu kształcenia jest postaci:

$$[\rho n_j^* (\bar{y}'_j - \alpha - \beta \bar{x}_j) - z_{1-\alpha/2} \sqrt{n_j^* (1 - \rho) \sigma_I^2 / n_j}, \rho n_j^* (\bar{y}'_j - \alpha - \beta \bar{x}_j) + z_{1-\alpha/2} \sqrt{n_j^* (1 - \rho) \sigma_I^2 / n_j}] \quad (17)$$

gdzie $z_{1-\alpha/2}$ kwantylem rozkładu $N(0,1)$ rzędu $1 - \alpha / 2$.

W Tabeli 3 prezentowane są główne charakterystyki statystyczne do oszacowanych modeli. Obliczenia zostały wykonane za pomocą programu R- Project.

W dodatku A oraz B zostały przedstawione sposoby estymacji wariancji σ_I^2 , σ^2 oraz opisany test mnożników Lagrange'a LM do badania istotności σ_I^2 (wzór B.1).

Otrzymane modele są dopasowane pod względem normalności reszt. Współczynniki nachylenia $\hat{\beta}$ charakteryzują się małym odchyleniem standardowym. Oszacowane wartości testu LM wskazują na to, że σ_I^2 jest statystycznie istotny na poziomie istotności 0,01. Zasadne jest więc stosowanie modelu efektów losowych (związanych z σ_I^2). Zauważmy, że w obu modelach ρ jest na podobnym poziomie 0,08. Wyniki gimnazjalne oraz maturalne (przeskalowane wzór (1)), są mocniej (dodatnio) skorelowane w przypadku danych humanistycznych, oraz charakteryzują się mniejszą wariancją.

Tab. 3. Podstawowe charakterystyki statystyczne modeli

Charakterystyki statystyczne	Model dla części humanistycznej	Model dla części matematycznej
Wsp. korelacji (Pearsona)	0,9515074	0,647615
σ^2	11,905508284	245,5981801
σ_I^2	0,783290	24,02437768
ρ	0,061730791	0,089104
P-value normalność	>0,01	>0,01
Współczynnik beta	0,9594070390	1,101871488
Odch. standardowe beta	0,003787	0,030512165
Współczynnik alfa	0,06551	-25,20932746
LM - p-value	<0,01	<0,01

Źródło: obliczenia własne za pomocą programu Excel oraz w środowisku R.

3.2. Miernik wykorzystujący współczynnik Giniego (indeks Sena)

Przedstawimy teraz inny sposób szacowania EWD. Jeżeli $U(x)$ jest użytecznością kształcenia danej szkoły uczniów o wiedzy, zdolności i zachowaniu x , a $f(x)$ jest gęstością tegoż rozkładu w danej populacji uczniów tejże szkoły, to średnią użyteczność kształcenia (SUK) dla uczniów tej szkoły można obliczyć ze wzoru:

$$SUK = \int_0^{\infty} U(x)f(x)dx,$$

który jest analogiem dobrobytu społecznego populacji o rozkładzie dochodów $f(x)$.

Do pomiaru nierównomierności rozkładu stosuje się współczynnik Giniego - G , który można zinterpretować geometrycznie za pomocą krzywej Lorenza (G jest równy podwojonemu polu ograniczonego przekątną jednostkowego kwadratu i krzywą Lorenza). Współczynnik Giniego jest relatywnym indeksem nierówności mierzącym skalę proporcji wiedzy (dochodów), a nie efektywną miarę nierówności. Krzywa Lorenza nie zmieni się, jeżeli wektor wiedzy (dochodów) pomnożymy przez dowolną dodatnią liczbę rzeczywistą. Zatem możemy powiedzieć, że współczynnik Giniego jest indeksem relatywnej nierównomierności.

Założmy, że mamy 2 krzywe Lorenza dla 2 różnych populacji A i B .

Definicja 1.

