

ZACHODNIOPOMORSKI UNIWERSYTET TECHNOLOGICZNY W SZCZECINIE
WYDZIAŁ INFORMATYKI

mgr inż. Jacek Klimaszewski

**Uczenie modeli liniowych z regularyzacją dla małych zbiorów
danych**

Rozprawa doktorska

Promotor: dr hab. inż. Marcin Korzeń, prof. ZUT

Szczecin 2024

Streszczenie

W niniejszej rozprawie poruszono zagadnienie uczenia modeli liniowych z regularyzacją, skupiając się na małych zbiorach danych (liczba atrybutów znacznie przewyższa liczbę obserwacji). Przeanalizowano i zbadano wydajność istniejących procedur uczenia na tych zbiorach oraz różnych form regularyzacji modeli. Sprawdzono aspekty numeryczne algorytmów, m.in. wpływ typu zmiennoprzecinkowego pojedynczej i podwójnej precyzji oraz rodzaju minimalizacji kierunkowej na szybkość uczenia oraz sposoby reprezentacji danych uczących. Wykazano teoretycznie i empirycznie, że w pewnym zakresie parametrów wykorzystanie formuły Woodbury’ego bądź rozkładu QR znacznie skraca czas uczenia modeli z regularyzacją ℓ^2 — sprawdzono analogiczne rozwiązanie dla modeli wyposażonych w rzadkie funkcje kary. Eksperymenty wykonano w środowisku Python wraz z fragmentami zaimplementowanymi w języku C++, biorąc za punkt odniesienia istniejące biblioteki uczenia maszynowego, jak `scikit-learn` czy `liblinear`. Przetestowano działanie proponowanych rozwiązań w zadaniach klasyfikacji na rzeczywistych i syntetycznych zbiorach danych oraz wskazano potencjalne zastosowania uczenia regularyzowanych modeli liniowych w problemach nieliniowych.

słowa kluczowe: uczenie maszynowe, modele liniowe, regularyzacja, optymalizacja, metoda Newtona,

Abstract

This PhD thesis deals with the issue of training regularized linear models on small datasets (i.e. datasets containing much less observations than features). Existing optimization algorithms were analysed and their behaviour on those datasets was examined. Some numerical issues were checked, such as influence of single precision or double precision floating-point type or line-search procedure on training speed. It was proved theoretically and empirically, that in some configurations use of Woodbury's formula or QR factorization leads to shorter time of training ℓ^2 -regularized models — the same approach was checked for sparsity-inducing penalizers. Experiments were conducted mainly in Python environment, with some parts written in C++, taking as a reference point existing machine learning libraries, i.e. `scikit-learn` or `liblinear`. Performance of proposed solutions were tested in the classification task for some real and sythetic datasets. Moreover, some potential use of training penalized linear models in non-linear problems was highlighted.

keywords: machine learning, linear models, regularization, optimization, Newton's method

Spis treści

1	Modele liniowe w uczeniu maszynowym	12
1.1	Sformułowanie problemu	12
1.2	Przejsie z modelu liniowego na nieliniowy	13
2	Uczenie modeli liniowych	14
2.1	Model regresji liniowej	14
2.2	Model regresji logistycznej	14
2.3	Idea regularyzacji	16
2.4	Ograniczenia wypukłe i regularyzacja ℓ^2	16
2.4.1	Zastosowanie rozkładu QR	21
2.4.2	Zastosowanie formuły Woodbury'ego	23
2.5	Regularyzacja ℓ^1	25
2.5.1	Subgradient	26
2.5.2	Spadek względem współrzędnych	26
2.5.3	Przekształcenie w problem gładkiej optymalizacji	29
2.6	Regularyzacja „elastic net”	34
2.6.1	Spadek względem współrzędnych	34
2.6.2	Przekształcenie w problem gładkiej optymalizacji	34
2.7	Aproksymacja normy ℓ^1	35
2.8	Regularyzacja z użyciem pseudonorm ℓ^q , $0 \leq q < 1$,	36
2.8.1	Spadek względem współrzędnych	37
2.8.2	Zastąpienie pseudonormy ℓ^q jej górnym ograniczeniem	39
2.9	Uogólniona regularyzacja	39

3	Realizacja modeli nieliniowych	41
3.1	Uczenie sieci neuronowych	41
3.2	Rzadki autoenkoder	42
3.3	Uczenie zrandomizowane, uczenie ekstremalne	45
4	Część eksperymentalna	47
4.1	Środowisko testowe i szczegóły numeryczne	47
4.2	Obliczanie iloczynu XX^T w przypadku macierzy rzadkiej	49
4.3	Porównanie implementacji metody gradientu sprzężonego	52
4.4	Przebieg eksperymentu dla modeli z regularyzacją ℓ^2 oraz ℓ^1	53
4.5	Wyniki eksperymentów dla regularyzacji ℓ^2	58
4.6	Wyniki eksperymentów dla regularyzacji ℓ^1	83
5	Zastosowanie regularyzacji w modelach nieliniowych	92
5.1	Rzadki autoenkoder	92
5.1.1	Przygotowanie i opis eksperymentu	92
5.1.2	Wyniki	94
5.2	Klasyfikacja nieliniowa — uczenie zrandomizowane	100
6	Podsumowanie i wnioski	105
	Spis literatury	108
	Załącznik A	118

Lista symboli

A, B, C oznaczają macierze.

$a_{:,j}$ oznaczają j -tą kolumnę macierzy A

$b_{i,:}$ oznaczają i -ty wiersz macierzy B

x, y, z oznaczają wektory.

$\langle \cdot, \cdot \rangle$ oznacza iloczyn skalarny

$\mathcal{D} = \{(x_i, y_i)\}_{i=1}^n$ dane uczące, gdzie $x_i \in \mathbb{R}^p$ oraz $y_i \in \{-1, +1\}$ w przypadku klasyfikacji binarnej lub $y_i \in \mathbb{R}$ w przypadku regresji.

a, b, c oznaczają skalary.

n liczba próbek

p liczba atrybutów

Wprowadzenie

Zagadnienie uczenia modelu na podstawie danych można podzielić na trzy grupy: zadanie uczenia, predykcji oraz testowania modelu. Zasadniczą część pracy dotyczy zadania uczenia modeli, głównie w sposób nadzorowany [55]. Celem zadania uczenia jest wybór modelu, który optymalizuje pewną miarę jakości predykcji na danych testowych. Taką miarą w przypadku zadania klasyfikacji może być błąd klasyfikacji, jednak zadanie uczenia zwykle nie jest rozwiązywane poprzez bezpośrednią optymalizację tej miary jakości na danych uczących — zwykle wykorzystuje się tu pewne miary, które są funkcjami ciągłymi lub gładkimi parametrów modelu, natomiast zadanie uczenia sprowadza się do rozwiązania pewnych zadań optymalizacyjnych. Jako funkcję celu wybiera się w najprostszym przypadku funkcję wiarygodności (zwaną czasami entropią krzyżową) lub błąd średniokwadratowy.

Zakres modeli analizowanych w pracy jest ograniczony do modeli liniowych. W przypadku dwóch podstawowych zadań (klasyfikacji i regresji) mamy dwie stosunkowo liczne grupy modeli, tj. klasyfikatory liniowe oraz liniowe modele regresyjne. Zasadniczą część metod przedstawionych i analizowanych w pracy odnosi się do obu grup modeli, z zastrzeżeniem, że niektóre zadania w przypadku modeli regresyjnych są nieco prostsze do rozwiązania. W pracy większy nacisk położono na zadania klasyfikacji. Praca ogranicza się do technik strojenia modeli liniowych, jednak to ograniczenie nie jest tak duże, jak mogłoby się wydawać. Zagadnienia optymalizacyjne, które pojawiają się przy strojeniu modeli liniowych, mogą być wykorzystywane w różnych modelach analizy danych, jak np. filtracja sygnałów. Przykłady pokazane w pracach [42, 62] wskazują, że omówione techniki strojenia modeli liniowych mogą być również efektywnie wykorzystane do uczenia modeli nieliniowych.

Modele liniowe mają stosunkowo prostą strukturę, tzn. wyjście modelu jest kombinacją liniową zmiennych i współczynników modelu z ewentualnym dodaniem pewnej skalarnej funkcji nieliniowej. Jednak ze względu na różny sposób uczenia (nawet dla tych samych danych) mogą prowadzić do różnych modeli, znacznie różniących się własnościami.

Mimo że współcześnie powszechne jest stosowanie znacznie bardziej złożonych modeli, to należy podkreślić, że istnieją zadania oraz zbiory danych, dla których stosowanie takich złożonych modeli napotyka przeszkody. Jednym z ograniczeń jest złożoność próbkowa¹ modelu [4] — potocznie mówiąc, jest to minimalny rozmiar próby, przy której model może być poprawnie uczony, a której zwiększanie nie prowadzi do wzrostu jakości predykcji na danych testowych. Złożone modele mają zwykle dużą złożoność próbkową, czyli do poprawnego uczenia wymagają dużych prób uczących. Istnieją jednak problemy analizy danych, w których zwiększanie rozmiaru próby jest kłopotliwe lub kosztochłonne (np. dane mikromacierzowe).

Prezentowane w pracy podejścia są zorientowane na przypadki, gdy próby są małe, co należy rozumieć jako małe w relacji do liczby atrybutów oraz gdy próba ucząca nie może być zwiększana w sposób dowolny. Należy podkreślić, że w przypadku stosunkowo prostych modeli liniowych ograniczenia typu Vapnika czy Rademachera [6, 80] gwarantują, że przy wzroście próby do nieskończoności wszystkie modele stają się takie same.

Zachowanie modelu może być kształtowane poprzez dobór odpowiedniej funkcji kryterialnej Q , która w ogólności może się składać z dwóch składników:

$$Q(\text{model}; \text{dane}) = L(\text{model}; \text{dane}) + \lambda \cdot R(\text{model}), \quad (1)$$

tj. składnika odpowiadającego za dopasowanie modelu do danych L oraz składnika regularyzującego model (lub funkcji kary) R . Taką postać funkcji kryterialnej mają główne procedury uczenia modeli liniowych. Przykładowo w modelu regresji liniowej jako funkcję L wykorzystuje się błąd średniokwadratowy, natomiast w zależności od składnika R wyróżniamy model standardowy (bez regularyzacji), regresję grzbietową² [33], gdy R jest normą ℓ^2 wektora wag, oraz modele LASSO lub LARS [18, 75, 83], gdy R jest normą ℓ^1 wektora wag. Podobnie w przypadku modelu SVM [13] funkcja L to margines separacji lub suma marginesów próbek, natomiast R oryginalnie jest normą ℓ^2 wektora wag. Historycznie, w przypadku strojenia sieci neuronowych, składnik kary w postaci normy ℓ^2 nazywa się butwieniem wag³ i jest stosowany od lat dziewięćdziesiątych XX wieku.

Wykorzystanie normy ℓ^2 rozwiązania jako składnika regularyzacyjnego historycznie pochodzi od A. N. Tchionva [78] i była ona użyta do rozwiązywania źle uwarunkowanych lub osobliwych problemów numerycznych. Podobne znaczenie ma dodanie innych składników regularyzacyjnych. Piewsze użycie regularyzacji ℓ^2 w analizie danych pojawiło się

¹ang. sample complexity

²ang. ridge regression

³ang. weight decay

z cytowanej pracy dot. regresji grzbietowej [33]. Na znaczenie i numeryczne trudności przy optymalizacji modeli z karą ℓ^1 wskazano po raz pierwszy w pracy [83], gdzie podkreślono między innymi, że w punkcie optimum część współczynników modelu zostaje wyzerowana. Oprócz aspektów numerycznych w modelu liniowym użycie regularyzacji typu ℓ^2 odpowiada założeniu, że parametry modelu mają rozkład gaussowski, natomiast analogicznie zastosowanie regularyzacji ℓ^1 odpowiada założeniu, że parametry modelu mają rozkład Laplace’a [75, 83, 89]. W pracach [75, 89] wskazano dodatkowo na: 1) grupujący efekt normy ℓ^2 , tj. współczynniki przy zmiennych silnie skorelowanych mają tendencję do zrównywania się (co w praktyce daje efekt uśredniania atrybutów) oraz 2) ponownie podkreślono fakt, że przy regularyzacji ℓ^1 wraz ze wzrostem parametru regularizacyjnego coraz więcej współczynników jest zerowanych. Prace te zapoczątkowały intensywne badania różnych form regularyzacji. Zwrócono również uwagę na możliwość stosowania pseudonorm niewypukłych. Niezależnie i nieco wcześniej na geometryczny sens regularyzacji typu ℓ^2 i związek z maksymalizacją marginesu zwrócił uwagę V. Vapnik wraz ze współautorami [13]. Regularyzacja ℓ^1 ze względu na efekt zerowania się współczynników ma związek z problemem selekcji zmiennych. Nawet współcześnie problem wyboru zmiennych w modelach liniowych uważany jest za jeden z istotniejszych w obrębie statystyki (por. [32]). Główne problemy obliczeniowe związane z uczeniem modeli liniowych z regularyzacją dotyczą uwarunkowania macierzy, zwłaszcza w przypadkach, gdy liczba atrybutów jest znaczna oraz gdy w danych jest dużo korelacji. Drugi problem dotyczy zadań uczenia z wykorzystaniem funkcji niegładkich, np. zawierającej normę ℓ^1 (np. modele lasso, fused lasso, elastic net), natomiast stosowanie pseudonorm $\ell^q, 0 \leq q < 1$ dodatkowo prowadzi do niewypukłych problemów optymalizacyjnych.

Należy podkreślić, że regularyzacja nie jest jedyną formą poprawy procesu uczenia oraz własności algorytmów uczących. Obecnie w dobie intensywnego rozwoju sieci głębokich powszechnie stosowane jest wzbogacanie danych⁴ [64, 71, 74] czy technika drop out [72].

Cel i teza pracy

Celem badań jest poszukiwanie efektywnych metod uczenia regularyzowanych modeli liniowych przy uwzględnieniu: rozmiaru danych (np. mała próba w stosunku do liczby atrybutów), zależności występujących w danych oraz własności regularyzatora (funkcji kary).

⁴ang. augmentation

Biorąc pod uwagę parametry próby uczącej (takie jak rozmiar danych, rozkład czy korelacje), możliwe jest opracowanie algorytmów uczących dla wybranych problemów uczenia z regularyzacją działających efektywniej w stosunku do istniejących rozwiązań. Pod pojęciem efektywności rozumie się uzyskanie krótszych czasów uczenia przy zachowaniu własności i parametrów modelu i referencyjnego algorytmu uczącego lub uzyskanie wyższej dokładności modeli przez zastosowanie odpowiednich technik regularyzacji.

Główne wyniki uzyskane w pracy

W pracy poruszono zagadnienie uczenia modeli liniowych z regularyzacją, skupiając się na „małych” zbiorach danych, to znaczy takich, w których liczba atrybutów znacznie przewyższa liczbę obserwacji. W takich przypadkach regularyzacja jest jednym z głównych narzędzi poprawy własności numerycznych algorytmów, natomiast różne formy regularyzacji umożliwiają kształtowanie własności maszyn uczących pożądaných z punktu widzenia analizy danych. Najmocniejszy wynik prezentowany w pracy dotyczy usprawnienia algorytmu uczenia modelu regresji z regularyzacją ℓ^2 , poprzez zastosowanie rozkładu QR, a w sposób szczególny formuły Woodbury’ego (sekcja 2.4). Pokazano formalnie oraz w eksperymentach numerycznych, że tablicując macierz $\mathbf{X}\mathbf{X}^T$ (gdzie $\mathbf{X}$ jest macierzą danych) można — zachowując równoważność algorytmów — istotnie przewyższyć standardową procedurę Newtona obliczania kierunku minimalizacji za pomocą metody gradientu sprzężonego. Ponadto ideę wykorzystania rozkładu QR oraz zastosowania formuły Woodbury’ego zaadaptowano do uczenia modeli z rzadkimi funkcjami kary (typu lasso, elastic net), jednak (co pokazały eksperymenty numeryczne) zysk w tym przypadku nie jest tak spektakularny, jak w przypadku regularyzacji ℓ^2 .

Całość eksperymentów wykonano w języku Python, biorąc za punkt odniesienia dostępne biblioteki uczenia maszynowego `scikit-learn` i `liblinear` oraz inne algorytmy opisane w literaturze. Większość modeli została samodzielnie zaimplementowana w języku C++ w ramach jednolitego środowiska testowego. W części numerycznej posłużono się bibliotekami numerycznej algebry liniowej, takimi jak: LAPACK, Armadillo, Intel® Math Kernel Library. Przeanalizowano szereg ustawień i własności numerycznych algorytmów, jak np. wpływ reprezentacji zmiennopozycyjnej, metody iteracyjnego rozwiązywania układów równań liniowych, sposoby realizacji minimalizacji kierunkowej, wpływ sposobu reprezentacji macierzy danych (gęste, rzadkie). W pracy przedstawiono również zastosowanie technik uczenia modeli liniowych do strojenia modeli nieliniowych.

W ramach badań przygotowano kilka prac dotyczących wybranych aspektów ucze-

nia modeli regularyzowanych oraz część prezentowanych wyników została opublikowana. W pracy [45] analizowano algorytmy polegające zastąpieniu niegładkiej funkcji wartości bezwzględnej jej gładką aproksymacją. W pracach [43, 44] przedstawiono i przeanalizowano zastosowanie niegładkich i niewypukłych „norm” ℓ^q dla $0 \leq q < 1$ w zagadnieniach uczenia modeli linowych i filtracji obrazów. W pracy [42] przeanalizowano zastosowanie rozkładu QR dla modeli z regularyzacją ℓ^2 oraz zaproponowano analogiczny algorytm dla innych form regularyzacji, w tym ograniczeń niegładkich. W pracy [62] przedstawiono zastosowanie technik regularyzacyjnych przy uczeniu sieci neuronowych oraz wskazano zastosowanie struktury sieci z regularyzacją typu ℓ^1 do strojenie rzadkiego autoenkodera. Proponowane rozwiązanie porównano z algorytmem uczenia słownikowego zaimplementowanego w pakiecie `scikit-learn`.

1. Modele liniowe w uczeniu maszynowym

1.1 Sformułowanie problemu

Model liniowy opiera się na liniowej kombinacji atrybutów obserwacji $\mathbf{x}$ oraz parametrów modelu $\mathbf{w}$, co można symbolicznie zapisać:

$$y = w_0 + \sum_{j=1}^p x_j \cdot w_j \equiv \mathbf{w}^T \cdot \begin{bmatrix} 1 \\ \mathbf{x} \end{bmatrix}. \quad (1.1)$$

Taką postać modelu wykorzystuje się w zadaniach klasyfikacji i regresji (w przypadku klasyfikacji patrzy się zwykle na znak wyrażenia y).

Problem uczenia modelu najczęściej sprowadza się do zadania optymalizacji pewnej funkcji straty (przykładowo może nią być błąd), która to określa stopień dopasowania modelu do danych uczących. Pożądaną własnością funkcji straty jest jej różniczkowalność, co daje możliwość zastosowania gradientowych metod optymalizacji.

Bezpośrednia optymalizacja funkcji straty może nie przynieść oczekiwanych rezultatów z kilku powodów:

- dane zwykle zawierają szum,
- wystarczająco złożony model (o zbyt wielu stopniach swobody) może się dopasować idealnie do próbek (tzw. przeuczenie),
- czasami liczba atrybutów przewyższa liczbę obserwacji (np. dane mikromacierzowe), co w wyniku utrudnia użycie pochodnych drugiego rzędu (niepełny rząd macierzy Hessego).

Jeden ze sposobów rozwiązujący powyższe problemy polega na dodaniu do funkcji straty pewnej funkcji kary — odpowiada to optymalizacji z ograniczeniami. Jej zadanie polega na preferowaniu takich modeli, które spełniają wskazane ograniczenia (np. mniejsza norma kwadratowa parametrów modelu). Czasami natura danych sugeruje rodzaj ograniczenia, które warto zastosować, aby zwiększyć jakość modelu — przykładowo w szeregu czasowym sąsiednie wartości są ze sobą silnie skorelowane, zatem nałożenie ograniczeń na różnice między nimi powinno skutkować nie tylko prostszym, ale i lepszym modelem.

Ostatecznie funkcja kosztu podlegająca optymalizacji ma postać:

$$Q(\mathbf{w}; \text{dane}) = L(\mathbf{w}; \text{dane}) + \lambda \cdot R(\mathbf{w}), \quad (1.2)$$

gdzie $L(\mathbf{w}; \text{dane})$ oznacza funkcję straty, $R(\mathbf{w})$ oznacza funkcję kary, zaś λ steruje siłą kary. Powyższy problem można także sformułować jako zadanie z ograniczeniami:

$$\begin{cases} L(\mathbf{w}; \text{dane}) \rightarrow \min \\ R(\mathbf{w}) \leq \gamma \\ \mathbf{w} \in \mathbb{R}^{p+1} \end{cases}. \quad (1.3)$$

Między parametrami λ i γ istnieje jednoznaczne odwzorowanie [29]. Wyraz wolny w_0 nie powinien podlegać karaniu.

1.2 Przejście z modelu liniowego na nieliniowy

Modele nieliniowe potrafią dla pewnych danych osiągnąć dużo lepsze wyniki niż modele liniowe, jednakże są sposoby, które umożliwiają modelom liniowym uzyskiwanie porównywalnych wyników. Literatura notuje co najmniej dwa takie sposoby.

Pierwszy z nich, stosowany powszechnie w kontekście Support Vector Machine, polega na zastosowaniu sztuczki z przekształceniem jądrowym [29]. Opiera się ona na obserwacji, że parametry modelu można wyrazić jako kombinację próbek:

$$\mathbf{w} = \sum_{i=1}^n \mathbf{x}_i \alpha_i = \mathbf{X}^T \boldsymbol{\alpha}, \quad (1.4)$$

a także na niejawnym podnoszeniu wymiarowości próbek. Efektem ubocznym tego podejścia jest redukcja wymiaru przestrzeni do n , co w przypadku $p \gg n$ powinno także skutkować skróceniem czasu uczenia.

Drugi z nich, znany jako Extreme Learning Machine [73], bazuje na dwuwarstwowej sieci neuronowej, w której strojeniu podlegają jedynie współczynniki warstwy wyjściowej, natomiast pierwsza warstwa ma ustalone wagi, a liczba neuronów w warstwie ukrytej jest dużo większa od liczby atrybutów. Dzięki temu zabiegowi dane, po przejściu przez pierwszą warstwę, jawnie trafiają do wyższej przestrzeni, w której separacja powinna być łatwiejsza.

2. Uczenie modeli liniowych

Niech dany będzie zbiór par uczących $\{(\mathbf{x}_i, y_i)\}_{i=1}^n$, gdzie $\mathbf{x}_i \in \mathbb{R}^p$ i $y_i \in \{-1, +1\}$ (w przypadku klasyfikacji binarnej) lub $y_i \in \mathbb{R}$ (w przypadku regresji) — przy tych oznaczeniach p to liczba atrybutów, a n to liczba próbek. Niech $p(\mathbf{x}) \equiv \Pr(y = 1)$ oznacza prawdopodobieństwo, że obserwacja $\mathbf{x}$ należy do klasy $+1$.

2.1 Model regresji liniowej

W celu wyznaczenia prostej regresji dokonuje się minimalizacji sumy kwadratów różnic prawdziwej i szacowanej wartości funkcji w punktach — funkcję straty można więc zapisać następująco:

$$L(\mathbf{w}) = \frac{1}{2} \sum_{i=1}^n (y_i - \langle \mathbf{w}, \mathbf{x}_i \rangle)^2 = \frac{1}{2} \|\mathbf{y} - \mathbf{X}\mathbf{w}\|_2^2. \quad (2.1)$$

Istnieje rozwiązanie analityczne powyższego problemu dane wzorem:

$$\mathbf{w} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y}. \quad (2.2)$$

W praktyce bezpośrednie użycie wzoru (2.2) (nieuwzględniającego regularyzacji) może albo dać słaby wynik (np. punkty odstające silnie zaburzają kierunek prostej), albo się nie powieść (z uwagi na niewystarczającą liczbę punktów) — w tym drugim przypadku nakładając dodatkowe ograniczenie można znaleźć za pomocą rozkładu QR takie rozwiązanie, które ma możliwie najmniejszą normę kwadratową [25].

2.2 Model regresji logistycznej

Do znalezienia postaci prawdopodobieństwa $p(\mathbf{x})$, aproksymowanego za pomocą kombinacji liniowej, wykorzystuje się logarytm szansy (logit):

$$\log \left(\frac{p(\mathbf{x})}{1 - p(\mathbf{x})} \right) = \langle \mathbf{w}, \mathbf{x} \rangle. \quad (2.3)$$

Odpowiednie przekształcenie prowadzi do funkcji sigmoidalnej:

$$p(\mathbf{x}) = \frac{e^{\langle \mathbf{w}, \mathbf{x} \rangle}}{1 + e^{\langle \mathbf{w}, \mathbf{x} \rangle}} = \frac{1}{1 + e^{-\langle \mathbf{w}, \mathbf{x} \rangle}}. \quad (2.4)$$

Korzystając z faktu, że zmienna decyzyjna przyjmuje wartość -1 lub $+1$, wyrażenie na prawdopodobieństwo przynależności próbki $\mathbf{x}_i$ do klasy y_i można zapisać w jednolitej formie:

$$P(y_i | \mathbf{x}_i) = \frac{1}{1 + e^{-y_i \cdot \langle \mathbf{w}, \mathbf{x}_i \rangle}}. \quad (2.5)$$

Współczynniki modelu regresji logistycznej znajduje się za pomocą metody największej wiarygodności. Funkcja wiarygodności modelu regresji logistycznej ma postać:

$$l(\mathbf{w}) = \prod_{i=1}^n \frac{1}{1 + e^{-y_i \cdot \langle \mathbf{w}, \mathbf{x}_i \rangle}}. \quad (2.6)$$

Jej maksymalizacja nie jest prosta, toteż w praktyce minimalizacji podlega ujemny logarytm funkcji wiarygodności:

$$L(\mathbf{w}) = - \sum_{i=1}^n \log \left(\frac{1}{1 + e^{-y_i \cdot \langle \mathbf{w}, \mathbf{x}_i \rangle}} \right) = \sum_{i=1}^n \log \left(1 + e^{-y_i \cdot \langle \mathbf{w}, \mathbf{x}_i \rangle} \right). \quad (2.7)$$

Jest to funkcja wypukła (kombinacja funkcji wypukłych), więc istnieje globalne minimum. Gradient i macierz Hessego powyższej funkcji względem $\mathbf{w}$ mają postać:

$$\frac{\partial L}{\partial \mathbf{w}} \equiv \mathbf{g} = -\mathbf{X}^T \mathbf{v}, \quad (2.8)$$

$$\frac{\partial^2 L}{\partial \mathbf{w} \partial \mathbf{w}^T} \equiv \mathbf{H} = \mathbf{X}^T \mathbf{D} \mathbf{X}, \quad (2.9)$$

gdzie $\mathbf{v}$ to wektor n -elementowy i:

$$v_i = y_i \cdot \frac{e^{-y_i \cdot \langle \mathbf{w}, \mathbf{x}_i \rangle}}{1 + e^{-y_i \cdot \langle \mathbf{w}, \mathbf{x}_i \rangle}} \equiv y_i \cdot \frac{1}{1 + e^{y_i \cdot \langle \mathbf{w}, \mathbf{x}_i \rangle}}, \quad (2.10)$$

a macierz $\mathbf{D}$ oznacza macierz diagonalną i $d_{i,i} = \frac{1}{1 + e^{-\langle \mathbf{w}, \mathbf{x}_i \rangle}} \cdot \left(1 - \frac{1}{1 + e^{-\langle \mathbf{w}, \mathbf{x}_i \rangle}} \right)$. Do rozwiązania układu równań $\frac{\partial L}{\partial \mathbf{w}} = \mathbf{0}$ (warunek konieczny istnienia ekstremum) trzeba użyć metod numerycznych, ponieważ nie istnieje analityczne rozwiązanie.

W toku obliczeń trzeba uważać na przepełnienie w potęgowaniu, gdy wyrażenie $-y_i \cdot \langle \mathbf{w}, \mathbf{x}_i \rangle$ osiąga dużą dodatnią wartość (dzieje się to w przypadku obserwacji niepoprawnie klasyfikowanych w danej chwili) — wtedy należy skorzystać z tożsamości:

$$\log \left(1 + e^{-y_i \cdot \langle \mathbf{w}, \mathbf{x}_i \rangle} \right) = -y_i \cdot \langle \mathbf{w}, \mathbf{x}_i \rangle + \log \left(1 + e^{y_i \cdot \langle \mathbf{w}, \mathbf{x}_i \rangle} \right). \quad (2.11)$$

Analogicznie w przypadku obliczania v_i w zależności od znaku tego wyrażenia należy wybrać pierwszą formułę (gdy $y_i \cdot \langle \mathbf{w}, \mathbf{x}_i \rangle > 0$) lub drugą (w przeciwnym wypadku).

2.3 Idea regularyzacji

W trakcie uczenia modelu zazwyczaj pojawiają się pewne problemy, z którymi trzeba sobie poradzić — należą do nich między innymi:

- szum danych — w przypadku pomiarów wyniki zwykle obarczone są pewnym błędem wynikającym z niedoskonałości urządzeń mierniczych bądź niedokładnego odczytania pomiaru;
- przeuczenie — jeśli model idealnie dopasuje się do danych (co objawia się bezbłędną klasyfikacją próbek uczących), może on utracić zdolność generalizowania (zwracania poprawnego wyniku dla próbek, które nie brały udziału w trakcie uczenia);
- brak jednoznacznego rozwiązania — gdy dane są separowalne lub gdy nie ma wystarczającej liczby obserwacji, granicę decyzyjną można poprowadzić na wiele sposobów.

Zastosowanie regularyzacji, polegającej na dodaniu do funkcji straty pewnej funkcji kary, która ogranicza zbiór wartości, jakie mogą przyjąć parametry modelu, w naturalny sposób niweluje wyżej wymienione problemy. Otrzymana funkcja nosi nazwę funkcji kosztu i można ją zapisać następująco:

$$Q(\mathbf{w}) = L(\mathbf{w}) + \lambda \cdot R(\mathbf{w}), \quad (2.12)$$

gdzie $L(\mathbf{w})$ to funkcja straty (zależna od danych i parametrów modelu), $R(\mathbf{w})$ to funkcja kary (zależna od parametrów modelu), zaś λ to parametr decydujący o sile regularyzacji ($\lambda > 0$). Funkcja kary może mieć różną postać, co przedstawiono w tabeli 2.1.

Ponadto dzięki regularyzacji można nałożyć pewną wiedzę a priori na parametry modelu, na przykład karząc różnice między sąsiednimi wartościami atrybutów, gdy spodziewamy się, że w jednowymiarowym sygnale występują „obszary płaskie”.

2.4 Ograniczenia wypukłe i regularyzacja ℓ^2

Po raz pierwszy o regularyzacji ℓ^2 wspomina się w kontekście rozwiązywania źle uwarunkowanych problemów [78]. W tym przypadku funkcja kary stanowi kwadrat normy euklidesowej (z pominięciem wyrazu wolnego w_0). Optymalizowana funkcja kosztu ma postać:

$$Q(\mathbf{w}) = L(\mathbf{w}) + \frac{\lambda}{2} \sum_{j=1}^p w_j^2. \quad (2.13)$$

Tabela 2.1: Przykładowe funkcje kary.

nazwa kary	postać funkcji $R(\mathbf{w})$
ridge (ℓ^2) [29, 78]	$\ \mathbf{w}\ _2^2$
LASSO (ℓ^1) [75]	$\ \mathbf{w}\ _1$
elastic net [89]	$\alpha\ \mathbf{w}\ _1 + \frac{1}{2}(1 - \alpha)\ \mathbf{w}\ _2^2$
Total Variation (TV) [65]	$\sum_{j=2}^p w_j - w_{j-1} $
uogólniony TV [10]	$\ \mathbf{F}\mathbf{w}\ _1$
fused LASSO [76]	$\alpha\ \mathbf{w}\ _1 + (1 - \alpha) \sum_{j=2}^p w_j - w_{j-1} $
SCAD [20]	$\sum f(w_j), f(x) = \begin{cases} x, & x \leq \lambda \\ \frac{x^2 - 2a\lambda x + \lambda^2}{2(1-a)\lambda}, & \lambda < x \leq a\lambda \\ \frac{\lambda(1+a)}{2}, & a\lambda < x \end{cases}$

W uczeniu maszynowym ten sposób regularyzacji, w kontekście regresji liniowej, funkcjonuje pod angielską nazwą ridge regression [29], a rozwiązanie dane jest wzorem:

$$\mathbf{w} = (\mathbf{X}^T \mathbf{X} + \lambda \mathbf{I})^{-1} \mathbf{X}^T \mathbf{y}, \quad (2.14)$$

gdzie $\mathbf{I}$ to macierz jednostkowa.

Do znalezienia minimum funkcji (2.13) w przypadku regresji logistycznej można posłużyć się dowolną gradientową metodą optymalizacji, zarówno pierwszego jak i drugiego rzędu [54]. W metodzie Newtona w postaci macierzowej wzór na kierunek poszukiwania minimum opisuje poniższe równanie:

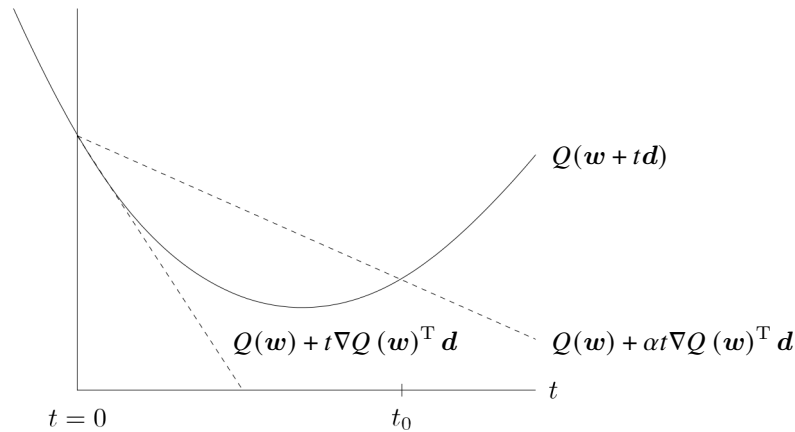
$$\mathbf{d} = (\mathbf{X}^T \mathbf{D} \mathbf{X} + \lambda \mathbf{I})^{-1} (\mathbf{X}^T \mathbf{v} - \lambda \mathbf{w}), \quad (2.15)$$

a nowy punkt to $\mathbf{w}^{(k+1)} = \mathbf{w}^{(k)} + t\mathbf{d}$, gdzie t stanowi krok znaleziony w wyniku minimalizacji kierunkowej. Samą metodę Newtona można przerwać, gdy $\sqrt{-\nabla Q(\mathbf{w}^{(k)})^T \mathbf{d}} \leq \epsilon$, gdzie $\nabla Q(\mathbf{w}^{(k)})^T \mathbf{d}$ to pochodna kierunkowa w k -tej iteracji, a ϵ to dokładność. Jeżeli pochodna kierunkowa jest dodatnia, to kierunek $\mathbf{d}$ nie wskazuje kierunku minimalizacji. W części eksperymentalnej powyższy model nazwano `l2_logreg_ordinary`.

Krok t można znajdować w taki sposób, by wyrażenie $Q(\mathbf{w} + t\mathbf{d})$ osiągało dokładne minimum (tzw. dokładna minimalizacja kierunkowa¹):

$$t = \arg \min_{s \geq 0} Q(\mathbf{w} + s\mathbf{d}). \quad (2.16)$$

¹ang. exact linesearch



Rysunek 2.1: Wrotna minimalizacja kierunkowa. Krzywa przedstawia funkcję kosztu Q ograniczoną do obszaru poszukiwań. Dolna przerywana linia reprezentuje liniową ekstrapolację funkcji, natomiast górna ma współczynnik nachylenia pomniejszony wprost proporcjonalnie do wartości α . Krok z przedziału $[0, t_0]$ spełnia warunek stopu wrotnej minimalizacji kierunkowej [9].

Podejście to można stosować szczególnie wtedy, gdy znalezienie minimum tejże funkcji jednej zmiennej ma niski koszt w porównaniu z kosztem wyznaczenia kierunku poszukiwań [9].

Gdy koszt obliczenia $Q(w + td)$ jest duży, w praktyce wystarczy jedynie przybliżona minimalizacja². Jedną z takich metod jest wrotna minimalizacja kierunkowa³, która zależy od dwóch stałych: $\alpha \in (0; 0,5)$ i $\beta \in (0, 1)$. Nazwa wywodzi się z faktu, że zaczynając od kroku jednostkowego zmniejszamy go proporcjonalnie do wartości czynnika β , aż będzie spełniony warunek stopu $Q(w + td) < Q(w) + \alpha t \nabla Q(w)^T d$. Z faktu, że d jest kierunkiem spadku, mamy $\nabla Q(w)^T d < 0$, więc dla dostatecznie małego t wynika

$$Q(w + td) \approx Q(w) + t \nabla Q(w)^T d < Q(w) + \alpha t \nabla Q(w)^T d,$$

co pokazuje, że wrotna minimalizacja kierunkowa w końcu się zatrzyma [9]. Podejście to podsumowuje algorytm 1. Wyliczając wcześniej wartość iloczynu Xd koszt obliczenia $Q(w + td)$ i $\nabla Q(w + td)^T d$ dla dowolnego t ma złożoność rzędu $O(n)$ [22]. W części eksperymentalnej modele, podczas uczenia których wykorzystano dokładną minimalizację kierunkową, opatrzone w nazwie symbolem e+, natomiast e- oznacza, że użyto przybliżoną minimalizację kierunkową w toku optymalizacji funkcji celu.

Do znalezienia kierunku d nie trzeba w sposób jawny odwracać macierz Hessego

²ang. inexact linesearch

³ang. backtracking line search

Algorytm 1 Wrotna minimalizacja kierunkowa.**Wejście:** kierunek spadku $\mathbf{d}$, $\alpha \in (0; 0,5)$, $\beta \in (0, 1)$. $t \leftarrow 1$ **while** $Q(\mathbf{w} + t\mathbf{d}) > Q(\mathbf{w}) + \alpha t \nabla Q(\mathbf{w})^T \mathbf{d}$ **do** $t \leftarrow \beta t$ **end while**

— zamiast tego stosuje się iteracyjną metodę gradientu sprzężonego w celu rozwiązania układu równań:

$$\left(\mathbf{X}^T \mathbf{D} \mathbf{X} + \lambda \mathbf{I} \right) \mathbf{d} = \mathbf{X}^T \mathbf{v} - \lambda \mathbf{w}, \quad (2.17)$$

ponieważ macierz układu równań jest dodatnio określona [5, 31]. Tempo zbieżności metody zależy od widma macierzy, zatem w celu zwiększenia tempa należy przekształcić układ równań do równoważnej postaci, w której poprawieniu uległo widmo — uzyskuje się to za pomocą pewnej macierzy⁴ $\mathbf{M}$, która w jakimś stopniu aproksymuje macierz $\mathbf{H}$, zaś rozwiązać należy układ równań

$$\mathbf{M}^{-1} \mathbf{H} \mathbf{d} = -\mathbf{M}^{-1} \mathbf{g}, \quad (2.18)$$

mający takie samo rozwiązanie, co $\mathbf{H} \mathbf{d} = -\mathbf{g}$, jednak w tym przypadku macierz $\mathbf{M}^{-1} \mathbf{H}$ posiada korzystniejsze spektrum. Im lepiej macierz $\mathbf{M}$ aproksymuje macierz $\mathbf{H}$, tym większe tempo zbieżności, ale koszt obliczenia $\mathbf{M}^{-1}$ wzrasta, dlatego znaleźć trzeba kompromis między dokładnością a złożonością. W najprostszym przypadku można za $\mathbf{M}$ przyjąć macierz diagonalną, która zawiera elementy z przekątnej macierzy $\mathbf{H}$ (tzw. preconditioner Jacobiego). Bardziej wyrafinowane postaci macierzy $\mathbf{M}$ opierają się na niekompletnej faktoryzacji [5, 66]. W części eksperymentalnej modele używające tego udoskonalenia mają w nazwie dopisek p+, a w przypadku $\mathbf{M} = \mathbf{I}$ do nazwy dopisano p-.

Szkic metody gradientu sprzężonego prezentuje algorytm 2. Metoda ta wymaga tylko procedury obliczającej iloczyn macierzy układu równań przez dowolny wektor $\mathbf{s}$ — opierając się na prawie łączności iloczynu można go rozpisać jako trzy iloczyny macierzowo-wektorowe w przypadku regresji logistycznej z regularyzacją ℓ^2 :

$$\mathbf{H} \mathbf{s} = \mathbf{X}^T (\mathbf{D} (\mathbf{X} \mathbf{s})) + \lambda \mathbf{s}. \quad (2.19)$$

Ponadto w początkowych iteracjach wystarczy jedynie zgrubne wyznaczanie kierunku — takie podejście prowadzi do „truncated Newton” [52].

⁴W języku angielskim o takiej macierzy mówi się preconditioner [25].

Algorytm 2 Metoda gradientu sprzężonego [5].

Wejście: $A, b, x^{(0)}$ (wybrany dowolnie, np. 0), $\epsilon, i_{\max}$

$r^{(0)} \leftarrow b - Ax^{(0)}$.

Rozwiąż układ równań $Mz^{(0)} = r^{(0)}$.

$\rho^{(0)} \leftarrow \langle r^{(0)}, z^{(0)} \rangle$.

$p^{(1)} \leftarrow z^{(0)}$.

for $i = 1, 2, \dots$ **do**

$\alpha^{(i)} \leftarrow \frac{\rho^{(i-1)}}{\langle p^{(i)}, Ap^{(i)} \rangle}$.

$x^{(i)} \leftarrow x^{(i-1)} + \alpha^{(i)} p^{(i)}$.

$r^{(i)} \leftarrow r^{(i-1)} - \alpha^{(i)} Ap^{(i)}$.

if $\|r^{(i)}\|_2 \leq \epsilon \|b\|_2 \vee i \geq i_{\max}$ **then**

 Zakończ.

end if

Rozwiąż układ równań $Mz^{(i)} = r^{(i)}$.

$\rho^{(i)} \leftarrow \langle r^{(i)}, z^{(i)} \rangle$.

$\beta^{(i)} \leftarrow \frac{\rho^{(i)}}{\rho^{(i-1)}}$.

$p^{(i+1)} \leftarrow z^{(i)} + \beta^{(i)} p^{(i)}$.

end for

Wyjście: rozwiązanie układu równań x .

Problem można także wyrazić w formie ważonych najmniejszych kwadratów, a następnie iteracyjnie rozwiązywać w celu uzyskania zbieżności [27]. Sprowadzenie do problemu najmniejszych kwadratów umożliwia też zastosowanie sztuczki z przekształceniem jądrowym [63].

Ponadto można z powodzeniem stosować metody gradientowe pierwszego rzędu do optymalizacji funkcji (2.13) [16, 68].

Włączając wyraz wolny w_0 równanie określające kierunek poszukiwań przybierze następującą postać:

$$\begin{bmatrix} d_0 \\ d \end{bmatrix} = \begin{bmatrix} \mathbf{1}^T D \mathbf{1} & \mathbf{1}^T D X \\ X^T D \mathbf{1} & X^T D X + \lambda I \end{bmatrix}^{-1} \cdot \begin{bmatrix} \mathbf{1}^T v \\ X^T v - \lambda w \end{bmatrix}. \quad (2.20)$$

Ponownie iloczyn powyższej macierzy i dowolnego wektora można rozisać jako trzy iloczyny macierzowo-wektorowe (opierając się na prawie łączności iloczynu), tym samym unikając jawnego formowania macierzy. W części eksperymentalnej modele zawierające wyraz wolny mają w nazwie b+, natomiast te bez wyrazu wolnego mają b-.

2.4.1 Zastosowanie rozkładu QR

W wyniku rozkładu QR macierzy rzeczywistej $\mathbf{X}_{n \times p}$ ($n \geq p$) o pełnym rzędzie powstaje macierz ortogonalna $\mathbf{Q}_{n \times n}$ i macierz górnotrójkątna $\mathbf{R}_{n \times p}$:

$$\mathbf{X} = \mathbf{Q}\mathbf{R} = \begin{bmatrix} \mathbf{Q}_1 & \mathbf{Q}_2 \end{bmatrix} \cdot \begin{bmatrix} \mathbf{R}_1 \\ \mathbf{0} \end{bmatrix} = \mathbf{Q}_1 \mathbf{R}_1, \quad (2.21)$$

a do wyznaczenia tego rozkładu można posłużyć się np. odbiciami Householdera. Rozkład $\mathbf{Q}_1 \mathbf{R}_1$ bywa nazywany ciekim rozkładem z uwagi na mniejsze wymiary wynikowych macierzy. W przypadku $n < p$ należy skorzystać z analogicznego rozkładu LQ bądź dokonać rozkład QR macierzy transponowanej i uwzględnić transpozycję w dalszych obliczeniach. Liczba operacji zmiennoprzecinkowych wymaganych do policzenia rozkładu QR za pomocą odbić Householdera wynosi $2p^2(n - \frac{p}{3})$ [25].

W przypadku regresji liniowej bez regularyzacji możemy znaleźć formułę na współczynniki za pomocą rozkładu QR [55]:

$$\mathbf{w} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y} = (\mathbf{R}^T \mathbf{Q}^T \mathbf{Q} \mathbf{R})^{-1} \mathbf{R}^T \mathbf{Q}^T \mathbf{y} = \mathbf{R}_1^{-1} \mathbf{Q}_1^T \mathbf{y}. \quad (2.22)$$

Jeśli $n < p$, można za pomocą rozkładu LQ znaleźć rozwiązanie, które będzie miało możliwie najmniejszą normę [25]:

$$\mathbf{w} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y} = (\mathbf{Q}^T \mathbf{L}^T \mathbf{L} \mathbf{Q})^{-1} \mathbf{Q}^T \mathbf{L}^T \mathbf{y} = \mathbf{Q}_1^T \mathbf{L}_1^{-1} \mathbf{y}. \quad (2.23)$$

Dla regresji liniowej z regularyzacją, po podstawieniu $\mathbf{LQ}$ w miejsce $\mathbf{X}$ w (2.14) uzyskamy:

$$\begin{aligned} \mathbf{w} &= (\mathbf{Q}^T \mathbf{L}^T \mathbf{L} \mathbf{Q} + \lambda \mathbf{I})^{-1} \mathbf{Q}^T \mathbf{L}^T \mathbf{y} = \mathbf{Q}^T (\mathbf{L}^T \mathbf{L} + \lambda \mathbf{I})^{-1} \mathbf{L}^T \mathbf{y} = \\ &= \begin{bmatrix} \mathbf{Q}_1^T & \mathbf{Q}_2^T \end{bmatrix} \begin{bmatrix} \mathbf{L}_1^T \mathbf{L}_1 + \lambda \mathbf{I} & \mathbf{0} \\ \mathbf{0} & \lambda \mathbf{I} \end{bmatrix}^{-1} \begin{bmatrix} \mathbf{L}_1^T \mathbf{y} \\ \mathbf{0} \end{bmatrix} = \mathbf{Q}_1^T (\mathbf{L}_1^T \mathbf{L}_1 + \lambda \mathbf{I})^{-1} \mathbf{L}_1^T \mathbf{y}. \end{aligned}$$

Bazując na fakcie, że regresja logistyczna z karą ℓ^2 jest obrotowo niezmiennicza⁵ [57], można w przypadku $p \gg n$ macierz danych $\mathbf{X}$ poddać rozkładowi LQ (bądź równoważnie macierz $\mathbf{X}^T$ rozkładowi QR) i tym samym zredukować wymiar problemu do n — powrót do pierwotnej przestrzeni uzyskuje się mnożąc wynik przez macierz $\mathbf{Q}^T$ (bądź $\mathbf{Q}$ w przypadku rozkładu QR) [28].

⁵Algorytm nazywamy obrotowo niezmienniczym, gdy jego wynik działania (np. wynik klasyfikacji czy prawdopodobieństwa posteriori) są identyczne, zarówno dla danych oryginalnych, jak i dla danych obróconych, tj. przemnożonych przez macierz ortogonalną.

Algorytm 3 Uczenie modelu regresji logistycznej z regularyzacją ℓ^2 za pomocą metody Newtona wykorzystującej rozkład LQ.

Dane wejściowe: $\mathbf{X} = \mathbb{R}^{n \times p}$, $\mathbf{y}_{n \times 1} : y_i \in \{-1, +1\}$, $n < p$

Inicjalizacja: $[\mathbf{L}_1, \mathbf{Q}_1] = \text{lq}(\mathbf{X})$, $\hat{\mathbf{w}}_{n \times 1}^{(0)} = \mathbf{0}$

for $k = 0, 1, \dots, k_{\max}$ **do**

Oblicz wartości $\mathbf{v}$ (zgodnie z (2.10)) i $\mathbf{D}$ dla bieżącego $\hat{\mathbf{w}}^{(k)}$.

Rozwiąż układ równań $(\mathbf{L}_1^T \mathbf{D} \mathbf{L}_1 + \lambda \mathbf{I}) \cdot \hat{\mathbf{d}} = \mathbf{L}_1^T \mathbf{v} - \lambda \hat{\mathbf{w}}$.

if $\sqrt{\langle \mathbf{L}_1^T \mathbf{v} - \lambda \hat{\mathbf{w}}^{(k)}, \hat{\mathbf{d}} \rangle} < \epsilon$ **then**

Zakończ pętlę i zwróć wynik.

end if

$\hat{\mathbf{w}}^{(k+1)} \leftarrow \hat{\mathbf{w}}^{(k)} + \hat{\mathbf{d}} \cdot \arg \min_t Q(\hat{\mathbf{w}}^{(k)} + t \hat{\mathbf{d}})$.

end for

Wyjście: $\mathbf{w} = \mathbf{Q}_1^T \hat{\mathbf{w}}$

Podstawiając $\mathbf{L}_1$ za $\mathbf{X}$ w równaniu (2.15) uzyskujemy na mocy twierdzenia pierwszego zawartego w [28] poniższą tożsamość:

$$\mathbf{d} = \mathbf{Q}_1^T \left(\mathbf{L}_1^T \mathbf{D} \mathbf{L}_1 + \lambda \mathbf{I} \right)^{-1} \left(\mathbf{L}_1^T \mathbf{v} - \lambda \mathbf{Q}_1 \mathbf{w} \right). \quad (2.24)$$

Przemnożenie przez $\mathbf{Q}_1$ rzutuje $\mathbf{w}$ do mniejszej przestrzeni, w której następuje rozwiązanie układu równań, po czym mnożenie przez $\mathbf{Q}_1^T$ przywraca wynik do pierwotnej przestrzeni. Nie trzeba za każdym razem wracać do oryginalnej przestrzeni — wszystkie obliczenia można dokonać w mniejszej przestrzeni, a dopiero końcowy rezultat cofnąć do pierwotnej przestrzeni. Powyższe podejście przedstawiono w algorytmie 3.

W celu włączenia wyrazu wolnego w_0 należy postąpić analogicznie, jak miało to miejsce w równaniu (2.20), z tą różnicą, że macierz $\mathbf{X}$ zastępujemy macierzą $\mathbf{L}_1$:

$$\begin{bmatrix} d_0 \\ \hat{\mathbf{d}} \end{bmatrix} = \begin{bmatrix} \mathbf{1}^T \mathbf{D} \mathbf{1} & \mathbf{1}^T \mathbf{D} \mathbf{L}_1 \\ \mathbf{L}_1^T \mathbf{D} \mathbf{1} & \mathbf{L}_1^T \mathbf{D} \mathbf{L}_1 + \lambda \mathbf{I} \end{bmatrix}^{-1} \cdot \begin{bmatrix} \mathbf{1}^T \mathbf{v} \\ \mathbf{L}_1^T \mathbf{v} - \lambda \hat{\mathbf{w}} \end{bmatrix}. \quad (2.25)$$

Wykorzystanie rozkładu QR nie jest czymś zupełnie nowym [28, 57], jednakże zastosowanie tej sztuczki wydaje się być mało powszechne w bibliotekach uczenia maszynowego — przykładowo technika ta nie jest wykorzystywana w pakiecie `scikit-learn` [61]. Możemy jednak w prosty sposób przerobić model regresji logistycznej z regularyzacją ℓ^2 , rozszerzając oryginalną klasę i przeciążając metodę `fit`. Ilustruje to poniższy fragment kodu:

```

1 import numpy as np
2 from sklearn.linear_model import LogisticRegression
3 class LR_with_lq(LogisticRegression):
4     def __init__(self, **kwargs):
5         assert kwargs['penalty'] == 'l2'
6         super().__init__(**kwargs)
7     def fit(self, X, y):
8         if X.shape[0] < X.shape[1]:
9             Q, R = np.linalg.qr(X.T, mode = 'reduced')
10            super().fit(R.T, y)
11            self.coef_ = self.coef_ @ Q.T
12        else:
13            super().fit(X, y)
14        return self

```

Takie usprawnienie powinno dawać zysk przy liczbie próbek mniejszej niż liczba atrybutów. Co więcej, ponieważ transformacja R do X za pomocą macierzy Q ($X = QR$), jest przekształceniem ortogonalnym, to może być ona wykorzystywana we wszystkich algorytmach, które są obrotowo niezmiennicze. Krótki eksperyment ze zbiorem danych `colon` o rozmiarze 62×2000 pokazuje zysk ok. dwuipółkrotny w czasie uczenia, uwzględniając koszt rozkładu QR.

W części eksperymentalnej (rodz. 4) model ten widnieje pod nazwą `l2_logreg_qr` (gdy użyto rozkładu QR) bądź `l2_logreg_lq` (gdy użyto rozkładu LQ).

2.4.2 Zastosowanie formuły Woodbury’ego

Dzięki rozkładowi QR rozmiar rozwiązywanego problemu redukuje się do liczby próbek, a także zyskujemy stabilność numeryczną procedury uczenia, jednak w przypadku macierzy rzadkich wyznaczenie rozkładu, który zachowa pierwotną rzadkość macierzy, wymaga większego nakładu pracy. Traktowanie macierzy rzadkiej jako gęstej może również nie przynieść zysku obliczeniowego. Mimo to efekt redukcji wielkości problemu można uzyskać po zastosowaniu formuły Woodbury’ego, pozwalającej obliczyć odwrotność sumy dwóch macierzy [30, 84]:

$$(A + UBV)^{-1} = A^{-1} - A^{-1}U \left(B^{-1} + VA^{-1}U \right)^{-1} VA^{-1}. \quad (2.26)$$

Podstawiając $A = \lambda I$, $U = X^T$, $B = D$, $V = X$ otrzymamy:

$$\left(\lambda I + X^T D X \right)^{-1} = \frac{1}{\lambda} I - \frac{1}{\lambda^2} X^T \left(D^{-1} + \frac{1}{\lambda} X X^T \right)^{-1} X. \quad (2.27)$$

Macierz $\frac{1}{\lambda} \mathbf{X} \mathbf{X}^T$ nie zależy w ogóle od parametrów modelu $\mathbf{w}$, więc przed rozpoczęciem optymalizacji należy ją policzyć, by skrócić czas rozwiązywania układu równań — z uwagi na jej symetryczność wystarczy przechowywać górny lub dolny trójkąt macierzy, co pozwoli dwukrotnie zmniejszyć zajętość w pamięci. Koszt wyznaczenia tej macierzy jest rzędu $O(n^2 p)$, natomiast w przypadku macierzy rzadkiej koszt ten będzie mniejszy proporcjonalnie do stopnia rzadkości. Trzeba też uważać na wartości bliskie zera na przekątnej macierzy $\mathbf{D}$, gdy prawdopodobieństwo przynależności i -tej próbki do klasy pozytywnej lub negatywnej dąży do zera lub jedności — można tego uniknąć prostym zabezpieczeniem: $\min(\max(d_{i,i}, \epsilon), 1 - \epsilon)$, przyjmując np. $\epsilon = 10^{-5}$ [21].

Wstawiając (2.27) do (2.15) i porządkując wszystko otrzymamy następujące równanie określające kierunek poszukiwań:

$$\mathbf{d} = \frac{1}{\lambda} \mathbf{X}^T (\mathbf{v} - \mathbf{z}) - \mathbf{w}, \quad (2.28)$$

gdzie $\mathbf{z}$ stanowi rozwiązanie poniższego układu równań:

$$\left(\mathbf{D}^{-1} + \frac{1}{\lambda} \mathbf{X} \mathbf{X}^T \right) \mathbf{z} = \frac{1}{\lambda} \mathbf{X} \mathbf{X}^T \mathbf{v} - \mathbf{X} \mathbf{w}. \quad (2.29)$$

Pomimo stosunkowo niedużego rozmiaru macierzy, metoda dokładna rozwiązywania układu równań może okazać się bardziej kosztowna od metody iteracyjnej — koszt policzenia rozkładu Cholesky’ego wymaga $\frac{n^3}{3}$ operacji zmiennoprzecinkowych, natomiast metoda iteracyjna wymaga w przybliżeniu kn^2 operacji, gdzie k oznacza liczbę iteracji. Uzyskanie satysfakcjonującego rozwiązania wymaga zazwyczaj mniejszej liczby iteracji niż n . Pewien problem ze zbieżnością metody iteracyjnej może tutaj powodować złe warunkowanie macierzy, więc nieodzowne może okazać się zastosowanie preconditioningu.

Aby znaleźć wzór na kierunek poszukiwań uwzględniający wyraz wolny, wprowadźmy $\mathbf{b} = \mathbf{X}^T \mathbf{D} \mathbf{1}$, $c = \mathbf{1}^T \mathbf{D} \mathbf{1}$ i zapiszmy (2.20) w postaci układu równań:

$$\begin{cases} c d_0 + \mathbf{b}^T \mathbf{d} &= \mathbf{1}^T \mathbf{v} \\ \mathbf{b} d_0 + (\mathbf{X}^T \mathbf{D} \mathbf{X} + \lambda \mathbf{I}) \mathbf{d} &= \mathbf{X}^T \mathbf{v} - \lambda \mathbf{w} \end{cases} \quad (2.30)$$

Wyznaczając z pierwszego równania d_0 i wstawiając do drugiego równania otrzymamy:

$$\begin{cases} d_0 &= \frac{\mathbf{1}^T \mathbf{v} - \mathbf{b}^T \mathbf{d}}{c} \\ \mathbf{d} &= (\mathbf{X}^T \mathbf{D} \mathbf{X} + \lambda \mathbf{I} - \frac{1}{c} \mathbf{b} \mathbf{b}^T)^{-1} \left(\mathbf{X}^T \mathbf{v} - \lambda \mathbf{w} - \frac{\mathbf{1}^T \mathbf{v}}{c} \mathbf{b} \right) \end{cases} \quad (2.31)$$

Podobnie jak w (2.27), podstawiając $\mathbf{A} = \lambda \mathbf{I} - \frac{1}{c} \mathbf{b} \mathbf{b}^T$ i korzystając z formuły Shermana-Morrisona do policzenia $\mathbf{A}^{-1}$ [70], a następnie z formuły Woodbury’ego do policzenia

$(\mathbf{X}^T \mathbf{D} \mathbf{X} + \lambda \mathbf{I} - \frac{1}{c} \mathbf{b} \mathbf{b}^T)^{-1}$ i wstawiając wynik do drugiego równania w (2.31) ostatecznie otrzymamy (dla czytelności wprowadzono zmienne pomocnicze):

$$\begin{aligned} \mathbf{s} &= \frac{1}{\lambda} \mathbf{X} \mathbf{X}^T \mathbf{D} \mathbf{1}, \\ \gamma &= \frac{\langle \mathbf{D} \mathbf{1}, \frac{1}{\lambda} \mathbf{X} \mathbf{X}^T \mathbf{v} - \mathbf{X} \mathbf{w} \rangle - \mathbf{1}^T \mathbf{v}}{c - \langle \mathbf{D} \mathbf{1}, \mathbf{s} \rangle}, \\ \mathbf{z} &= \left(\mathbf{D}^{-1} + \frac{1}{\lambda} \mathbf{X} \mathbf{X}^T + \frac{\mathbf{s} \mathbf{s}^T}{c - \langle \mathbf{D} \mathbf{1}, \mathbf{s} \rangle} \right)^{-1} \left(\frac{1}{\lambda} \mathbf{X} \mathbf{X}^T \mathbf{v} - \mathbf{X} \mathbf{w} + \gamma \mathbf{s} \right), \\ \mathbf{d} &= \frac{1}{\lambda} \mathbf{X}^T \left(\mathbf{v} - \mathbf{z} + \left(\gamma - \frac{\langle \mathbf{z}, \mathbf{s} \rangle}{c - \langle \mathbf{D} \mathbf{1}, \mathbf{s} \rangle} \right) \mathbf{D} \mathbf{1} \right) - \mathbf{w}. \end{aligned} \quad (2.32)$$

Mając $\mathbf{d}$ można z łatwością policzyć d_0 zgodnie z (2.31). Należy pamiętać, że do iloczynu $\mathbf{X} \mathbf{w}$ nie dodano w_0 . Włączenie wyrazu wolnego wiąże się z dodatkowym kosztem policzenia wektora $\mathbf{s}$ i kilkoma iloczynami skalarnymi. Dodatkowo macierz $\mathbf{X}^T \mathbf{D} \mathbf{X} + \lambda \mathbf{I} - \frac{1}{c} \mathbf{b} \mathbf{b}^T$ może być nieokreślona, dlatego zamiast metody gradientu sprzężonego do rozwiązania powyższego układu równań należy użyć metodę, która działa poprawnie w przypadku macierzy nieokreślonych, np. metoda MINRES bądź SYMMLQ [60].

W części eksperymentalnej zbadano dwa modele — w pierwszym obliczono i zapamiętano iloczyn $\mathbf{X} \mathbf{X}^T$, by wykorzystać go w kolejnych iteracjach (ten model nazwano `l2_logreg_wdbr`), a w drugim nie skorzystano z tej optymalizacji (ten z kolei nazwano `l2_logreg_woodbury`).

2.5 Regularyzacja ℓ^1

Idea zastosowania normy ℓ^1 w zadaniu regularyzacji i przycinania pojawiła się w latach 90. [83], a później ukuto dla niej nazwę „lasso” w kontekście regresji liniowej [75]. W przypadku regresji liniowej optymalizowana funkcja kosztu wygląda następująco:

$$Q(\mathbf{w}) = \frac{1}{2} \sum_{i=1}^n (y_i - \langle \mathbf{w}, \mathbf{x}_i \rangle)^2 + \lambda \sum_{j=1}^p |w_j|, \quad (2.33)$$

natomiast w przypadku regresji logistycznej funkcja kosztu ma postać:

$$Q(\mathbf{w}) = \sum_{i=1}^n \log \left(1 + e^{-y_i \cdot \langle \mathbf{w}, \mathbf{x}_i \rangle} \right) + \lambda \sum_{j=1}^p |w_j|. \quad (2.34)$$

Funkcja w punkcie $w_j = 0$ posiada ostrze (pochodna nie istnieje), co przyczynia się do powstawania rzadkich modeli (część parametrów modelu zostaje wyzerowana). Brak gładkości uniemożliwia bezpośrednie zastosowanie metod gradientowych, jednakże wprowadzenie pewnych modyfikacji niweluje to ograniczenie.

2.5.1 Subgradient

Zrózniczkowanie funkcji kosztu zawierającej karę ℓ^1 względem zmiennej w_j prowadzi do następujących przypadków:

$$\frac{\partial Q(\mathbf{w})}{\partial w_j} = \begin{cases} \frac{\partial L(\mathbf{w})}{\partial w_j} + \lambda, & w_j > 0 \\ \frac{\partial L(\mathbf{w})}{\partial w_j} - \lambda, & w_j < 0 \\ \frac{\partial L(\mathbf{w})}{\partial w_j} - \lambda, & w_j = 0 \wedge \frac{\partial L(\mathbf{w})}{\partial w_j} > \lambda \\ \frac{\partial L(\mathbf{w})}{\partial w_j} + \lambda, & w_j = 0 \wedge \frac{\partial L(\mathbf{w})}{\partial w_j} < -\lambda \\ 0, & \text{wpr.} \end{cases} \quad (2.35)$$

Jeśli $L(\mathbf{w})$ jest funkcją kwadratową, to pochodna funkcji $Q(\mathbf{w})$ jest funkcją kawałkami liniową (ang. piecewise linear). Dzięki tej obserwacji można skonstruować wydajny algorytm znajdujący minimum funkcji $Q(\mathbf{w})$, którego złożoność odpowiada złożoności zwyczajnej (bez regularyzacji) metody najmniejszych kwadratów [18]. Analogiczne podejście można zastosować w przypadku regresji logistycznej, formułując funkcję straty w postaci ważonej funkcji kwadratowej [49].

Wykorzystując fakt, że pochodna istnieje poza zerem (czyli poza osiami układu współrzędnych), można znaleźć minimum lokalnej aproksymacji kwadratowej, a następnie dokonać projekcji znalezionego rozwiązania na osie układu współrzędnych, aby część parametrów, które mają być rzeczywiście wyzerowane, zostały faktycznie wyzerowane [3, 26]. Alternatywne metody iteracyjnego ściągania-progowania można znaleźć w pracach [7, 11].

2.5.2 Spadek względem współrzędnych

Modyfikacja polega na optymalizacji jednego parametru (lub bloku parametrów) przy ustalonych pozostałych parametrach [21, 79]. Optymalizacji podlega albo bezpośrednio funkcja kosztu, albo jej lokalna aproksymacja kwadratowa [87]. W przypadku regresji liniowej optymalizacja pojedynczego parametru sprowadza się do miękkiego progowania⁶:

$$w_j^* \leftarrow \frac{\mathcal{S}_\lambda \left(\sum_i x_{i,j} \cdot \left(y_i - \sum_{k \neq j} w_k \cdot x_{i,k} \right) \right)}{\sum_i x_{i,j}^2}, \quad (2.36)$$

gdzie $\mathcal{S}_\lambda(z)$ jest operatorem miękkiego progowania:

$$\mathcal{S}_\lambda(z) = \text{sgn}(z) \cdot \max(|z| - \lambda, 0). \quad (2.37)$$

⁶ang. soft thresholding

W tej postaci koszt korekty pojedynczego współczynnika wynosi $O(np)$. Można zauważyć, że:

$$\begin{aligned} y_i - \sum_{k \neq j} w_k \cdot x_{i,k} &= y_i - \langle \mathbf{w}, \mathbf{x}_i \rangle + w_j \cdot x_{i,j} \\ &= r_i + w_j \cdot x_{i,j}, \end{aligned} \quad (2.38)$$

gdzie $\mathbf{r}$ jest wektorem reszt [21]. Tym sposobem formuła (2.36) upraszcza się:

$$w_j^* \leftarrow \frac{\mathcal{S}_\lambda \left(\sum_i x_{i,j} \cdot r_i + w_j \cdot \sum_i x_{i,j}^2 \right)}{\sum_i x_{i,j}^2}, \quad (2.39)$$

redukując złożoność do $O(n)$. Wartości $\sum_i x_{i,j}^2$ można policzyć raz i zapamiętać. Korekta i -tej współrzędnej wektora reszt wygląda następująco:

$$r_i \leftarrow r_i + (w_j - w_j^*) \cdot x_{i,j}. \quad (2.40)$$

W przypadku regresji logistycznej, aproksymowanej lokalnie funkcją kwadratową, formuła na korektę j -tego parametru ma postać:

$$w_j^* \leftarrow \frac{\mathcal{S}_\lambda \left(\sum_i x_{i,j} \cdot r_i + w_j \cdot \sum_i d_{i,i} \cdot x_{i,j}^2 \right)}{\sum_i d_{i,i} \cdot x_{i,j}^2}. \quad (2.41)$$

Na początku $\mathbf{r} = \mathbf{v}$ (2.8), $d_{i,i}$ to i -ty element macierzy diagonalnej $\mathbf{D}$ (2.9), zaś korekta i -tej współrzędnej wektora reszt wygląda następująco:

$$r_i \leftarrow r_i + (w_j - w_j^*) \cdot d_{i,i} \cdot x_{i,j}. \quad (2.42)$$

Aby dalej usprawnić proces minimalizacji, należy na początku wyznaczyć zbiór aktywnych (niezerowych) parametrów, a następnie operować jedynie na nich. Jeśli w trakcie obliczeń dany parametr osiągnie wartość 0, to jest on usuwany ze zbioru [21, 86]. Ponadto mieszanie kolejności optymalizowania poszczególnych parametrów modelu również zazwyczaj sprzyja skracaniu czasu uczenia [86].

Po znalezieniu minimum trzeba udać się do nowego punktu, obliczyć nowe wartości wag $\mathbf{v}$ i $\mathbf{D}$ i powtórzyć cały proces — podsumowuje to algorytm 4. Jako nowy punkt można wziąć znalezione minimum $\mathbf{w}^*$ bądź punkt leżący między poprzednim minimum $\mathbf{w}^{(k)}$ a znalezionym, minimalizujący wyrażenie (2.34). Eksperymenty pokazują, że to drugie podejście przeważa nad pierwszym [86].

W początkowych iteracjach nie trzeba dokładnie rozwiązywać problemu, toteż warunek stopu dla wewnętrznej pętli metody spadku względem współrzędnych powinien

Algorytm 4 Uczenie modelu regresji logistycznej z regularyzacją ℓ^1 za pomocą metody spadku względem współrzędnych.

Dane wejściowe: $X = \mathbb{R}^{n \times p}$, $y_{n \times 1}$

Inicjalizacja: $w^{(0)} \leftarrow 0$, $\delta \leftarrow 1$, $k \leftarrow 0$

repeat

Oblicz wartości g , D , r dla bieżącego $w^{(k)}$.

Wyznacz zbiór $\mathcal{J}$ indeksów aktywnych (niezerowych) parametrów.

$w^* \leftarrow w^{(k)}$.

$n_{\text{it}} \leftarrow 0$.

repeat

$\Delta_{\max} \leftarrow 0$.

Wymieszaj kolejność indeksów w zbiorze $\mathcal{J}$.

for all $j \in \mathcal{J}$ **do**

$w_{\text{opt}} \leftarrow \frac{S_{\lambda}(\sum_i x_{i,j} \cdot r_i + w_j^* \cdot \sum_i d_{i,i} \cdot x_{i,j}^2)}{\sum_i d_{i,i} \cdot x_{i,j}^2}$.

$r \leftarrow r + (w_j^* - w_{\text{opt}}) \cdot D \cdot x_{:,j}$.

$\Delta_{\max} \leftarrow \max(\Delta_{\max}, |w_j^* - w_{\text{opt}}|)$.

$w_j^* \leftarrow w_{\text{opt}}$.

end for

$n_{\text{it}} \leftarrow n_{\text{it}} + 1$.

until $\Delta_{\max} < \delta$

if $n_{\text{it}} = 1$ **then**

$\delta \leftarrow \frac{\delta}{4}$.

end if

$\beta \leftarrow \arg \min_t Q(t w^* + (1 - t) w^{(k)})$.

$w^{(k+1)} = \beta w^* + (1 - \beta) w^{(k)}$.

$k \leftarrow k + 1$.

until $\|\partial Q\|_2 < \epsilon$

się zmieniać w razie potrzeby (gdy po jednym obiegu pętli maksymalna zmiana $|\Delta_{\max}|$ okazuje się mniejsza niż początkowy próg δ , to jest on zmniejszany czterokrotnie [86]).

W części eksperymentalnej przetestowano wersję algorytmu znajdującą się w bibliotece `liblinear` (nazwano ją `l1_logreg_liblinear`).

2.5.3 Przekształcenie w problem gładkiej optymalizacji

Inne podejście opiera się na zastąpieniu normy ℓ^1 liniowymi ograniczeniami. Zrealizować można to dwojako:

- zastępując zmienną nieograniczoną w_j parą zmiennych ograniczonych $w_j^+ - w_j^-$ [69]:

$$\begin{cases} \sum_{i=1}^n \log \left(1 + e^{-y_i \langle \mathbf{w}^+ - \mathbf{w}^-, \mathbf{x}_i \rangle} \right) + \lambda \sum_{j=1}^p w_j^+ + w_j^- \rightarrow \min \\ w_j^+ \geq 0, w_j^- \geq 0, j = 1, \dots, p \end{cases} \quad (2.43)$$

Powyższy problem z ograniczeniami można rozwiązać np. za pomocą algorytmu L-BFGS-B [88]. W części eksperymentalnej model ten umieszczono pod nazwą `l1_logreg_lbfsgb`.

- wprowadzając zmienną u_j stanowiącą ograniczenie zmiennej w_j [46]:

$$\begin{cases} \sum_{i=1}^n \log \left(1 + e^{-y_i \langle \mathbf{w}, \mathbf{x}_i \rangle} \right) + \lambda \sum_{j=1}^p u_j \rightarrow \min \\ -u_j \leq w_j \leq u_j, j = 1, \dots, p \end{cases} \quad (2.44)$$

Niegładkie ograniczenie $-u_j \leq w_j \leq u_j, j = 1, \dots, p$ można zastąpić gładką i wypukłą funkcją logarytmiczną, uzyskując w ten sposób problem bez ograniczeń, który następnie można rozwiązać za pomocą metody Newtona [46]. W części eksperymentalnej pod nazwą `l1_logreg_ipm` kryje się ten model.

Rozbicie pierwotnego problemu na dwa podproblemy

Kolejny sposób oparty na optymalizacji z ograniczeniami polega na sformułowaniu równoważnego problemu, w którym część gładka i niegładka są rozdzielone — mianowicie problemowi pierwotnemu

$$\min_{\mathbf{w}} L(\mathbf{w}) + \lambda \|\mathbf{w}\|_1 \quad (2.45)$$

odpowiada problem z ograniczeniem równościowym

$$\begin{cases} L(\mathbf{w}) + \lambda \|\mathbf{z}\|_1 \rightarrow \min \\ \mathbf{w} - \mathbf{z} = 0 \end{cases}, \quad (2.46)$$

który można rozwiązać przy użyciu algorytmu ADMM⁷ [10, 85]:

$$\mathbf{w}^{(k+1)} \leftarrow \arg \min_{\mathbf{w}} \left(L(\mathbf{w}) + \frac{\rho}{2} \left\| \mathbf{w} - \mathbf{z}^{(k)} + \mathbf{u}^{(k)} \right\|_2^2 \right), \quad (2.47)$$

$$\mathbf{z}^{(k+1)} \leftarrow S_{\frac{\lambda}{\rho}} \left(\mathbf{w}^{(k+1)} + \mathbf{u}^{(k)} \right), \quad (2.48)$$

$$\mathbf{u}^{(k+1)} \leftarrow \mathbf{u}^{(k)} + \mathbf{w}^{(k+1)} - \mathbf{z}^{(k+1)}. \quad (2.49)$$

⁷ang. alternating direction method of multipliers

gdzie $\rho > 0$. Minimalizacja funkcjonału wyrażonego za pomocą operatora bliskości⁸ we wzorze (2.47) odpowiada rozwiązywaniu problemu z regularyzacją ℓ^2 . Formuła (2.48) minimalizuje wyrażenie $\lambda \|z\|_1 + \frac{\rho}{2} \|w^{(k+1)} - z + u^{(k)}\|_2^2$. Na końcu korekcie ulega przeskalowana zmienna dualna u i całość powtarza się aż do spełnienia warunku stopu [10].

W celu rozwiązania (2.47) można ponownie posłużyć się rozkładem QR lub formułą Woodbury'ego. Dla rozkładu QR równanie (2.24) przybierze formę

$$d = Q_1^T \left(L_1^T D L_1 + \rho I \right)^{-1} \left(L_1^T v - \rho Q_1 \left(w - z^{(k)} + u^{(k)} \right) \right), \quad (2.50)$$

a po transformacji wektorów w , $z^{(k)}$ i $u^{(k)}$ do mniejszej przestrzeni można je zapisać następująco:

$$\hat{d} = \left(L_1^T D L_1 + \rho I \right)^{-1} \left(L_1^T v - \rho \left(\hat{w} - \hat{z}^{(k)} + \hat{u}^{(k)} \right) \right). \quad (2.51)$$

Z kolei dla formuły Woodbury'ego równanie (2.28) przyjmie postać

$$d = \frac{1}{\rho} X^T (v - \gamma) - \left(w - z^{(k)} + u^{(k)} \right), \quad (2.52)$$

gdzie γ stanowi rozwiązanie poniższego układu równań:

$$\left(D^{-1} + \frac{1}{\rho} X X^T \right) \gamma = \frac{1}{\rho} X X^T v - X \left(w - z^{(k)} + u^{(k)} \right). \quad (2.53)$$

Z uwagi na konieczność rzutowania do pierwotnej przestrzeni, rozwiązanie oparte na rozkładzie QR będzie mniej satysfakcjonujące pod względem złożoności obliczeniowej niż to oparte na formule Woodbury'ego, dlatego w części eksperymentalnej skupiono się na tym drugim podejściu.

Zastąpienie normy ℓ^1 jej górnym ograniczeniem

Jeszcze inne podejście opiera się na zastąpieniu normy ℓ^1 jej górnym ograniczeniem za pomocą funkcji kwadratowej [20, 47]:

$$\|w\|_1 \leq \frac{1}{2} \sum_{j=1}^p \left(\frac{w_j^2}{|w'_j|} + |w'_j| \right), \quad (2.54)$$

gdzie powyższa nierówność jest prawdziwa dla dowolnego w' , a równość zachodzi wtedy i tylko wtedy, gdy $w' = w$. Zastępując czynnik karzący w równaniu (2.34) otrzymujemy:

$$Q(w) = \sum_{i=1}^n \log \left(1 + e^{-y_i \langle w, x_i \rangle} \right) + \frac{\lambda}{2} \sum_{j=1}^p \left(\frac{w_j^2}{|w'_j|} + |w'_j| \right). \quad (2.55)$$

⁸ang. proximity operator

Kierunek poszukiwań d określa poniższe równanie:

$$d = \left(X^T D X + E \right)^{-1} \left(X^T v - E w \right), \quad (2.56)$$

gdzie E jest macierzą diagonalną i $e_{j,j} = \frac{\lambda}{|w_j|}$ (przyjmując, że $w' = w$). Początkowy wektor w musi mieć niezerowe składowe, ponieważ raz wyzerowana składowa pozostanie zerem, a jeśli w trakcie optymalizacji wartość składowej $|w_j|$ spadnie poniżej precyzji maszynowej, to można ją bezpiecznie wyzerować. Dobry punkt startowy stanowi rozwiązanie problemu (2.13) [36]. Można także zacząć od wektora zerowego i znaleźć minimum w kierunku wskazanym przez antysubgradient, dzięki temu część współczynników może już na początku zostać wyeliminowana, co przyspieszy początkowe iteracje — ten sposób inicjalizacji wykorzystano w eksperymencie.

Formuła (2.56) nie jest stabilna numerycznie, ponieważ w mianowniku mogą pojawić się wartości bliskie zera, jednakże istnieją warianty formuły stabilnej numerycznie [30, 47]:

$$d = -E^{-\frac{1}{2}} \left(E^{-\frac{1}{2}} X^T D X E^{-\frac{1}{2}} + I \right)^{-1} E^{-\frac{1}{2}} \cdot \partial Q \quad (2.57)$$

$$= \left(I - E^{-1} X^T \left(X E^{-1} X^T + D^{-1} \right)^{-1} X \right) \cdot \left(E^{-1} X^T v - w \right), \quad (2.58)$$

gdzie ∂Q oznacza subgradient policzony zgodnie ze wzorem (2.35) z tą różnicą, że dla $w_j = 0$ mamy $\frac{\partial Q(w)}{w_j} = 0$. W przypadku (2.57) macierz układu równań ma wymiary $p \times p$, natomiast w przypadku (2.58) macierz układu równań ma wymiary $n \times n$ (skorzystano z zależności

$$\left(X^T D X + E \right)^{-1} = E^{-1} - E^{-1} X^T \left(X E^{-1} X^T + D^{-1} \right)^{-1} X E^{-1}$$

i macierz E^{-1} wyłączono z prawej strony, po czym wmnóżono ją w antygradient funkcji). W obu wypadkach mamy do czynienia z macierzą dodatnio określoną, więc odwracanie macierzy można zastąpić rozwiązywaniem układu równań metodą gradientu sprzężonego, unikając przy tym jawnego formowania macierzy. Pochodną kierunkową dla przypadku (2.57) opisuje poniższa formuła:

$$\nabla Q(w)^T d = \langle \partial Q, d \rangle \equiv \langle -X^T v, d \rangle + \lambda \sum_{j=1}^p \text{sgn}(w_j) \cdot d_j. \quad (2.59)$$

W części eksperymentalnej model opierający się na rozwiązywaniu (2.57) nazwano `l1_logreg_ordinary`, natomiast ten oparty na (2.58) nazwano `l1_logreg_woodbury`.

Porządkując wyrazy w równaniu (2.58) można je zapisać w następującej formie:

$$\mathbf{d} = \mathbf{E}^{-1} \mathbf{X}^T \left(\mathbf{v} - \left(\mathbf{X} \mathbf{E}^{-1} \mathbf{X}^T + \mathbf{D}^{-1} \right)^{-1} \left(\mathbf{X} \mathbf{E}^{-1} \mathbf{X}^T \mathbf{v} - \mathbf{X} \mathbf{w} \right) \right) - \mathbf{w}. \quad (2.60)$$

Do rozwiązania jest układ równań:

$$\left(\mathbf{X} \mathbf{E}^{-1} \mathbf{X}^T + \mathbf{D}^{-1} \right) \mathbf{z} = \mathbf{X} \mathbf{E}^{-1} \mathbf{X}^T \mathbf{v} - \mathbf{X} \mathbf{w}. \quad (2.61)$$

Gdy $\mathbf{w}$ jest rzadki, należy ten fakt uwzględnić w obliczeniach, żeby nie wykonywać niepotrzebnych mnożeń przez 0 — iloczyn macierzowo-wektorowy $\mathbf{X} \mathbf{E}^{-1} \mathbf{X}^T \mathbf{v}$ można rozpisać następująco:

$$\mathbf{X} \mathbf{E}^{-1} \mathbf{X}^T \mathbf{v} \equiv \sum_{j=1}^p e_{j,j}^{-1} \cdot \langle \mathbf{x}_{:,j}, \mathbf{v} \rangle \cdot \mathbf{x}_{:,j}, \quad (2.62)$$

z kolei działanie $\mathbf{a} = \mathbf{E}^{-1} \mathbf{X}^T \mathbf{v}$ można zapisać następująco:

$$a_j = e_{j,j}^{-1} \cdot \langle \mathbf{x}_{:,j}, \mathbf{v} \rangle, \quad (2.63)$$

gdzie $\mathbf{x}_{:,j}$ jest j -tą kolumną macierzy $\mathbf{X}$. Tym samym dla zwiększenia wydajności obliczeń macierz danych powinna być przechowywana w porządku kolumnowym.

Formalnie pochodną kierunkową funkcji $Q(\mathbf{w})$ względem kierunku $\mathbf{d}$ opisuje następujące równanie:

$$\nabla Q(\mathbf{w})^T \mathbf{d} = -\mathbf{v}^T (\mathbf{X} \mathbf{d}) + (\mathbf{v} - \mathbf{z})^T (\mathbf{X} \mathbf{w}) - \lambda \sum_j |w_j|, \quad (2.64)$$

jednakże można z powodzeniem korzystać z równania (2.59), zaś samą optymalizację funkcji (2.34) można zakończyć w momencie spełnienia poniższego warunku:

$$\sqrt{-\nabla Q(\mathbf{w})^T \mathbf{d}} \leq \epsilon \quad (2.65)$$

W przypadku uwzględnienia wyrazu wolnego w_0 , układ równań przybierze następującą postać:

$$\begin{bmatrix} c & \mathbf{b}^T \\ \mathbf{b} & \mathbf{X}^T \mathbf{D} \mathbf{X} + \mathbf{E} \end{bmatrix} \cdot \begin{bmatrix} d_0 \\ \mathbf{d} \end{bmatrix} = \begin{bmatrix} \mathbf{1}^T \mathbf{v} \\ \mathbf{X}^T \mathbf{v} - \mathbf{E} \mathbf{w} \end{bmatrix}, \quad (2.66)$$

gdzie $\mathbf{b} = \mathbf{X}^T \mathbf{D} \mathbf{1}$, $c = \mathbf{1}^T \mathbf{D} \mathbf{1}$. Przekształcając odpowiednio otrzymamy:

$$\begin{cases} d_0 = \frac{\mathbf{1}^T \mathbf{v} - \mathbf{b}^T \mathbf{d}}{c} \\ (\mathbf{X}^T \mathbf{D} \mathbf{X} + \mathbf{E} - \frac{1}{c} \mathbf{b} \mathbf{b}^T) \mathbf{d} = \mathbf{X}^T \mathbf{v} - \mathbf{E} \mathbf{w} - \frac{\mathbf{1}^T \mathbf{v}}{c} \mathbf{b} \end{cases} \quad (2.67)$$

W tym przypadku macierz symetryczna $\mathbf{X}^T \mathbf{D} \mathbf{X} + \mathbf{E} - \frac{1}{c} \mathbf{b} \mathbf{b}^T$ może być nieokreślona, dlatego do rozwiązania powyższego układu równań należy posłużyć się np. metodą MINRES bądź SYMMLQ [60].