Będziemy mówić, że rozkład populacji A dominuje nad rozkładem populacji B w sensie Lorenza $F_A(x) \geq F_B(x) \Leftrightarrow$ krzywa Lorenza dla populacji A jest nad krzywą Lorenza dla populacji B (jeżeli się przecinają, to są nieporównywalne).

Łatwo zauważyć, że jeżeli rozkład populacji A dominuje nad rozkładem populacji B w sensie Lorenza, to pole koncentracji dla A jest mniejsze niż pole koncentracji dla B i tym samym w A jest mniejsza nierównomierność rozkładu dochodów niż w B . Zatem na zbiorze rozkładów wiedzy porządek dominacji w sensie Lorenza jest w pewien sposób równoważny z porządkiem nierównomierności rozkładu wiedzy.

Twierdzenie Atkinsona. Niech $F(x)$ i $G(x)$ będą dystrybuantami dwóch rozkładów dochodów o takich samych dochodach przeciętnych $\mu_F = \mu_G$. Wtedy:

$$L_F(p) \geq L_G(p) \Leftrightarrow \int_0^{\infty} U(x)f(x)dx \geq \int_0^{\infty} U(x)g(x)dx, \quad (18)$$

dla każdej $U(x)$ takiej, że $U(x)$ rośnie wklęsłe ($U'(x) > 0$ i $U''(x) < 0$).

Czyli SUK populacji A (o dystrybuancie wiedzy $F(x)$) jest nie mniejsza niż SUK populacji B (przy takiej samej wiedzy przeciętnej) wtedy i tylko wtedy, gdy rozkład populacji A dominuje w sensie Lorenza nad rozkładem populacji B , czyli gdy w grupie A jest mniejsza nierówność wiedzy i zachowania niż w grupie B .

Twierdzenie Atkinsona uogólnił Shorrocks [1983], wprowadzając uogólnioną krzywą Lorenza $GL(x)$:

$$GL(x) = \mu \cdot L(x), \quad (19)$$

gdzie $L(x)$ jest krzywą Lorenza, μ jest dochodem przeciętnym.

Twierdzenie Shorrocks. Niech $F(x)$ i $G(x)$ będą dystrybuantami 2 rozkładów dochodów ($f(x)$ i $g(x)$ odpowiednio do ich gęstości). Wówczas:

$$GL_F(p) \geq GL_G(p) \Leftrightarrow \int_0^{\infty} U(x)f(x)dx \geq \int_0^{\infty} U(x)g(x)dx, \quad (20)$$

dla wszystkich $U(x)$, takich że $U'(x) > 0$ i $U''(x) < 0$, oraz dla każdego $p \in [0, 1]$.

Czyli SUK populacji A jest nie mniejszy niż SUK populacji B wtedy i tylko wtedy, gdy rozkład populacji A dominuje w sensie uogólnionych krzywych Lorenza nad rozkładem populacji B. Uogólniona krzywa Lorenza dla populacji o rozkładzie F jest definiowana następująco:

$$GL(F; p) = \int_0^p F^{-1}(q) dq, \text{ dla } p \in [0, 1], \quad (21)$$

i stowarzyszony z nią częściowy porządek GL jest zdefiniowany następująco:

$FGLG \Leftrightarrow GL(F; p) \geq GL(G; p), \forall p \in [0, 1]$ oraz $GL(F; p) > GL(G; p)$ dla pewnego $p \in [0, 1]$.

Sen [1973] wprowadził następującą skróconą miarę dobrobytu (indeks Sena), którą wykorzystamy w pomiarze wiedzy danej grupy:

$$IS_F = m_F(1 - G_F). \quad (22)$$

Uwaga. Porządek uogólnionych krzywych Lorenza implikuje porządek według skróconej miary (dobrobytu) wiedzy Sena:

$$\begin{aligned} FGLG &\Rightarrow \int_0^1 GL_F(p) dp \geq \int_0^1 GL_G(p) dp \Leftrightarrow \mu_F \int_0^1 L_F(p) dp \geq \mu_G \int_0^1 L_G(p) dp \Leftrightarrow \\ &\Leftrightarrow \mu_F \left(\frac{1 - 2S_1}{2} \right) \geq \mu_G \left(\frac{1 - 2S_2}{2} \right) \Leftrightarrow \frac{1}{2} \mu_F (1 - G_F) \geq \frac{1}{2} \mu_G (1 - G_G) \Leftrightarrow IS_F \geq IS_G, \end{aligned}$$

gdzie S_1 jest polem pod krzywą Lorenza dla F , a S_2 jest polem pod krzywą Lorenza dla G .