Formułę stabilną numerycznie można uzyskać w podobny sposób, jak miało to miejsce w przypadku (2.57):

$$\mathbf{d} = -\mathbf{E}^{-\frac{1}{2}} \left(\mathbf{E}^{-\frac{1}{2}} \mathbf{X}^T \left(\mathbf{D} - \frac{1}{c} \mathbf{D} \mathbf{1} \mathbf{1}^T \mathbf{D} \right) \mathbf{X} \mathbf{E}^{-\frac{1}{2}} + \mathbf{I} \right)^{-1} \mathbf{E}^{-\frac{1}{2}} \left(\partial \mathbf{Q} + \frac{1}{c} \mathbf{v}^T \mathbf{b} \right), \quad (2.68)$$

a do rozwiązania jest układ równań o wymiarach $p \times p$. Macierz $\mathbf{D} - \frac{1}{c} \mathbf{D} \mathbf{1} \mathbf{1}^T \mathbf{D}$ jest osobliwa, więc w celu użycia formuły Woodbury'ego wprowadźmy macierz $\hat{\mathbf{E}} = \mathbf{E} - \frac{1}{c} \mathbf{b} \mathbf{b}^T$. Korzystając z formuły Shermana-Morrisona obliczymy $\hat{\mathbf{E}}^{-1}$ [70]:

$$\hat{\mathbf{E}}^{-1} = \mathbf{E}^{-1} + \frac{\mathbf{E}^{-1} \mathbf{b} \mathbf{b}^T \mathbf{E}^{-1}}{c - \mathbf{b}^T \mathbf{E}^{-1} \mathbf{b}}. \quad (2.69)$$

Po uporządkowaniu do rozwiązania jest układ równań podobny do (2.61):

$$\left(\mathbf{X} \hat{\mathbf{E}}^{-1} \mathbf{X}^T + \mathbf{D}^{-1} \right) \mathbf{z} = \mathbf{X} \mathbf{E}^{-1} \mathbf{X}^T \mathbf{v} - \mathbf{X} \mathbf{w} + \gamma \mathbf{X} \mathbf{E}^{-1} \mathbf{b}. \quad (2.70)$$

gdzie $\gamma = \frac{\langle \mathbf{v}, \mathbf{X} \mathbf{E}^{-1} \mathbf{b} \rangle - \langle \mathbf{D} \mathbf{1}, \mathbf{X} \mathbf{w} \rangle - \mathbf{1}^T \mathbf{v}}{c - \langle \mathbf{D} \mathbf{1}, \mathbf{X} \mathbf{E}^{-1} \mathbf{b} \rangle}$. Iloczyn macierzowo-wektorowy $\mathbf{X} \hat{\mathbf{E}}^{-1} \mathbf{X}^T \mathbf{s}$ można rozpisać analogicznie:

$$\mathbf{X} \hat{\mathbf{E}}^{-1} \mathbf{X}^T \mathbf{s} \equiv \mathbf{X} \mathbf{E}^{-1} \mathbf{X}^T \mathbf{s} + \frac{\langle \mathbf{s}, \mathbf{X} \mathbf{E}^{-1} \mathbf{b} \rangle}{c - \langle \mathbf{D} \mathbf{1}, \mathbf{X} \mathbf{E}^{-1} \mathbf{b} \rangle} \mathbf{X} \mathbf{E}^{-1} \mathbf{b}, \quad (2.71)$$

zaś rozwiązanie po uproszczeniu opisuje równanie:

$$\mathbf{d} = \mathbf{E}^{-1} \mathbf{X}^T \left(\mathbf{v} - \mathbf{z} + \left(\gamma - \frac{\langle \mathbf{z}, \mathbf{X} \mathbf{E}^{-1} \mathbf{b} \rangle}{c - \langle \mathbf{D} \mathbf{1}, \mathbf{X} \mathbf{E}^{-1} \mathbf{b} \rangle} \right) \mathbf{D} \mathbf{1} \right) - \mathbf{w}. \quad (2.72)$$

Pochodną kierunkową można policzyć zgodnie z równaniem (2.59), pamiętając o dodaniu składnika $g_0 \cdot d_0$. Włączenie wyrazu wolnego wiąże się z dodatkowym kosztem $O(\#_{\text{nnz}} \cdot p)$ obliczenia iloczynu $\mathbf{X} \mathbf{E}^{-1} \mathbf{b}$ i kilkoma iloczynami skalarnymi, gdzie $\#_{\text{nnz}}$ oznacza liczbę niezerowych parametrów

Zysk z przechowywania macierzy $\mathbf{X} \mathbf{E}^{-1} \mathbf{X}^T$ (analogicznie jak miało to miejsce w przypadku macierzy $\frac{1}{\lambda} \mathbf{X} \mathbf{X}^T$) wydaje się wątpliwy, ponieważ po każdej zmianie parametrów trzeba wyliczyć korektę (zmienia się przecież macierz $\mathbf{E}^{-1}$), co wiąże się z kosztem rzędu $O(\#_{\text{nnz}} \cdot n)$. Dodatkowo sama macierz może mieć niepełny rząd (gdy liczba parametrów aktywnych będzie mniejsza niż n), zatem koszt iloczynu macierzowo-wektorowego zapewne będzie większy niż użycie formuły (2.62).

2.6 Regularyzacja „elastic net”

Liczba niezerowych parametrów w modelu z regularyzacją ℓ^1 zawsze nie przekroczy $\min(n, p)$, a gdy w zbiorze znajdują się grupy silnie skorelowanych atrybutów, to spośród nich zostanie zazwyczaj arbitralnie wybranych po jednym (z kolei regularyzacja ℓ^2 ściąga równomiernie całą grupę), co w pewnym stopniu ogranicza tę technikę jako narzędzie automatycznej selekcji istotnych cech. Powyższe ograniczenia likwiduje połączenie regularyzacji ℓ^2 i ℓ^1 , co szerzej znane jest jako regularyzacja elastic net [89]. W przypadku regresji liniowej optymalizowana funkcja kosztu wygląda następująco:

$$Q(\mathbf{w}) = \frac{1}{2} \sum_{i=1}^n (y_i - \langle \mathbf{w}, \mathbf{x}_i \rangle)^2 + \lambda \sum_{j=1}^p \alpha \cdot |w_j| + \frac{1-\alpha}{2} w_j^2, \quad (2.73)$$

natomiast w przypadku regresji logistycznej funkcja kosztu ma postać:

$$Q(\mathbf{w}) = \sum_{i=1}^n \log \left(1 + e^{-y_i \cdot \langle \mathbf{w}, \mathbf{x}_i \rangle} \right) + \lambda \sum_{j=1}^p \alpha \cdot |w_j| + \frac{1-\alpha}{2} w_j^2, \quad (2.74)$$

gdzie $\alpha \in (0, 1)$ steruje wpływem poszczególnych kar.

2.6.1 Spadek względem współrzędnych

W przypadku regresji liniowej zmianie ulega jedynie równanie (2.39) [21]:

$$w_j^* \leftarrow \frac{\mathcal{S}_{\lambda\alpha} \left(\sum_i x_{i,j} \cdot r_i + w_j \cdot \sum_i x_{i,j}^2 \right)}{\lambda \cdot (1 - \alpha) + \sum_i x_{i,j}^2}, \quad (2.75)$$

a w przypadku regresji logistycznej równanie (2.41) przybiera postać:

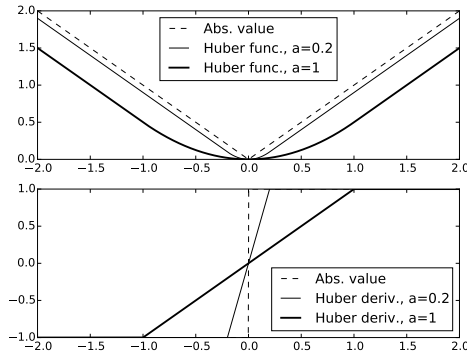
$$w_j^* \leftarrow \frac{\mathcal{S}_{\lambda\alpha} \left(\sum_i x_{i,j} \cdot r_i + w_j^* \cdot \sum_i d_{i,i} \cdot x_{i,j}^2 \right)}{\lambda \cdot (1 - \alpha) + \sum_i d_{i,i} \cdot x_{i,j}^2}. \quad (2.76)$$

2.6.2 Przekształcenie w problem gładkiej optymalizacji

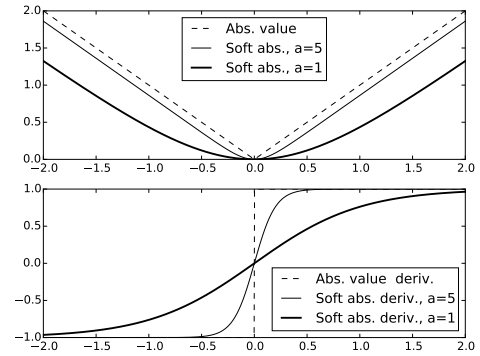
Analogicznie do układu (2.43) można zapisać układ równań w tym przypadku:

$$\begin{cases} \sum_{i=1}^n \log \left(1 + e^{-y_i \cdot \langle \mathbf{w}^+ - \mathbf{w}^-, \mathbf{x}_i \rangle} \right) + \lambda \frac{1-\alpha}{2} \|\mathbf{w}^+ - \mathbf{w}^-\|_2^2 + \lambda\alpha \sum_{j=1}^p w_j^+ + w_j^- \rightarrow \min \\ w_j^+ \geq 0, w_j^- \geq 0, j = 1, \dots, p \end{cases} \quad (2.77)$$

i ponownie rozwiązać za pomocą algorytmu L-BFGS-B.



Rysunek 2.2: Funkcja Hubera [35] — gładka aproksymacja wartości bezwzględnej oraz jej pochodną w porównaniu do wartości bezwzględnej wraz z jej pochodną.



Rysunek 2.3: Funkcja soft_abs — gładka aproksymacja wartości bezwzględnej oraz jej pochodną w porównaniu do funkcji wartości bezwzględnej wraz z jej pochodną.

Zastąpienie normy ℓ^1 jej górnym ograniczeniem

Ideę przedstawioną w sekcji 2.5.3 można z powodzeniem zastosować przy tej regularyzacji, zmieniając jedynie elementy macierzy E^{-1} :

$$e_{j,j}^{-1} = \frac{|w_j|}{\lambda\alpha + \lambda \cdot (1 - \alpha) |w_j|} \quad (2.78)$$

i równanie opisujące pochodną kierunkową:

$$\nabla Q(\mathbf{w})^T \mathbf{d} = \langle -\mathbf{X}^T \mathbf{v}, \mathbf{d} \rangle + \lambda (1 - \alpha) \langle \mathbf{w}, \mathbf{d} \rangle + \lambda \alpha \sum_{j=1}^p \text{sgn}(w_j) \cdot d_j. \quad (2.79)$$

2.7 Aproksymacja normy ℓ^1

Kolejny sposób radzenia sobie z nieróżniczkowalnością normy ℓ^1 polega na zastąpieniu jej gładką aproksymacją. Powoduje to utratę rzadkości rozwiązania, jednakże efekt grupowania (ściągnięcia w otoczenie zera) pozostaje zachowany. Jeżeli chcemy uzyskać rozwiązanie w pełni rzadkie, musimy po strojeniu wykonać progowanie współczynników poniżej arbitralnie przyjętej wartości.

W literaturze można natknąć się na przynajmniej trzy sposoby aproksymacji:

- w pierwszym funkcję w otoczeniu $[-\alpha, \alpha]$ przybliża się za pomocą funkcji kwadratowej (metoda Hubera [35]):

$$f_\alpha(x) = \begin{cases} \frac{1}{2\alpha} x^2, & |x| \leq \alpha \\ |x| - \frac{1}{2}\alpha, & \text{w.p.r.} \end{cases} \quad (2.80)$$

Im α bliższe zeru, tym aproksymacja staje się dokładniejsza (rys. 2.2). Funkcja nie posiada ciągłej drugiej pochodnej.

- w drugim korzysta się z obserwacji, że tangens hiperboliczny przypomina funkcję signum, więc jej całka przybliża wartość bezwzględną [50] (funkcja oznaczona jako `soft_abs`, por. rys. 2.3):

$$f_\alpha(x) = \frac{1}{\alpha} \log(\cosh(\alpha x)) \equiv |x| + \frac{1}{\alpha} \log\left(\frac{1 + e^{-2\alpha|x|}}{2}\right). \quad (2.81)$$

Zwiększenie wartości α poprawia dokładność przybliżenia. Funkcja ta jest ściśle wypukła oraz posiada ciągłe i gładkie pochodne, co pozwala na użycie gradientowych metod optymalizacji:

$$\frac{\partial Q}{\partial w_j} \equiv g_j = \frac{\partial L}{\partial w_j} + \lambda \tanh(\alpha w_j), \quad (2.82)$$

$$\frac{\partial^2 Q}{\partial \mathbf{w} \partial \mathbf{w}^T} = \frac{\partial^2 L}{\partial \mathbf{w} \partial \mathbf{w}^T} + \mathbf{E}, \quad (2.83)$$

gdzie $\mathbf{E}$ jest macierzą diagonalną i $e_{j,j} = \lambda \alpha (1 - \tanh^2(\alpha w_j))$. W przypadku $n < p$ można skorzystać z formuły Woodbury'ego w celu redukcji wymiarowości problemu — dla regresji logistycznej wzór na kierunek wygląda następująco:

$$\mathbf{d} = -\mathbf{E}^{-1} \mathbf{g} + \mathbf{E}^{-1} \mathbf{X}^T \left(\mathbf{X} \mathbf{E}^{-1} \mathbf{X}^T + \mathbf{D}^{-1} \right)^{-1} \mathbf{X} \mathbf{E}^{-1} \mathbf{g}. \quad (2.84)$$

W analogiczny sposób można dołączyć wyraz wolny w_0 .

W przypadku tego sposobu trzeba również zachować ostrożność, ponieważ na przekątnej macierzy $\mathbf{E}$ mogą pojawić się wartości bliskie zera.

- w trzecim przypadku wartość bezwzględną przybliża się za pomocą pierwiastka z funkcji kwadratowej:

$$f_\alpha(x) = \sqrt{x^2 + \alpha}, \quad (2.85)$$

gdzie α jest pewną małą dodatnią liczbą.

2.8 Regularyzacja z użyciem pseudonorm ℓ^q , $0 \leq q < 1$,

Norma ℓ^1 znajduje się na granicy wypukłości — obniżając dalej wykładnik uzyskamy funkcję niewypukłą. Funkcja kosztu ma postać:

$$Q(\mathbf{w}) = L(\mathbf{w}) + \lambda \sum_{j=1}^p |w_j|^q. \quad (2.86)$$

Z uwagi na niewypukłość funkcji celu znalezione minimum będzie w ogólności lokalnym, a nie globalnym minimum. W dodatku pochodna kierunkowa nie istnieje, dlatego jako kryterium stopu można posłużyć się różnicą względną wartości funkcji między kolejnymi iteracjami.

2.8.1 Spadek względem współrzędnych

Rozpatrzmy najpierw ogólny przypadek jednowymiarowy:

$$f(x) = \frac{\mu}{2} (y - x)^2 + \lambda \cdot |x|^q, \quad \mu > 0, \quad q \in (0, 1). \quad (2.87)$$

Dla pewnej wartości $\lambda = \lambda_{\text{critical}}$ wartość funkcji w zerze równa się wartości funkcji w punkcie x_{critical} — innymi słowy istnieją dwa minima globalne. Aby znaleźć wartości $\lambda_{\text{critical}}$ i x_{critical} należy rozwiązać poniższy układ równań:

$$\begin{cases} f(0) = f(x_{\text{critical}}) \\ f'(x_{\text{critical}}) = 0 \end{cases}. \quad (2.88)$$

Rozwiązując otrzymamy $x_{\text{critical}} = \frac{2-2q}{2-q} \cdot y$ i $\lambda_{\text{critical}} = \mu \cdot (2 - 2q)^{1-q} \cdot \left(\frac{|y|}{2-q}\right)^{2-q}$. Gdy $\lambda > \lambda_{\text{critical}}$, to minimum globalne funkcji (2.87) znajduje się w zerze, a w przeciwnym razie należy skorzystać z numerycznej metody poszukiwania minimum, ponieważ nie istnieje analityczne rozwiązanie. W tym celu można użyć jednowymiarowej metody Newtona-Raphsona dzięki różniczkowalności funkcji dla $x \neq 0$:

$$\begin{cases} f(x)' = \mu \cdot (x - y) + \lambda \cdot q \cdot \text{sgn}(x) \cdot |x|^{q-1}, \\ f(x)'' = \mu + \lambda \cdot q \cdot (q - 1) \cdot |x|^{q-2}, \end{cases} \quad (2.89)$$

przyjmując $x_0 = y$. W przypadku $\lambda = \lambda_{\text{critical}}$ należy arbitralnie wskazać któreś z minimumów.

Podstawiając (2.1) do (2.86) i różniczkując po zmiennej w_j otrzymamy:

$$\frac{\partial Q}{\partial w_j} = \sum_i x_{i,j}^2 \cdot \left(w_j - \frac{\sum_i x_{i,j} \cdot (y_i - \sum_{k \neq j} x_{i,k} w_k)}{\sum_i x_{i,j}^2} \right) + \lambda \cdot q \cdot \text{sgn}(w_j) \cdot |w_j|^{q-1}. \quad (2.90)$$

Powyższe wyrażenie przypomina (2.89): $\mu = \sum_i x_{i,j}^2$, $x = w_j$, $y = \frac{\sum_i x_{i,j} \cdot (y_i - \sum_{k \neq j} x_{i,k} w_k)}{\sum_i x_{i,j}^2}$. Ponownie można skorzystać ze sztuczki (2.38) obniżającej koszt obliczenia wartości y . Z kolei dla regresji logistycznej w algorytmie 4 wystarczy jedynie podmienić $x = w_{\text{opt}}$, $y = \frac{\sum_i x_{i,j} \cdot r_i}{\sum_i d_{i,i} \cdot x_{i,j}^2} + w_j^*$ i wprowadzić $\mu = \sum_i d_{i,i} \cdot x_{i,j}^2$, po czym policzyć nową wartość w_{opt} dla j -tego parametru modelu.

Ponieważ mamy do czynienia z problemem niewypukłym, algorytm może nie zatrzymać się w skończonej liczbie kroków, a minimalizacja kierunkowa wymaga większej uwagi, aby znaleźć odpowiedni krok.

Przypadek szczególny $q = 0$

Przypadek ten sprowadza się do znalezienia najlepszego podzbioru atrybutów⁹ [29]. Bywa on także nazywany twardym progowaniem przez analogię do miękkiego progowania w przypadku $q = 1$ [20]. Funkcja (2.87) w tym przypadku traci ciągłość w punkcie $x = 0$ i wygląda następująco:

$$f(x) = \begin{cases} \frac{\mu}{2} (y - x)^2 + \lambda, & x \neq 0, \\ \frac{\mu}{2} y^2, & x = 0. \end{cases} \quad (2.91)$$

Przyjmuje ona dwa minima: jedno w punkcie $x = 0$, a drugie w punkcie $x = y$. Wartość $\lambda_{\text{critical}}$ wynosi $\frac{\mu}{2} y^2$.

Przypadek szczególny $q = \frac{1}{2}$

W pracy [44] przeanalizowano dokładnie przypadek $q = \frac{1}{2}$, w którym minimalizację (2.87) można sprowadzić do minimalizacji wielomianu czwartego stopnia. Niech $x \geq 0$ i $t = x^{\frac{1}{2}}$ (przypadek $x < 0$ jest symetryczny):

$$f(t) = \frac{\mu}{2} (y - t^2)^2 + \lambda t. \quad (2.92)$$

Różniczkując względem t , a następnie dzieląc przez 2μ i przyrównując wynik do zera uzyskamy:

$$t^3 - yt + \frac{\lambda}{2\mu} = 0. \quad (2.93)$$

Powyższe równanie posiada trzy pierwiastki [58]:

$$\begin{cases} \alpha &= 2 \cdot \sqrt{\frac{y}{3}} \cdot \cos \theta, \\ \beta &= 2 \cdot \sqrt{\frac{y}{3}} \cdot \cos \left(\theta + \frac{2\pi}{3} \right), \\ \gamma &= 2 \cdot \sqrt{\frac{y}{3}} \cdot \cos \left(\theta + \frac{4\pi}{3} \right), \end{cases} \quad (2.94)$$

gdzie

$$\theta = \frac{1}{3} \cdot \arccos \left(-\frac{\lambda}{4\mu \left(\frac{y}{3} \right)^{\frac{3}{2}}} \right). \quad (2.95)$$

Interpretacja rozwiązań: α to lokalne minimum (2.92), z kolei γ to lokalne maksimum, natomiast $\beta < 0$. Pamiętając o podstawieniu $t = x^{\frac{1}{2}}$, wartość α należy podnieść do kwadratu, by otrzymać x .

⁹ang. best subset selection

2.8.2 Zastąpienie pseudonormy ℓ^q jej górnym ograniczeniem

Podobnie jak miało to miejsce w przypadku normy ℓ^1 , możemy zastąpić pseudonormę ℓ^q funkcją kwadratową [40]:

$$Q(\mathbf{w}) = L(\mathbf{w}) + \frac{\lambda}{2} \sum_{j=1}^p \left(\frac{q \cdot w_j^2}{|w'_j|^{2-q}} + (2-q) |w'_j|^q \right) \quad (2.96)$$

i zapisać rozwiązanie dla regresji logistycznej w analogiczny sposób, jak w (2.56), zmieniając jedynie elementy macierzy diagonalnej $\mathbf{E}$, tj. $e_{j,j} = \frac{\lambda q}{|w'_j|^{2-q}}$. Dalej należy postąpić analogicznie, jak przedstawiono to dla normy ℓ^1 .

2.9 Uogólniona regularyzacja

Karze nie muszą podlegać wyłącznie pojedyncze współczynniki modelu — jeżeli np. podejrzewamy, że szukany model zawiera wiele płaskich obszarów, to możemy nałożyć karę ℓ^1 na różnice sąsiednich współczynników, co zminimalizuje ich wahanie się. Niech dana będzie macierz różnicy $\mathbf{F}_{(p-1) \times p}$ zbudowana następująco:

$$f_{i,j} = \begin{cases} -1, & j = i, \\ 1, & j = i + 1, \\ 0, & \text{wpr.} \end{cases} \quad (2.97)$$

Powyższą macierz należy umieścić w funkcji kary (np. ℓ^1 lub ℓ^2) danej w tabeli 2.1, zaś samą funkcję w (2.12). Przyjmując następującą postać funkcji kosztu:

$$Q(\mathbf{w}) = \frac{1}{2} \|\mathbf{w} - \mathbf{y}\|_2^2 + \|\mathbf{F}\mathbf{w}\|_1, \quad (2.98)$$

otrzymamy zagadnienie odszumiania¹⁰, gdzie $\mathbf{y}$ to sygnał źródłowy, który chcemy odzyskać [65].

Normę ℓ^1 można z powodzeniem zastąpić normą ℓ^2 , a różnice pierwszego rzędu różnicami wyższych rzędów [19], np. macierz $\mathbf{F}_{(p-2) \times p}$ obliczająca różnice drugiego rzędu ma postać:

$$f_{i,j} = \begin{cases} 1, & j = i \vee j = i + 2, \\ -2, & j = i + 1, \\ 0, & \text{wpr.} \end{cases} \quad (2.99)$$

W przypadku sygnałów dwuwymiarowych (np. obrazów) karać należy współczynniki w poszczególnych wierszach i kolumnach. Postać macierzy różnic $\mathbf{F}$ zależy od porządku,

¹⁰ang. total variation denoising

w jakim trzymano pierwotny sygnał, zanim przekształcono go w wektor, łącząc kolejno wiersze lub kolumny z sobą. Niech $(\hat{n}, \hat{p})$ oznacza pierwotną liczbę wierszy i kolumn sygnału (po wektoryzacji mamy $p = \hat{n}\hat{p}$):

- W porządku wierszowym $\mathbf{H} = \mathbf{I}_{\hat{n}} \otimes \hat{\mathbf{F}}_{(\hat{p}-1) \times \hat{p}}$, gdzie $\mathbf{H}$ to macierz obliczająca różnice elementów w wierszach, zaś $\hat{\mathbf{F}}_{(\hat{p}-1) \times \hat{p}}$ to macierz obliczająca różnice sygnału jednowymiarowego. Z kolei macierz $\mathbf{V}$ obliczającą różnice pierwszego rzędu w kolumnach buduje się następująco:

$$v_{i,j} = \begin{cases} -1, & j = i, \\ 1, & j = i + \hat{p}, \\ 0, & \text{wpr.} \end{cases} \quad (2.100)$$

- W porządku kolumnowym $\mathbf{V} = \mathbf{I}_{\hat{p}} \otimes \hat{\mathbf{F}}_{(\hat{n}-1) \times \hat{n}}$, natomiast macierz $\mathbf{H}$ definiuje się następująco:

$$h_{i,j} = \begin{cases} -1, & j = i, \\ 1, & j = i + \hat{n}, \\ 0, & \text{wpr.} \end{cases} \quad (2.101)$$

Ostatecznie macierz różnic $\mathbf{F}$ w przypadku dwuwymiarowym stanowi złożenie obu macierzy różnic, co można symbolicznie zapisać:

$$\mathbf{F} = \begin{bmatrix} \mathbf{H} \\ \mathbf{V} \end{bmatrix}. \quad (2.102)$$

Do regularyzacji różnic można dodać jednocześnie karanie pojedynczych współczynników modelu¹¹ [76], co w przypadku normy ℓ^1 pozwoli wychwycić pojedyncze „piki” tudzież równiny.

¹¹ang. fused lasso

3. Realizacja modeli nieliniowych

3.1 Uczenie sieci neuronowych

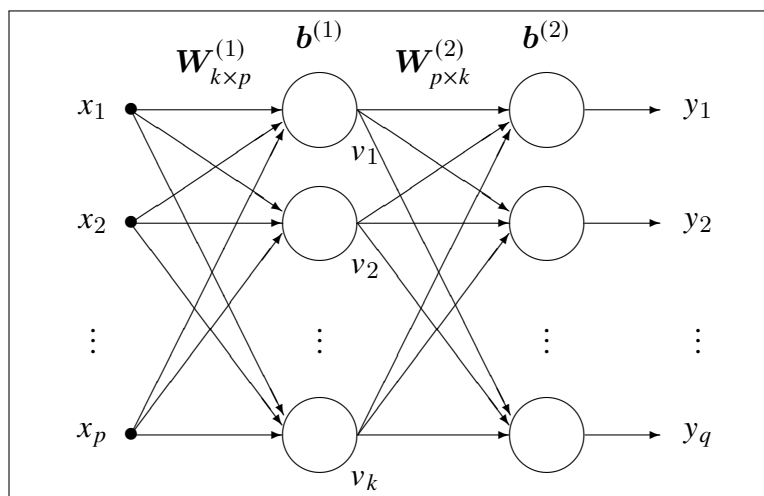
Rozpatrzmy prosty model sieci neuronowej z pojedynczą warstwą ukrytą złożoną z k neuronów ($k < p$), przedstawiony na rysunku 3.1. Wyjście z sieci opisuje wyrażenie:

$$y(x) = \sigma_2 \left(\mathbf{W}^2 \cdot \sigma_1 \left(\mathbf{W}^1 \cdot \mathbf{x} + \mathbf{b}^1 \right) + \mathbf{b}^2 \right), \quad (3.1)$$

natomiast wyjście z warstwy pośredniej dane jest wzorem:

$$v(x) = \sigma_1 \left(\mathbf{W}^1 \cdot \mathbf{x} + \mathbf{b}^1 \right), \quad (3.2)$$

gdzie $\sigma_1(\cdot)$ to funkcja aktywacji w warstwie ukrytej, a $\sigma_2(\cdot)$ to funkcja aktywacji w warstwie wyjściowej oraz $\mathbf{W}_{k \times p}^1$, $\mathbf{b}_{k \times 1}^1$, $\mathbf{W}_{q \times k}^2$ i $\mathbf{b}_{q \times 1}^2$ są odpowiednio wagami i wartościami progowymi dla pierwszej i drugiej warstwy. Rozważane w pracy funkcje aktywacji σ to identyczność oraz tangens hiperboliczny.



Rysunek 3.1: Sieć neuronowa z jedną warstwą ukrytą.

Sieć o takiej architekturze może spełniać różnorodne funkcje, począwszy od uniwersalnego klasyfikatora i regresora. W pracy przedstawione zostaną dwa zagadnienia wykorzystujące prezentowane techniki regularyzacji:

uczenie zrandomizowane — uczenie sieci wykorzystujące jedynie techniki strojenia klasyfikatorów/regresorów liniowych, w tym techniki opisane w poprzednich punktach.

rzadki autoenkoder — autoenkoder bazujący na dwuwarstwowej sieci neuronowej z zerowaną pewną liczbą wag w warstwie pośredniej i warstwie wyjściowej; taki autoenkoder jest pewną formą rzadkiego wariantu metody PCA lub pewną formą faktoryzacji macierzy danych.

3.2 Rzadki autoenkoder

Sformułowanie zagadnienia

Niech $\mathbf{X}$ oznacza macierz danych. Rozważmy procedurę PCA¹:

$$\begin{aligned} \mathbf{C} &= \text{cov}(\mathbf{X}) \\ [\mathbf{P}, \boldsymbol{\lambda}] &= \text{eig}(\mathbf{C}), \end{aligned} \tag{3.3}$$

gdzie $\mathbf{P}$ jest macierzą złożoną kolumnowo z wektorów własnych, a $\boldsymbol{\lambda}$ jest wektorem odpowiadających im wartości własnych. Wtedy mamy:

$$\begin{aligned} \mathbf{X} &= \mathbf{P} \mathbf{P}^T \mathbf{X} \\ \mathbf{X} &\approx \mathbf{P}_{:,1:k} \mathbf{P}_{:,1:k}^T \mathbf{X} \end{aligned} \tag{3.4}$$

a dla $k < p$ model ten może być traktowany jako najprostsza forma autoenkodera.

Metoda PCA nie jest jednak wolna od pewnych wad, np. minimalizacja wariancji składowych jest zależna od wartości odstających [15] oraz w standardowym podejściu większość zmiennych ma wpływ na każdą główną składową, a ich bezpośrednia interpretacja jest często kłopotliwa.

Przeciwieństwem tego podejścia są procedury rzadkie, które produkują składowe zależne tylko od małego podzbioru zmiennych wejściowych. Rzadkość komponentów składowych może być atrakcyjna z wielu punktów widzenia. Przykładowo stosując taką reprezentację składników można zaoszczędzić miejsce w pamięci i skrócić czas obliczeń. W przypadku specjalnych typów sygnałów o określonej strukturze zależności (jak obrazy, dźwięki lub dane chemometryczne), takie podejście zapewnia rozwiązania, które często

¹ang. principal component analysis

można zinterpretować w sposób naturalny. Przykładowo jako rozwiązanie można uzyskać oddzielne obszary w danych, np. usta i oczy w przypadku obrazów twarzy.

Rzadkość jest zazwyczaj rozumiana jako liczba współczynników zerowych w rozwiązaniu lub jako względna liczba współczynników zerowych w odniesieniu do rozmiaru danych. Rzadkość rozwiązania może być zapewniona przez dodanie do funkcji błędu składnika regularyzacyjnego z karą ℓ^1 na współczynniki [29, 83]. Istnieje wiele różnych technik zapewniających rzadkość komponentów [15, 56]. Jednym z takich podejść jest uczenie słownikowe, które zaimplementowano w klasie SparsePCA w pakiecie `scikit-learn` [39, 53].

Sieć o architekturze przedstawionej na rysunku 3.1, gdy $p = q$ może pełnić funkcję autoenkodera, tj. struktury, której celem jest odtwarzanie danych wejściowych, uczonej zestawem danych $\{(\mathbf{x}_i, \mathbf{x}_i)\}_{i=1}^n$, gdzie n to liczba obserwacji, a p to liczba atrybutów. Sieć o takiej strukturze może być rozpatrywana jako wariant rzadkiej metody PCA [8, 56]. Przyjmijmy funkcje aktywacji identycznościowe, tzn. $\sigma_1(z) = z$, $\sigma_2(z) = z$. Rozważmy funkcję kosztu zawierającą składniki z karą ℓ^1 na współczynniki warstwy pośredniej i wyjściowej:

$$L(\mathbf{X}, \mathbf{Y}, \mathbf{W}^1, \mathbf{W}^2) = \frac{1}{2} \sum_{i=1}^n \|\mathbf{y}_i - f(\mathbf{x}_i)\|_2^2 + \lambda \cdot (\|\mathbf{W}^1\|_1 + \|\mathbf{W}^2\|_1). \quad (3.5)$$

Procedura uczenia

Wyniki przedstawione w tym rozdziale zostały opublikowane w pracy [62]. Rozważamy ogólne zadanie uczenia nadzorowanego: niech $\mathbf{X}_{p \times n} = \{\mathbf{x}_1, \dots, \mathbf{x}_n\}$ będzie częścią wejściową zbioru danych (obserwacje znajdują się w kolumnach), a $\mathbf{Y}_{p \times n} = \{\mathbf{y}_1, \dots, \mathbf{y}_n\}$ wyjściową. Rozważmy następującą funkcję (część dodatnia, `relu`, por. sekcja 2.5.3):

$$(a_i)^+ = \max(a_i, 0) = \begin{cases} a_i, & \text{dla } a_i > 0 \\ 0, & \text{wpr.j} \end{cases},$$

podobnie możemy zdefiniować część ujemną: $(a_i)^- = (-a_i)^+$ oraz podobnie dla wektora $\mathbf{a} = [a_1, \dots, a_p]$, używamy notacji: $\mathbf{a}^+ = [(a_1)^+, \dots, (a_p)^+]$ i $\mathbf{a}^- = [(a_1)^-, \dots, (a_p)^-]$. Korzystając z definicji otrzymamy (por. sekcja 2.5.3) $\mathbf{a} = \mathbf{a}^+ - \mathbf{a}^-$, czyli rozbito zmienną na różnicę części dodatniej i ujemnej, przy czym zarówno część ujemna, jak i dodatnia są nieujemne.

Rozważamy dwuwarstwową sieć neuronową z k neuronami w warstwie ukrytej, jak pokazano na rys. 3.1, wyjściem z sieci i wyjściem z warstwy pośredniej danymi wzorami (3.1)

oraz (3.2). Celem uczenia jest optymalizacja (minimalizacja) następującej funkcji:

$$\arg \min_{\mathbf{W}^1, \mathbf{b}^1, \mathbf{W}^2, \mathbf{b}^2} \frac{1}{2n} \sum_{i=1}^n \left\| \sigma \left(\mathbf{W}^2 \sigma \left(\mathbf{W}^1 \mathbf{x}_i + \mathbf{b}^1 \right) + \mathbf{b}^2 \right) - \mathbf{y}_i \right\|_2^2 + \alpha \cdot \left(\|\mathbf{W}^1\|_1 + \|\mathbf{W}^2\|_1 \right), \quad (3.6)$$

gdzie α jest współczynnikiem regularyzacji. Ze względu na nieróżniczkowalność normy $\|\cdot\|_1$, zamieniamy problem optymalizacyjny (3.6) na równoważny problem optymalizacji funkcji gładkiej z ograniczeniami, używając tożsamości $|x| \equiv x^+ - x^-$ dla $x^+, x^- \geq 0$:

$$\begin{aligned} \arg \min_{\mathbf{W}^{1+}, \mathbf{W}^{1-}, \mathbf{b}^1, \mathbf{W}^{2+}, \mathbf{W}^{2-}, \mathbf{b}^2} & \frac{1}{2n} \sum_{i=1}^n \left\| \sigma \left(\mathbf{W}_*^2 \cdot \sigma \left(\mathbf{W}_*^1 \mathbf{x}_i + \mathbf{b}^1 \right) + \mathbf{b}^2 \right) - \mathbf{y}_i \right\|_2^2 + \\ & + \alpha \cdot \left(\mathbf{1}^T \mathbf{W}^{1+} \mathbf{1} + \mathbf{1}^T \mathbf{W}^{1-} \mathbf{1} + \mathbf{1}^T \mathbf{W}^{2+} \mathbf{1} + \mathbf{1}^T \mathbf{W}^{2-} \mathbf{1} \right), \\ & \mathbf{W}^{1+} \geq 0, \mathbf{W}^{1-} \geq 0, \mathbf{W}^{2+} \geq 0, \mathbf{W}^{2-} \geq 0, \end{aligned} \quad (3.7)$$

gdzie $\mathbf{W}_*^1 = \mathbf{W}^{1+} - \mathbf{W}^{1-}$, $\mathbf{W}_*^2 = \mathbf{W}^{2+} - \mathbf{W}^{2-}$, a operator wektorowy „ $\geq$ ” jest stosowany do każdego elementu wektora/macierzy.

Do rozwiązania zadania (3.7) można użyć algorytmu L-BFGS-B [88]. Podobne podejście rozważano np. w [52, 69] w przypadku uczenia regresji logistycznej z karą ℓ^1 (por. sekcja 2.5.3).

Wykorzystując procedurę L-BFGS-B sieć można uczyć globalnie lub użyć następującej metody krokowej: zaczynając od jednego neuronu w warstwie ukrytej ($k = 1$) sieć neuronowa jest trenowana, a następnie poszerzeniu ulega bieżące rozwiązanie poprzez dodanie kolejnej jednostki — ta operacja wstawia do macierzy $\mathbf{W}^{1+}$ nowy wiersz i nową kolumnę do macierzy $\mathbf{W}^{2+}$ oraz nowy element do wektora $\mathbf{b}^1$, podobnie należy zrobić z macierzami $\mathbf{W}^{1-}$ i $\mathbf{W}^{2-}$. Wartości nowych współczynników można zainicjalizować losowo. Po dodaniu neuronu procedura uczenia jest ponawiana.

Początkowe wagi sieci neuronowej mogą być ustawione losowo, jednak w pracy inicjalizowano wagi $\mathbf{W}^1$ macierzą transformacji wyznaczoną za pomocą procedury PCA, a wagi $\mathbf{W}^2$ macierzą transponowaną, z kolei wektory $\mathbf{b}^1$ i $\mathbf{b}^2$ wykluczono z modelu.

Przedstawione rozwiązanie jest rozwiązaniem ogólnym, które można wykorzystać do uczenia dowolnych sieci neuronowych z rzadkimi ograniczeniami typu ℓ^1 . Chcąc użyć taką sieć jako rzadki autokoder równoważny procedurze PCA, ustawiamy σ jako funkcję identycznościową (co oznacza liniową funkcję aktywacji) oraz dane wejściowe ustawiamy jako cel predykcji. k -ty wiersz $\mathbf{W}^1$ można interpretować jako k -tą główną składową, a $\mathbf{W}^2$ jest macierzą odwrotnej transformacji do wymiaru p .

Proponowane rozwiązanie może być porównane z procedurą SparsePCA z pakietu `scikit-learn` [39, 53], wykorzystującą uczenie słownikowe. Podejście to rozwiązuje następujący problem optymalizacji:

$$(U^*, V^*) = \arg \min_{U, V} \frac{1}{2} \|X - VU\|_2^2 + \alpha \|V\|_1 \quad (3.8)$$

przy ograniczeniach: $\|u_{j,:}\|_2 = 1$ dla $1 \leq j \leq k$

Patrząc na (3.6) oraz (3.4) mamy $\|X - W^2 W^1 X\|_2^2$, gdzie W^2 odpowiada V i jest traktowane jako słownik, a $W^1 X$ odpowiada macierzy U i jest częścią kodującą (zakodowaną reprezentacją). W prezentowanym rzadkim autoenkoderze wymaga się, aby obie macierze wag były rzadkie oraz bardziej naturalne jest użycie W^1 jako macierzy transformacji i W^2 jako transformacji odwrotnej. W uczeniu słownikowym wymagane jest, aby słownik był rzadki (ograniczenia z normą ℓ^1)

3.3 Uczenie zrandomizowane, uczenie ekstremalne

Sieć neuronowa z jedną warstwą ukrytą i nieliniową funkcją aktywacji umożliwia wygenerowanie nieliniowej granicy decyzyjnej. Istnieje cały szereg wariantów strojenia takich struktur, głównie bazujących na idei wstecznej propagacji błędu.

Alternatywne podejście do uczenia polega na ustaleniu wag warstwy wejściowej w sposób losowy. Przepuszczenie próby uczącej przez tę warstwę skutkować będzie podniesieniem wymiarowości, o ile liczba neuronów będzie dużo większa od liczby atrybutów w danych uczących. Z kolei w przestrzeni o wyższym wymiarze dużo łatwiej znaleźć hiperpłaszczyznę dobrze separującą klasy decyzyjne (por. np. [80]). Przy zamrożonych losowych wagach warstwy ukrytej problem uczenia sprowadza się do strojenia tylko warstwy wyjściowej. Dzięki temu zabiegowi otrzymujemy liniowy klasyfikator wraz z losową warstwą transformującą, który generuje nieliniową granicę decyzyjną w oryginalnej przestrzeni.

Uczenie takie współcześnie spotykane jest pod nazwą uczenia ekstremalnego (ang. *extreme learning*) [17, 73, 90], jednak idea losowego wyboru parametrów warstw pośrednich sieci jest znana od lat dziewięćdziesiątych np. pod nazwami RVFL (Random Vector Functional Link) lub RFHNs (Randomly Fixed Hidden Neurons) [34, 37]². Wykorzystując techniki regularyzacji możemy bardzo efektywnie nastroić sieć, a korzystając z rzadkich regularyzacji możemy również znacznie zredukować rozmiar sieci i ustalić optymalną liczbę neuronów w warstwie pośredniej.

²por. ciekawy komentarz [82] zawierający historię rozwoju tej koncepcji.

Uczenie zrandomizowane wymaga jedynie prostego kroku losowania wag warstwy pośredniej, które ustala losowe odwzorowanie pomiędzy wejściem $\mathbf{X}$ a wyjściem z warstwy pośredniej $\mathbf{V}$ oraz nastrojenia warstwy wyjściowej na danych $\{(\mathbf{v}_i, \mathbf{y}_i)\}_{i=1}^n$, zamiast na parach $\{(\mathbf{x}_i, \mathbf{y}_i)\}_{i=1}^n$. Dobór wag warstwy pośredniej najlepiej zrealizować rozmieszczając je na tle danych wejściowych lub wybrać z danych wejściowych, biorąc pod uwagę geometryczną interpretację neuronu składającego się z części liniowej oraz nieliniowej funkcji aktywacji. Na podobnej idei bazują przekształcenia jądrowe (kernel trick), gdzie standardowo jako parametry przekształcenia wykorzystuje się cały zbiór uczący [80].

4. Część eksperymentalna

4.1 Środowisko testowe i szczegóły numeryczne

Na środowisko testowe składa się komputer z procesorem Intel® Core™ i7-10700 i pamięcią operacyjną o pojemności 128 GiB, pracujący pod kontrolą systemu operacyjnego Ubuntu 22.04.4 LTS. Zasadniczo eksperymenty przygotowano w języku Python, z poziomu którego wywoływane są fragmenty kodu związane z uczeniem modeli, napisane w C++ dla zachowania maksymalnej wydajności. Wszystkie programy skompilowano przy użyciu narzędzi dostępnych w systemie Ubuntu, tj. `g++` lub `gfortran` (dla tych napisanych w języku Fortran) w wersji 11.4, z włączoną optymalizacją trzeciego stopnia za pomocą przełącznika `-O3`.

Wykorzystano bibliotekę Armadillo v14.0 [67] do realizacji obliczeń numerycznych w C++, ponieważ dostarcza jednolity interfejs dla macierzy gęstej (w porządku kolumnowym) i macierzy rzadkiej (w formacie skompresowanym kolumnowym CSC).

Implementacje niskopoziomowe procedur BLAS i LAPACK zaczerpnięto z biblioteki Intel® Math Kernel Library [38], którą zainstalowano za pomocą narzędzia `apt`. Biblioteka ta dla zwiększenia wydajności obliczeń korzysta z wielu rdzeni procesora, dlatego podczas eksperymentów wymuszono korzystanie tylko z jednego rdzenia, aby umożliwić uczciwe porównanie algorytmów.

Z biblioteki `liblinear` v2.47, która oryginalnie operuje na macierzach rzadkich w formacie skompresowanym wierszowym CSR, wzięto implementację przyciętej metody Newtona [23], do której dodano obsługę macierzy gęstej oraz typu zmiennoprzecinkowego pojedynczej precyzji. Z tej samej biblioteki wybrano implementację metody spadku względem współrzędnych dla regresji logistycznej z regularyzacją ℓ^1 [86], dodając również obsługę macierzy gęstych oraz typu zmiennoprzecinkowego pojedynczej precyzji.

Klasą bazową każdego modelu jest klasa abstrakcyjna `function` reprezentująca funkcję celu, zawierająca interfejs metod obliczających np. gradient funkcji celu bądź iloczyn macierzy Hessego funkcji celu przez dany wektor. Przygotowano także klasę `functionf`,

w której podmieniono typ zmiennoprzecinkowy podwójnej precyzji na typ zmiennoprzecinkowy pojedynczej precyzji. W pierwotnej implementacji metody Newtona z biblioteki `liblinear` należy do funkcji przekazać wskaźnik na obiekt klasy `function`, chcąc znaleźć optimum danej funkcji celu — korzystając z mechanizmu szablonów w C++ umożliwiono przekazywanie obiektów obu klas.

Właściwości modelu regresji logistycznej (spójrz na poniższy listing) należy ustawić za pomocą parametrów szablonu:

- pierwszy parametr włącza wyraz wolny (jeśli ustawiono go na `true`),
- drugi parametr aktywuje preconditioning,
- trzeci parametr umożliwia użycie dokładnej minimalizacji kierunkowej,
- czwarty parametr określa typ macierzy (albo `arma::Mat<>` dla macierzy gęstej, albo `arma::SpMat<>` dla macierzy rzadkiej).

Na podstawie użytego typu zmiennoprzecinkowego kompilator decyduje, którą wersję interfejsu odziedziczyć:

```

1  template <
2      bool use_intercept ,
3      bool use_precon ,
4      bool use_exact_ls ,
5      typename matrix_class ,
6      typename el_type = std::remove_cvref_t<matrix_class>::elem_type
7  >
8  class logistic_regression : public std::conditional_t<
9      std::is_same_v<double , el_type> ,
10     function ,
11     std::conditional_t<
12         std::is_same_v<float , el_type> ,
13         functionf ,
14         void>>
15 {
16     // ...
17 };

```

Następnie w klasie `logistic_regression` dostarczono implementacje abstrakcyjnych metod, tak by reprezentowała podstawowy model regresji logistycznej bez regularyzacji. W kolejnych krokach wykorzystano za pośrednictwem dziedziczenia tę klasę do utworzenia pozostałych klas reprezentujących poszczególne modele.

Dokonano modyfikacji procedury `SYMMLQ`¹, dodając wsparcie dla typu zmiennoprzecinkowego pojedynczej precyzji oraz usuwając ograniczenie związane z procedurami `Aprod` i `Msolve` (przyjmowały one tylko trzy parametry: rozmiar macierzy `n`, wektor

¹Kod źródłowy pochodzi z <https://web.stanford.edu/group/SOL/software/symmlq/symmlq-f77.zip>

wejściowy x i wektor wynikowy y), aby móc przekazywać obiekty reprezentujące funkcje celu w analogiczny sposób, jak miało to miejsce w implementacji metody Newtona w bibliotece `liblinear`.

W implementacji algorytmu L-BFGS-B v3.0² [88] pozbyto się zależności od biblioteki LINPACK na rzecz jej następcy, tj. biblioteki LAPACK, tzn. zastąpiono procedury DPOFA i DTRSL odpowiednio procedurami DPOTRF i DTRTRS [2] — w ten sposób całość obliczeń oparto na bibliotece Intel® Math Kernel Library. Istniejącą implementację niskim kosztem przerobiono w ten sposób, by działała również dla typu zmiennoprzecinkowego pojedynczej precyzji.

W implementacji metody wewnętrznego punktu³ dodano opcję logowania postępu do pliku zgodnie z przyjętym schematem. Autorzy implementacji przyjęli wierszowy porządek macierzy danych, dlatego do pomiaru włączono czas potrzebny na transpozycję macierzy (transpozycja nie jest dokonywana w miejscu). Porzucono pomysł adaptacji metody do typu pojedynczej precyzji, ponieważ wiązało się to ze zmianami w wielu plikach.

4.2 Obliczanie iloczynu XX^T w przypadku macierzy rzadkiej

W modelu regresji logistycznej z regularyzacją ℓ^2 po zastosowaniu formuły Woodbury’ego w celu obliczenia macierzy odwrotnej do macierzy Hessego pojawił się składnik XX^T o wymiarach $n \times n$. Intuicyjnie w pewnym zakresie powinna pojawić się korzyść z jednorazowego obliczenia tego iloczynu, by później móc się do niego odwoływać podczas rozwiązywania układu równań. Dla macierzy gęstej iloczyn ten oblicza procedura `?SYRK` z biblioteki BLAS, gdzie „?” należy zastąpić literą D dla podwójnej precyzji lub S dla pojedynczej precyzji, a w wyniku powstaje macierz dolno- bądź górnotrójkątna (macierz kwadratowa z wypełnionym dolnym bądź górnym trójkątem), ponieważ mamy do czynienia z macierzą symetryczną, więc nie ma potrzeby obliczania całej macierzy. By maksymalnie wykorzystać pamięć podręczną autorzy procedury⁴ budują wynik kolumna po kolumnie.

W przypadku macierzy rzadkiej sposób obliczania się komplikuje z uwagi na różne

²Kod źródłowy pochodzi z <https://users.iems.northwestern.edu/~nocedal/Software/Lbfgsb.3.0.tar.gz>

³Kod źródłowy pochodzi z https://web.stanford.edu/~boyd/l1_logreg/download/l1_logreg-0.8.2.tar.gz

⁴Kod źródłowy znajduje się pod adresem https://www.netlib.org/lapack/explore-html/dc/d05/dsyrrk_8f_source.html

sposoby przechowywania macierzy, dlatego przetestowano dwa podejścia dla kilku losowych macierzy rzadkich w porządku kolumnowym o różnym stopniu wypełnienia⁵, przyjąwszy wynik jako macierz gęstą dolnotrójkątną:

- w pierwszym przypadku wykorzystano „naiwny” sposób obliczenia wyniku, polegający na sumowaniu iloczynów zewnętrznych⁶ kolejnych kolumn macierzy X :

$$XX^T \equiv \sum_{j=1}^p \mathbf{x}_{:,j} \mathbf{x}_{:,j}^T, \quad (4.1)$$

a podsumowuje to poniższy kod (XXT to tablica o długość $n \cdot n$):

```
1 for (arma::uword i = 0; i < X.n_cols; ++i)
2     for (auto it = X.begin_col(i); it != X.end_col(i); ++it)
3         for (auto jt = it; jt != X.end_col(i); ++jt)
4             XXT[it.row()*X.n_rows + jt.row()] += *it * *jt;
```

Dla wysokiego stopnia wypełnienia macierzy X czas obliczeń może się wydłużyć z powodu wielokrotnego przechodzenia po macierzy wynikowej. W tabeli wynikowej temu sposobowi przypisano nr 1.

- w drugim przypadku zaadaptowano procedurę ?SYRK zapisującą wynik w dolnym trójkącie, gdzie wynik liczony jest kolumna po kolumnie (col_it to alias dla typu iteratora pozwalającego przejść po kolumnie macierzy):

```
1 auto cmp = [](const auto &a, const auto &b) {
2     if (a.row() == b.row())
3         return a.col() > b.col();
4     return a.row() > b.row();
5 };
6 std::priority_queue<col_it, std::vector<col_it>, decltype(cmp)> nzi(cmp);
7 for (arma::uword j = 0; j < X.n_cols; j++)
8     if (X.begin_col(j).col() == j)
9         nzi.push(X.begin_col(j));
10 while (!nzi.empty()) {
11     auto col = nzi.top();
12     nzi.pop();
13     auto j = col.row(), k = col.col();
14     for (auto it = col; it != X.end_col(k); ++it)
15         XXT[it.row() + j*X.n_rows] += *col * (*it);
16     ++col;
17     if (col != X.end_col(k))
18         nzi.push(col);
19 }
```

W tabeli wynikowej temu sposobowi przypisano nr 2.

⁵Stopień wypełnienia to iloraz liczby niezerowych elementów i iloczynu wymiarów macierzy $n \cdot p$.

⁶Działanie to w bibliotece BLAS realizuje procedura ?SYR.

Tabela 4.1: Czas wyznaczania iloczynu $\mathbf{X}\mathbf{X}^T$ dla macierzy rzadkiej w podwójnej precyzji dla trzech opisanych sposobów (nr), przy zadanym wypełnieniu (sp), n — l. wierszy, p — l. kolumn.

sp	nr	p=100000			p=1000000		
		n=100	n=1000	n=10000	n=100	n=1000	n=10000
10^{-3}	1	0,001246	0,005947	0,444052	0,011596	0,036657	2,44171
	2	0,002971	0,045212	1,50472	0,039607	1,1049	18,1561
	3	0,000806	0,010594	0,895082	0,010979	0,233275	7,56482
10^{-2}	1	0,001857	0,045611	13,3153	0,017567	0,417333	134,202
	2	0,026391	0,633123	13,8742	0,567729	11,1309	181,413
	3	0,0046	0,215629	13,7634	0,113054	2,49794	100,096
10^{-1}	1	0,019872	2,67712	395,156	0,199497	23,97	4053,38
	2	0,389611	8,9242	292,341	5,07064	101,59	3181,89
	3	0,112088	6,10633	311,866	1,48863	62,9627	2946,42

Tabela 4.2: Czas wyznaczania iloczynu $\mathbf{X}\mathbf{X}^T$ dla macierzy rzadkiej w pojedynczej precyzji dla trzech opisanych sposobów (nr), przy zadanym wypełnieniu (sp), n — l. wierszy, p — l. kolumn.

sp	nr	p=100000			p=1000000		
		n=100	n=1000	n=10000	n=100	n=1000	n=10000
10^{-3}	1	0,001206	0,004404	0,301092	0,011767	0,033401	2,01659
	2	0,002877	0,042853	1,22071	0,039498	1,08033	16,8522
	3	0,000741	0,007922	0,71066	0,010352	0,234735	6,96401
10^{-2}	1	0,001849	0,03921	12,8991	0,017701	0,362959	130,754
	2	0,026887	0,616852	13,1467	0,56545	10,8668	179,486
	3	0,004169	0,191002	12,63	0,088646	2,41417	94,5575
10^{-1}	1	0,020864	1,89437	243,705	0,206076	17,9507	2477,67
	2	0,380083	8,25811	275,335	5,05641	97,327	3002,67
	3	0,095689	5,5231	263,007	1,32293	55,2086	2566,24

- sprawdzono też czas obliczenia wyrażenia $\mathbf{X}^T\mathbf{X}$, choć iloczyn macierzy rzadkich biblioteka Armadillo traktuje jako macierz rzadką. W tabeli wynikowej temu sposobowi przypisano nr 3.

W eksperymencie zmierzono czas (w sekundach) jednokrotnego policzenia wyniku dla losowo wygenerowanej macierzy liczb z przedziału $[0, 1)$ o podanych wymiarach i stopniu wypełnienia. Wynik eksperymentu przedstawiają tabele 4.1 i 4.2.

W większości przypadków sposób pierwszy („naiwny”) przewyższa pozostałe sposoby pod kątem czasu obliczeń. Mimo zaadaptowania kodu z biblioteki BLAS na potrzeby

macierzy rzadkiej, tylko w jednym przypadku ($n=10000$, $p=100000$, $sp=10^{-1}$, podwójna precyzja) udało się przewyższyć pozostałe sposoby — testowane macierze okazały się zbyt rzadkie, by wykorzystać potencjał proponowanego algorytmu.

Na podstawie wyników powyższego eksperymentu wybrano sposób pierwszy do realizacji iloczynu $\mathbf{X}\mathbf{X}^T$ w przypadku macierzy rzadkiej.

4.3 Porównanie implementacji metody gradientu sprzężonego

W pakiecie `scikit-learn` znajduje się klasa `LogisticRegression`, która trenuje model regresji logistycznej z regularyzacją ℓ^2 (odpowiada on modelowi `l2_logreg_ordinary`, por. tab. 4.5) m.in. za pomocą metody Newtona, w której kierunek poszukiwań wyznacza algorytm gradientu sprzężonego bez preconditioningu (parametr `solver='newton-cg'` w konstruktorze klasy). Analiza kodu źródłowego metody Newtona użytej do uczenia modelu `LogisticRegression` wykazała następujące różnice w stosunku do implementacji metody Newtona znajdującej się w bibliotece `liblinear`:

- warunek stopu metody Newtona to $\|\mathbf{g}\|_\infty \leq \epsilon$, gdzie $\mathbf{g}$ to gradient funkcji kosztu w bieżącej iteracji, a ϵ to zadana przez użytkownika dokładność — w pracy obserwowano wartość pochodnej kierunkowej, natomiast pierwotnie w `liblinear` optymalizacja kończy się w momencie wystąpienia warunku $\|\mathbf{g}\|_2 \leq tol \cdot \|\mathbf{g}_0\|_2$, gdzie $\mathbf{g}_0$ to gradient funkcji kosztu w punkcie 0,
- warunek stopu metody gradientu sprzężonego to $\|\mathbf{r}^{(i)}\|_1 \leq \min\left(\frac{1}{2}, \sqrt{\|\mathbf{g}\|_1}\right) \cdot \|\mathbf{g}\|_1$, gdzie $\mathbf{r}^{(i)}$ to wektor różnic zdefiniowany tak samo, jak w algorytmie 2 — w `liblinear` warunek stopu opiera się na badaniu stosunku redukcji wartości funkcji kwadratowej (aproksymującej funkcję celu) w bieżącym kroku do średniej redukcji, a próg wynosi $\min\left(\epsilon_{cg}, \sqrt[4]{\mathbf{g}^T \mathbf{M}^{-1} \mathbf{g}}\right)$ [23], gdzie ϵ_{cg} to zadana przez użytkownika dokładność,
- w minimalizacji kierunkowej, oprócz kryterium dostatecznej redukcji wartości funkcji⁷ (tak jak w algorytmie 1), autorzy skorzystali dodatkowo z kryterium wystarczającej krzywizny⁸, a procedurę minimalizacji zaczerpnęli z implementacji algorytmu L-BFGS-B w języku Fortran — w `liblinear` użyto natomiast wrotnej minimalizacji kierunkowej ze współczynnikami $\alpha = 0,01$ i $\beta = 0,5$.

W celach porównawczych obu implementacji przygotowano eksperyment, w którym trenowano dwa modele regresji logistycznej:

⁷znanego również jako kryterium Armijo [59]

⁸w połączeniu z kryterium Armijo są to tzw. (silne) kryteria Wolfe’a [59]

- `LogisticRegression` wykorzystujący do uczenia implementację metody Newtona z pakietu `scikit-learn` (w tabeli oznaczono ją `newton-cg`),
- `l2_logreg_ordinary` wykorzystujący do uczenia implementację metody Newtona z biblioteki `liblinear` (w tabeli oznaczono ją `liblinear`),

a uczenia dokonano na zbiorach danych wymienionych w tabeli 4.4 przy ustalonej wartości $\lambda = 10^{-4}$. Jakościowo obie procedury uczenia i oba modele są matematycznie równoważne, tzn. zbiegają do tego samego optimum. Przetestowano dwa ustawienia modeli: uwzględniający wyraz wolny (b+) i bez wyrazu wolnego (b-). Sprawdzono zachowanie obu metod dla dwóch precyzji typu zmiennoprzecinkowego: podwójnej (oznaczono ją literą d) i pojedynczej (oznaczono ją literą s). Dla typu podwójnej precyzji użyto $\epsilon = 10^{-8}$, natomiast dla typu pojedynczej precyzji wybrano $\epsilon = 10^{-4}$. We własnej implementacji modelu (`l2_logreg_ordinary`) wyłączono preconditioning. Wyniki przedstawiono w tabeli 4.3.

Na pierwszy rzut oka widać, że dla zbiorów rzadkich metoda `newton-cg` z pakietu `scikit-learn` pod względem czasu wykonania działała znacznie szybciej od metody Newtona z biblioteki `liblinear`, choć w większości przypadków wykonała więcej iteracji. Przyczyną tego stanu rzeczy są prawdopodobnie różnice w obsłudze macierzy rzadkiej, a także różnice wymienione wcześniej. Z kolei dla macierzy gęstych widać uderzającą różnicę na korzyść implementacji wziętej z biblioteki `liblinear`. Zauważono również skrócenie czasu uczenia po przejściu na pojedynczą precyzję, jednak w przypadku metody `newton-cg` dla testowanych macierzy rzadkich nie zaobserwowano dwukrotnego przyspieszenia.

4.4 Przebieg eksperymentu dla modeli z regularyzacją ℓ^2 oraz ℓ^1

Na podstawie zaprezentowanych wyników, w dalszym toku badań, jako punkt referencyjny przy porównywaniu modeli regresji logistycznej z regularyzacją ℓ^2 przyjęto model `l2_logreg_ordinary`. Dzięki temu uzyskano jednolity szkielet pętli uczącej z jednolitym warunkiem stopu oraz z jednolitymi bibliotekami numerycznej algebry liniowej. W pętli tej wygodnie można podmieniać i badać poszczególne elementy mające wpływ na zbieżność, np. sposób minimalizacji kierunkowej, dodanie preconditioningu czy podmianę metody optymalizacji. W przypadku modeli z regularyzacją ℓ^1 różnice pomiędzy modelami są znaczne, toteż każdy model stanowi osobną klasę.

Do przetestowania algorytmów wykorzystano cztery zbiory testowe: dwa zestawy danych gęstych (`extreme` i `PCam16k`) oraz dwa zestawy danych w postaci rzadkiej (`RCV1`

Tabela 4.3: Porównanie czasów optymalizacji dla wybranych zbiorów danych. k oznacza liczbę iteracji metody Newtona, n_{cg} to skumulowana liczba iteracji algorytmu gradientu sprzężonego, f^* to końcowa wartość funkcji, a t to czas wyrażony w sekundach.

zbiór	model	metoda	k	n_{cg}	f^*	t
RCV1	b+d	liblinear	10	41	7059,81	2,68
		newton-cg	15	51	7059,81	1,22
	b+s	liblinear	7	27	7059,75	1,34
		newton-cg	10	30	7059,81	0,99
	b-d	liblinear	9	31	7070,56	2,18
		newton-cg	16	48	7070,56	0,83
	b-s	liblinear	9	28	7070,50	1,26
		newton-cg	10	28	7070,56	0,84
news20	b+d	liblinear	11	66	6482,03	6,76
		newton-cg	17	94	6482,03	2,23
	b+s	liblinear	12	57	6481,78	3,95
		newton-cg	11	65	6482,02	2,71
	b-d	liblinear	13	62	6493,09	6,59
		newton-cg	17	87	6493,09	2,03
	b-s	liblinear	10	51	6492,83	3,80
		newton-cg	10	47	6493,08	1,88
extreme	b+d	liblinear	13	88	1970,53	12,92
		newton-cg	18	129	1970,53	16,19
	b+s	liblinear	17	89	1970,53	7,77
		newton-cg	13	89	1970,53	10,10
	b-d	liblinear	12	80	1970,53	11,82
		newton-cg	17	101	1970,53	12,95
	b-s	liblinear	7	38	1970,53	3,24
		newton-cg	14	75	1970,53	11,27
PCam16k	b+d	liblinear	6	31	9678,65	60,27
		newton-cg	16	56	9678,65	103,15
	b+s	liblinear	11	42	9678,63	20,53
		newton-cg	13	52	9678,65	46,27
	b-d	liblinear	8	33	10419,44	30,30
		newton-cg	12	42	10419,44	109,83
	b-s	liblinear	7	28	10419,43	13,36
		newton-cg	7	25	10419,44	45,08

Tabela 4.4: Zestawy danych wykorzystane w eksperymentach.

Dane	Opis	Rozmiar $n \times p$	Typ
extreme	Dane syntetyczne, losowo wygenerowane 10000 cech w problemie dwóch spiral	10000×10000	gęste
PCam16k	Dane rzeczywiste pobrane z repozytorium OpenML komendą: <code>X, y = fetch_openml(data_id=42811, return_X_y= True)</code>	16000×27649	gęste
RCV1	Baza textminingowa Reuters RCV1 [51] po wektoryzacji, wybrana klasa „33”	23149×47236	rzadkie
news20	dane textminingowe po wektoryzacji zebrane z 20 grup dyskusyjnych zawierające ponad 18000 wpisów, pobrane komendą: <code>data = fetch_20newsgroups_vectorized(subset='all', return_X_y= False, normalize=True)</code> wybrana klasa „1”	18846×130107	rzadkie

i news20). Charakterystyka zestawów danych wykorzystywanych w eksperymentach została przedstawiona w tabeli 4.4. W trakcie eksperymentu pokazany będzie wpływ liczby próbek n oraz wartości parametru regularizacyjnego λ na przebieg procesu uczenia. Na wykresach i w tabelach liczba próbek przedstawiona będzie w sposób znormalizowany, tj. w odniesieniu do liczby atrybutów (n/p). Wszystkie algorytmy zaprojektowane są pod przypadek $n \leq p$, natomiast zysk z proponowanych podejść widoczny jest dla mniejszych licznosci prób $n \ll p$. W eksperymentach przyjęto minimalną licznosc $n = 20$ próbek, natomiast wartość końcową uzależniono od rozmiaru dostępnych danych lub ograniczeń pamięciowych. Eksperymenty wykonano dla 10 wartości n/p rozmieszczonych w skali logarytmicznej. Oprócz liczby próbek na przebieg procesu uczenia ma wpływ parametr regularizacyjny λ . Część eksperymentów wykonano przy arbitralnie ustalonej wartości $\lambda = 10^{-4}$ przy zmianie parametru n . Przedstawiono również eksperyment ukazujący wpływ parametru λ na przebieg procesu uczenia.

W eksperymentach wykorzystano algorytmy opisane szczegółowo w rozdziale 2. W przypadku regularyzacji ℓ^2 główny model referencyjny to `l2_logreg_ordinary`, zaś

Tabela 4.5: Modele wykorzystane w eksperymentach wraz z szacowaną złożonością obliczeniową pojedynczej iteracji metody Newtona dla przypadku gęstego w wariancie optymistycznym i pesymistycznym; n — liczba próbek, p — liczba atrybutów, k_i — liczba iteracji procedury gradientu sprzężonego (zwykle mniejsza od liczby wierszy/kolumn macierzy układu).

Model	Opis	Złożoność iteracji
l2_logreg_ordinary	metoda Newtona wraz liniowym gradientem sprzężonym, odpowiednik procedury newton-cg pakietu scikit-learn	$O(npk_1)$, $O(np^2)$
l2_logreg_woodbury	odwrotność macierzy $\lambda I + X^T D X$ wyznaczana jest w każdej iteracji za pomocą formuły Woodbury’ego, por. sekcja 2.4.2	$O(npk_2)$, $O(n^2 p)$
l2_logreg_wdbr	odwrotność macierzy $\lambda I + X^T D X$ wyznaczana jest w każdej iteracji za pomocą formuły Woodbury’ego, przy wcześniejszym stabilizowaniu macierzy $X X^T$, por. sekcja 2.4.2; koszt stabilizowania iloczynu $O(n^2 p)$	$O(n^2 k_3)$
l2_logreg_qr	transformacja macierzy X^T za pomocą rozkładu QR, por. sekcja 2.4.1; sposób bardziej naturalny uwzględniając porządek kolumnowy macierzy X^T ; koszt wyznaczenia rozkładu $O(n^2 p)$	$O(n^2 k_4)$
l2_logreg_lq	transformacja macierzy X za pomocą rozkładu LQ, por. sekcja 2.4.1; sposób bardziej naturalny, biorąc pod uwagę rozmiar macierzy ($n < p$); koszt wyznaczenia rozkładu $O(n^2 p)$	$O(n^2 k_5)$

pozostałe to własne implementacje autora: l2_logreg_woodbury, l2_logreg_wdbr, l2_logreg_qr, l2_logreg_lq, podsumowane w tabeli 4.5. Wszystkie procedury mają podobną strukturę: zewnętrzna pętla (kroki algorytmu uczenia) i wewnętrzna (iteracyjne rozwiązywanie odpowiedniego układu równań do wyznaczenia kolejnego kierunku) oraz skalarna optymalizacja kierunkowa. Wszystkie porównywane algorytmy mają ustawione identycznie warunek zatrzymania pętli iteracyjnej rozwiązującej układ równań liniowych

oraz pętli zewnętrznej zatrzymującej algorytm uczenia. Każdy algorytm zaczynał optymalizację z punktu początkowego $w = 0$.

Od strony formalno-teoretycznej wszystkie te procedury powinny mieć równoważne iteracje pętli zewnętrznej w tym sensie, że w każdym kroku spadek wartości funkcji celu powinien następować o tę samą wartość. Rozwiązywane w każdym kroku układy równań są różne, stąd liczby iteracji algorytmów rozwiązujących układ równań liniowych mogą od siebie odbiegać, bo inne są także wymiarowości macierzy układów równań, jednakże wektory wag w każdej iteracji mogą być sprowadzone do postaci równoważnej (wspólnej przestrzeni oryginalnej), np. w przypadku metody z rozkładem QR/LQ konieczne jest przemnożenie wag w przestrzeni transformowanej przez macierz Q .

Przed uruchomieniem eksperymentów zweryfikowano empirycznie poprawność algorytmów na mniejszych zbiorach danych i wszystkie algorytmy dawały identyczne wyniki, zarówno jeżeli idzie o liczbę iteracji pętli zewnętrznej, jak i wartości funkcji celu w kolejnych iteracjach oraz wynik końcowy w postaci wag.

Od strony numerycznej na wykonanie algorytmu mają wpływ następujące elementy (por. sekcja 2.4):

procedura rozwiązywania układu równań liniowych: wykorzystywane są gradient sprzężony wraz z przyciętą metodą Newtona [23] oraz SYMMLQ [60];

bias użycie wyrazu wolnego (b+) lub model bez wyrazu wolnego (b-);

użycie dokładnej minimalizacji kierunkowej (e+) lub przybliżonej (e-);

preconditioner wykorzystanie tzw. preconditioningu (macierzy ściskającej) celem poprawy uwarunkowania macierzy układu i zbieżności procedury iteracyjnej rozwiązującej układ równań; możliwe wartości: p+ (użycie preconditionera), p- (brak, układ oryginalny);

typ zmiennopozycyjny d (podwójna precyzja, 64 bity), s (pojedyncza precyzja, 32 bity) [14].

Jeden z eksperymentów porównawczych przedstawiono w tabeli 4.6. Tabela zawiera wyniki uzyskane na zbiorze `lsvt` (por. repozytorium OpenML, zbiór `id=1484`) o rozmiarze $n \times p = 126 \times 309$ przy ustawieniach $\lambda = 10^{-4}$ oraz identycznych pozostałych ustawieniach numerycznych. Jak można zauważyć, w pierwszej iteracji wyniki wszystkich algorytmów zgadzają się z dokładnością do wyświetlanej precyzji. W kolejnych iteracjach widać, że identyczne numerycznie są algorytmy `l2_logreg_ordinary` oraz `l2_logreg_QR`, mimo że drugi z algorytmów pracuje w przetransformowanej przestrzeni, tj. na danych macierzy R z rozkładu QR. Dla obu algorytmów bazujących na formule Woodbury’ego widać drobną rozbieżność numeryczną względem dwóch pozostałych algo-

Tabela 4.6: Porównanie zbieżności testowanych algorytmów na zbiorze lsvt (por. repozytorium OpenML zbiór id=1484) o rozmiarze $n \times p = 126 \times 309$ przy ustawieniach $\lambda = 10^{-4}$ oraz ustawieniach: b-p-e-d (bez wyrazu wolnego, bez preconditioningu, bez dokładnej minimalizacji kierunkowej w typie podwójnej precyzji); it oznacza numer iteracji, gTd rzut gradientu na nowy kierunek (pochodna kierunkowa), f wartość funkcji celu, cg liczba iteracji procedury rozwiązującej układ równań liniowych.

it	12_logreg_ordinary			12_logreg_woodbury			12_logreg_wdbr			12_logreg_qr		
	gTd	f	cg	gTd	f	cg	gTd	f	cg	gTd	f	cg
1	-15,7842	79,1355	4	-15,7842	79,1355	3	-15,7842	79,1355	3	-15,7842	79,1355	4
2	-0,287993	78,9777	4	-0,286014	78,9779	5	-0,287309	78,9778	5	-0,287993	78,9777	4
3	-7,879e-03	78,9737	4	-8,449e-03	78,9737	6	-7,731e-03	78,9738	6	-7,879e-03	78,9737	4
4	-6,324e-06	78,9737	4	-7,944e-05	78,9737	3	-2,655e-04	78,9737	2	-6,324e-06	78,9737	4
5	-3,981e-12	78,9737	3	-6,309e-08	78,9737	7	-7,165e-07	78,9737	7	-3,981e-12	78,9737	3
6				-1,811e-13	78,9737	3	-6,599e-14	78,9737	6			

rytmów, zwłaszcza w pochodnej kierunkowej $\mathbf{g}^T \mathbf{d}$ (wyraźniejszą począwszy od czwartej iteracji), oraz oba algorytmy wykonały o jedną iterację więcej, gdyż warunek stopu był ustalony na $|\mathbf{g}^T \mathbf{d}| < 10^{-8}$. Algorytmy bazujące na formule Woodbury’ego mają do rozwiązania nieco inny układ równań liniowych, stąd liczby iteracji rozwiązujących dany układ (kolumna cg) mogą być różne. Ponadto eksperymenty prezentowane w dalszych sekcjach wskazują, że oba algorytmy działają nieco stabilniej przy włączonym preconditioningu. Patrząc na wartość funkcji celu (f) wszystkie algorytmy zachowują się identycznie z dokładnością do wyświetlanej precyzji.

4.5 Wyniki eksperymentów dla regularyzacji ℓ^2

W kolejnych tabelach i na wykresach przedstawiono 8 (tylko double) lub 16 wariantów algorytmu uczenia kodowanych odpowiednią konfiguracją parametrów. Przykładowo wyrażenie b+p-e-d oznacza model z wyrazem wolnym, bez użycia preconditioningu, bez dokładnej minimalizacji kierunkowej, a obliczenia w typie podwójnej precyzji. Podstawowym typem numerycznym jest 64 bitowy typ zmiennopozycyjny (double) [14]. Wszystkie eksperymenty wykonano również dla typu 32 bitowego (single), głównie po to, aby przebadać zachowanie algorytmów przy obniżonej precyzji, gdzie problemy numeryczne różnych wariantów są bardziej widoczne.

Tabele 4.7, 4.8, 4.9, 4.10 przedstawiają podsumowanie wyników eksperymentów dla różnej liczby próbek uczących (n) oraz wskazują obszary parametrów z przewagą proponowanego algorytmu 12_logreg_wdbr nad pozostałymi algorytmami.

Tabela 4.7: Podsumowanie wyników dla danych *extreme*, dla różnych rozmiarów zbioru danych (n). Przewagę algorytmu *l2_logreg_wdbr* nad pozostałymi oznaczono przez „+”. Typ zmien-nopozycyjny (d/s), b (+/-) model z/bez wyrazu wolnego, p (+/-) model z/bez preconditionerem, e (+/-) model z/bez dokładnej minimalizacji kierunkowej.

type	b	p	e	20	36	68	125	232	429	793	1465	2707	4999
d	+	+	+	-	-	+	+	+	+	+	+	+	+
			-	-	-	+	+	+	+	+	+	+	+
		-	+	-	-	+	+	+	+	+	+	+	+
			-	-	-	+	+	+	+	+	+	+	+
	-	+	+	-	-	-	+	+	+	+	+	+	+
			-	-	-	-	+	+	+	+	+	+	+
		-	+	-	-	-	+	+	+	+	+	+	+
			-	-	-	-	+	+	+	+	+	+	+
s	+	+	+	-	-	-	+	+	+	+	+	+	+
			-	-	-	-	+	+	+	+	+	+	+
		-	+	-	-	-	+	+	+	+	+	+	+
			-	-	-	-	+	+	+	+	+	+	+
	-	+	+	-	-	-	+	+	+	+	+	+	+
			-	-	-	-	+	+	+	+	+	+	+
		-	+	-	-	-	-	+	+	+	+	+	+
			-	-	-	-	+	+	+	+	+	+	+

Na rysunkach 4.1, 4.3, 4.5, 4.7 przedstawiono czasy wykonania poszczególnych algorytmów wraz ze wzrostem liczby próbek, uśrednione po dziesięciu powtórzeniach z losowym wyborem zestawu uczącego. Na wykresach wąsami zaznaczono również odchylenia standardowe wyników. Na wykresach na osi OX przedstawiono unormowaną liczbę próbek n/p . Wyniki przedstawiono tylko dla typu *double*. Ponadto dla danych rzadkich nie wykonano algorytmów z rozkładami QR/LQ, gdyż nie gwarantują one rzadkich macierzy Q i R , a zastąpienie rzadkiej macierzy danych X przez gęste macierze R i Q wydaje się dalece nieefektywne pamięciowo.

Na rysunkach 4.2, 4.4, 4.6, 4.8 zwizualizowano tempo zbieżności algorytmów dla

Tabela 4.8: Podsumowanie wyników dla danych PCam16k, dla różnych rozmiarów zbioru danych (n). Przewagę algorytmu 12_logreg_wdbr nad pozostałymi oznaczono przez „+”. Typ zmien-nopozycyjny (d/s), b (+/-) model z/bez wyrazu wolnego, p (+/-) model z/bez preconditionerem, e (+/-) model z/bez dokładnej minimalizacji kierunkowej.

typ	b	p	e	n	20	38	75	147	286	557	1085	2112	4111	8000
d	+	+	+	-	-	+	+	+	+	+	+	+	+	-
			-	+	-	+	+	+	+	+	+	+	+	-
		-	+	+	+	+	+	+	+	+	+	+	+	-
			-	+	+	+	+	+	+	+	+	+	+	-
	-	+	+	+	-	+	+	+	+	+	+	+	+	+
			-	+	+	+	+	+	+	+	+	+	+	+
		-	+	-	-	+	+	+	+	+	+	+	+	-
			-	+	-	+	+	+	+	+	+	+	+	-
s	+	+	+	-	-	-	+	+	+	+	+	+	+	-
			-	-	-	-	+	+	+	+	+	+	+	-
		-	+	-	-	+	+	+	+	+	+	+	+	-
			-	-	-	+	+	+	+	+	+	+	+	+
	-	+	+	-	-	-	+	+	+	+	+	+	+	-
			-	-	-	-	+	+	+	+	+	+	+	-
		-	+	-	-	-	+	+	+	+	+	+	+	-
			-	-	-	-	+	+	+	+	+	+	+	-

wybranych rozmiarów danych uczących (n). Wyniki przedstawiono dla wszystkich 16 wariantów algorytmów. Mimo że wszystkie algorytmy powinny mieć zbliżone iteracje, widać pewne odstępstwa spowodowane błędami numerycznymi oraz uwarunkowaniem macierzy. Widać przykładowo, że wpływ preconditionera Jacobiego jest niekorzystny dla algorytmu 12_logreg_ordinary, natomiast korzystnie wpływa na czas wykonania pozostałych algorytmów (por. 4.4). Implementacja 12_logreg_ordinary znacznie lepiej zachowuje się dla modelu bez wyrazu wolnego. Ogólnie zgodnie z oczekiwaniami widać zysk zarówno z zastosowania rozkładu QR, jak i formuły Woodbury’ego z tablicowaniem macierzy XX^T w stosunku do algorytmu referencyjnego. Ponadto patrząc

Tabela 4.9: Podsumowanie wyników dla danych news20, dla różnych rozmiarów zbioru danych (n). Przewagę algorytmu 12_logreg_wdbr nad pozostałymi oznaczono przez „+”. Typ zmien-nopozycyjny (d/s), b (+/-) model z/bez wyrazu wolnego, p (+/-) model z/bez preconditionerem, e (+/-) model z/bez dokładnej minimalizacji kierunkowej.

typ	b	p	e	n	20	42	91	196	419	898	1922	4114	8805	18846
d	+	+	+	+	+	+	+	+	+	+	+	-	-	-
			-	+	+	+	+	+	+	+	+	-	-	-
		-	+	+	+	+	+	+	+	+	+	-	-	-
			-	+	+	+	+	+	+	+	+	-	-	-
		-	+	+	+	+	+	+	+	+	-	-	-	-
			-	+	+	+	+	+	+	+	+	-	-	-
	-	+	+	+	+	+	+	+	+	+	-	-	-	-
			-	+	+	+	+	+	+	+	+	-	-	-
		-	+	-	-	-	-	-	-	-	-	-	-	-
			-	-	-	-	+	-	-	-	-	-	-	-
		-	+	+	+	+	+	+	+	+	+	-	-	-
			-	+	+	+	+	+	+	+	+	-	-	-
s	+	+	+	+	+	+	+	+	+	+	+	-	-	-
			-	+	+	+	+	+	+	+	+	-	-	-
		-	+	-	-	-	-	-	-	-	-	-	-	-
			-	-	-	-	-	-	-	-	-	-	-	-
		-	+	-	-	-	-	-	-	-	-	-	-	-
			-	-	-	-	-	-	-	-	-	-	-	-
	-	+	+	-	-	-	-	-	-	-	-	-	-	-
			-	-	-	-	-	-	-	-	-	-	-	-
		-	+	-	-	-	-	-	-	-	-	-	-	-
			-	-	-	-	-	-	-	-	-	-	-	-
		-	+	-	-	-	-	-	-	-	-	-	-	-
			-	-	-	-	-	-	-	-	-	-	-	-

na tabele i wykresy można zobaczyć, że dla małych prób procedura 12_logreg_wdbr działała najkorzystniej, z wyjątkiem pojedynczych przypadków bardzo małych prób, gdzie korzystniejszy był algorytm z rozkładem QR. Dla danych gęstych widać przewagę proponowanej procedury praktycznie w całym badanym zakresie (do $n/p = 0,3$ lub $n/p = 0,5$). Zgodnie z oczekiwaniami, rozkład QR był korzystniejszy niż rozkład LQ, ponieważ jest lepiej dostosowany do porządku kolumnowego, w jakim przechowywane były dane po transpozycji. Dla danych rzadkich zysk nie jest tak wyraźny, jak dla danych gęstych. Dokładna analiza złożoności obliczeniowej jest w tym przypadku zagadnieniem złożonym, gdyż zależy ona nie tylko od stopnia wypełnienia macierzy, ale i od stopnia wypełnienia

Tabela 4.10: Podsumowanie wyników dla danych RCV1, dla różnych rozmiarów zbioru danych (n). Przewagę algorytmu `l2_logreg_wdbr` nad pozostałymi oznaczono przez „+”. Typ zmien-nopozycyjny (d/s), b (+/-) model z/bez wyrazu wolnego, p (+/-) model z/bez preconditionerem, e (+/-) model z/bez dokładnej minimalizacji kierunkowej.

typ	b	p	e	n	19	42	95	210	462	1017	2233	4902	10760	23617
d	+	+	+	+	+	+	+	+	+	+	-	-	-	-
			-	+	+	+	+	+	+	+	-	-	-	-
		-	+	+	+	+	+	+	+	+	-	-	-	-
			-	-	+	+	+	+	+	+	-	-	-	-
	-	+	+	+	+	+	+	+	+	+	-	-	-	-
			-	+	+	+	+	+	+	+	-	-	-	-
		-	+	+	+	+	+	+	+	+	-	-	-	-
			-	+	+	+	+	+	+	+	-	-	-	-
s	+	+	+	+	+	+	+	+	+	+	+	-	-	-
			-	+	+	+	+	+	+	+	+	-	-	-
		-	+	-	-	-	-	-	-	+	-	-	-	-
			-	-	-	-	-	-	-	+	-	-	-	-
	-	+	+	+	+	-	-	-	-	+	+	-	-	-
			-	-	+	+	+	-	-	+	+	-	-	-
		-	+	-	-	-	-	-	-	+	-	-	-	-
			-	-	-	-	-	-	-	+	-	-	-	-

kolumn/wierszy oraz rzadkości wyniku (por. [1]).

Kolejny eksperyment pokazuje wpływ parametru λ na przebieg procesu uczenia. Eksperyment przeprowadzono zaczynając od dużej wartości parametru regularizacyjnego i zmniejszając go kolejno $\lambda = 10^3, 10^2, \dots, 10^{-4}$. Eksperyment był prowadzony w ten sposób, że wagi z poprzednio wyuczonego modelu stanowiły punkt startowy dla kolejnego modelu⁹.

Dla dużych wartości λ w funkcji celu $Q(w; dane) = L(w; dane) + R(w)$ dominuje składnik kwadratowy R , dzięki czemu optymalizowana funkcja jest dobrze przybliżana za

⁹ang. warm start

pomocą lokalnej aproksymacji kwadratowej, co prowadzi do szybkiej zbieżności. Wraz ze spadkiem wartości λ zaczyna dominować składnik związany z funkcją wiarygodności L i problem optymalizacyjny staje się bardziej złożony. Ponieważ koszt iteracji w proponowanych algorytmach `l2_logreg_wdbr` oraz `l2_logreg_qr` jest mniejszy, więc zysk z zastosowania rozkładu QR bądź formuły Woodbury’ego z tablicowaniem macierzy $\mathbf{X}\mathbf{X}^T$ będzie bardziej widoczny, zwłaszcza w sytuacji, gdy algorytm uczenia do zbieżności wymaga większej liczby iteracji. Należy podkreślić, że oba algorytmy przed wejściem do pętli algorytmu uczenia mają pewien narzut związany z policzeniem rozkładu QR lub wyznaczeniem macierzy $\mathbf{X}\mathbf{X}^T$.

Na rysunkach 4.9, 4.10, 4.11 4.12 przedstawiono wyniki eksperymentu wyznaczania rodziny modeli dla parametru regularyzacyjnego zmieniającego się od dużych wartości do małych: $\lambda = 10^3, 10^2, \dots, 10^{-4}$. Rysunki przedstawiają skumulowany czas eksperymentu w funkcji malejącej λ . Widać, że dla dużych wartości λ zgodnie z oczekiwaniami przewaga algorytmów `l2_logreg_wdbr` oraz `l2_logreg_qr` jest mniejsza, natomiast staje się ona wyraźniejsza dla coraz mniejszych wartości λ . Ponadto widać w większości obszarów parametru λ przewagę proponowanego algorytmu wykorzystującego formułę Woodbury’ego. Na gęstych danych przewaga obu algorytmów (zwłaszcza dla małych prób) jest wyraźna — algorytm bazujący na rozkładzie QR ma nieznaczną przewagę tylko dla bardzo małych prób oraz dużej wartości λ . Jedynie na rzadkim zbiorze RCV1 widać przewagę referencyjnego modelu `l2_logreg_ordinary`, która maleje wraz ze spadkiem wartości λ . Widać również, że przewaga proponowanych algorytmów maleje wraz z podnoszeniem rozmiaru próby. Można również zauważyć, że dla coraz większych prób na zbiorach gęstych zauważalna jest nieznaczną przewagę algorytmu `l2_logreg_woodbury` nad algorytmem `l2_logreg_ordinary` i również pozostałymi. Patrząc na złożoność obliczeniową kroku iteracji wewnętrznej (rozwiązywanie układu równań liniowych przy wyznaczaniu nowego kierunku) trudno wytłumaczyć jednoznacznie to zachowanie. Oba kroki mają podobną złożoność obliczeniową ze wskazaniem na korzyść modelu `l2_logreg_ordinary` — jedyna różnica jest taka, że w przypadku formuły Woodbury’ego rozwiązywany jest układ o mniejszej wymiarowości ($n \times n$ zamiast $p \times p$), co może przełożyć się na mniejszą liczbę iteracji procedury iteracyjnej rozwiązującej układ równań.

Na rysunkach 4.13, 4.14, 4.15 4.16 przedstawiono wyniki eksperymentu wyznaczania rodziny modeli dla parametru regularyzacyjnego zmieniającego się od dużych wartości do małych: $\lambda = 10^3, 10^2, \dots, 10^{-4}$. W odróżnieniu od poprzedniego eksperymentu, ten był prowadzony w ten sposób, że kolejny model startował zawsze od wag $\mathbf{w} = \mathbf{0}$. Rysunki przedstawiają czas uśredniony jednorazowego uczenia modelu dla zadanej wartości λ .

Tak jak w poprzednim przypadku, dla dużych wartości λ przewaga algorytmów `l2_logreg_wdbr` oraz `l2_logreg_qr` jest mniejsza, a staje się ona wyraźniejsza dla coraz mniejszych wartości λ . Ponadto można zaobserwować w większości obszarów parametru λ przewagę proponowanego algorytmu wykorzystującego formułę Woodbury’ego. Dla gęstych danych przewaga obu algorytmów jest wyraźniejsza — algorytm bazujący na rozkładzie QR ma nieznaczną przewagę jedynie dla bardzo małych prób. Dla rzadkich zbiorów danych pewien zysk był widoczny dla małych prób na zbiorze `news20`, zaś na zbiorze `RCV1` za każdym razem krócej trwało uczenie modelu referencyjnego. Pokazuje to, że rodzaj rzadkości danych ma istotny wpływ na zachowanie algorytmów.

Przykładowe zagregowane wyniki ilościowe przedstawiono w tabelach 4.11 i 4.12. W tabeli 4.11 umieszczono wyniki podsumowujące eksperyment dla zbioru `PCam16k` (por. wykres 4.14). Wyniki dotyczą przypadku modelu z wyrazem wolnym (b+): dla każdego z modeli wybrano najlepsze ustawienia parametrów p i e , po czym uśredniono czasy z dziesięciu powtórzeń eksperymentu. Analogiczne wyniki dla przypadku rzadkiego zbioru `news20` (por. wykres 4.15) przedstawia tabela 4.12.