Uwaga. Indeks Sena (IS) w większym stopniu zależy od μ niż od G dla $G \leq 0,5$; a dla $G > 0,5$ odwrotnie. Stąd, aby zmierzyć skuteczność danej szkoły wystarczy wziąć np. stosunek miary wiedzy Sena liczonej dla wyników egzaminu maturalnego ($IS_M(y)$) i wyników egzaminu gimnazjalnego ($IS_G(x)$). Tak wyrażony współczynnik będziemy oznaczali przez EWS (edukacyjna wartość dodana opierająca się na idei indeksu Sena):

$$EWS(\vec{x}, \vec{y}) = \frac{IS_M(\vec{y})}{IS_G(\vec{x})}.$$

Z trzeciego wniosku Lazeara wynika, że im większa nierówność wiedzy w klasie, tym trudniej zoptymalizować wynik kształcenia. Zatem w pomiarze efektywności kształcenia powinno brać się pod uwagę wielkość nierównomierności (G) rozkładu wiedzy.

4. REZULTATY I WNIOSKI

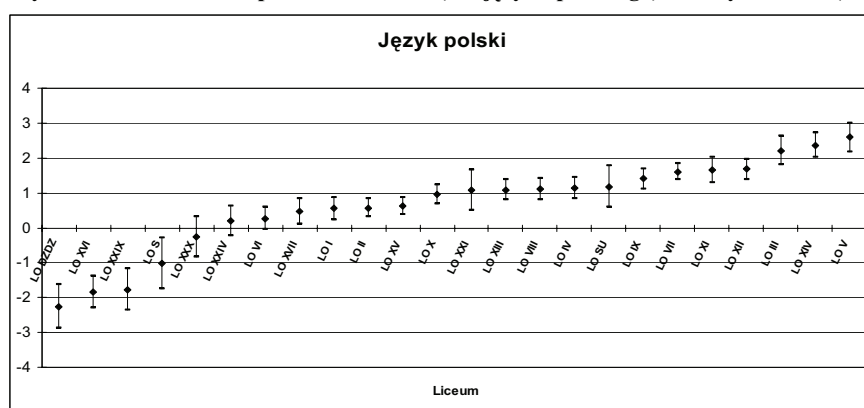
Przedstawimy otrzymane rezultaty. W tabeli 4 pokazana jest próba poszukiwania zależności efektywności kształcenia od charakterystyki szkoły, a dokładniej od stażu pracy nauczycieli w nich uczących. Okazuje się, że dodatnia zależność w obydwu modelach jest tylko dla nauczania matematyki. Im więcej nauczycieli matematyki z powyżej pięcio – letnim stażem jest w danej szkole, tym dana szkoła ma lepszy wynik kształcenia

Tab. 4. Współczynniki korelacji Pearsona rankingu z charakterystykami nauczycieli

Język polski			Matematyka		
Charakterystyki szkół*	EWD	EWS	Charakterystyki szkół*	EWD	EWS
SP5	0,0193	0,0552	SM5	0,4748	0,5090
SP	-0,1328	-0,0611	SM	0,2246	0,2305
SW	0,3067	0,3373	SW	0,1523	0,1397
SP5 przez SP	0,3760	0,2696	SM5 przez SM	0,5271	0,5866
SP przez SW	-0,2439	-0,2271	SM przez SW	0,1454	0,1602
SP5 przez SW	-0,1110	-0,1283	SM5 przez SW	0,3783	0,4171
SumaP5 przez SumaW	-0,2953	-0,2642	SumaM5 przez SumaW	0,1838	0,1827
SumaP przez SumaW	-0,2991	-0,2358	SumaM przez SumaW	0,1860	0,1911
Udział mężczyzn	-0,0244	0,0099	Udział mężczyzn	-0,5253	-0,4914