Na podstawie przedstawionych wyników widać również, że przewaga algorytmów `l2_logreg_wdbr` oraz `l2_logreg_qr` maleje wraz z podnoszeniem rozmiaru próby, natomiast dla dużych prób na zbiorze gęstym obserwowana jest pewna drobna przewaga algorytmu `l2_logreg_woodbury` względem `l2_logreg_ordinary`. Porównując rozwiązywane układy równań dla modelu `l2_logreg_ordinary`:

$$\left(\mathbf{X}^T \mathbf{D} \mathbf{X} + \lambda \mathbf{I} \right) \mathbf{d} = \mathbf{X}^T \mathbf{v} - \lambda \mathbf{w}, \quad (4.2)$$

oraz proponowanego modelu `l2_logreg_woodbury` (wersja bez tablicowania):

$$\left(\mathbf{D}^{-1} + \frac{1}{\lambda} \mathbf{X} \mathbf{X}^T \right) \mathbf{z} = \frac{1}{\lambda} \mathbf{X} \mathbf{X}^T \mathbf{v} - \mathbf{X} \mathbf{w}. \quad (4.3)$$

widać, że złożoność obliczeniowa jest podobna $O(np)$, z przewagą dla formuły klasycznej — jedyne nasuwające się wyjaśnienie jest takie, że układy różnią się wymiarowością na korzyść obliczeń z użyciem formuły Woodbury’ego, tj. macierz układu ma wymiar $n \times n$ zamiast $p \times p$, jak w przypadku standardowym.

W eksperymentach przebadano również wpływ typu zmiennopozycyjnego na dokładność i czas obliczeń. Eksperymenty pokazują, że wszystkie algorytmy przy odpowiednich ustawieniach zachowują się generalnie stabilnie przy obniżonej precyzji, czasami wymagając większej liczby iteracji dla uzyskania zbieżności. Widoczny jest również krótszy czas obliczeń. Wniosek ten może mieć znaczenie przy wykonywaniu obliczeń z użyciem kart

Tabela 4.11: Podsumowanie wyników dla zbioru danych PCam16k — tabela przedstawia czas obliczeń dla najlepszych ustawień modelu, zagregowane po n , λ oraz rodzaju modelu.

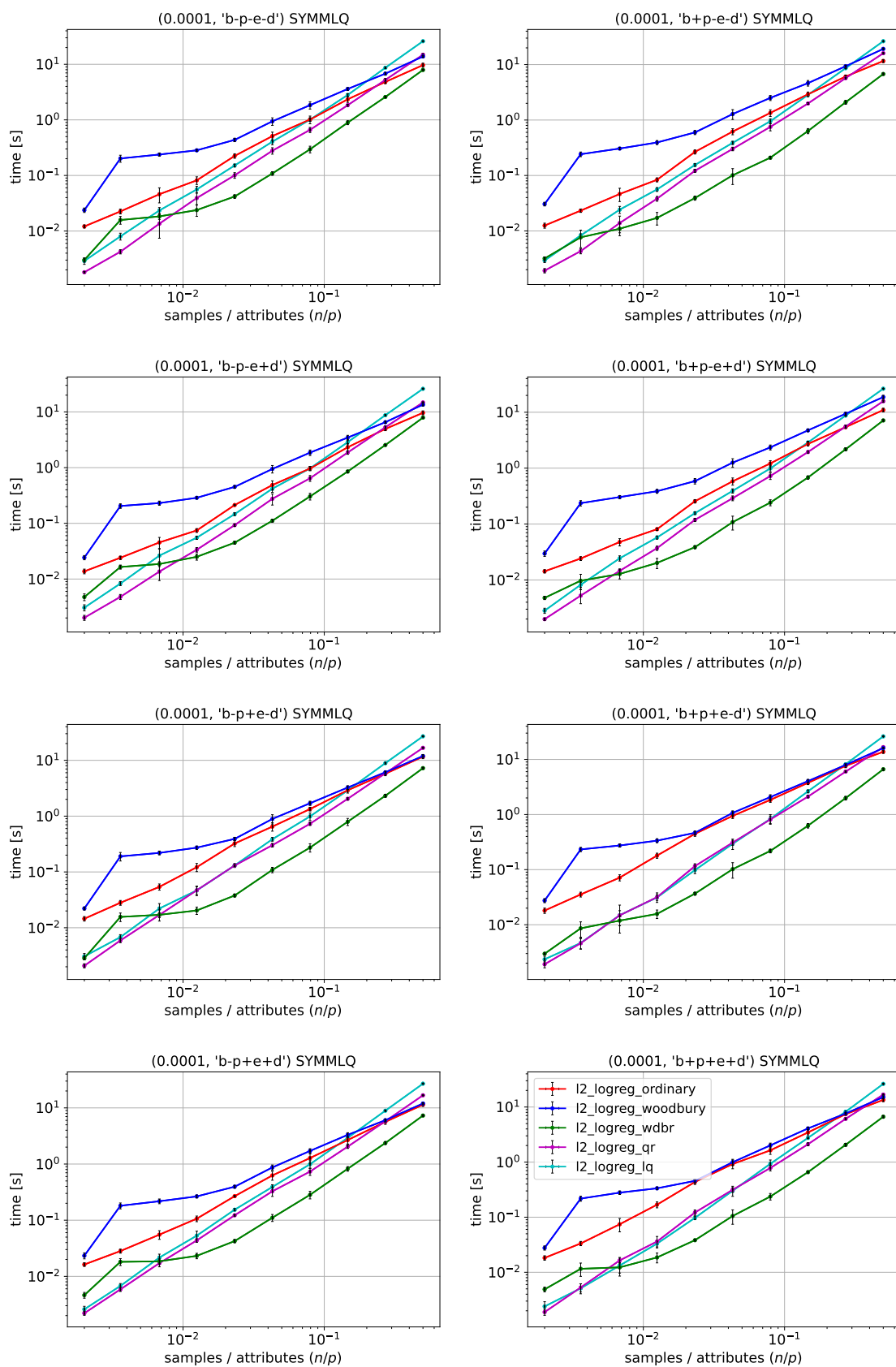
n	λ model	10^{-4}	10^{-3}	10^{-2}	10^{-1}	10^0	10^1	10^2	10^3
20	l2_logreg_ordinary	0,0987	0,068	0,0471	0,0407	0,0298	0,0233	0,0169	0,0091
	l2_logreg_qr	0,0149	0,0191	0,0171	0,0188	0,0201	0,02	0,0181	0,0182
	l2_logreg_wdbr	0,0165	0,00979	0,00729	0,00756	0,00452	0,00373	0,00273	0,0019
	l2_logreg_woodbury	0,0975	0,0607	0,0401	0,0312	0,0245	0,0186	0,0135	0,00796
38	l2_logreg_ordinary	0,195	0,134	0,0856	0,0745	0,0525	0,0436	0,0311	0,0161
	l2_logreg_qr	0,0331	0,0279	0,0378	0,0352	0,0325	0,0402	0,0212	0,0339
	l2_logreg_wdbr	0,0746	0,0449	0,0507	0,034	0,0345	0,017	0,0148	0,00941
	l2_logreg_woodbury	0,509	0,262	0,194	0,142	0,111	0,0856	0,0612	0,0361
75	l2_logreg_ordinary	0,458	0,318	0,214	0,178	0,122	0,0964	0,0696	0,0364
	l2_logreg_qr	0,075	0,0826	0,0906	0,0913	0,0837	0,0871	0,0777	0,0942
	l2_logreg_wdbr	0,0934	0,0501	0,0618	0,0345	0,0237	0,0282	0,0177	0,0145
	l2_logreg_woodbury	0,596	0,335	0,291	0,218	0,165	0,126	0,0931	0,0479
147	l2_logreg_ordinary	1,17	0,738	0,51	0,453	0,297	0,224	0,161	0,0808
	l2_logreg_qr	0,18	0,254	0,253	0,262	0,264	0,262	0,264	0,264
	l2_logreg_wdbr	0,165	0,124	0,1	0,0797	0,0648	0,0505	0,0432	0,0284
	l2_logreg_woodbury	1,16	0,664	0,568	0,426	0,298	0,227	0,159	0,092
286	l2_logreg_ordinary	3,25	2,19	1,43	1,04	0,891	0,649	0,453	0,211
	l2_logreg_qr	0,666	0,693	0,678	0,695	0,679	0,688	0,692	0,688
	l2_logreg_wdbr	0,269	0,212	0,181	0,143	0,132	0,113	0,0901	0,0681
	l2_logreg_woodbury	2,93	1,67	1,35	0,881	0,706	0,551	0,384	0,211
557	l2_logreg_ordinary	4,96	3,41	2,25	1,76	1,26	0,934	0,659	0,335
	l2_logreg_qr	1,26	1,35	1,32	1,35	1,43	1,45	1,52	1,54
	l2_logreg_wdbr	0,566	0,392	0,391	0,276	0,296	0,269	0,235	0,197
	l2_logreg_woodbury	4,08	2,44	2,07	1,58	1,05	0,898	0,611	0,336
1085	l2_logreg_ordinary	9,96	5,99	4,4	3,84	2,67	1,92	1,32	0,647
	l2_logreg_qr	2,78	3,48	3,37	3,54	3,6	3,73	3,78	3,77
	l2_logreg_wdbr	1,56	1,02	1,1	1,05	0,906	0,899	0,781	0,704
	l2_logreg_woodbury	7,47	3,76	4,09	2,17	2,07	1,68	1,23	0,666
2112	l2_logreg_ordinary	20	12,8	8,59	6,59	4,93	3,4	2,53	1,31
	l2_logreg_qr	10	9,6	9,54	9,48	9,52	9,58	9,57	9,49
	l2_logreg_wdbr	4,67	4,01	3,71	3,49	3,26	3,13	2,76	2,62
	l2_logreg_woodbury	11,8	9,08	6,47	5,17	4,55	3,53	2,34	1,31
4111	l2_logreg_ordinary	41,6	27,7	17,7	13,9	11	7,74	5,46	2,64
	l2_logreg_qr	29,9	28,7	28,4	28,1	27,9	27,9	27,9	27,8
	l2_logreg_wdbr	15,8	13,4	13,4	12	11,9	11,1	11	10,6
	l2_logreg_woodbury	29,7	17	12,6	11,5	8,29	6,64	4,45	2,68
8000	l2_logreg_ordinary	87,2	52,4	36,3	31,3	20,9	15,4	11	5,33
	l2_logreg_qr	94,8	89,7	87,7	86,6	86,8	86,2	85,9	84,8
	l2_logreg_wdbr	45,9	44,7	44,1	41,7	42	40,4	40,8	39,4
	l2_logreg_woodbury	76,3	32,4	30,8	21	14,2	12,3	8,78	5,38

graficznych GP-GPU. W obliczeniach na kartach chętnie korzysta się z typów obniżonej precyzji, gdyż oszczędzamy na narzutach związanych z przesyłaniem danych, przesyłając

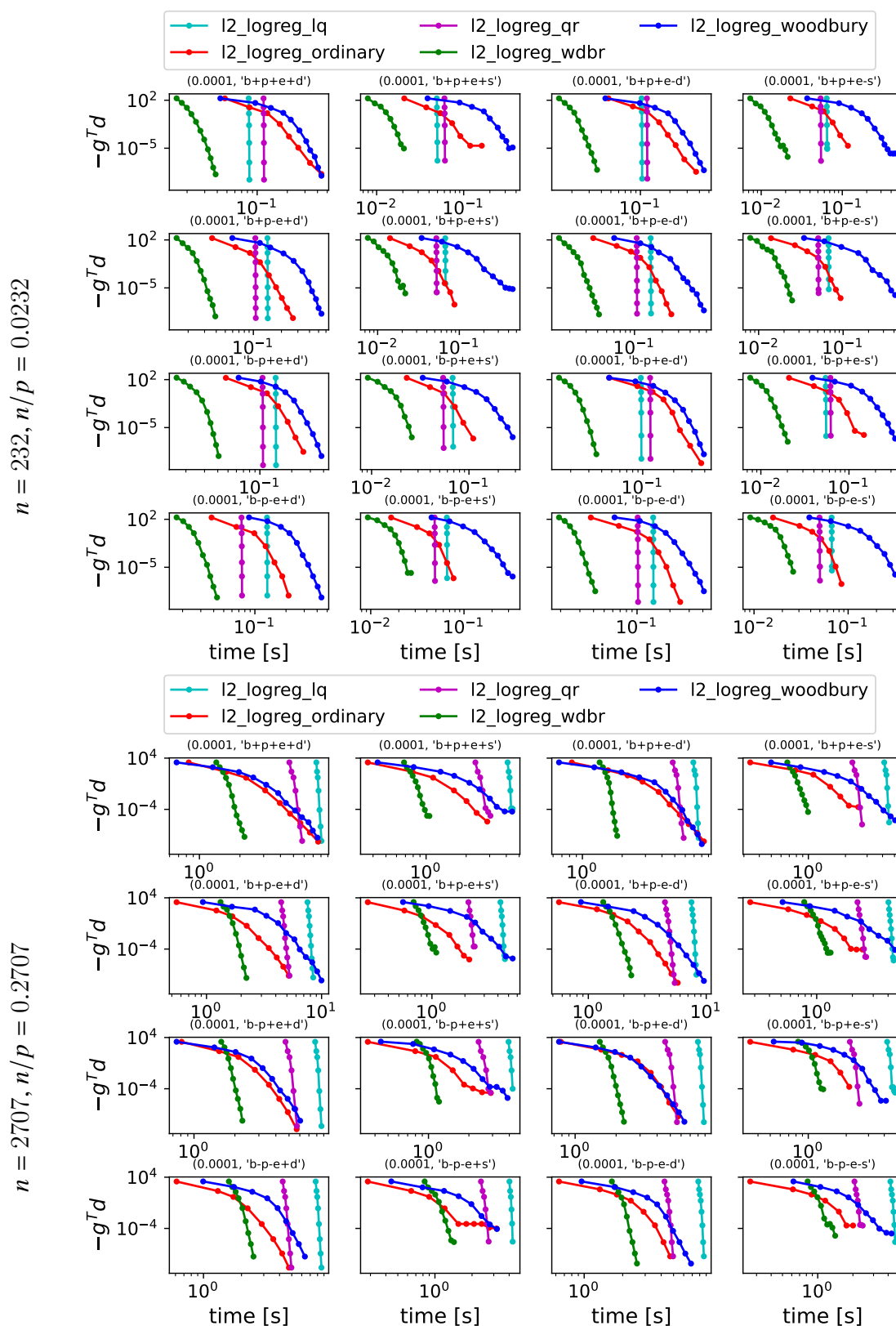
Tabela 4.12: Podsumowanie wyników dla zbioru danych news20 — tabela przedstawia czas obliczeń dla najlepszych ustawień modelu, zagregowane po n , λ oraz rodzaju modelu.

n	λ model	10^{-4}	10^{-3}	10^{-2}	10^{-1}	10^0	10^1	10^2	10^3
20	12_logreg_ordinary	0,586	0,406	0,303	0,217	0,168	0,12	0,0968	0,0478
	12_logreg_wdbr	0,321	0,271	0,211	0,155	0,122	0,0927	0,071	0,0419
	12_logreg_woodbury	1,97	1,48	0,998	0,622	0,439	0,342	0,246	0,149
42	12_logreg_ordinary	0,546	0,443	0,369	0,274	0,198	0,129	0,0955	0,054
	12_logreg_wdbr	0,322	0,262	0,201	0,157	0,119	0,0927	0,0668	0,041
	12_logreg_woodbury	2,22	1,66	0,998	0,63	0,444	0,346	0,246	0,142
91	12_logreg_ordinary	0,625	0,46	0,315	0,227	0,173	0,121	0,0832	0,0468
	12_logreg_wdbr	0,33	0,256	0,192	0,151	0,112	0,0855	0,0598	0,0343
	12_logreg_woodbury	2,35	1,63	0,947	0,583	0,426	0,319	0,22	0,114
196	12_logreg_ordinary	0,74	0,516	0,358	0,26	0,193	0,127	0,0833	0,0452
	12_logreg_wdbr	0,334	0,259	0,198	0,156	0,116	0,0882	0,0621	0,0355
	12_logreg_woodbury	2,7	1,65	0,953	0,637	0,433	0,327	0,22	0,112
419	12_logreg_ordinary	0,795	0,549	0,375	0,269	0,196	0,13	0,0847	0,0458
	12_logreg_wdbr	0,37	0,276	0,214	0,168	0,127	0,098	0,0696	0,0421
	12_logreg_woodbury	2,92	1,72	0,95	0,654	0,447	0,321	0,214	0,108
898	12_logreg_ordinary	0,912	0,604	0,405	0,284	0,207	0,134	0,0899	0,049
	12_logreg_wdbr	0,432	0,328	0,258	0,208	0,163	0,132	0,102	0,0722
	12_logreg_woodbury	3,15	1,79	1,03	0,696	0,478	0,35	0,234	0,122
1922	12_logreg_ordinary	1,12	0,72	0,491	0,35	0,254	0,163	0,105	0,0527
	12_logreg_wdbr	0,793	0,6	0,495	0,422	0,36	0,319	0,282	0,24
	12_logreg_woodbury	3,72	2,01	1,14	0,769	0,53	0,393	0,261	0,131
4114	12_logreg_ordinary	1,54	1	0,708	0,523	0,392	0,263	0,161	0,0995
	12_logreg_wdbr	3,67	2,86	2,45	2,19	2,04	1,95	1,85	1,76
	12_logreg_woodbury	4,58	2,48	1,41	0,951	0,668	0,509	0,351	0,19
8805	12_logreg_ordinary	2,75	1,81	1,26	0,936	0,698	0,468	0,321	0,143
	12_logreg_wdbr	21,4	17,6	15	13	12,1	11,3	10,9	10,1
	12_logreg_woodbury	6,25	3,27	1,89	1,33	0,962	0,733	0,503	0,266
18846	12_logreg_ordinary	5,88	4,03	2,81	2,11	1,57	1,05	0,698	0,37
	12_logreg_wdbr	106	84,6	67,3	57,8	53	50,4	47,8	45,3
	12_logreg_woodbury	11,2	5,82	3,5	2,45	1,73	1,32	0,895	0,453

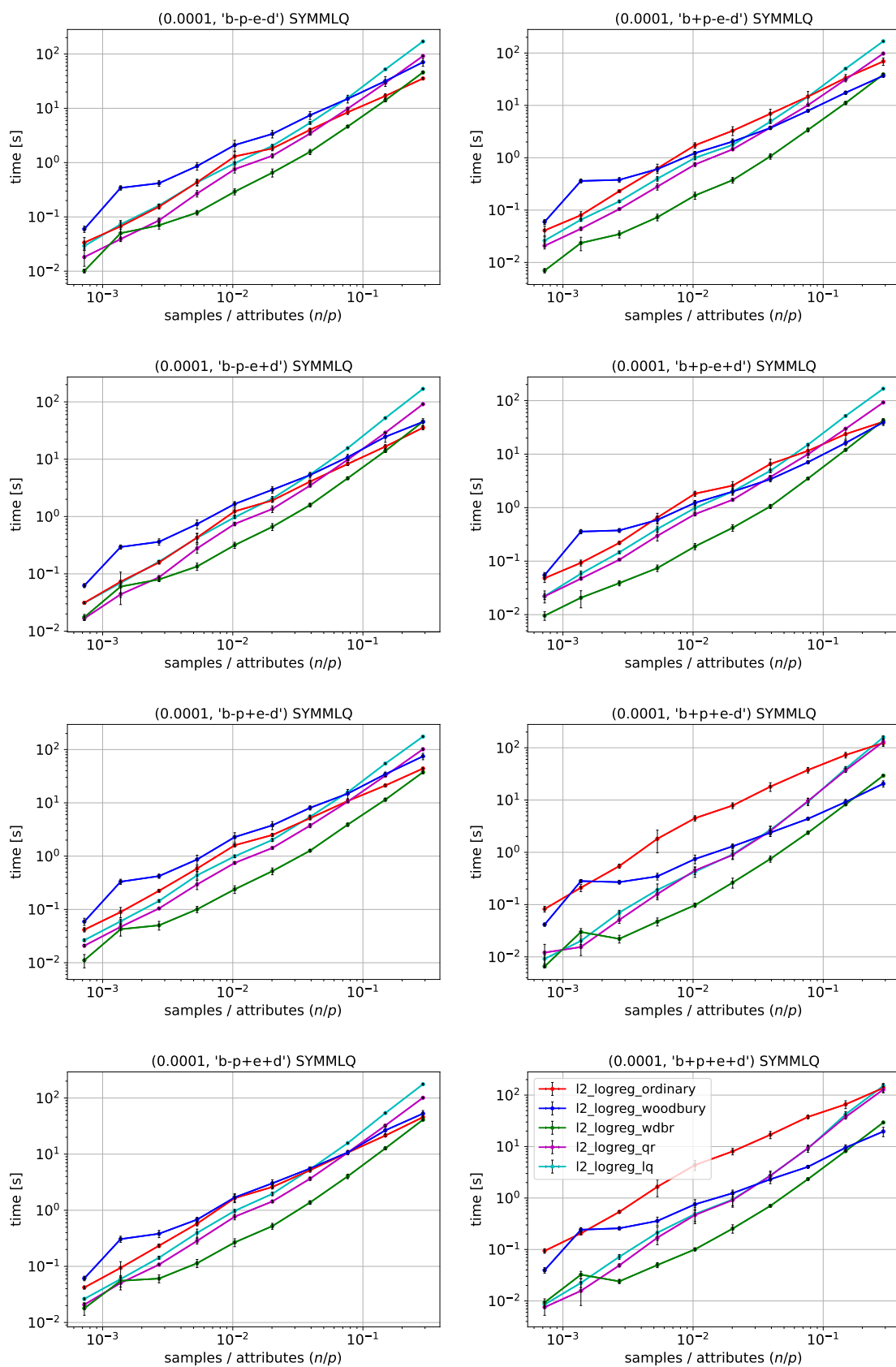
na kartę dane o mniejszym rozmiarze.



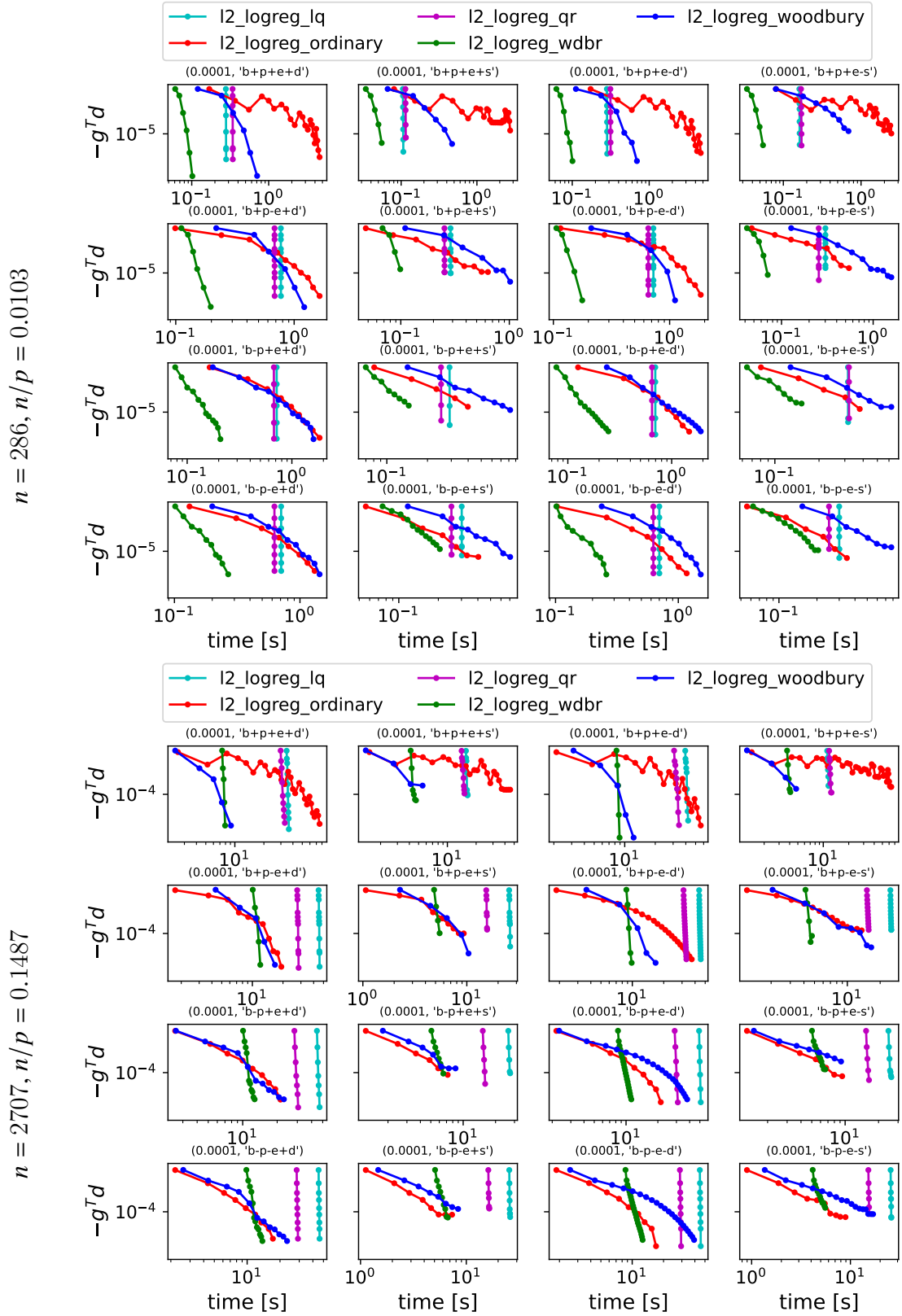
Rysunek 4.1: Porównanie czasów strojenia modeli w funkcji n/p . Dane gęste: extreme.



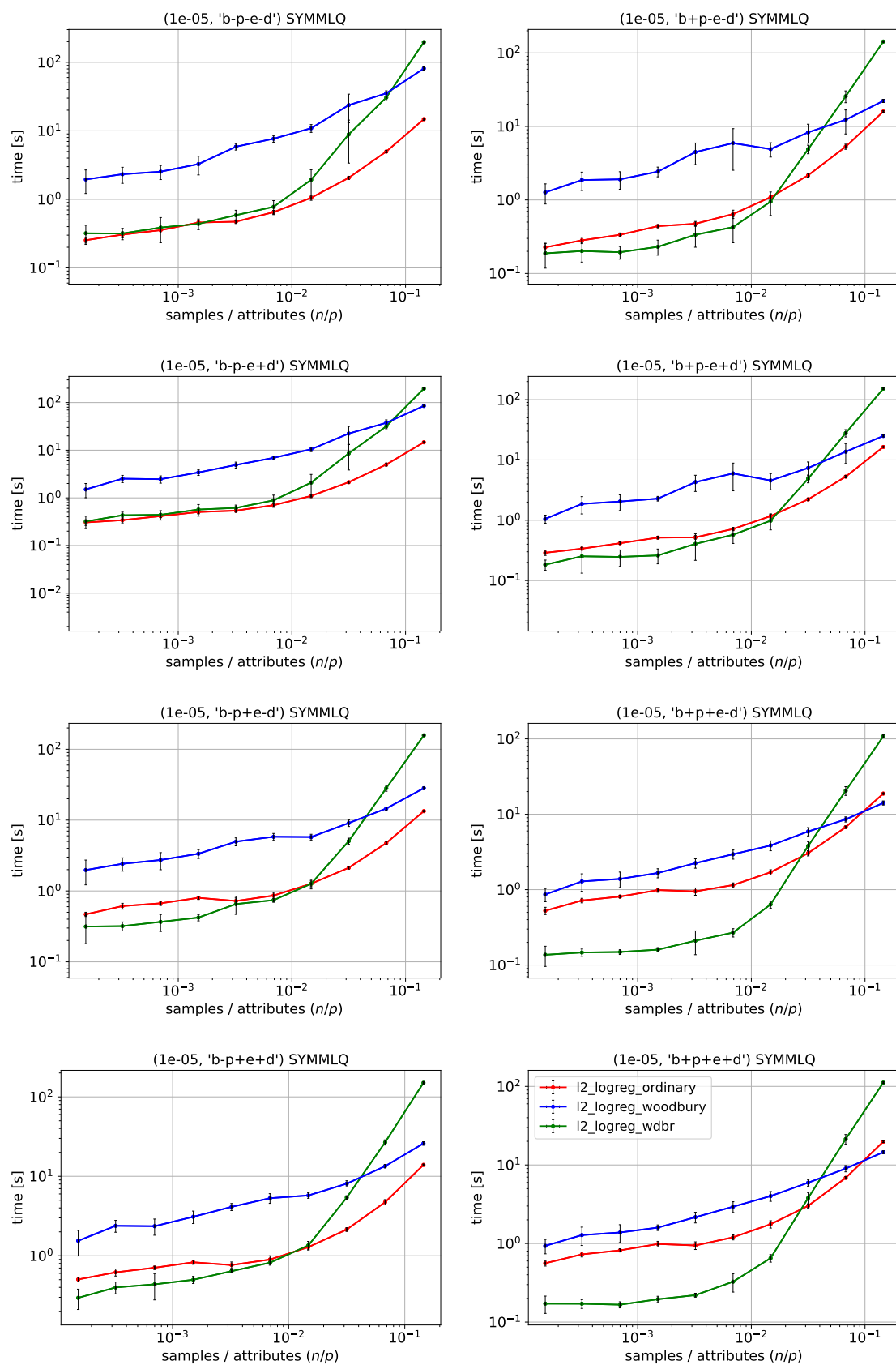
Rysunek 4.2: Porównanie zbieżności procedur dla różnej liczby próbek. Dane: extreme.



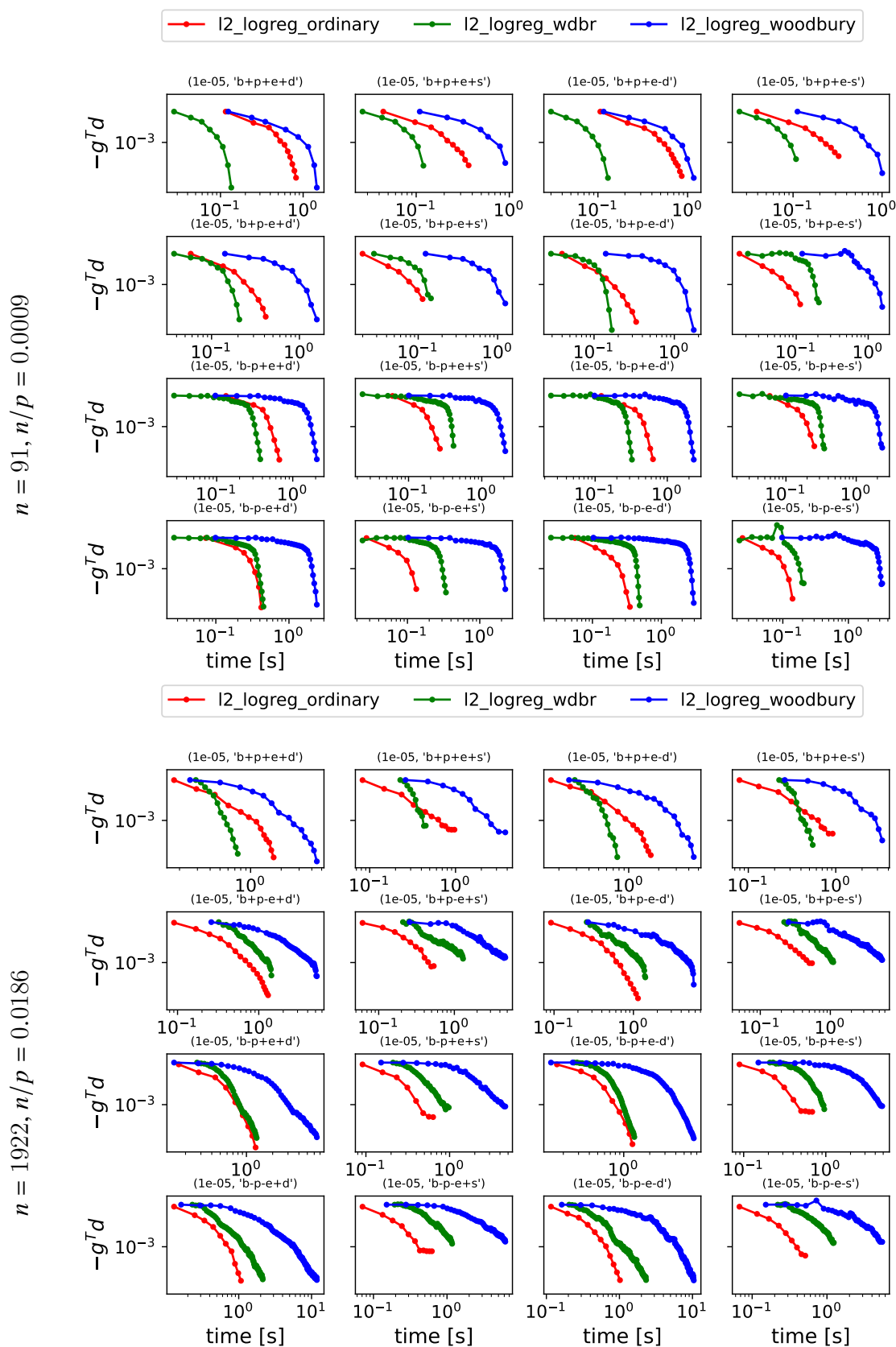
Rysunek 4.3: Porównanie czasów strojenia modeli w funkcji n/p . Dane gęste: PCam16k.



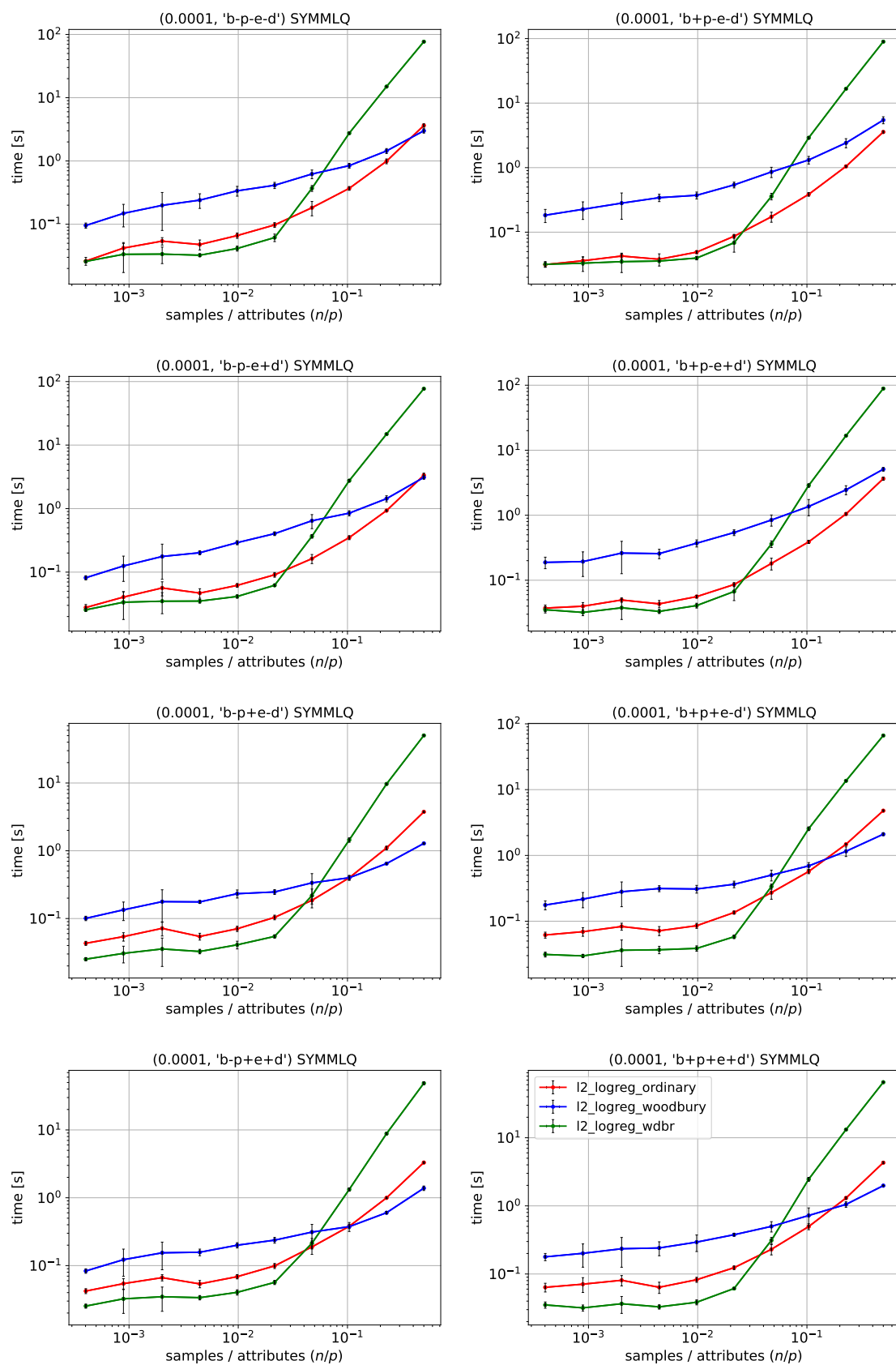
Rysunek 4.4: Porównanie zbieżności procedur dla różnej liczby próbek. Dane: PCam16k.



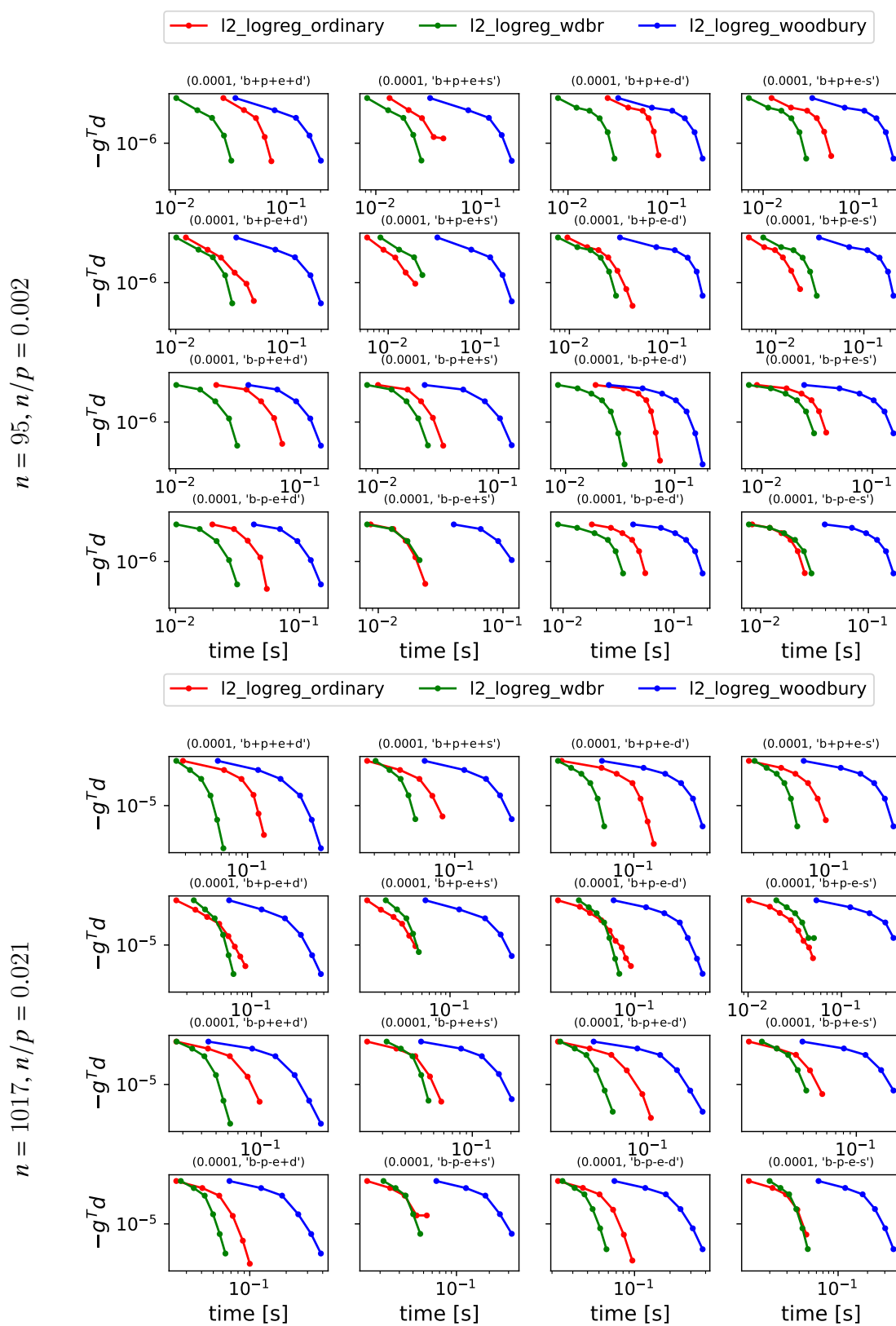
Rysunek 4.5: Porównanie czasów strojenia modeli w funkcji n/p . Dane rzadkie: news20.



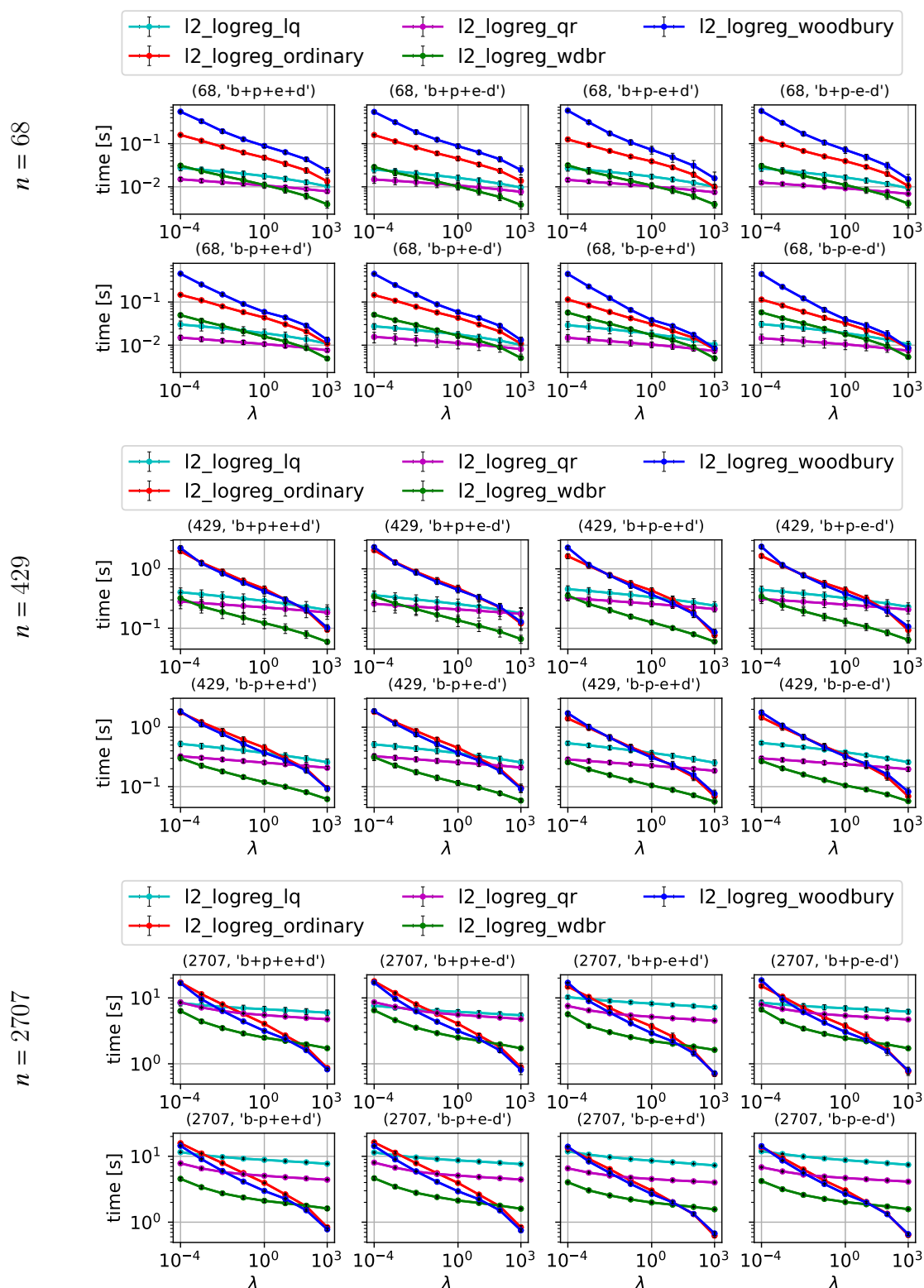
Rysunek 4.6: Porównanie zbieżności procedur dla różnej liczby próbek. Dane: news20.