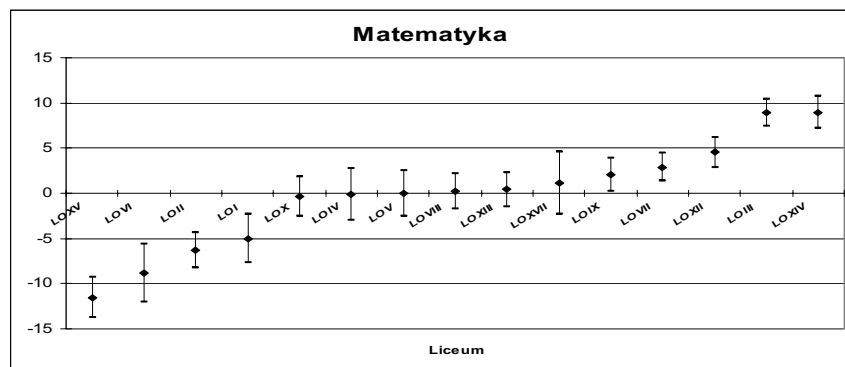
Źródło: Kuratorium Oświaty we Wrocławiu.

Rys. 1. EWD oraz 95% przedział ufności (dla języka polskiego) zadany wzorem (17)



Źródło: obliczenia własne za pomocą programu Excel.

Rys. 2. EWD oraz 95% przedział ufności (dla matematyki) zadany wzorem (17)



Źródło: obliczenia własne za pomocą programu Excel.

Tab. 5. Rankingi szkół (RS)

	Model VC				Miernik wykorzystujący indeks Sena								GAZ. WROC.
	Język polski		Matematyka		Język polski				Matematyka				
Szkoła	EWD	RS	EWD	RS	Indeks Sena Gim.	Indeks Sena matura	EWS	RS	Indeks Sena Gim.	Indeks Sena matura	EWS	RS	RS
LO DZDZ	-2,2566	24			55,1270	36,1870	0,6564	24					18
LO I	0,5527	16	-4,9920	12	66,9870	52,2670	0,7803	7	65,7450	33,1630	0,5044	12	-
LO II	0,5771	15	-6,2677	13	70,5870	52,9200	0,7497	18	75,1490	37,8610	0,5038	13	9
LO III	2,2117	3	8,9073	2	82,9620	63,5240	0,7657	12	87,9260	65,5510	0,7455	1	2
LO IV	1,1336	9	-0,0951	10	68,1280	53,5390	0,7859	5	71,2320	41,0560	0,5764	11	12
LO IX	1,4022	7	2,1015	5	74,2710	57,6280	0,7759	8	74,4150	46,3650	0,6231	5	6
LO S	-1,0305	21			66,9090	46,6690	0,6975	22					11
LO SU	1,1815	8			71,6370	57,4220	0,8016	3					10
LO V	2,5901	1	0,0292	9	78,9630	63,2280	0,8007	4	78,1530	46,1510	0,5905	9	3
LO VI	0,2704	18	-8,7870	14	65,6710	50,7110	0,7722	11	69,0320	32,3130	0,4681	14	15
LO VII	1,6086	6	2,8972	4	78,4570	59,4900	0,7582	16	81,0050	51,4350	0,6350	4	5
LO VIII	1,1056	10	0,2417	8	74,2130	55,2990	0,7451	20	75,0800	44,5480	0,5933	8	7
LO X	0,9455	13	-0,3661	11	69,8720	54,0190	0,7731	9	73,7500	42,9910	0,5829	10	13
LO XI	1,6647	5			66,3990	55,2390	0,8319	2					14
LO XII	1,6820	4	4,5479	3	78,5230	58,8200	0,7491	19	79,5200	52,2450	0,6570	3	8
LO XIII	1,0837	11	0,4406	7	72,8870	55,6960	0,7641	13	73,8090	43,8620	0,5943	7	4
LO XIV	2,3633	2	8,9580	1	79,9340	62,4490	0,7813	6	83,1080	61,6340	0,7416	2	1
LO XV	0,6338	14	-11,5375	15	68,0050	52,5500	0,7727	10	69,0720	30,0000	0,4343	15	16
LO XVI	-1,8428	23			58,2270	41,3700	0,7105	21					-
LO XVII	0,4754	17	1,1386	6	67,7200	50,7900	0,7500	17	68,6400	41,7990	0,6090	6	-
LO XXI	1,0741	12			60,5710	51,9640	0,8579	1					-
LO XXIV	0,2093	19			64,7110	49,3800	0,7631	15					12
LO XXIX	-1,7674	22			61,7840	42,5190	0,6882	23					17
LO XXX	-0,2645	20			58,5430	44,6790	0,7632	14					-