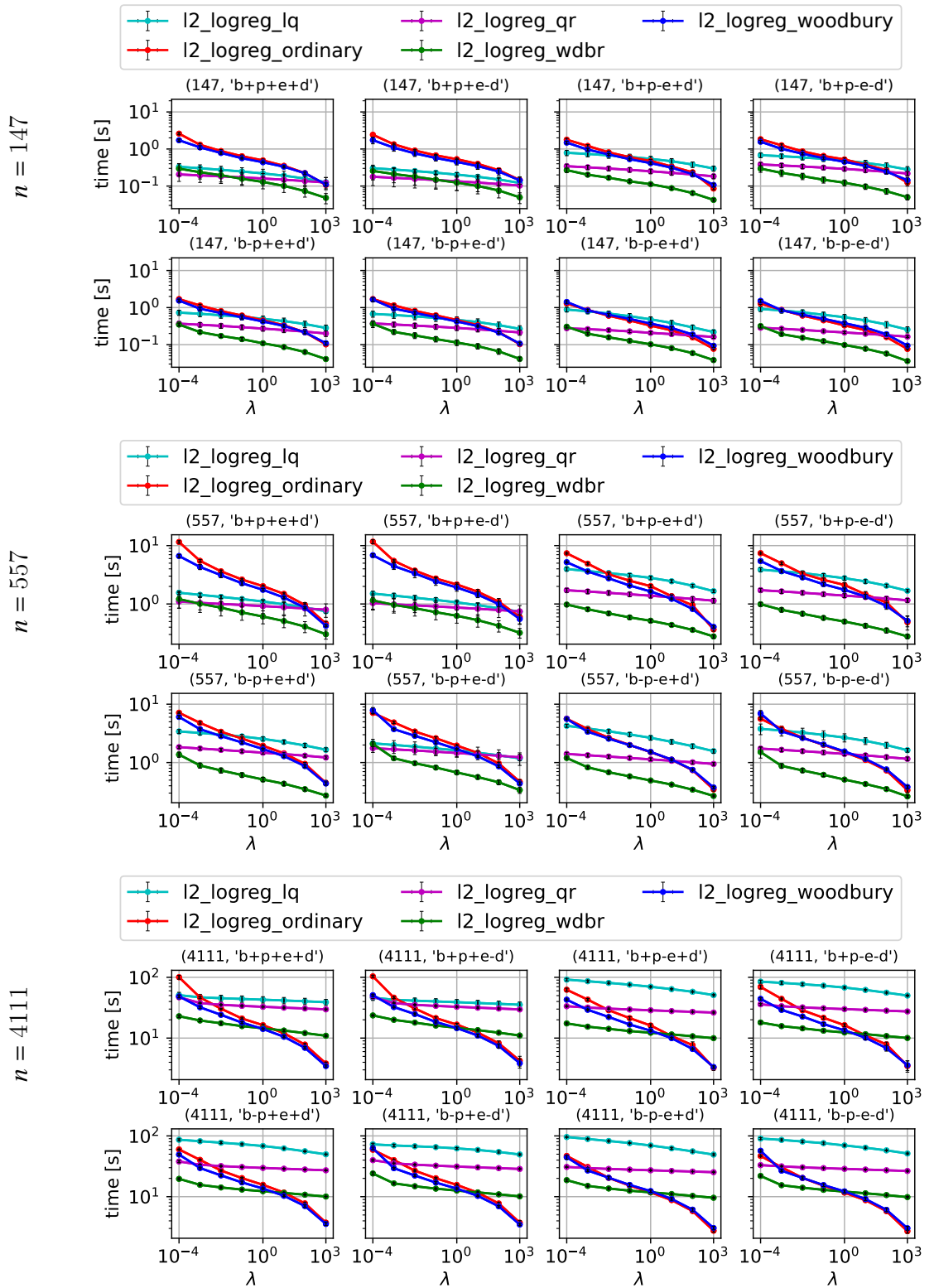
Rysunek 4.7: Porównanie czasów strojenia modeli w funkcji n/p . Dane rzadkie: RCV1.



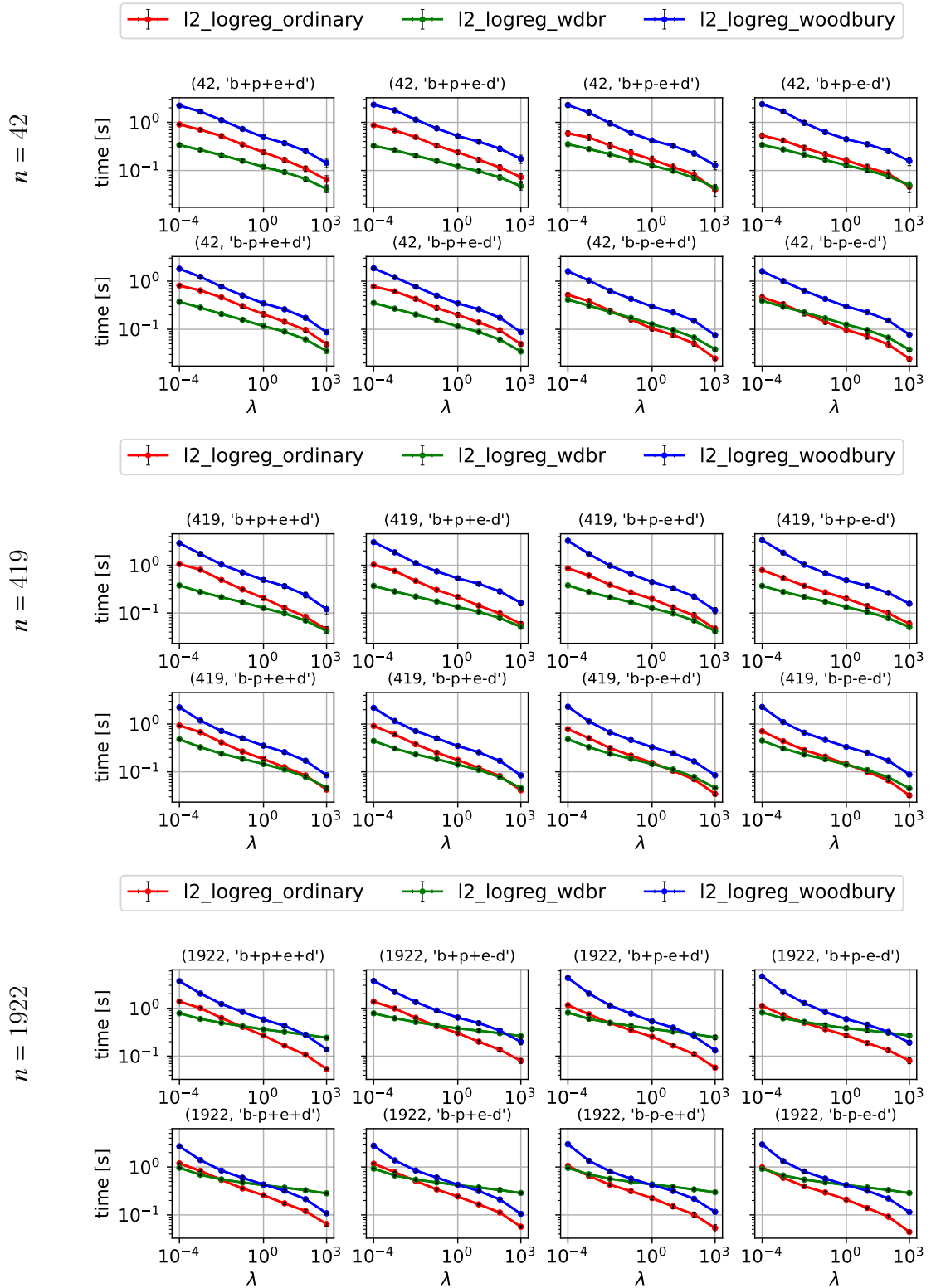
Rysunek 4.8: Porównanie zbieżności procedur dla różnej liczby próbek. Dane: RCV1.



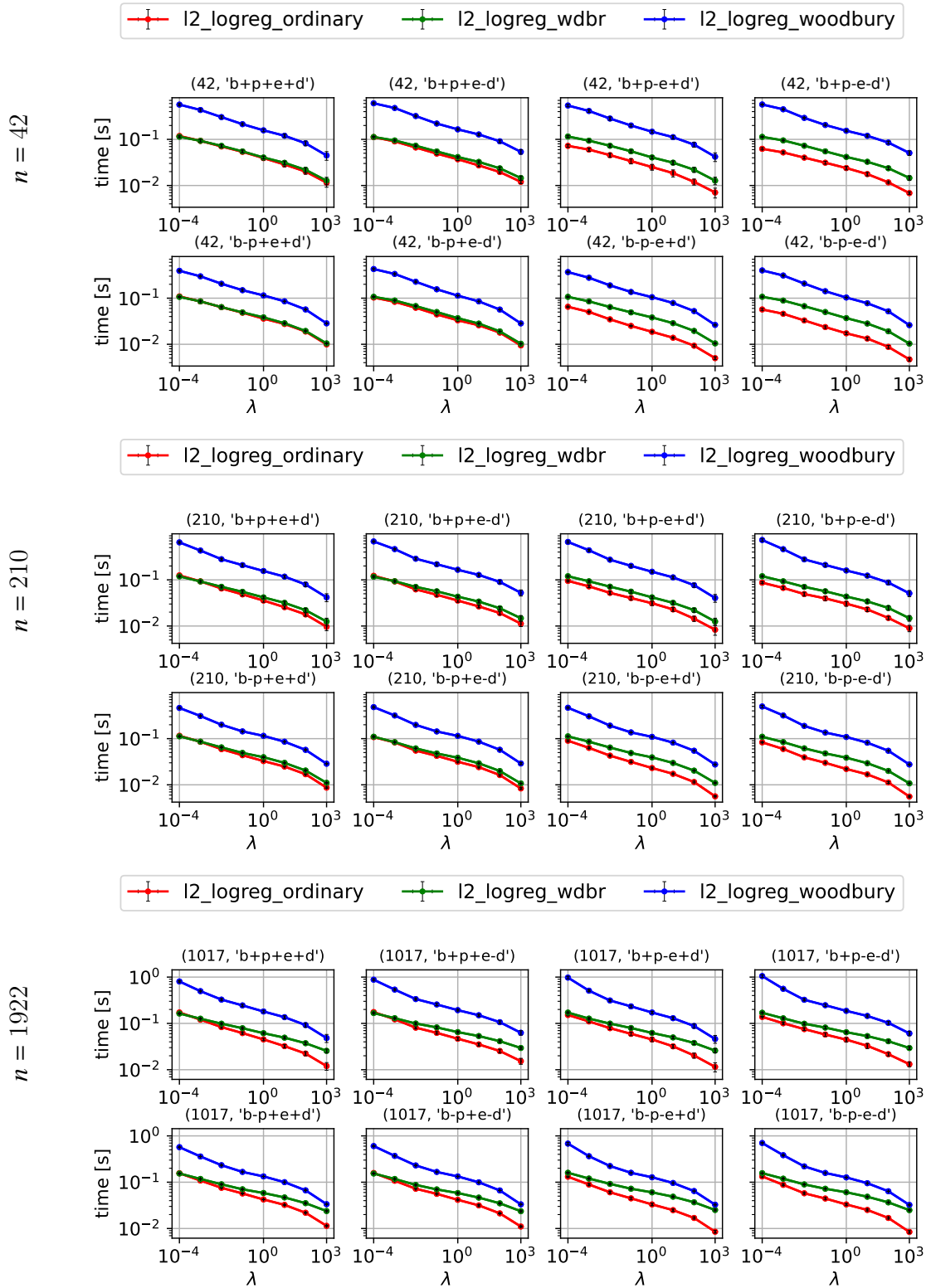
Rysunek 4.9: Porównanie czasów wyznaczania ścieżki dla $\lambda = 10^3, 10^2, \dots, 10^{-4}$; kolejny model startuje od wag modelu poprzedniego. Liczba próbek kolejno od góry: $n = 68, 429, 2707$. Wykres przedstawia skumulowany czas obliczeń. Dane: extreme.



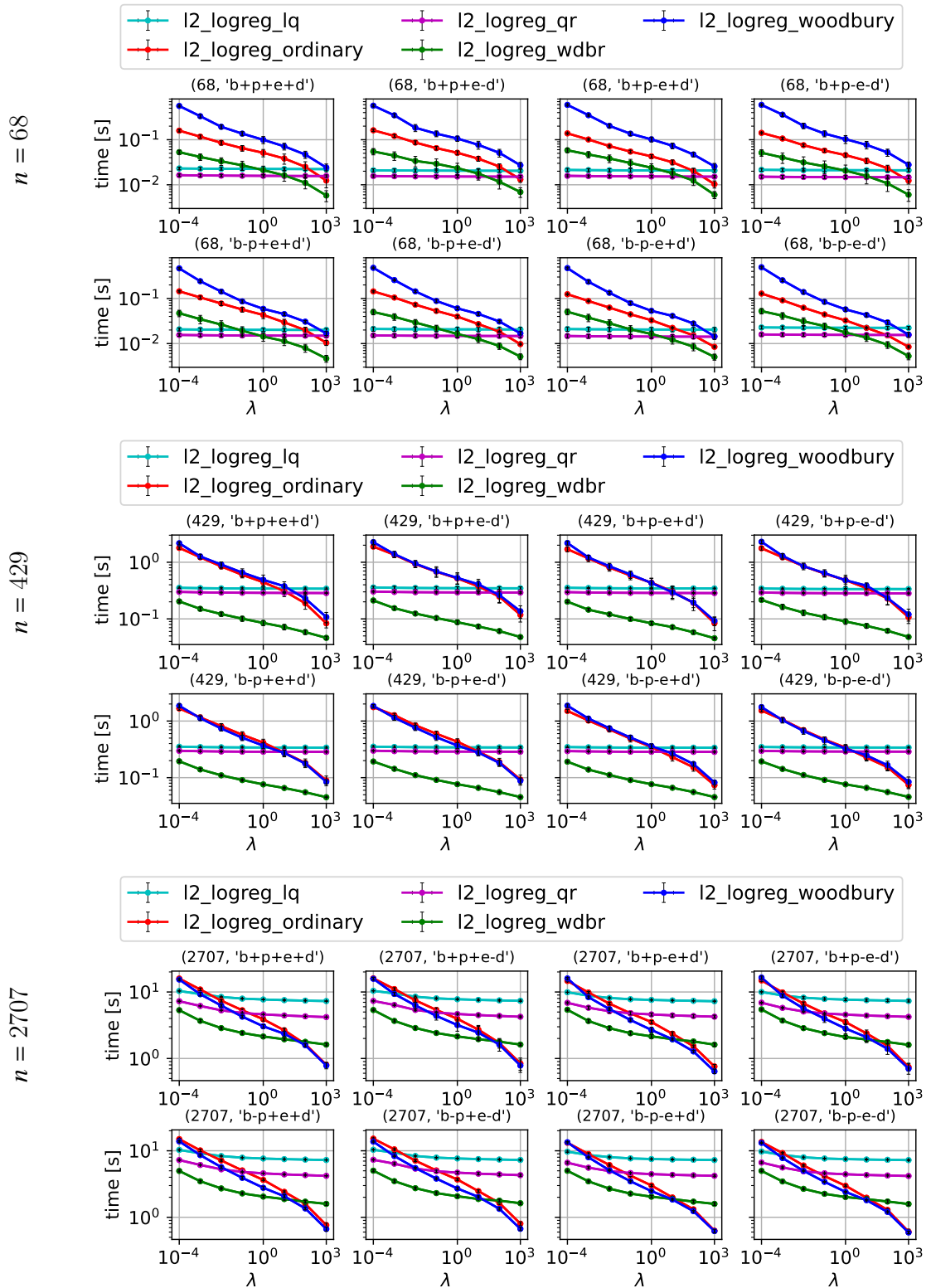
Rysunek 4.10: Porównanie czasów wyznaczania ścieżki dla $\lambda = 10^3, 10^2, \dots, 10^{-4}$; kolejny model startuje od wag modelu poprzedniego. Liczba próbek kolejno od góry: $n = 147, 557, 4111$. Wykres przedstawia skumulowany czas obliczeń. Dane: PCam16k.



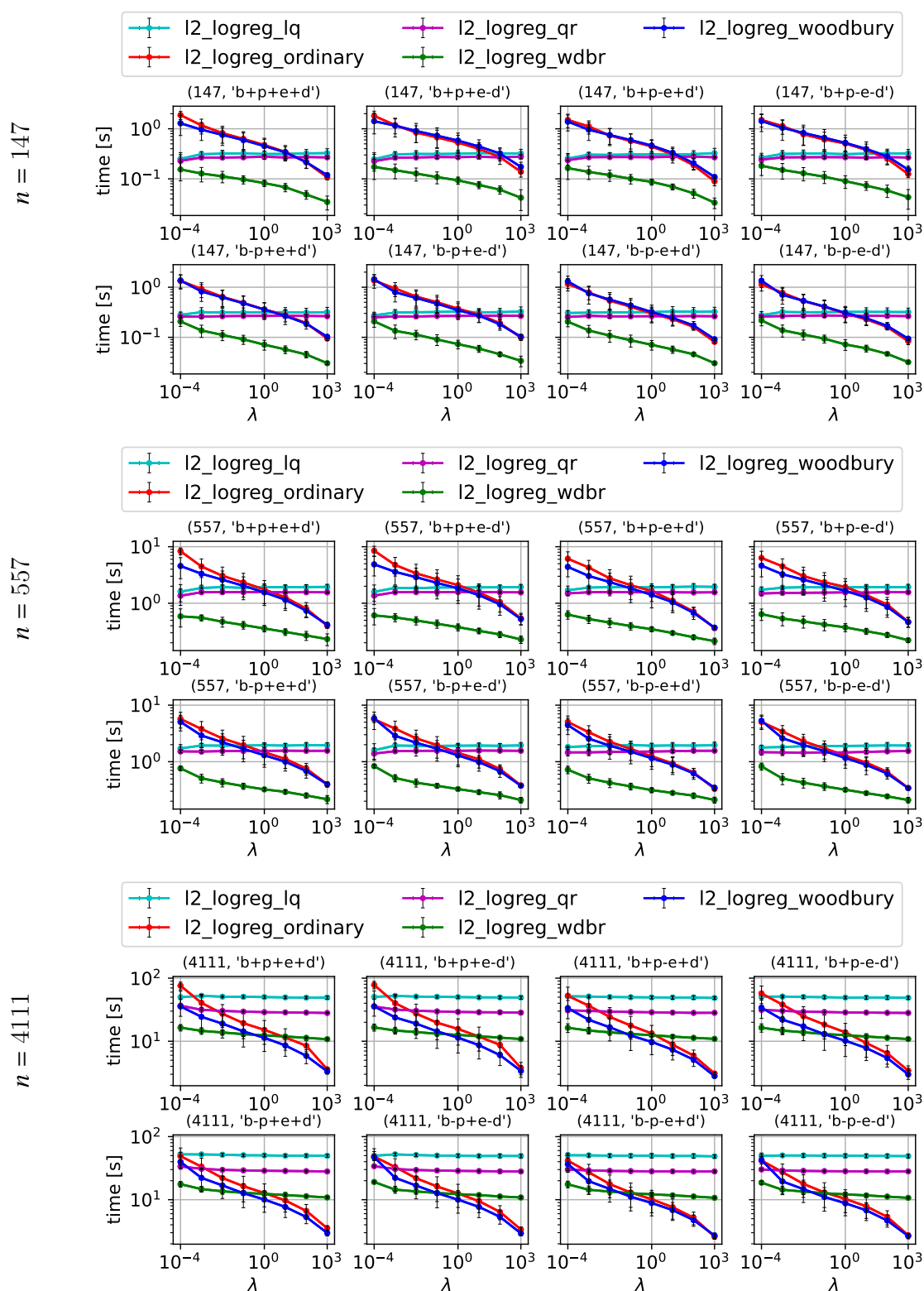
Rysunek 4.11: Porównanie czasów wyznaczania ścieżki dla $\lambda = 10^3, 10^2, \dots, 10^{-4}$; kolejny model startuje od wag modelu poprzedniego. Liczba próbek kolejno od góry: $n = 42, 419, 1922$. Wykres przedstawia skumulowany czas obliczeń. Dane: news20.



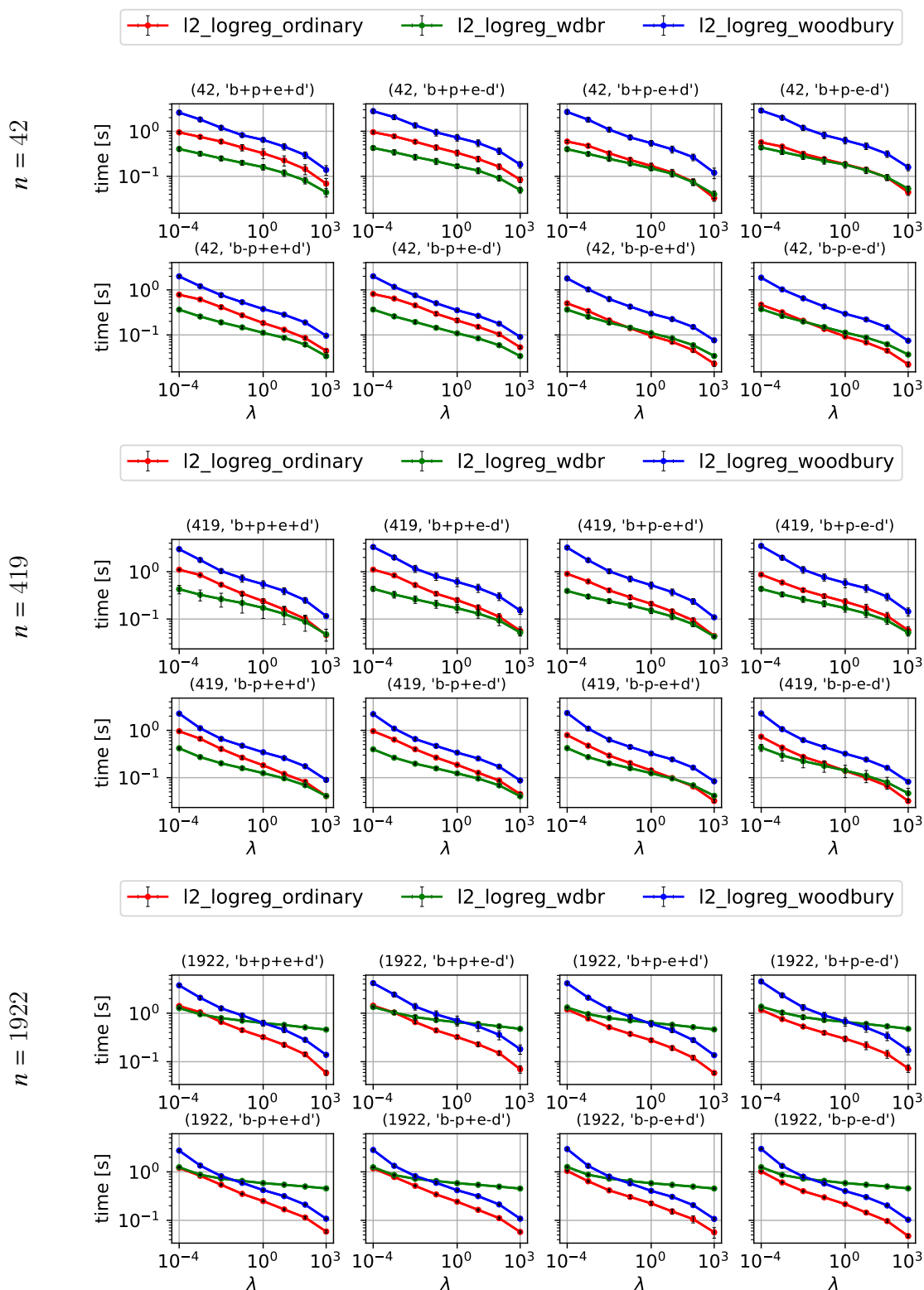
Rysunek 4.12: Porównanie czasów wyznaczania ścieżki dla $\lambda = 10^3, 10^2, \dots, 10^{-4}$; kolejny model startuje od wag modelu poprzedniego. Liczba próbek kolejno od góry: $n = 42, 210, 1922$. Wykres przedstawia skumulowany czas obliczeń. Dane: RCV1.



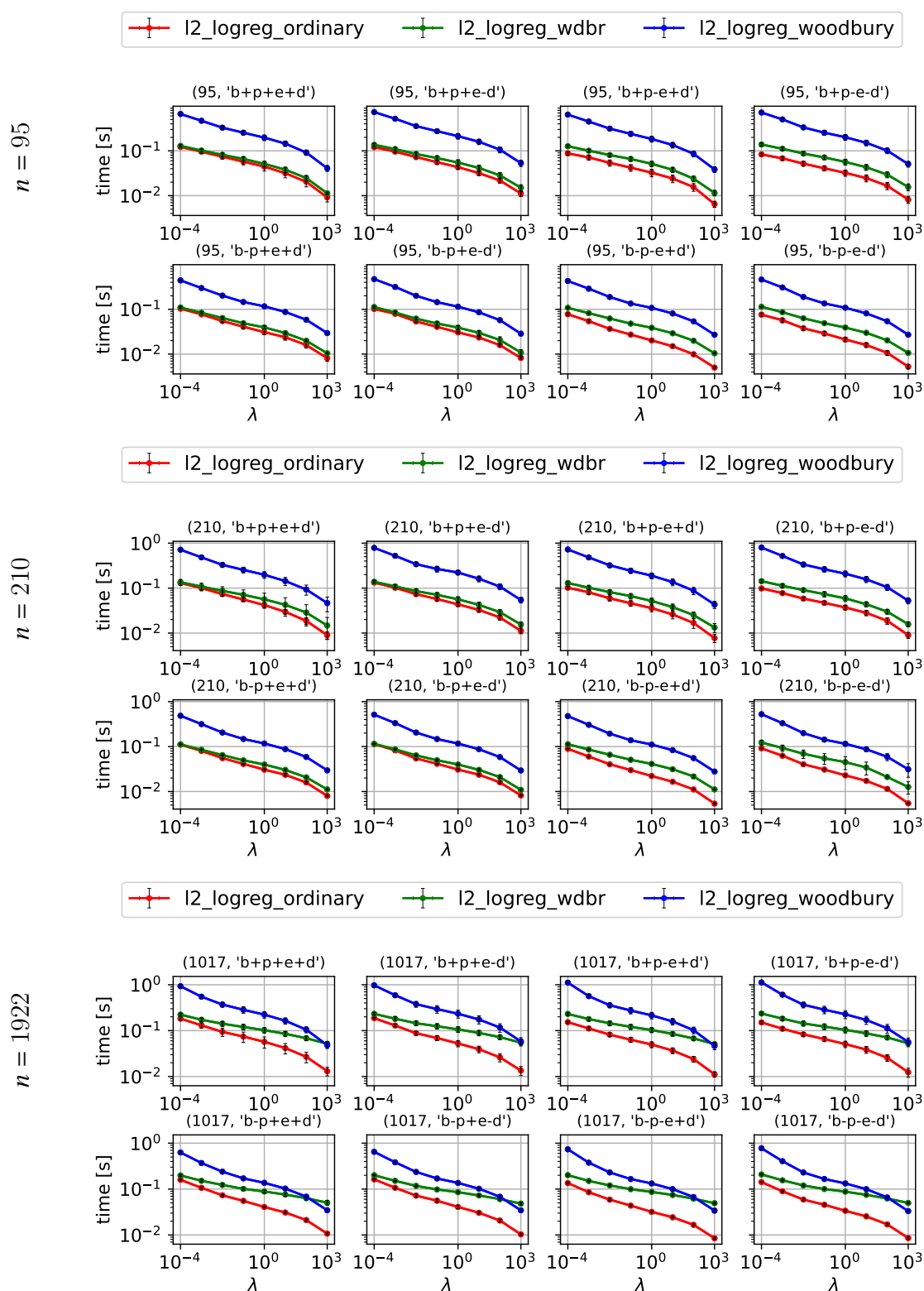
Rysunek 4.13: Porównanie czasów wyznaczania ścieżki dla $\lambda = 10^3, 10^2, \dots, 10^{-4}$; kolejny model startuje od $w = 0$. Liczba próbek kolejno od góry: $n = 68, 429, 2707$. Dane: extreme.



Rysunek 4.14: Porównanie czasów wyznaczania ścieżki dla $\lambda = 10^3, 10^2, \dots, 10^{-4}$; kolejny model startuje od $w = 0$. Liczba próbek kolejno od góry: $n = 147, 557, 4111$. Dane: PCam16k.



Rysunek 4.15: Porównanie czasów wyznaczania ścieżki dla $\lambda = 10^3, 10^2, \dots, 10^{-4}$; kolejny model startuje od $w = 0$. Liczba próbek kolejno od góry: $n = 42, 419, 1922$. Dane: news20.



Rysunek 4.16: Porównanie czasów wyznaczania ścieżki dla $\lambda = 10^3, 10^2, \dots, 10^{-4}$; kolejny model startuje od $w = 0$. Liczba próbek kolejno od góry: $n = 95, 210, 1922$. Dane: RCV1.

4.6 Wyniki eksperymentów dla regularyzacji ℓ^1

Eksperymenty z modelami z regularyzacją ℓ^1 wykonano w identycznym układzie, jak w przypadku opisanym we wcześniejszej sekcji. W eksperymentach porównywano wybrane procedury opisane w rozdziale 2 w sekcji 2.5, skrótowo podsumowane w tabeli 4.13 oraz wykorzystano te same zbiory danych.

Tabela 4.13: Procedury wykorzystane w eksperymentach z uczeniem modeli regresji logistycznej z regularyzacją ℓ^1 .

Procedura	Opis
<code>l1_logreg_ipm</code>	metoda Newtona z ograniczeniami, rozwiązywany za pomocą metody wewnętrznego punktu (por. sekcja 2.5.3)
<code>l1_logreg_liblinear</code>	metoda spadku względem współrzędnych, procedura zaczerpnięta z pakietu <code>liblinear</code> (por. sekcja 2.5.2)
<code>l1_logreg_cd</code>	metoda spadku względem współrzędnych, własna implementacja w C++, nieznaczne różnice względem <code>liblinear</code> (por. sekcja 2.5.2)
<code>l1_logreg_lbfgsb</code>	metoda z wykorzystaniem gładkiej optymalizacji z ograniczeniami za pomocą algorytmu L-BFGS-B (por. sekcja 2.5.3)
<code>l1_logreg_ordinary</code>	metoda wykorzystująca górne ograniczenie normy ℓ^1 przez normę ℓ^2 , kierunek znajdowany przez rozwiązywanie układu (2.57)
<code>l1_logreg_woodbury</code>	podobnie jak wyżej, kierunek znajdowany przez rozwiązywanie układu (2.58)
<code>elnet_lr_lbfgsb</code>	zaimplementowana metoda elastic net z wykorzystaniem gładkiej optymalizacji z ograniczeniami za pomocą algorytmu L-BFGS-B (por. sekcja 2.6), w eksperymentach z wyzerowaną karą typu ℓ^2 — przy tych ustawieniach model równoważny modelowi <code>l1_logreg_lbfgsb</code>

Tabela 4.14: Porównanie zbieżności testowanych algorytmów dla regularyzacji ℓ^1 na zbiorze lsvt (por. repozytorium OpenML zbiór id=1484) o rozmiarze $n \times p = 126 \times 309$ przy ustawieniach $\lambda = 10^{-4}$ oraz ustawieniach: b-p-e-d (bez wyrazu wolnego, bez preconditioningu, bez dokładnej minimalizacji kierunkowej w typie double); it oznacza numer iteracji, gTd rzut gradientu na nowy kierunek (pochodna kierunkowa), f wartość funkcji celu, cg liczba iteracji procedury rozwiązującej układ równań liniowych.

	l1_logreg_ordinary			l1_logreg_woodbury		
it	gTd	f	cg	gTd	f	cg
1	-0.170364	84.1824	1	-0.170364	84.1824	2
2	-0.0719646	84.1045	1	-0.0719593	84.1045	1
3	-0.0189374	84.0836	1	-0.0189363	84.0836	2
4	-0.0167302	84.0731	1	-0.0167284	84.0731	2
5	-0.00174617	84.0713	1	-0.00174603	84.0713	2
6	-0.00162302	84.0703	1	-0.00162284	84.0703	2
7	-0.000161771	84.0701	1	-0.000161757	84.0701	2
8	-6.70121e-05	84.07	1	-6.70074e-05	84.07	2
9	-5.8941e-05	84.07	1	-5.89344e-05	84.07	2
10	-6.1486e-06	84.07	1	-6.14812e-06	84.07	2
11	-5.74207e-06	84.07	1	-5.74141e-06	84.07	2
12	-5.70253e-07	84.07	1	-5.70205e-07	84.07	2
13	-2.35855e-07	84.07	1	-2.35838e-07	84.07	2
14	-2.08508e-07	84.07	1	-2.08484e-07	84.07	2
15	-2.16595e-08	84.07	1	-2.16578e-08	84.07	2
16	-2.0317e-08	84.07	1	-2.03147e-08	84.07	2
17	-2.01045e-09	84.07	1	-2.01028e-09	84.07	2
Wagi	$W_{186} = 2.73 \cdot 10^{-10}$, $W_{187} = 2.82 \cdot 10^{-10}$, $W_{188} = 0.569745$			$W_{186} = 2.73 \cdot 10^{-10}$, $W_{187} = 2.82 \cdot 10^{-10}$, $W_{188} = 0.569745$		

	l1_logreg_lbfsgb			l1_logreg_liblinear			l1_logreg_cd		
it	gTd	f	cg	gTd	f	cg	gTd	f	cg
1	-0.0932117	84.4821	3	-4.58855	-	1	-5.00323	85.5457	1
2	-0.114826	84.3672	1	-1.89621	-	3	-1.47852	84.07	3
3	-0.290918	84.0741	1	-1.88509e-08	-	1	-9.16914e-09	84.07	1
4	-0.00189289	84.0717	1						
5	-6.23783e-06	84.07	1						
Wagi	$W_{188} = 0.569721$			$W_{188} = 0.569745$			$W_{188} = 0.569745$		

Ponownie przed uruchomieniem eksperymentów zweryfikowano empirycznie poprawność algorytmów. Wyniki przykładowego eksperymentu porównawczego przedstawiono w tabeli 4.14. Tabela zawiera wyniki uzyskane na zbiorze `lsvt` (por. repozytorium OpenML zbiór `id=1484`) o rozmiarze $n \times p = 126 \times 309$ przy ustawieniach $\lambda = 10^{-4}$ oraz identycznych pozostałych ustawieniach. Z prezentowanych algorytmów porównywalne są dwa algorytmy bazujące na spadku względem współrzędnych, tj. `l1_logreg_liblinear` i `l1_logreg_cd`. W tym przypadku rozwiązanie posiada tylko jeden niezerowy współczynnik W_{188} , pozostałe zmienne są wyzerowane. Oba algorytmy dały dokładnie taki sam wynik z dokładnością do wyświetlanej precyzji. Bardzo podobny wynik (z dokładnością do 4 cyfr znaczących) oraz podobną liczbę iteracji wykonał algorytm `l1_logreg_lbfgsb`. Wszystkie z nich startowały z punktu początkowego $w = 0$.

Dwa pozostałe algorytmy (`l1_logreg_woodbury` i `l1_logreg_ordinary`) zatrzymały się na takiej samej wartości funkcji celu, natomiast dodatkowe dwa współczynniki W_{186} i W_{187} okazały się niezerowe, chociaż ich wartość jest właściwie na poziomie zera numerycznego rzędu 10^{-10} — gdyby algorytmy działały dłużej, to te dwa współczynniki zostałyby wyzerowane. Oba algorytmy są równoważne (co można zauważyć śledząc iteracje w tabeli 4.14). Te algorytmy startowały z punktu początkowego $w = d \cdot \arg \min_t Q(td)$, gdzie d to antysubgradient policzony w zerze. Długa ścieżka uczenia już na wstępie dyskwalifikuje te dwa podejścia, gdyż dla dużo większych zbiorów danych czas znacząco się wydłuży.

Pozostałe algorytmy, pomimo zbieżności do tej samej wartości funkcji celu oraz dawania numerycznie identycznych rozwiązań, nie są równoważne (jeżeli chodzi sposób działania i iteracje).

Na rysunku 4.17 przedstawiono czasy wykonania poszczególnych algorytmów wraz ze wzrostem liczby próbek, uśrednione po dziesięciu powtórzeniach z losowym wyborem zestawu uczącego. Na wykresach wąsami zaznaczono również odchylenia standardowe wyników. Na wykresach na osi OX przedstawiono unormowaną liczbę próbek n/p . Wyniki przedstawiono tylko dla typu `double`. Metoda `l1_logreg_ipm` wspiera jedynie konfigurację `b+p+e-d`, natomiast dwie pozostałe metody mogą uczyć modele z wyrazem wolnym lub bez niego, korzystając z przybliżonej minimalizacji kierunkowej (`b-`) i bez możliwości wykorzystania preconditioningu (`p-`).

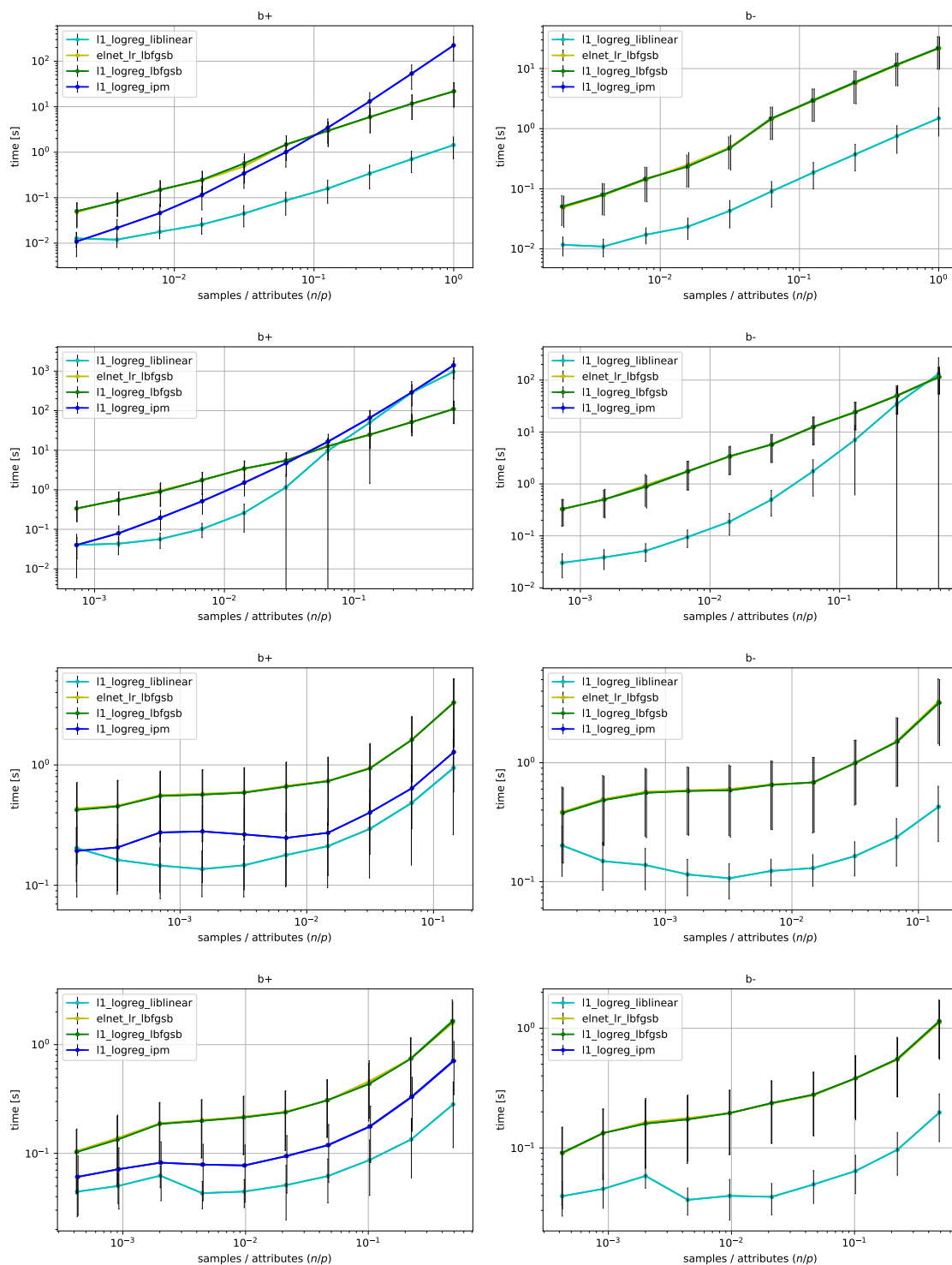
Na rysunkach 4.18, 4.19, 4.20, 4.21 zwizualizowano tempo zbieżności algorytmów dla wybranych rozmiarów danych uczących (n) i parametru $\lambda = 10^{-2} \cdot \lambda_{\max}$, gdzie $\lambda_{\max}$ to najmniejsza wartość, dla której wagi w modelu wynikowym wynoszą zero (z wyjątkiem wyrazu wolnego, który nie podlega regularyzacji). Wyniki przedstawiono dla wszyst-

kich dwóch wariantów algorytmów (z wyrazem wolnym i bez wyrazu wolnego). Każda z metod inaczej podchodzi do rozwiązywania problemu, zatem mają różny koszt związany z wyznaczaniem nowego punktu. Kroki w metodzie spadku względem współrzędnych są szybkie, jednak trzeba ich wykonać więcej, by uzyskać satysfakcjonujący wynik. Z kolei pojedyncza iteracja w metodzie wewnętrznego punktu wymaga większego nakładu pracy, jednakże prowadzi to do szybszej zbieżności. Metoda oparta na algorytmie L-BFGS-B balansuje między tymi dwiema metodami pod kątem kosztu pojedynczej iteracji.