Źródło: obliczenia własne za pomocą programu Excel.

Model VC – język polski

Pod względem humanistycznym w jakości kształcenia najwyższej sklasyfikowane zostało LO V. Przewaga nad drugim w rankingu humanistycznym LO XIV wynosi około 0,22. Na trzeciej pozycji uplasowało się LO III. Zdecydowanie najgorzej w rankingu humanistycznym wypada LO DZDZ (co jest zgodne z rankingiem zaprezentowanym przez Gazetę Wrocławską). Rysunek 1 pokazuje jakie licea wykazują istotne różnice pod względem oszacowanych EWD. Zauważmy, że wyznaczone przedziały ufności dla pierwszych dwóch liceów są wyraźnie lepsze od pozostałych, z wyjątkiem LO III (z którym mają część wspólną). Dużą niepewnością oszacowanych EWD charakteryzuje się LO SU którego przedział ufności ma część wspólną z przedziałami trzynastu innych liceów. Podobna sytuacja dotyczy LO XXI.

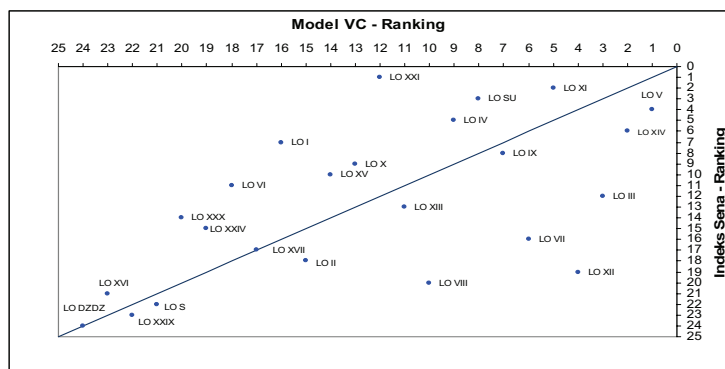
Model VC – matematyka

Pod kątem matematycznym najlepiej wypadło LO XIV, aczkolwiek należy podkreślić, że przewaga nad LO III jest minimalna (około 0,05), oraz przedział ufności dla LO III jest całkowicie zwarty w przedziale LO XIV. Ponadto jak to wynika z rysunku 2 istnieje istotna przewaga tych dwóch liceów nad pozostałymi. Zauważmy również, że lider poprzedniego rankingu (humanistycznego) zajął dziewiątą pozycję. Podobnie jak w poprzednim modelu dużą niepewnością jest obciążony wynik LO XVII, gdyż oszacowany przedział ufności ma część wspólną z większością przedziałów zaznaczonych na rysunku 2. Zauważmy również, że oszacowane EWD dla liceów LO X, LO IV, LO V, LO VIII, LO XIII są bardzo zbliżone, zaś ich przedziały ufności zawierają się w sobie (biorąc parami) stąd ryzyko sformułowania nietrafnej oceny jest w tym przypadku duże. Najgorszym liceum w tym rankingu jest LO XV, dla którego wyznaczony przedział ufności ma część wspólna tylko z LO VI zaś od pozostałych jest wyraźnie gorszy.