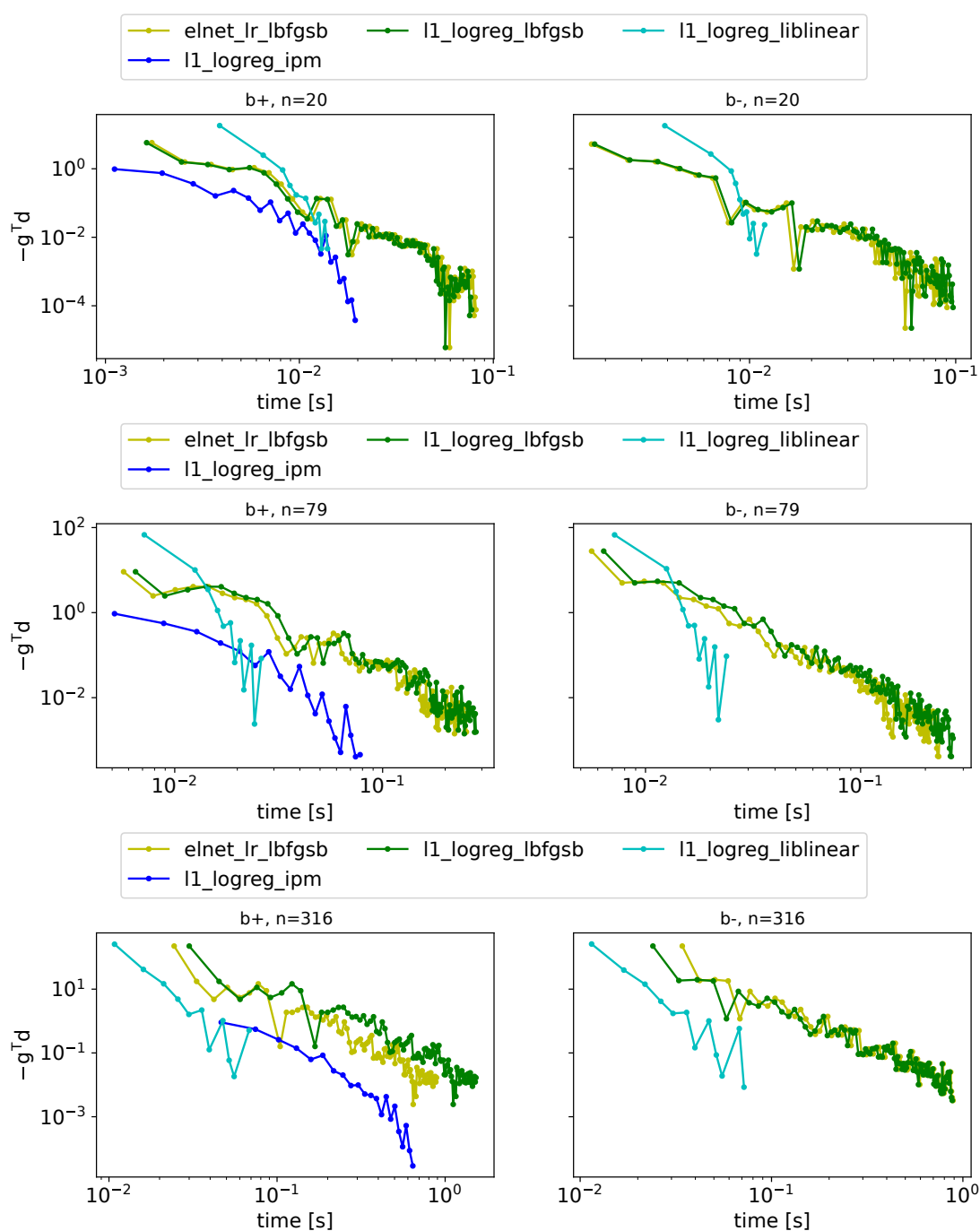
Patrząc na rysunki 4.18, 4.19, 4.20, 4.21 widać, że dla dostatecznie małych zbiorów `l1_logreg_ipm` może być konkurencyjny w stosunku do `l1_logreg_liblinear`. Ogólnie różnica między tymi algorytmami (zwłaszcza na zbiorach rzadkich) nie jest duża (nieznacznie na korzyść `l1_logreg_liblinear`), natomiast dla większych wartości n/p na zbiorach gęstych staje się ona wyraźniejsza. Rysunek 4.17 nie oddaje dokładnie różnic między algorytmami, ponieważ porównywane metody różnie definiują warunek stopu i zatrzymanie wspomnianych algorytmów następuje często później (bliżej optimum, przy niższej wartości funkcji celu), natomiast dla algorytmu `l1_logreg_liblinear` ustalony jest on domyślnie wcześniej. Skutkuje to również tym, że algorytmy te zatrzymują się czasami w różniących się nieznacznie punktach końcowych (jeżeli chodzi o rzadkość modeli).

Patrząc na rysunek 4.17 można zauważyć, że dla zbioru gęstego PCam16k uczenie modelu z wyrazem wolnym oparte na algorytmie L-BFGS-B okazało się efektywniejsze dla większej próby (przy stosunku liczby próbek do liczby atrybutów na poziomie 10%). Zysk wytłumaczyć można następująco — przy większej liczbie próbek zależności między atrybutami (korelacje) w tym zbiorze zaczynają się ujawniać, przez co optymalizacja parametrów pojedynczo nie prowadzi do szybkiej zbieżności. Obniżając mocniej współczynnik λ spodziewać należy się wydłużenia czasu uczenia w metodzie opartej na spadku względem współrzędnych

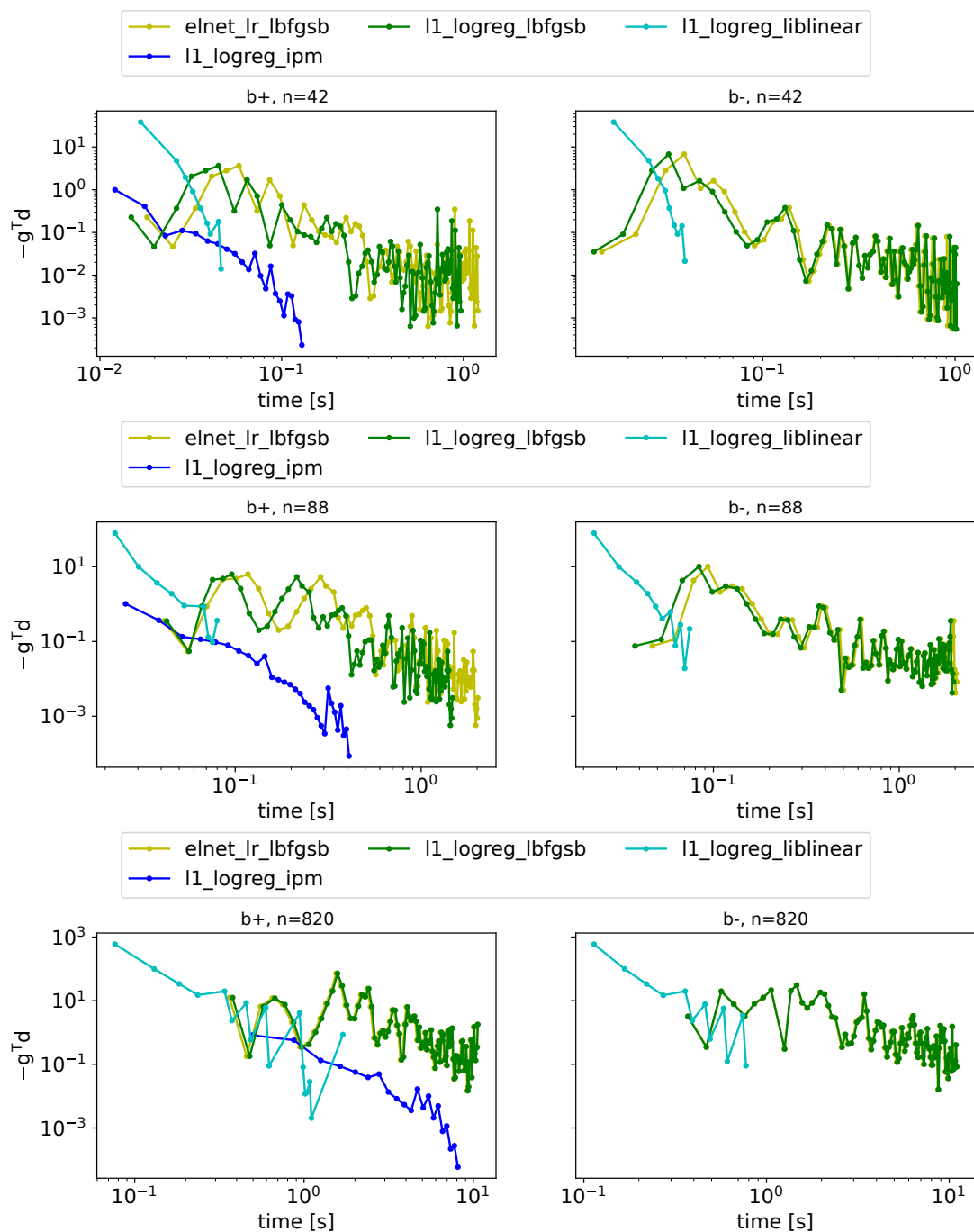
Z rysunku można odczytać także, że dla rzadkich zbiorów danych metoda IPM radzi sobie nieznacznie gorzej względem metody spadku względem współrzędnych — przy mniejszej wartości λ powinien być obserwowany zysk w stosunku do metody spadku względem współrzędnych.



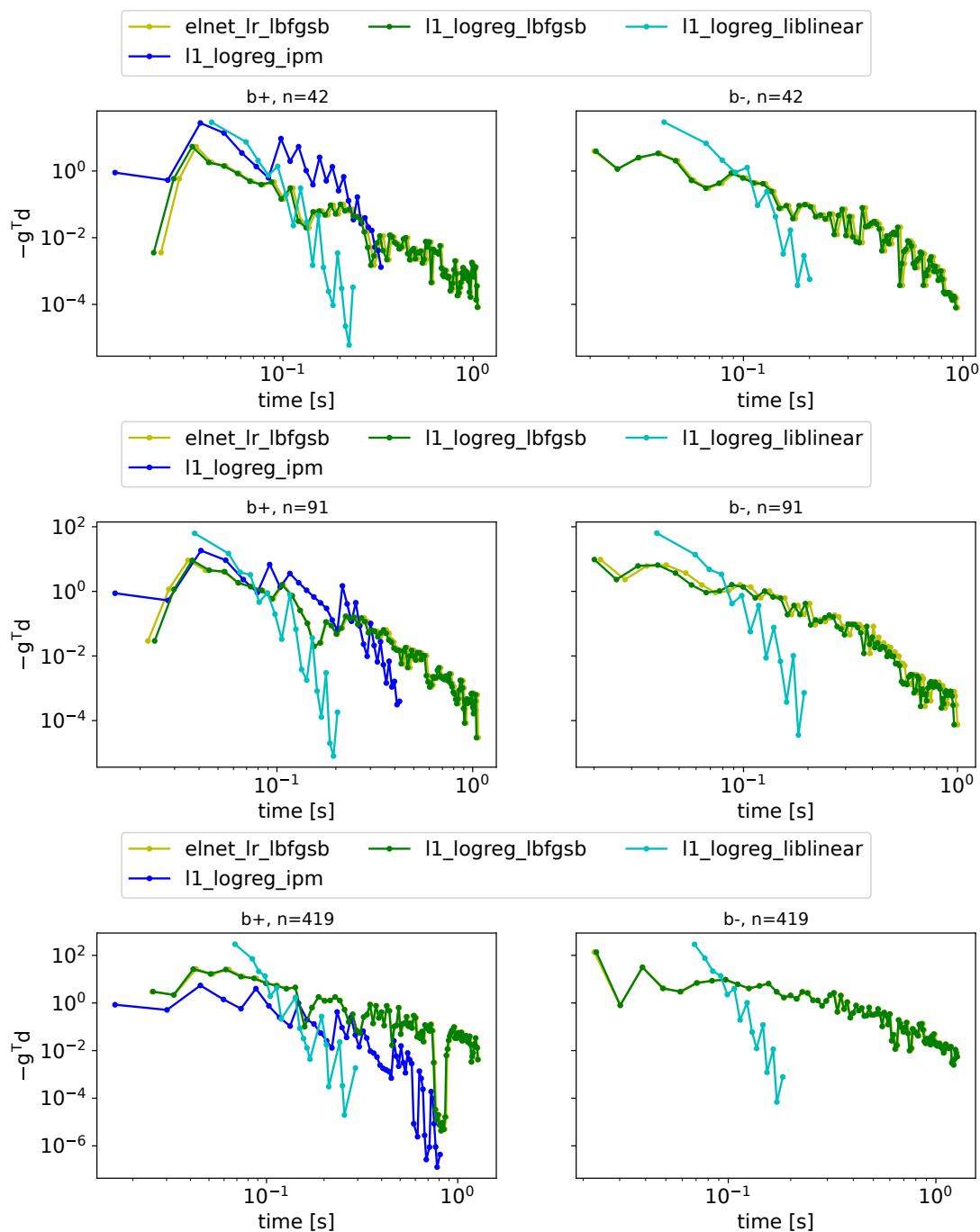
Rysunek 4.17: Porównanie czasów strojenia modeli w funkcji n/p . Kolejność zbiorów danych od góry: extreme, PCam16k, news20, RCV1.



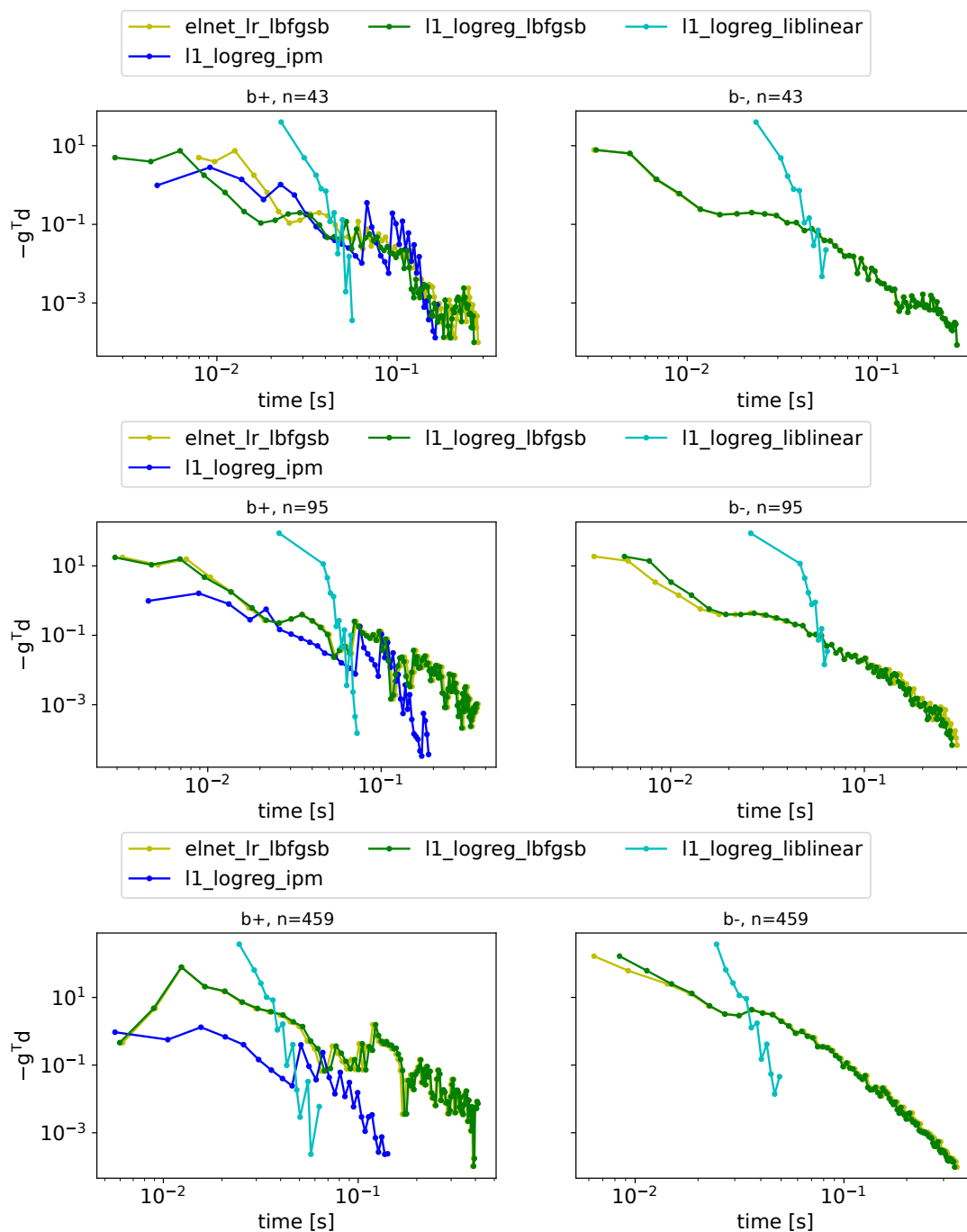
Rysunek 4.18: Porównanie zbieżności modeli na zbiorze *extreme* dla różnych rozmiarów danych. Lewa kolumna to model z wyrazem wolnym, a prawa kolumna to model bez wyrazu wolnego.



Rysunek 4.19: Porównanie zbieżności modeli na zbiorze PCam16k. Lewa kolumna to model z wyrazem wolnym, a prawa kolumna to model bez wyrazu wolnego.



Rysunek 4.20: Porównanie zbieżności modeli na zbiorze news20. Lewa kolumna to model z wyrazem wolnym, a prawa kolumna to model bez wyrazu wolnego.



Rysunek 4.21: Porównanie zbieżności modeli na zbiorze RCV1. Lewa kolumna to model z wyrazem wolnym, a prawa kolumna to model bez wyrazu wolnego.

5. Zastosowanie regularyzacji w modelach nieliniowych

5.1 Rzadki autoenkoder

5.1.1 Przygotowanie i opis eksperymentu

Model dwuwarstwowej sieci neuronowej uczoney z regularyzacją ℓ^1 o strukturze autoenkodera opisany w sekcji 3.2 przetestowano na zbiorach: Gasoline, MNIST oraz Olivetti faces. Eksperymenty przeprowadzono w środowisku Python/scikit-learn.

Głównym modelem odniesienia jest procedura SparsePCA z pakietu scikit-learn [61] w języku Python. Własna oryginalna struktura rzadkiego autoenkodera w eksperymentach oznaczona jako `neural_network` została zaimplementowana w języku C++ i podpięta do środowiska Python, w którym przeprowadzono całość eksperymentów. Jako drugą metodę referencyjną wykorzystano metodę PCA, która jest oczywiście najszybszą z rozważanych metod, jednak generuje rozwiązania jakościowo znacznie odbiegające od pozostałych.

W eksperymencie porównywano czas obliczeń i jakość uzyskiwanych wyników (głównie rzadkość rozwiązania oraz dokładność odtworzenia) dla różnych parametrów regularyzacji α . W tym eksperymencie liczba dopasowanych komponentów (w przypadku sieci — neuronów warstwy ukrytej) została ustawiona arbitralnie na 50 dla faces i MNIST oraz na 10 dla gasoline. Metoda SparsePCA wykorzystuje odmienną technikę uczenia

Tabela 5.1: Podsumowanie parametrów zbiorów danych wykorzystanych w eksperymentach.

Dane	l. próbek (n)	liczba atrybutów (p)	Rozstęp atrybutów
Gasoline	60	401	$[-0.084, 1.33]$
MNIST	60000	784	$[0, 255]$
Olivetti faces	400	4096	$[0, 1]$

(uczenie słownikowe), przez co znaczenie parametru regularyzacyjnego α jest inne niż w proponowanej procedurze, stąd zamiast wartości parametru regularyzacyjnego punktem odniesienia w porównaniu wyników będzie rzadkość rozwiązania.

W eksperymentach porównawczych wykorzystano następujące zestawy danych:

1. `gasoline` — ten zestaw danych pochodzi z pakietu `pls` [41] języka R i zawiera 60 sygnałów widmowych bliskiej podczerwieni dla próbek benzyny, opisanych przez 401 atrybutów, odpowiadających długościom fal od 900 nm do 1700 nm. W zbiorze tym celem jest predykcja liczby oktanowej paliwa.
2. `MNIST` — zestaw danych MNIST [48] zawierający odręcznie pisane cyfry 0–9, który zawiera 60 000 obrazów treningowych i 10 000 obrazów testowych zapisanych na macierzach 28×28 , co daje łącznie 784 atrybuty.
3. `Olivetti faces` — zestaw danych [12] pobrany z pakietu `scikit-learn` [61] zawiera 400 obrazów twarzy — 10 różnych obrazów dla każdej klasy. Obrazy znajdują się w macierzach o wymiarach 64×64 , co daje łącznie 4096 atrybutów. Obrazy zostały wykonane w różnych warunkach, przy różnym oświetleniu, mimice i szczegółach twarzy (np. obecność okularów).

Dla wszystkich modeli przy walidacji wyników użyto osobnych zestawów testowych. Zbiór MNIST zawiera oryginalny zestaw testowy, pozostałe zestawy danych podzielono na część uczącą i testową w sposób losowy za pomocą funkcji `train_test_split` (z pakietu `scikit-learn`) w proporcji 1:1. Liczba iteracji i tolerancja dla algorytmów uczących zostały ustawione arbitralnie odpowiednio jako 1000 oraz 10^{-3} .

Oba rozwiązania były testowane w tych samych warunkach w modelu jednowątkowym. Jedyną równoległą częścią była n -krotna walidacja z jednym wątkiem na iterację. Eksperymenty zostały przeprowadzone na maszynie z procesorem Xeon E5-2699 v4 2,20 GHz i 128 GB pamięci RAM. Wykorzystano następującą procedurę testowania dopasowania:

```
1 for trial_num in [1, ..., number_of_trials]:
2     x_train, x_test = train_test_split(X, test_size=0.5)
3     for model in [SparsePCA, neural_network]:
4         for alpha in list_of_alphas:
5             model.fit(x_train)
6             mse_train[trial_num] = score(model, x_train)
7             mse_test[trial_num] = score(model, x_test)
8 score_train = mean(mse_train)
9 score_test = mean(mse_test)
```

Tabela 5.2: Szczegółowe porównanie proponowanego autoenkodera (`neural_network`) z procedurą `SparsePCA/scikit-learn` oraz standardową procedurą PCA. Rzadkość oznacza średnią liczbę zerowych elementów na komponent.

Data	neural-network					SparsePCA/scikit-learn					PCA		
	α	MSE train	MSE test	time	spar-sity	α	MSE train	MSE test	time	spar-sity	MSE train	MSE test	time
Gasoline	$1.4e-5$	$4.4e-6$	$7.3e-6$	0.073	0.09	$2.3e-4$	$1.2e-6$	$1.1e-5$	10.57	65.65	0.016	0.016	0.064
	$2.7e-5$	$4.5e-6$	$7.3e-6$	0.071	0.62	$5.2e-4$	$1.3e-6$	$1.2e-5$	11.08	86.63			
	$5.2e-5$	$7.4e-6$	$1.2e-5$	0.469	115.1	0.001	$1.6e-6$	$1.3e-5$	7.915	121.3			
	$1.0e-4$	$1.7e-5$	$2.4e-5$	1.205	339.7	0.003	$2.4e-6$	$1.4e-5$	8.058	174.5			
	$1.9e-4$	$2.3e-5$	$3.0e-5$	1.268	374.8	0.006	$4.6e-6$	$2.1e-5$	7.459	225.9			
Olivetti faces	$1.0e-4$	$2.0e-3$	$3.8e-3$	392.6	564.8	0.014	$2.0e-3$	$3.9e-3$	163.5	344.0	$6.4e-3$	$8.2e-3$	2.118
	$3.7e-4$	$2.0e-3$	$3.9e-3$	281.5	2212	0.032	$2.3e-3$	$3.8e-3$	2699	1332			
	$7.2e-4$	$2.1e-3$	$4.0e-3$	268.3	2927	0.072	$2.8e-3$	$3.9e-3$	1542	2190			
	$1.4e-3$	$2.3e-3$	$4.1e-3$	248.4	3371	0.164	$3.3e-3$	$4.2e-3$	485.4	2817			
	$2.7e-3$	$2.9e-3$	$4.4e-3$	228.9	3375	0.373	$4.3e-3$	$5.2e-3$	180.5	3347			
MNIST	0.019	766.1	770.9	2969	204.4	5.179	766.1	771.0	46.83	191.2	778.9	783.7	0.864
	2.683	766.3	771.0	1271	362.5	13.89	766.3	771.1	46.08	230.7			
	19.31	769.2	773.7	1264	521.9	100.0	775.7	780.3	85.00	351.9			
	138.9	825.9	829.5	1268	746.1	268.3	831.7	835.2	755.6	616.8			
	372.8	901.9	904.4	1250	770.7	1931	1118	1119	127.7	762.6			

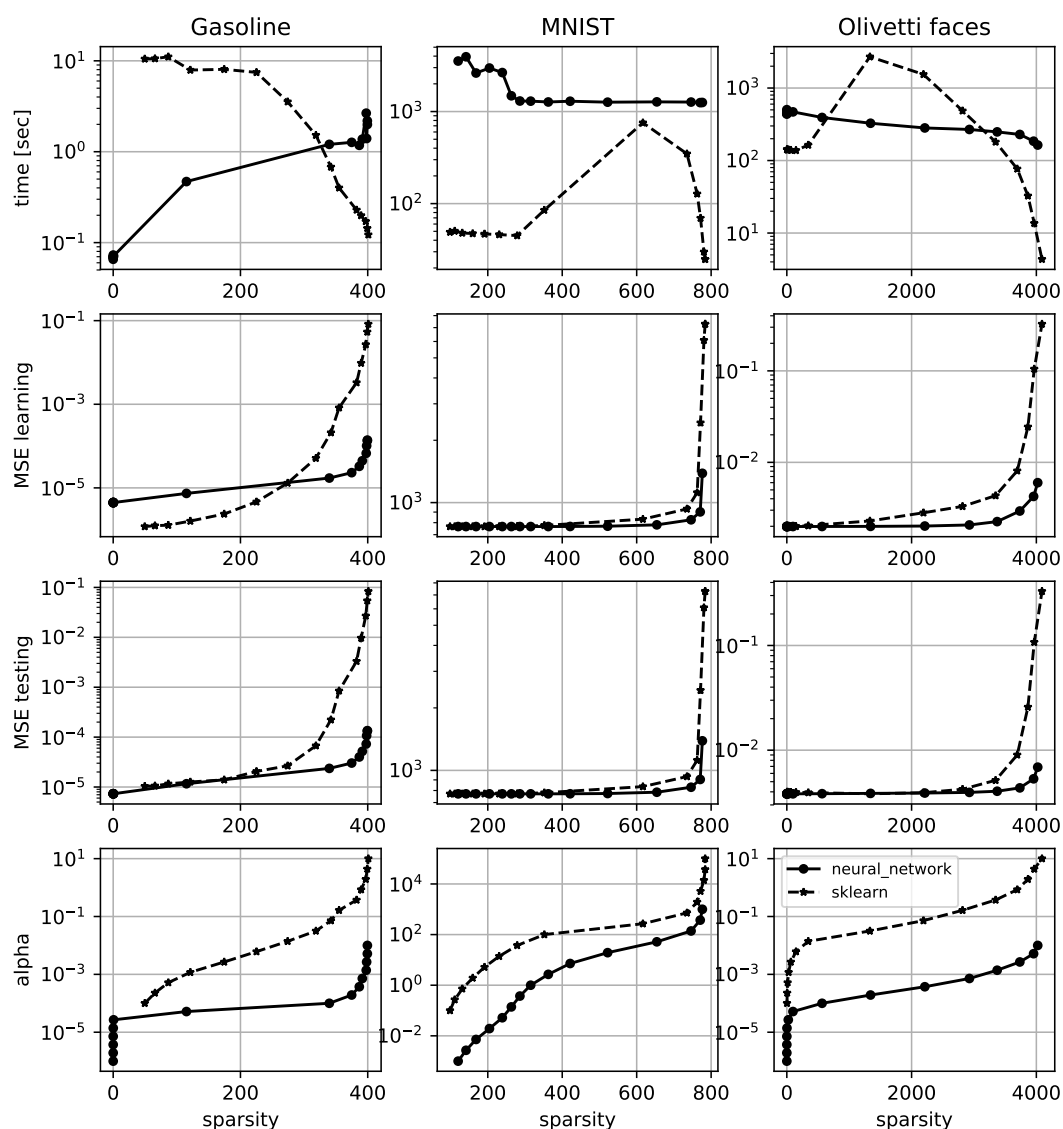
Wyniki prezentują odpowiednie miary jakości modeli (czas strojenia i jakość rekonstrukcji) w odniesieniu do rzadkości rozwiązań. Oznacza to, że dwa modele o tej samej rzadkości zostały porównane pod względem czasów obliczeń i jakości rekonstrukcji. Jako miarę jakości zastosowano standardowy średni błąd kwadratowy rekonstrukcji (`mse`). Rzadkość w części eksperymentalnej oznacza średnią liczbę zerowych współczynników na komponent.

5.1.2 Wyniki

Tabela 5.2 przedstawia porównanie proponowanego autoenkodera (`neural_network`) z procedurą `SparsePCA/scikit-learn` oraz standardową procedurą PCA (też z pakietu `scikit-learn`). Na rysunku 5.1 przedstawiono czasy obliczeń, jakość rekonstrukcji (tj. błąd MSE na danych uczących i testowych) oraz wartość parametru regularyzacyjnego α przy uzyskanej rzadkości.

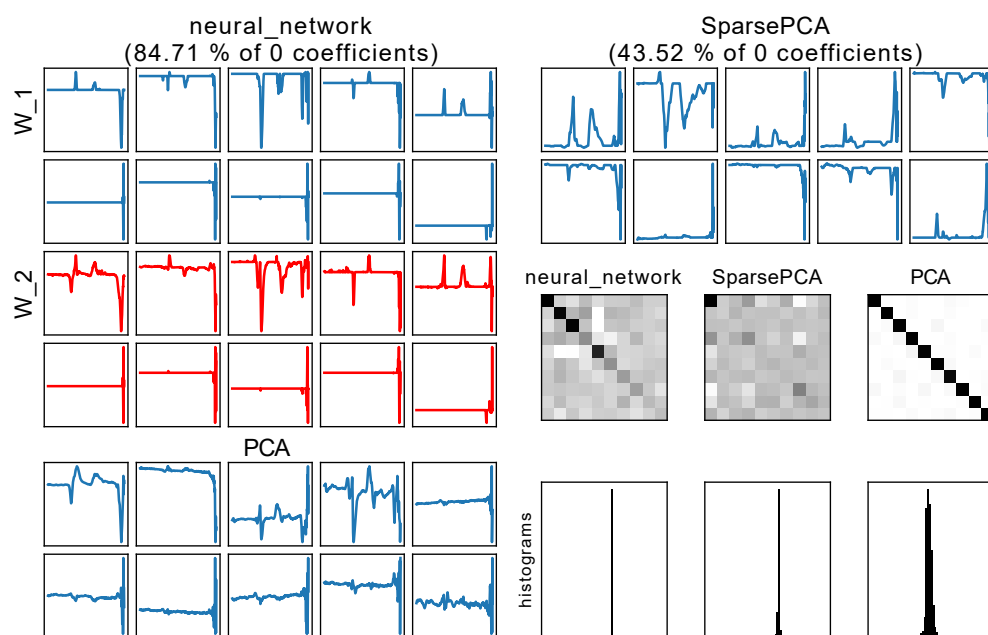
Jak można zauważyć na rysunku 5.1, pomimo że `SparsePCA` jest szybsze dla większej rzadkości, proponowane rozwiązanie (`neural_network`) zachowuje wyższą dokładność rekonstrukcji dla większej rzadkości, zarówno na zbiorze treningowym, jak i testowym. Ponadto istnieją obszary rzadkości, dla których proponowana procedura działa szybciej.

Porównanie komponentów (współczynników neuronów warstwy ukrytej) w zestawie danych `faces` jest przedstawione na rysunku: 5.6. Różne kolory czerwony/niebieski

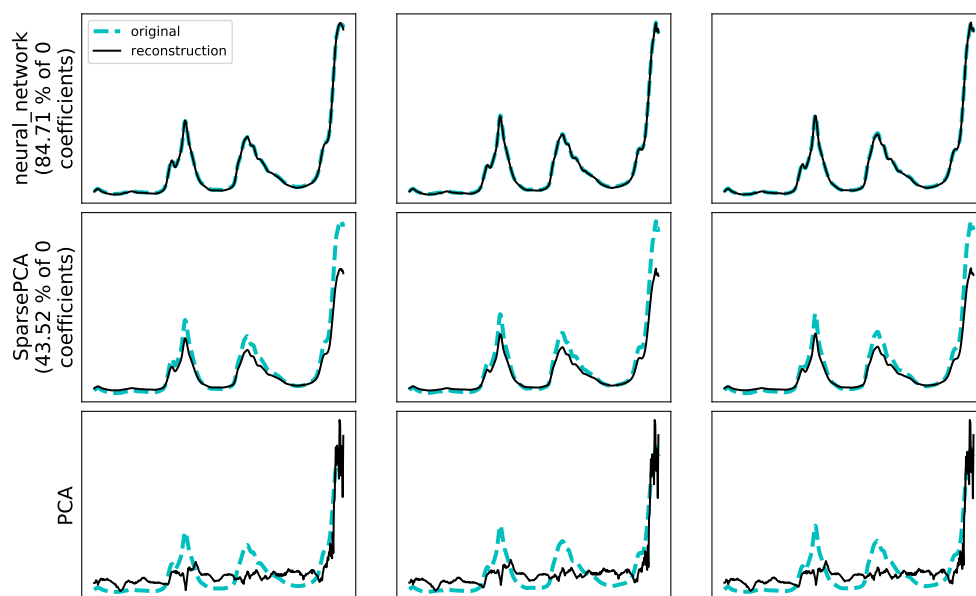


Rysunek 5.1: Porównanie proponowanego podejścia `neural_network` z procedurą SparsePCA/scikit-learn. Rzadkość oznacza średnią liczbę zerowych współczynników na składnik, α (alpha) jest parametrem regularyzacji. Od góry czas uczenia, błąd rekonstrukcji na danych uczących i testowych, parametr regularyzacyjny α .

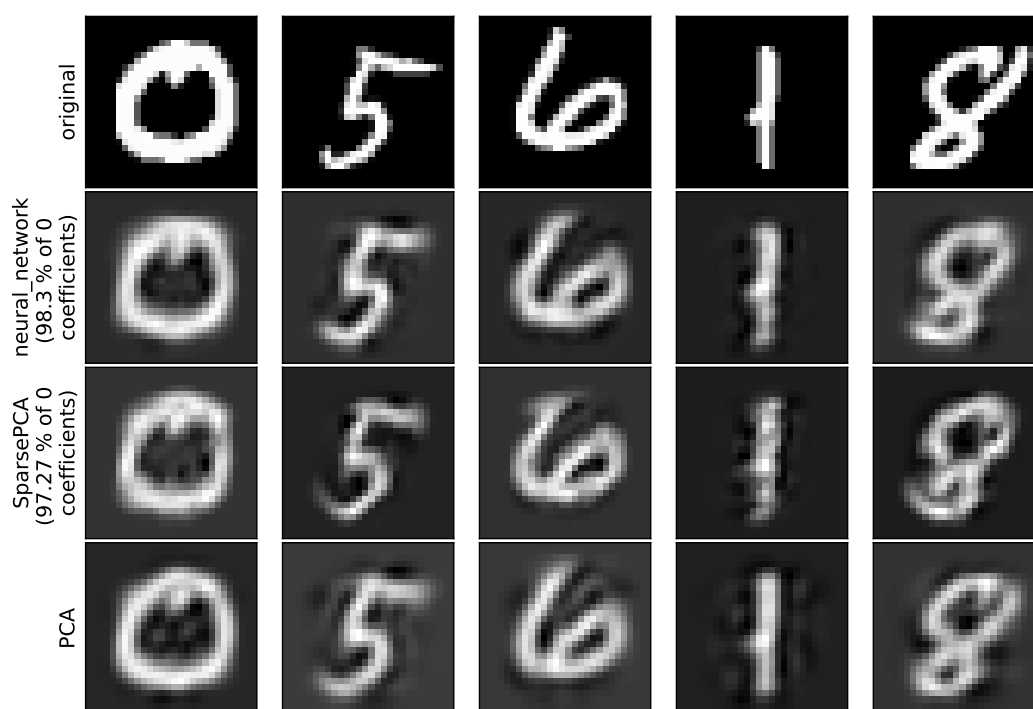
ski oznaczają współczynniki ujemne/dodatnie, a kolor biały oznacza dokładną wartość równą zero. Dodatkowo, dla porównania przedstawiono wyniki dla zwykłej procedury PCA oraz macierze kowariancji dla komponentów. Ortogonalność komponentów zwykłej procedury PCA jest gwarantowana teoretycznie, co potwierdzają również eksperymenty. Patrząc na wizualizacje macierzy kowariancji widzimy, że komponenty znalezione przez `neural_network` są bliskie ortogonalnym (wyraźniejsza przekątna), znacznie bardziej



Rysunek 5.2: Porównanie 10 komponentów wyznaczonych na zbiorze gasoline. Niebieski kolor oznacza komponenty (W^1), czerwony to wektory używane do rekonstrukcji (W^2). W dole po lewej składowe główne wyznaczone za pomocą procedury PCA. W dole po prawej wizualizacja macierzy kowariancji komponentów, a poniżej histogram wag dla poszczególnych modeli; uzyskana rzadkość: neural_network (93,47%), SparsePCA/scikit-learn (56,33%), PCA (0%).



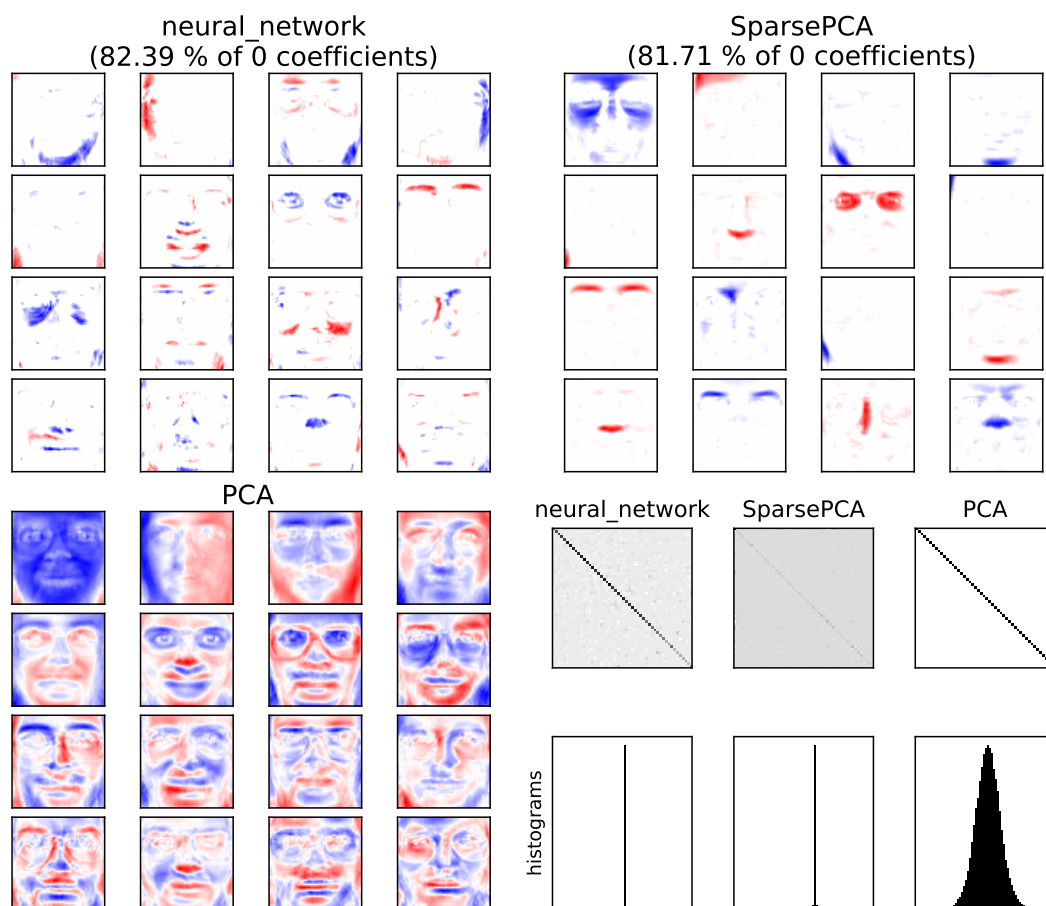
Rysunek 5.3: Dokładność rekonstrukcji dla zbioru gasoline — porównano modele neural_network, SparsePCA oraz model PCA. Oryginalne sygnały zaznaczono linią przerywaną, a rekonstrukcję linią ciągłą.



Rysunek 5.4: Porównanie jakości rekonstrukcji na zbiorze MNIST. Pierwszy wiersz zawiera oryginalne obrazy, poniżej rekonstrukcja przy użyciu kolejno modeli: `neural_network`, `SparsePCA` i `PCA`.



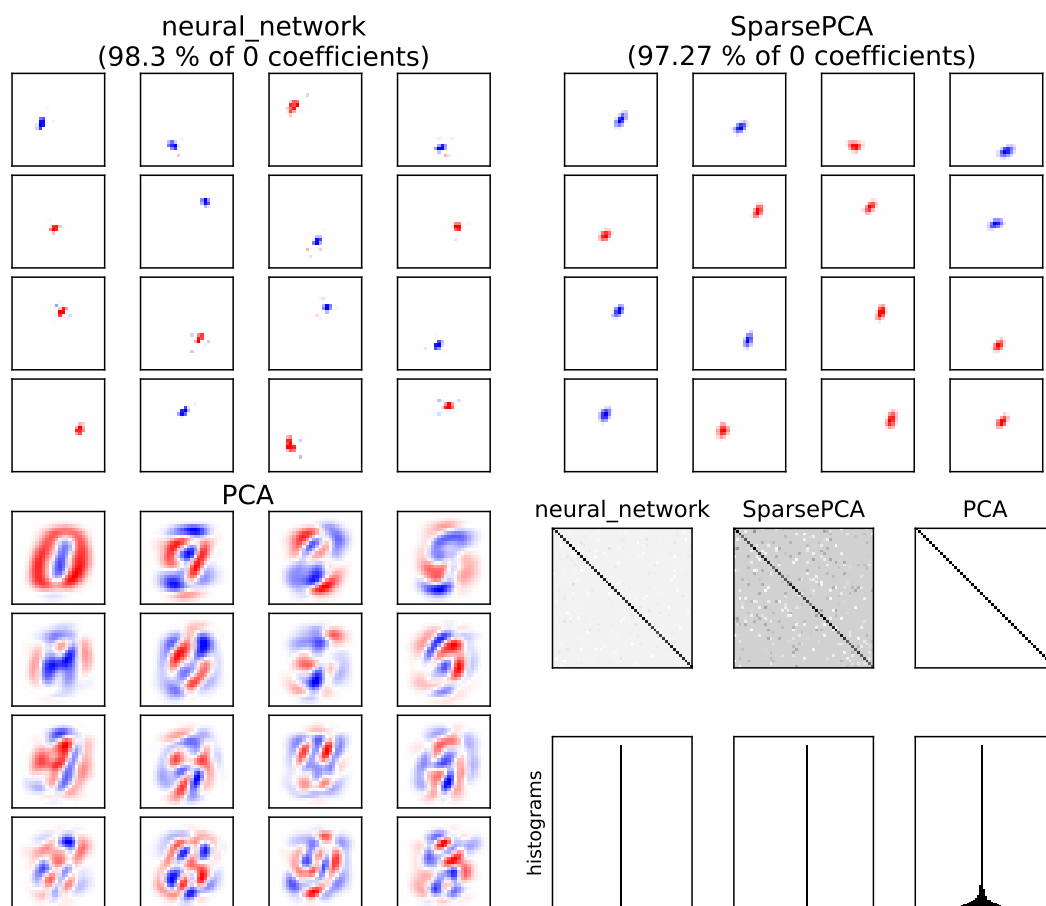
Rysunek 5.5: Porównanie jakości rekonstrukcji na zbiorze Olivetti faces. Pierwszy wiersz zawiera oryginalne obrazy, poniżej rekonstrukcja przy użyciu kolejno modeli: `neural_network`, `SparsePCA` i `PCA`.



Rysunek 5.6: Przykładowe 16 komponentów (spośród wszystkich 50) w zbiorze danych Olivetti faces. Kolor czerwony oznacza wartości ujemne, niebieski dodatnie, biały jest dokładnym zerem; lewy dolny róg: wyniki zwykłej procedury PCA; prawy dolny róg: macierze kowariancji dla komponentów oraz histogramy współczynników modeli. Uzyskana rzadkość: neural_network (82,39%), SparsePCA/scikit-learn (81,71%), PCA (0%).

niż te znalezione przez SparsePCA. Wyniki takie obserwowane są dla wszystkich trzech zbiorów: gasoline, MNIST i Olivetti faces.