Model opierający się na indeksie Sena – język polski

Zauważmy, że w tym przypadku otrzymaliśmy ranking, który jest znacząco różny od rankingu wyliczonego dla modelu VC. Korelacja pomiędzy tymi dwoma rankingami wynosi około 0,593. Rysunek 3 pokazuje zależność pomiędzy rankingami otrzymanymi z pomocą modelu VC oraz opierającego się na idei Sena. Jeżeli liceum otrzymało w rankingu opierającym się na indeksie Sena wyższą lokatę niż w rankingu otrzymanym za pomocą modelu VC wówczas leży nad przekątną, jeżeli gorszą wówczas leży pod przekątną (w przypadku równości leżą na prostej). Zaznaczone punkty należy interpretować tak, że im większa odległość od prostej tym większa rozbieżność w pozycji rankingowej. Na pierwszej pozycji uplasowało się LO XXI. Położenie tego liceum na rysunku 3 wskazuje na istotną różnicę pomiędzy otrzymanymi rangami (model VC pozycja 12). Z rysunku 12 zauważamy również, że LO XII ma największą odległość od prostej co powoduje spadek z pozycji 4 (model VC) na pozycję 19. Na drugiej pozycji uplasowało się LO XI, na trzeciej zaś LO SU. Przypomnijmy, że model Sena dowartościuje te szkoły, które musiały dokonać większego wysiłku dydaktycznego powiększając wzrost wiedzy i umiejętności każdego ucznia.

Rys. 3. Zależność pomiędzy rankingami otrzymanymi przy wykorzystaniu dwóch modeli rozpatrywanych liceów dla języka polskiego



Źródło: obliczenia własne za pomocą programu Excel.

Model opierający się na indeksie Sena – matematyka

Otrzymane rezultaty w tym przypadku są niemalże idealnie zgodne z rankingiem otrzymanym za pomocą modelu VC, różnice możemy zauważyć tylko dla liderów tzn. LO III jest na pierwszej pozycji zaś LO XIV na pozycji drugiej.

Przeprowadzone analizy pokazały jak zmieniają się rankingi jeżeli weźmiemy pod uwagę zróżnicowanie wyników gimnazjalnych oraz maturalnych (zmienimy kryteria). Otrzymane rezultaty w każdym modelu wykazały, że efektywność kształcenia języka polskiego jest dużo bardziej zróżnicowana niż nauczanie matematyki. Matematyka jako przedmiot obowiązkowy pojawiła się na maturze w 2010 r. po bardzo długiej przerwie, zatem pełniejszy obraz efektywności kształcenia uzyskamy biorąc pod uwagę wyniki tegorocznych matur.

LITERATURA

- Aitkin M., Anderson D. Hinde J., [1981], *Statistical modeling of data on teaching styles*, Statistic Soc Vol. 144, s. 419-461.
- Aitkin M., Longford N., [1986], *Statistical Modelling Issues in School Effectiveness Studies*, Journal of the Royal Statistical Society Vol. 149, No. 1, s. 1-43.
- Atkinson A. B., [1970], *On the measurement of inequality*, Journal of Economic Theory. Vol. 2, s. 244-263.
- Balestra P, Nerlove M. [1966], *Pooling Cross Section and Time Series Data in the Estimation of a Dynamic Model: The Demand for Natural Gas* Econometrica, Vol. 34, No. 3, s. 585-612.
- Baltagi B., [1995], *Econometric Analysis of Panel Data*, Willey.
- Biernacki M., [2006], *Porządki generowane krzywą Lorenza*, Mathematical Economics 10, Wydawnictwo Akademii Ekonomicznej, Wrocław, s. 11-16.
- Biernacki M., [2009], *Kilka uwag do pomiaru jakości kształcenia*, Quality of Life Improvement through Social Cohesion, Publishing House of Wrocław University of Economics, s. 111-123.
- Croissant Y., Milla G., [2008], *Panel Data Econometrics in R: The plm Package*, Journal of Statistical Software Vol. 27, No. 2, s. 1-48.

- Dempster A., Rubin D., Tsutakawa R., [1981], *Estimation in Covariance Components Models*, Journal of the American Statistical Association Vol. 76, No. 374, s. 341-353.
- Hasio C., [1999], *Analysis of Panel Data*, Cambridge University Press.
- Jakubowski J., Sztencel R., [2004], *Wstęp do teorii prawdopodobieństwa*, SCRIPT.
- Laird N., [1982], *Computation of variance components using the em algorithm*, Journal of Statistical Computation and Simulation, Vol. 14, No. 3, s. 295-303.
- Lazear E., [2001], *Educational production*, Quarterly Journal of Economics, vol. 116, s. 777-803.
- Lindley D. Smith A., [1972], *Bayes Estimates for the Linear Model*, Journal of the Royal Statistical Society. Series B (Methodological), Vol. 34, No. 1, s. 1-4.
- Marks D. Fogelman K. et al. (general discussion), [1984], *Assessment of Examination Performance in Different Types of Schools*, Journal of the Royal Statistical Society Vol. 147, No. 4, s. 569-581.
- Sen A.K., [1973], *On Ignorance and equal distribution*, American Economic Review. Vol. 63. No. 5, s. 1022-1024.
- Shorrocks A.F., [1983], *Ranking Income Distributions*. Economica, Vol. 50, s. 3-17.

Dodatek A – estymacja komponentów wariancji

Znając wariancje obu części składnika losowego modelu bylibyśmy w stanie przeprowadzić estymację parametrów α oraz β . Jednak w praktyce σ^2 oraz σ_i^2 nie są znane.

Konieczne jest więc oszacowanie ich wartości. W literaturze możemy spotkać wiele sposobów estymacji tych parametrów (np. Baltagi [1995] przedstawia cztery różne sposoby). Aby uzyskać ocenę wariancji σ^2 można użyć estymatora postaci:

$$\sigma^2 = \frac{\sum_{j=1}^k \sum_{i=1}^{n_j} (y'_{ij} - \hat{\alpha}_j - \hat{\beta}_2 x_{ij})^2}{n - k - 1}, \quad (\text{A.1})$$

gdzie $\hat{\beta}_2$ jest opisany wzorem (13), oraz:

$$\hat{\alpha}_j = \bar{y}'_j - \hat{\beta}_2 \bar{x}_j. \quad (\text{A.2})$$

Wyjaśnię ideę konstrukcji powyższego estymatora. Odejmując stronami równania zapisane wzorami (2) i (14) otrzymuję:

$$y'_{ij} - \bar{y}'_j = \beta(x_{ij} - \bar{x}_{ij}) + e_{ij} - \bar{e}_j. \quad (\text{A.3})$$

Szacując współczynnik β metodą najmniejszych kwadratów otrzymamy współczynnik zapisany we wzorze (13). Metodę najmniejszych kwadratów zastosowaliśmy, gdyż reszty mają rozkład normalny ze średnią zero oraz wariancją $\text{var}(e_{ij} - \bar{e}_j) = \sigma^2(n_j - 1) / n_j (\approx \sigma^2)$. Stąd jeżeli wariancje składnika losowego obliczymy jako drugi moment centralny otrzymamy:

$$E \frac{1}{n_j} \sum_{i=1}^{n_j} (e_{ij} - \bar{e}_j)^2 = \sigma^2 \frac{(n_j - 1)}{n_j}.$$

Zastępując odpowiednie wyrażenia ich estymatorami mamy $(n_j - 1)\hat{\sigma}^2 = \sum_{i=1}^{n_j} (\hat{e}_{ij} - \bar{\hat{e}}_j)^2$.

Dla k szkół jest k takich estymatorów. Sumując ostatnie równanie (po wszystkich szkołach j od 1 do k) otrzymujemy:

$$\hat{\sigma}^2 = \frac{\sum_{j=1}^k \sum_{i=1}^{n_j} (\hat{e}_{ij} - \bar{\hat{e}}_j)^2}{n-k}. \quad (\text{A.4})$$

Ze wzoru (A.3) widzimy, że $(\hat{e}_{ij} - \bar{\hat{e}}_j)^2 = y'_{ij} - \bar{y}'_j - \hat{\beta}_2(x_{ij} - \bar{x}_{ij}) = y'_{ij} - \hat{\alpha}_j - \hat{\beta}_2 x_{ij}$. Wstawiając to równanie do wzoru (A.4) oraz biorąc poprawkę na liczbę szacowanych parametrów otrzymujemy estymator opisany wzorem (A.1). Myśląc o liczbie szacowanych parametrów mam na myśli $\hat{\beta}_2$ oraz $\hat{\alpha}_j$ $j \in \{1, \dots, k\}$, co wyjaśnia skąd we wzorze (A.1) jest $n-k-1$.

Przejdźmy do oszacowania σ_I^2 . Ze wzoru (2) widzimy, że $y'_{ij} - \alpha - \beta x_{ij} = \xi_j + e_{ij}$ stąd mamy $\text{var}(y'_{ij} - \alpha - \beta x_{ij}) = \sigma^2 + \sigma_I^2$. Jako ocenę estymatora wariancji całkowitej modelu (2) możemy przyjąć wariancję modelu klasycznej regresji linowej (dopasowanych do wszystkich danych). Zapisana powyżej idea sprowadza się do wzoru:

$$\hat{\sigma}_I^2 = \frac{\sum_{j=1}^k \sum_{i=1}^{n_j} (y'_{ij} - \hat{\alpha}' - \hat{\beta}_1 x_{ij})^2}{n-2} - \hat{\sigma}^2, \quad (\text{A.5})$$

gdzie: $\hat{\beta}_1$ jest opisany wzorem (12), zaś $\hat{\alpha}' = \bar{y}' - \hat{\beta}_1 \bar{x}$.

Dodatek B– Test Breuscha-Pagana

W celu przeprowadzenia testu, czy uzyskane efekty losowe są istotne, użyję testu Breusch'a-Pagana (np. Baltagi [1995]). Jest to test mnożników Lagrange'a, w którym stawiam hipotezy: $H_0: \sigma_I^2 = 0$ hipoteza alternatywna: $H_1: \sigma_I^2 \neq 0$.

Testuję za pomocą statystyki:

$$LM = \frac{\left(\sum_{j=1}^k n_j \right)^2}{\left(2 \sum_{j=1}^k n_j (n_j - 1) \right)} \left[\frac{\sum_{j=1}^k \left(\sum_{i=1}^{n_j} e'_{ij} \right)^2}{\sum_{j=1}^k \sum_{i=1}^{n_j} e'^2_{ij}} - 1 \right] \sim \chi^2(1), \quad (\text{B.1})$$

gdzie e'_{ij} są to reszty otrzymane w wyniku zastosowania metody MNK do wszystkich danych (nie zależnie od szkół). Powyższy wzór mówi, że statystyka testowa LM ma asymptotyczny rozkład chi-kwadrat (przy założeniu hipotezy zerowej) z jednym stopniem swobody. Hipotezę zerową odrzucamy, jeżeli wartość statystyki LM należy do prawostronnego obszaru krytycznego.

EFFICIENCY OF LEARNING IN SECONDARY SCHOOLS IN WROCLAW

In the article authors try to measure the effectiveness of learning in secondary schools in Wrocław. For this purpose, they use models for panel data. The use of panel models to measure the effectiveness have been launched by Aitkins and Longford in 1986. Our work focuses on the use of the ideas presented by these authors, which we apply to data describing the results of secondary schools and matriculation exams, and school characteristics.