Jakość rekonstrukcji dla danych widmowych gasoline jest przedstawiona na rysunku 5.3 (użyto 10 komponentów). Z kolei dla obrazów (dane MNIST i Olivetti faces) jakość rekonstrukcji przedstawiono na rysunkach 5.4 i 5.5 (w tym przypadku użyto 50 komponentów). Jak widać, nawet dla tak względnie małej liczby komponentów jakość jest całkiem dobra. Przy bardzo zbliżonej rzadkości proponowane podejście daje lepszą jakość rekonstrukcji w stosunku do metody SparsePCA. W przypadku faces patrząc na komponenty można zaobserwować nieco wyraźniejsze szczegóły, takie jak



Rysunek 5.7: Przykładowe 16 komponentów (spośród wszystkich 50) w zbiorze danych MNIST. Kolor czerwony oznacza wartości ujemne, niebieski dodatnie, biały jest dokładnym zerem. Lewy dolny róg: wyniki zwykłej procedury PCA, a w prawym dolnym rogu: macierze kowariancji dla komponentów oraz histogramy współczynników modeli. Uzyskana rzadkość: neural_network (98.3%), SparsePCA/scikit-learn (97.27%), PCA (0%).

okulary lub wąsy, broda czy obrys twarzy. W przypadku danych gasoline przy tej samej rzadkości (około 80%, tj. 320 wyzerowanych współczynników) jakość proponowanego podejścia neural_network jest znacznie lepsza. W przypadku zestawów danych gasoline i faces, posiadających stosunkowo mało próbek (ok. 10% liczby atrybutów), jakość rekonstrukcji obu rzadkich metod jest znacznie lepsza niż w przypadku standardowej procedury PCA, a przy większym zbiorze danych (jak MNIST) wpływ rzadkości na jakość rekonstrukcji jest mniejszy.

5.2 Klasyfikacja nieliniowa — uczenie zrandomizowane

Wyniki opisane w tej sekcji zostały opublikowane w pracy [42].

Zgodnie z opisem w sekcji 3.3 w tym eksperymencie wybieramy losowo współczynniki ukrytej warstwy $\mathbf{W}^1$ i $\mathbf{b}^1$, a liczbę neuronów w warstwie ukrytej ustalono na $k = 1000$. Mając ustalone wagi warstwy ukrytej wystarczy nauczyć współczynniki warstwy wyjściowej $\mathbf{W}^2$ i $\mathbf{b}^2$, wykorzystując zamiast macierzy $\mathbf{X}_{n \times p}$ nową macierz przekształconych danych $\mathbf{V}_{n \times k}$, w której nowa obserwacja $\mathbf{v}_i$ odpowiada wynikowi $\tanh(\mathbf{W}^1 \cdot \mathbf{x}_i + \mathbf{b}^1)$. Eksperymenty wykonano w języku Python. Na potrzeby eksperymentów przygotowano klasę `ExtremeClassifier` (w paradygmacie `fit-predict` pakietu `scikit-learn`), która zależy od liczby neuronów w warstwie ukrytej k , rodzaju liniowego klasyfikatora wykorzystanego do strojenia warstwy wyjściowej i jego parametrów. W eksperymentach wszystkie klasyfikatory liniowe mają jednakowo inicjalizowaną warstwę uczącą. Wykorzystano osiem modeli liniowych:

L1-Liblinear model regresji logistycznej z karą ℓ^1 — klasa `LogisticRegression` z ustawieniami `penalty='l1'` oraz `solver='liblinear'`

L1-QR model z regularyzacją ℓ^1 opisany w sekcji 2.5.3 wykorzystujący wcześniejszy rozkład QR/LQ (podobnie do (2.60), por. [42])

Lq-QR model analogiczny do L1-QR wykorzystujący wcześniejszy rozkład QR/LQ, z regularyzacją w postaci pseudonormy ℓ^q , przy $q = 0,8$, por. [43, 44]

L1-QR-soft model zastępujący niegładką normę ℓ^1 jej gładkim przybliżeniem, opisany w sekcji 2.7, por. [45].

L2-QR model regresji logistycznej z regularyzacją ℓ^2 z wcześniejszym wykonaniem rozkładu QR, opisany w sekcji 2.4.1.

L2-lbfgs model regresji logistycznej z karą ℓ^2 — klasa `LogisticRegression` z ustawieniami `penalty='l2'` oraz `solver='lbfgs'`

L2-NewtonCG model regresji logistycznej z karą ℓ^2 — klasa `LogisticRegression` z ustawieniami `penalty='l2'` oraz `solver='newton-cg'`

MLP-lbfgs dodatkowo dla porównania dodano również nieliniowy model — perceptron wielowarstwowy (klasa `MLPClassifier`) z dwoma warstwami oraz z k neuronami w warstwie ukrytej. Model ten jest uczony w standardowy sposób przy użyciu dostępnego w pakiecie algorytmu L-BFGS (oznaczony jako `MLP-lbfgs`). Jest to model o strukturze równoważnej z pozostałymi modelami, z tą różnicą, że model ten stroi współczynniki obu warstw w procesie wstecznej propagacji — poprzednie modele losują i zamrażają współczynniki warstwy pośredniej. Model ten posiada

domyślnie ustawiony parametr $\alpha = 0,0001$ odpowiadający mechanizmowi „weight decay”, czyli regularyzacji ℓ^2 .

Model L1-Liblinear stanowi punkt odniesienia dla modeli rzadkich, a modele L2-lbfgs, l2-netwoncg, MLP-lbfgs stanowią punkt odniesienia dla modeli z regularyzacją ℓ^2 .

W przypadku rzadkich funkcji kary tylko algorytmy L1-Liblinear i L1-QR dają ilościowo te same wyniki, jednak L1-Liblinear może działać kilkukrotnie szybciej. Inne modele (L1-QR-soft, Lq-QR) dają jakościowo różne wyniki. Algorytm Lq-QR dla $q = 0,8$ uzyskał najlepszą rzadkość i był również nieznacznie szybszy niż L1-QR. Pakiet scikit-learn udostępnia rozwiązania dla kar typu ℓ^2 i ℓ^1 , lecz nie posiada takowego dla ogólnego przypadku ℓ^q .

Jako dane testowe wykorzystano trzy sztuczne zbiory:

1. Chessboard3 Szachownica 3×3
2. Chessboard4 Szachownica 4×4
3. Spirals Dwie zagnieżdżone w sobie spirale

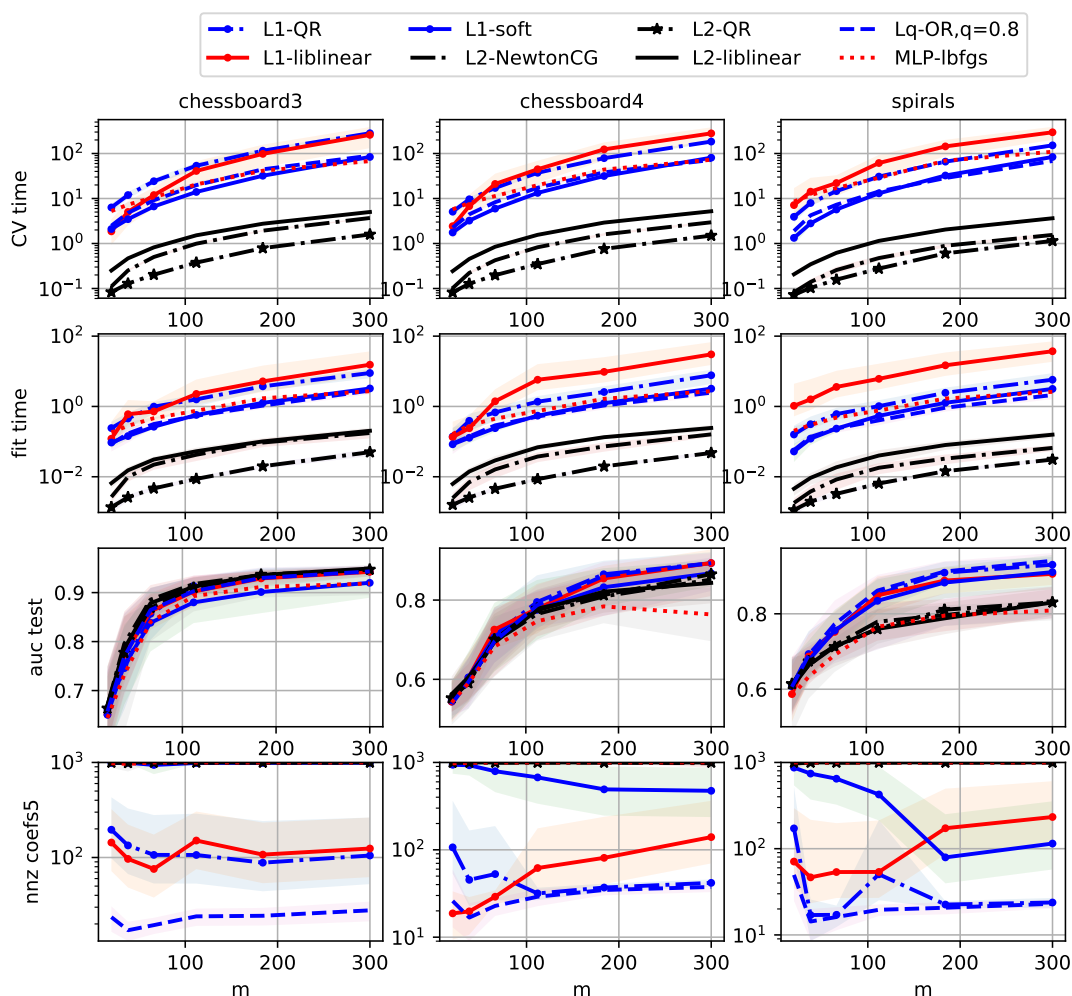
Eksperyment przebiegał następująco: zwiększano rozmiar próby uczącej od 20 do 300 próbek, dla każdego rozmiaru danych uczących i dla każdego klasyfikatora wybierano optymalną wartość parametru $C = 1/\lambda$ za pomocą walidacji krzyżowej, a w trakcie kolejnych iteracji notowano czas uczenia i testowania, błędy na zbiorze testowym oraz liczbę niezerowych współczynników (dotyczy modeli z karą $\ell^{0 < q \leq 1}$). Wyniki zbiorcze eksperymentu przedstawiono na rysunku 5.8.

Jak widać w przypadku regularyzacji ℓ^2 rozwiązanie oparte na rozkładzie QR zawsze daje lepsze czasy dopasowania niż zwykłe „rozwiązujące” dostępne w pakiecie scikit-learn. Czas dopasowania L1-QR jest 2–5 razy krótszy niż L1-Liblinear, szczególnie w przypadku szachownicy 4×4 i dwóch spirali — wynika to z faktu niedostosowania biblioteki do danych gęstych. Patrząc na jakość, widzimy, że modele rzadkie są zbliżonej jakości i okazały się znacznie lepsze niż modele z regularyzacją ℓ^2 , zwłaszcza wyraźnie widać to dla danych Chessboard4 i Spirals.

Dla zbioru Spirals najlepszy okazał się model Lq-QR i był to również model najrzadszy. Model MLP-lbfgs uzyskał najgorszą jakość dopasowania oraz porównywalny czas uczenia modeli rzadkich.

Eksperyment pokazuje, że użycie rozkładu QR można skutecznie połączyć z uczeniem zrandomizowanym (RVFL) z różnymi regularyzatorami. Co więcej, eksperymenty potwierdzają, że takie uczenie działa stabilniej niż zwykły algorytm uczenia sieci neuronowych, szczególnie w przypadku dużej liczby ukrytych neuronów. Przykładowe granice

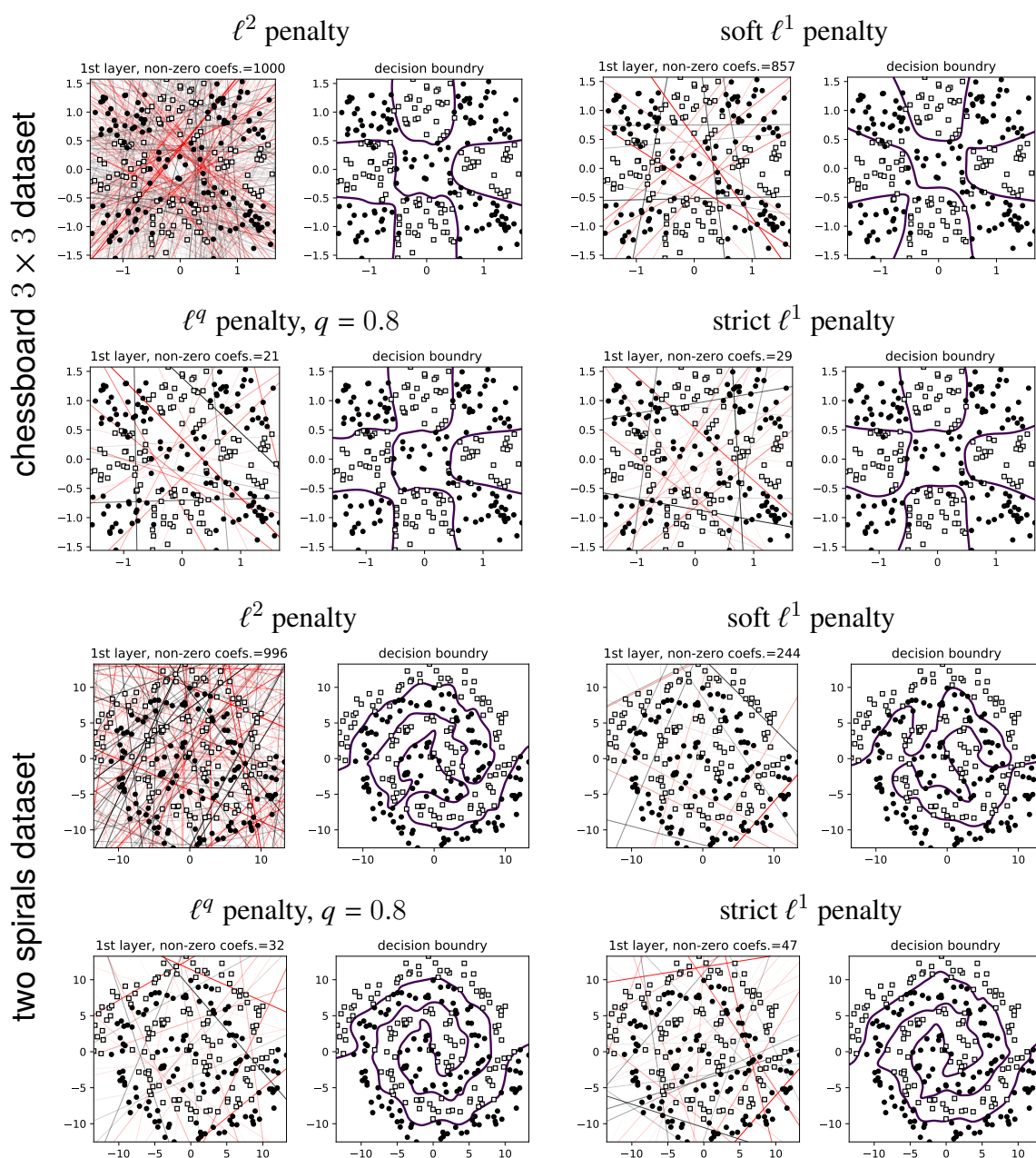
decyzyjne, rzadkość rozwiązania i wizualizacje wag neuronów ukrytych na tle danych wejściowych przedstawiono na rysunku 5.9, z którego można np. odczytać, że modele z rzadką regularyzacją redukują liczbę neuronów w warstwie ukrytej z 1000 do ok. 20–50.



Rysunek 5.8: Wyniki eksperymentalne dla uczenia zrandomizowanego. Porównanie na sztucznych zbiorach danych. CV time to czas procedury walidacji krzyżowej wyboru parametru regularyzacyjnego, fit time to czas dopasowania dla znalezionej λ , auc test to miara AUC (pole pod krzywą ROC) na danych testowych, a nnz coefs to liczba niezerowych współczynników.

Przedstawione zastosowanie rozkładu QR macierzy w celu poprawy procedury Newtona uczenia modeli regresji logistycznej może być efektywnie wykorzystane w obu przypadkach, tj. w przypadku obrotowo niezmienniczym (z karą ℓ^2) oraz w przypadku modeli obrotowo zmiennych (z rzadkimi funkcjami kary)¹. Wykonane eksperymenty wskazują, że użycie rozkładu QR może istotnie poprawić klasyczny algorytm Newton-CG dla regresji

¹W tym drugim podejściu zastosowanie formuły Woodbury’ego będzie miało mniejszą złożoność obliczeniową, bo znika konieczność policzenia rozkładu. Artykuł powstał zanim podjęto bezpośrednią próbę



Rysunek 5.9: Przykładowe granice decyzyjne dla różnych funkcji składników regularyzacyjnych (ℓ^2 , ℓ^1 z gładką aproksymacją funkcji wartości bezwzględnej, $\ell^{q=0.8}$, ℓ^1) dla różnych zbiorów testowych. Na rysunku (lewa strona) pokazano wizualizację współczynników pierwszej warstwy sieci neuronowej przedstawione jako linie — intensywność i kolor reprezentują wielkość i znak danego współczynnika.

z karą ℓ^2 w przypadku danych o rozmiarze $n \ll p$.

zastosowania formuły Woodbury’ego do uczenia modelu regresji logistycznej z karą ℓ^1 z góry ograniczoną przez karą ℓ^2 .

Eksperymenty pokazały, że stosując rozkład QR i formułę Woodbury’ego możemy rozwiązać problem uczenia modelu regresji również z rzadkimi funkcjami kar. Zysk numeryczny z zastosowania rozkładu QR nie w tym przypadku tak spektakularny, jak w przypadku kary ℓ^2 . W pracy [42] pokazano, że stosując rozkład QR uzyskujemy teoretyczną górną granicę lepszą niż w przypadku ogólnej procedury Newton-CG. Eksperymenty numeryczne wykazały, że w przypadku trudniejszych i skorelowanych danych (np. w przypadku uczenia ekstremalnego) oraz większych wartości parametru λ takie podejście może działać efektywniej niż metoda spadku względem współrzędnych (L1-Liblinear). Jednak należy przyznać, że w typowych i prostszych przypadkach metody oparte na spadku względem współrzędnych (w tym biblioteka Liblinear) są szybsze.

6. Podsumowanie i wnioski

W pracy przebadano rozmaite konfiguracje modelu regresji logistycznej, próbując przyspieszyć procedurę uczenia. Dla małych zbiorów danych wykazano teoretycznie i empirycznie, że w pewnym zakresie podejście zwyczajne (`l2_logreg_ordinary`, odpowiednik modelu `LogisticRegression` z ustawionym parametrem `solver='newton-cg'`) okazuje się mniej wydajne pod względem czasu uczenia od rozwiązania uzyskanego za pomocą formuły Woodbury’ego, co szczególnie dobrze widać dla zbiorów gęstych. O ile koncepcja wykorzystania rozkładu QR w zagadnieniu uczenia modeli liniowych (głównie w kontekście regresji) pojawia się w literaturze, o tyle analogiczne zastosowanie formuły Woodbury’ego wydaje się jeszcze niespotykane.

W przypadku rzadkich zbiorów danych analogiczny zysk dla modeli z karą ℓ^2 okazuje się trudniejszy do uzyskania. W eksperymentach było to obserwowane dla bardzo małych zbiorów. Dokładna analiza teoretyczna złożoności procedur w przypadku danych rzadkich wydaje się trudna, gdyż bardzo to zależy od rodzaju rzadkości macierzy wejściowych, ale również od rzadkości macierzy wynikowych powstających w wyniku kolejnych mnożeń, a także sposobów ich reprezentacji [1]. W przeprowadzanych eksperymentach korzystne okazało się przechowywanie wyników mimo wszystko w macierzach gęstych.

Pewna przewaga proponowanego podejścia może się ujawnić wtedy, gdy trzeba wielokrotnie uczyć model na tym samym zbiorze danych (np. podczas walidacji krzyżowej przy poszukiwaniu optymalnego parametru λ) — wtedy istotne znaczenie ma również kolejność parametrów λ_k , dla których chcemy nauczyć model. W takiej procedurze wartości λ powinny tworzyć malejący ciąg, zaś przy uczeniu kolejnego modelu rozsądnym założeniem jest skorzystanie z parametrów modelu nauczonego dla czynnika λ_{k-1} .

W ramach badań dotyczących modeli z regularyzacją ℓ^1 również zaproponowano model bazujący na rozkładzie QR oraz formule Woodbury’ego. Jednak zysk z tego podejścia, mimo teoretycznej atrakcyjności, nie okazał się lepszy od metod referencyjnych. W przypadku regresji logistycznej z regularyzacją ℓ^1 uważna implementacja metody spadku względem współrzędnych, korzystająca ze wszystkich udoskonaleń, radzi sobie zdecydo-

wanie najlepiej pomimo swej prostoty. Wszystkim metodom uczącym dodatkowo pomóc może wykluczenie na wstępie atrybutów, które przy zadanej wartości λ na pewno nie wejdą do wynikowego modelu [24, 77, 81]. W przypadku metody opartej na zastąpieniu normy ℓ^1 przez normę ℓ^2 największym mankamentem jest względnie długi czas potrzebny do uzyskania w pełni rzadkiego rozwiązania, ponieważ zerowanie współczynników zmierzających do zera powoduje chwilowe zaburzenie.

Dla części z procedur konieczna była również adaptacja do danych rzadkich, co ujawniło w eksperymentach pewne zalety metod wyższego rzędu (właśnie dla danych rzadkich) w stosunku metody bazującej na spadku względem współrzędnych. Jednym z wniosków dotyczących adaptacji algorytmów do danych rzadkich dotyczy rozkładu QR. Rozkład ten dla dużych danych jest problematyczny pamięciowo, a zwłaszcza konieczność wyznaczenia macierzy Q , o której należy założyć, że jest gęsta, a jej rozmiar jest równy rozmiarowi danych. Problem ten nie występuje w przypadku zastosowania formuły Woodbury’ego.

W ostatnich wersjach pakietu `scikit-learn` (>1.2) wprowadzono dla regresji z karą ℓ^2 nową metodę strojenia `newton-cholesky`, jednak jest ona dostosowana do przypadku $n \gg p$ i zupełnie nieprzydatna w przypadkach rozważanych w pracy. Eksperymenty pokazują, że nawet drobna modyfikacja istniejącego kodu pakietu `scikit-learn` (jak ta przedstawiona na stronie 22) może dać zysk obliczeniowy dla małych zbiorów danych.

W pracy przebadano szereg własności numerycznych algorytmów i nasuwają się tu pewne wnioski odnośnie ich konfiguracji. W przypadku algorytmów rozwiązujących zadanie uczenia w przestrzeni zredukowanej (metoda z rozkładem QR/LQ czy z formułą Woodbury’ego) zastosowanie preconditionera ma istotne znaczenie, natomiast w przypadku klasycznej procedury Newton-CG nie jest to kluczowe. Z rozważanych procedur gradientu sprzężonego i SYMMLQ, ta pierwsza, mimo że w pewnych konfiguracjach bywa szybsza, nie radzi sobie z problemami nieokreślonymi, a takie macierze czasami pojawiają się np. po zastosowaniu formuły Woodbury’ego uwzględniającego wyraz wolny.

Korzystając z paradygmatu extreme machine learning (zrandomizowanego ustalania wag warstwy wejściowej) można w prosty sposób na bazie istniejącego zbioru danych przygotować nowy, z dużo większą liczbą cech, gdzie granica separacji będzie łatwiejsza do znalezienia. W takim podejściu rzadkie modele liniowe są naturalnym i bardzo atrakcyjnym narzędziem selekcji cech przekształcenia zrandomizowanego. Przeprowadzone eksperymenty pokazały duże zalety takiego podejścia w porównaniu do standardowych algorytmów uczenia sieci wielowarstwowych (np. zwykła wsteczna propagacja czy L-BFGS). W przypadku przewymiarowanych sieci, standardowe metody bazujące na gradiencie prostym mają tendencję utykania w minimach lokalnych (np. część neuronów

warstwy ukrytej nie podlega procesowi uczenia). Ponadto uczenie takiego klasyfikatora z losową warstwą pośrednią jest znacznie szybsze, ponieważ wymaga tylko strojenia warstwy wyjściowej. Takie podejście może być łatwo zaadaptowane do głębokiego uczenia, np. wybierając parametry jąder w sposób losowy i strojąc warstwę wyjściową procedurami uczenia klasyfikatorów liniowych. Podejście to może stanowić punkt startowy przy douczaniu sieci.

Zastosowanie metod regularyzowanych, a zwłaszcza rzadkich regularyzacji (typu lasso czy fused lasso) może być wykorzystane do ekstrakcji cech posiadających pewne własności, atrakcyjne z punktu widzenia interpretacji wyników, np. wykryte piki w spektrogramie czy oczy, broda, okulary na zdjęciach twarzy. Mimo że nie jest to nowa obserwacja, to w ramach pracy udało się zaimplementować strukturę rzadkiego autoenkodera, który może służyć jako ekstraktor cech rzadkich. Opracowana implementacja okazała się w pewnych przypadkach konkurencyjna w stosunku do zaimplementowanego w pakiecie `scikit-learn` uczenia słownikowego.

W pracy duży nacisk położono na metody wyższego rzędu, jednak należy mieć świadomość, że współczesne uczenie maszynowe opiera się głównie na metodach pierwszego rzędu, tj. metodach gradientowych. Pakiety automatycznego różniczkowania (typu Tensorflow czy Torch), na których opierają się biblioteki głębokiego uczenia, mają zdefiniowaną funkcję wartości bezwzględnej (`abs`) wraz z jej pochodną w postaci funkcji `signum` (`sign`). Za pomocą tej funkcji możemy z łatwością zbudować model wyposażony w regularyzację typu ℓ^1 , jednak stosując typową metodę gradientu prostego (bez wyrafinowanej zmiany współczynnika uczenia) nie uzyskamy efektu zerowania się współczynników, a proces zbieżności w pobliże rozwiązania jest bardzo powolny. Ogólnie bezpośrednie stosowanie rzadkich regularyzacji w pakietach wydaje się kłopotliwe, choć jest możliwe.

Spis literatury

- [1] Amir Abboud i in. “The Time Complexity of Fully Sparse Matrix Multiplication”. W: Proceedings of the 2024 Annual ACM-SIAM Symposium on Discrete Algorithms (SODA), s. 4670–4703. DOI: 10.1137/1.9781611977912.167. URL: <https://epubs.siam.org/doi/abs/10.1137/1.9781611977912.167>.
- [2] E. Anderson i in. LAPACK Users’ Guide. Third. Society for Industrial i Applied Mathematics, 1999. DOI: 10.1137/1.9780898719604. eprint: <https://epubs.siam.org/doi/pdf/10.1137/1.9780898719604>. URL: <https://epubs.siam.org/doi/abs/10.1137/1.9780898719604>.
- [3] Galen Andrew i Jianfeng Gao. “Scalable Training of L_1 -Regularized Log-Linear Models”. W: Proceedings of the 24th International Conference on Machine Learning. ICML ’07. Corvallis, Oregon, USA: Association for Computing Machinery, 2007, s. 33–40. ISBN: 9781595937933. DOI: 10.1145/1273496.1273501. URL: <https://doi.org/10.1145/1273496.1273501>.
- [4] Martin Anthony i Peter L. Bartlett. Neural Network Learning: Theoretical Foundations. USA: Cambridge University Press, 2009.
- [5] Richard Barrett i in. Templates for the Solution of Linear Systems: Building Blocks for Iterative Methods. Society for Industrial i Applied Mathematics, 1994. DOI: 10.1137/1.9781611971538. eprint: <https://epubs.siam.org/doi/pdf/10.1137/1.9781611971538>. URL: <https://epubs.siam.org/doi/abs/10.1137/1.9781611971538>.
- [6] Peter L. Bartlett i Shahar Mendelson. “Rademacher and gaussian complexities: risk bounds and structural results”. W: J. Mach. Learn. Res. 3 (2003), 463–482.
- [7] Amir Beck i Marc Teboulle. “A Fast Iterative Shrinkage-Thresholding Algorithm for Linear Inverse Problems.” W: SIAM Journal on Imaging Sciences 2.1 (2009),

- s. 183–202. URL: <http://dblp.uni-trier.de/db/journals/siamis/siamis2.html#BeckT09>.
- [8] H. Bourlard i Y. Kamp. “Auto-association by multilayer perceptrons and singular value decomposition”. W: *Biological Cybernetics* 59.4 (1988), s. 291–294. DOI: 10.1007/BF00332918. URL: <https://doi.org/10.1007/BF00332918>.
 - [9] Stephen Boyd i Lieven Vandenberghe. *Convex Optimization*. New York, NY, USA: Cambridge University Press, 2004. ISBN: 0521833787.
 - [10] Stephen Boyd i in. “Distributed Optimization and Statistical Learning via the Alternating Direction Method of Multipliers”. W: *Foundations and Trends in Machine Learning* 3.1 (sty. 2011), s. 1–122. ISSN: 1935-8237. DOI: 10.1561/22000000016. URL: <http://dx.doi.org/10.1561/22000000016>.
 - [11] Kristian Bredies, Dirk A. Lorenz i Stefan Reiterer. “Minimization of Non-Smooth, Non-Convex Functionals by Iterative Thresholding”. W: *Journal of Optimization Theory and Applications* 165.1 (kw. 2015), s. 78–112. ISSN: 0022-3239. DOI: 10.1007/s10957-014-0614-7. URL: <https://doi.org/10.1007/s10957-014-0614-7>.
 - [12] AT&T Laboratories Cambridge. “The Database of Faces”. W: (1994). URL: <https://www.cl.cam.ac.uk/research/dtg/attarchive/facedatabase.html>.
 - [13] C. Cortes i V. Vapnik. “Support Vector Networks”. W: *Machine Learning* 20 (1995), s. 273–297.
 - [14] IEEE Standard for Floating-Point Arithmetic. *Standard IEEE Std 754-2008*. New York, NY, USA: IEEE Computer Society, 2008.
 - [15] Alexandre d’Aspremont i in. “A Direct Formulation for Sparse PCA Using Semi-definite Programming”. W: *CoRR cs.CE/0406021* (czer. 2004). DOI: 10.2139/ssrn.563524.
 - [16] Aaron Defazio, Francis Bach i Simon Lacoste-Julien. “SAGA: A Fast Incremental Gradient Method with Support for Non-Strongly Convex Composite Objectives”. W: *Proceedings of the 27th International Conference on Neural Information Processing Systems - Volume 1. NIPS’14*. Montreal, Canada: MIT Press, 2014, s. 1646–1654.

- [17] B. Deng i in. “An Overview of Extreme Learning Machine”. W: 2019 4th International Conference on Control, Robotics and Cybernetics (CRC). IEEE, 2019, s. 189–195.
- [18] Bradley Efron i in. “Least Angle Regression”. English. W: The Annals of Statistics 32.2 (2004), s. 407–451. ISSN: 00905364.
- [19] Paul H. C. Eilers. “A Perfect Smoother”. W: Analytical Chemistry 75.14 (2003), s. 3631–3636. DOI: <https://doi.org/10.1021/ac034173t>.
- [20] Jianqing Fan i Runze Li. “Variable Selection via Nonconcave Penalized Likelihood and its Oracle Properties”. W: Journal of the American Statistical Association 96.456 (2001), s. 1348–1360. DOI: 10.1198/016214501753382273.
- [21] Jerome Friedman, Trevor Hastie i Rob Tibshirani. “Regularization Paths for Generalized Linear Models via Coordinate Descent”. W: Journal of Statistical Software, Articles 33.1 (2010), s. 1–22. ISSN: 1548-7660. DOI: 10.18637/jss.v033.i01. URL: <https://www.jstatsoft.org/v033/i01>.
- [22] Leonardo Galli i Chih-Jen Lin. Supplementary materials for “A Study on Truncated Newton Methods for Linear Classification”. 2020. URL: "<https://www.csie.ntu.edu.tw/~cjlin/papers/tncg/supplement.pdf>".
- [23] Leonardo Galli i Chih-Jen Lin. “A Study on Truncated Newton Methods for Linear Classification”. W: IEEE Transactions on Neural Networks and Learning Systems 33.7 (2022), s. 2828–2841. DOI: 10.1109/TNNLS.2020.3045836.
- [24] Laurent Ghaoui, Vivian Viallon i Tarek Rabbani. “Safe Feature Elimination for the LASSO and Sparse Supervised Learning Problems”. W: CoRR abs/1009.3515 (wrz. 2010).
- [25] Gene H. Golub i Charles F. Van Loan. Matrix Computations. Johns Hopkins Studies in the Mathematical Sciences. Johns Hopkins University Press, 2013. ISBN: 9781421407944.
- [26] Pinghua Gong i Jieping Ye. “A Modified Orthant-Wise Limited Memory Quasi-Newton Method with Convergence Analysis”. W: Proceedings of the 32nd International Conference on Machine Learning. Red. Francis Bach i David Blei. T. 37. Proceedings of Machine Learning Research. Lille, France: PMLR, 2015, s. 276–284. URL: <https://proceedings.mlr.press/v37/gonga15.html>.

- [27] P. J. Green. “Iteratively reweighted least squares for maximum likelihood estimation, and some robust and resistant alternatives (with discussion)”. W: Journal of the Royal Statistical Society, Series B, Methodological 46 (1984), s. 149–192.
- [28] T. Hastie i R. Tibshirani. Expression arrays and the $p \gg n$ problem. 2003.
- [29] Trevor Hastie, Robert Tibshirani i Jerome Friedman. The Elements of Statistical Learning. Springer Series in Statistics. New York, NY, USA: Springer New York Inc., 2001.
- [30] H. V. Henderson i S. R. Searle. “On Deriving the Inverse of a Sum of Matrices”. W: SIAM Review 23.1 (1981), s. 53–60. DOI: 10.1137/1023004.
- [31] M. R. Hestenes i E. Stiefel. “Methods of Conjugate Gradients for Solving Linear Systems”. W: Journal of research of the National Bureau of Standards 49 (1952), s. 409–436.
- [32] Tim Hesterberg i in. “Least angle and ℓ_1 penalized regression: A review”. W: Statistics Surveys 2.none (2008), s. 61 –93. DOI: 10.1214/08-SS035. URL: <https://doi.org/10.1214/08-SS035>.
- [33] Arthur E. Hoerl i Robert W. Kennard. “Ridge Regression: Biased Estimation for Nonorthogonal Problems”. W: Technometrics 12.1 (1970), s. 55–67. DOI: 10.1080/00401706.1970.10488634.
- [34] G.B. Huang, L. Chen i C.K. Siew. “Universal Approximation Using Incremental Constructive Feedforward Networks with Random Hidden Nodes”. W: IEEE Transactions on Neural Networks 17.4 (2006), 879—892.
- [35] Peter J. Huber. “Robust estimation of a location parameter”. W: Annals of Mathematical Statistics 35.1 (mar. 1964), s. 73–101. ISSN: 0003-4851. DOI: 10.1214/aoms/1177703732. URL: <http://dx.doi.org/10.1214/aoms/1177703732>.
- [36] R. Hunter i Runze Li. “Variable Selection Using MM Algorithm”. W: Annals of Statistics (2005), s. 1617–1642.
- [37] B. Igel'nik i Yoh-Han Pao. “Stochastic choice of basis functions in adaptive function approximation and the functional-link net”. W: Trans. Neur. Netw. 6.6 (1995), 1320–1329. ISSN: 1045-9227. DOI: 10.1109/72.471375. URL: <https://doi.org/10.1109/72.471375>.

- [38] Intel. Intel® oneAPI Math Kernel Library Documentation. 2023. URL: <https://www.intel.com/content/www/us/en/developer/tools/oneapi/onemkl-documentation.html>.
- [39] Rodolphe Jenatton, Guillaume Obozinski i Francis Bach. Structured Sparse Principal Component Analysis. 2009. URL: <https://hal.archives-ouvertes.fr/hal-00414158>.
- [40] Ata Kabán i Robert J. Durrant. “Learning with $L_{q<1}$ vs L_1 -Norm Regularisation with Exponentially Many Irrelevant Features”. W: Machine Learning and Knowledge Discovery in Databases. Red. Walter Daelemans, Bart Goethals i Katharina Morik. Berlin, Heidelberg: Springer Berlin Heidelberg, 2008, s. 580–596.
- [41] John H. Kalivas. “Two Data Sets of Near Infrared Spectra”. W: Chemometrics and Intelligent Laboratory Systems 37 (1997), s. 255–259.
- [42] Jacek Klimaszewski i Marcin Korzeń. “Fitting Penalized Logistic Regression Models Using QR Factorization”. W: Computational Science – ICCS 2020. Red. Valeria V. Krzhizhanovskaya i in. Cham: Springer International Publishing, 2020, s. 44–57. ISBN: 978-3-030-50417-5.
- [43] Jacek Klimaszewski i Marcin Korzeń. “Image Smoothing Using ℓ^p Penalty for $0 \leq p \leq 1$ with Use of Alternating Minimization Algorithm”. W: Progress in Computer Recognition Systems. Red. Robert Burduk, Marek Kurzynski i Michał Wozniak. Cham: Springer International Publishing, 2020, s. 224–234. ISBN: 978-3-030-19738-4.
- [44] Jacek Klimaszewski i Marcin Korzeń. “Optimization of ℓ^p regularized Linear Models via Coordinate Descent”. W: Schedae Informaticae 25 (2016), s. 61–72.
- [45] Jacek Klimaszewski, Michał Sklyar i Marcin Korzeń. “Learning ℓ^1 -penalized logistic regressions with smooth approximation”. W: 2017 IEEE International Conference on INnovations in Intelligent SysTems and Applications (INISTA). 2017, s. 126–130. DOI: 10.1109/INISTA.2017.8001144.
- [46] Kwangmoo Koh, Seung-Jean Kim i Stephen Boyd. “An Interior-Point Method for LargeScale ℓ_1 -Regularized Logistic Regression”. W: Journal of Machine Learning Research 8 (grud. 2007), s. 1519–1555. ISSN: 1532-4435. URL: <http://dl.acm.org/citation.cfm?id=1314498.1314550>.

- [47] B. Krishnapuram i in. “Sparse Multinomial Logistic Regression: Fast Algorithms and Generalization Bounds”. W: IEEE Transactions on Pattern Analysis and Machine Intelligence 27.6 (2005), s. 957–968.
- [48] Yann LeCun i Corinna Cortes. “MNIST handwritten digit database”. W: (2010). URL: <http://yann.lecun.com/exdb/mnist/>.
- [49] Su-In Lee i in. “Efficient L_1 Regularized Logistic Regression”. W: AAAI. AAAI Press, 2006, s. 401–408. URL: <http://dblp.uni-trier.de/db/conf/aaai/aaai2006.html#LeeLAN06>.
- [50] Yuh-Jye Lee i O.L. Mangasarian. “SSVM: A Smooth Support Vector Machine for Classification”. W: Computational Optimization and Applications 20.1 (2001), s. 5–22. ISSN: 1573-2894. DOI: 10.1023/A:1011215321374. URL: <https://doi.org/10.1023/A:1011215321374>.
- [51] David D. Lewis i in. “RCV1: A New Benchmark Collection for Text Categorization Research”. W: J. Mach. Learn. Res. 5 (2004), 361–397. ISSN: 1532-4435.
- [52] Chih-Jen Lin, Ruby C. Weng i S. Sathiya Keerthi. “Trust Region Newton Method for Logistic Regression.” W: Journal of Machine Learning Research 9 (22 kw. 2010), s. 627–650.
- [53] Julien Mairal i in. “Online Dictionary Learning for Sparse Coding”. W: Proceedings of the 26th Annual International Conference on Machine Learning. ICML '09. Montreal, Quebec, Canada: ACM, 2009, s. 689–696. ISBN: 978-1-60558-516-1. DOI: 10.1145/1553374.1553463. URL: <http://doi.acm.org/10.1145/1553374.1553463>.
- [54] T. P. Minka. A comparison of numerical optimizers for logistic regression. 2003. URL: <https://tminka.github.io/papers/logreg/minka-logreg.pdf>.
- [55] Kevin P. Murphy. Machine learning: a probabilistic perspective. Cambridge, Mass.: MIT Press, 2013.
- [56] Andrew Ng. CS294A lecture notes: Sparse autoencoder. 2019. URL: <https://web.stanford.edu/class/cs294a/sparseAutoencoder.pdf>.
- [57] Andrew Y. Ng. “Feature Selection, L_1 vs. L_2 Regularization, and Rotational Invariance”. W: Proceedings of the Twenty-first International Conference on Machine Learning. ICML '04. Banff, Alberta, Canada: ACM, 2004, s. 78–. ISBN: 1-58113-

838-5. DOI: 10.1145/1015330.1015435. URL: <http://doi.acm.org/10.1145/1015330.1015435>.

- [58] R. W. D. Nickalls. “A New Approach to Solving the Cubic: Cardan’s Solution Revealed”. W: *The Mathematical Gazette* 77.480 (1993), s. 354–359. ISSN: 00255572. URL: <http://www.jstor.org/stable/3619777>.
- [59] Jorge Nocedal i Stephen J. Wright. *Numerical Optimization*. 2e. New York, NY, USA: Springer, 2006.
- [60] C. C. Paige i M. A. Saunders. “Solution of Sparse Indefinite Systems of Linear Equations”. W: *SIAM Journal on Numerical Analysis* 12.4 (1975), s. 617–629. DOI: 10.1137/0712047. eprint: <https://doi.org/10.1137/0712047>. URL: <https://doi.org/10.1137/0712047>.
- [61] F. Pedregosa i in. “Scikit-learn: Machine Learning in Python”. W: *Journal of Machine Learning Research* 12 (2011), s. 2825–2830.
- [62] Jan Rodziewicz-Bielewicz, Jacek Klimaszewski i Marcin Korzeń. “Regularized Learning of Neural Network with Application to Sparse PCA”. W: *Artificial Intelligence and Soft Computing*. Red. Leszek Rutkowski i in. Cham: Springer International Publishing, 2019, s. 203–214. ISBN: 978-3-030-20912-4.
- [63] Volker Roth. “Probabilistic Discriminative Kernel Classifiers for Multi-Class Problems”. W: *Pattern Recognition*. Red. Bernd Radig i Stefan Florczyk. Berlin, Heidelberg: Springer Berlin Heidelberg, 2001, s. 246–253. ISBN: 978-3-540-45404-5.
- [64] Donald B. Rubin. “The Calculation of Posterior Distributions by Data Augmentation: Comment: A Noniterative Sampling/Importance Resampling Alternative to the Data Augmentation Algorithm for Creating a Few Imputations When Fractions of Missing Information Are Modest: The SIR Algorithm”. W: *Journal of the American Statistical Association* 82.398 (1987), s. 543–546. ISSN: 01621459. URL: <http://www.jstor.org/stable/2289460>.
- [65] Leonid I. Rudin, Stanley Osher i Emad Fatemi. “Nonlinear Total Variation Based Noise Removal Algorithms”. W: *Proceedings of the Eleventh Annual International Conference of the Center for Nonlinear Studies on Experimental Mathematics: Computational Issues in Nonlinear Science*. Los Alamos, New Mexico, USA: Elsevier North-Holland, Inc., 1992, s. 259–268.

- [66] Yousef Saad. Iterative Methods for Sparse Linear Systems. Second. Other Titles in Applied Mathematics. SIAM, 2003. ISBN: 978-0-89871-534-7. DOI: 10.1137/1.9780898718003.
- [67] Conrad Sanderson i Ryan Curtin. “Armadillo: a template-based C++ library for linear algebra”. W: Journal of Open Source Software 1.2 (2016), s. 26. DOI: 10.21105/joss.00026. URL: <https://doi.org/10.21105/joss.00026>.
- [68] Mark Schmidt, Nicolas Le Roux i Francis Bach. “Minimizing Finite Sums with the Stochastic Average Gradient”. W: Mathematical Programming (2013).
- [69] Shai Shalev-Shwartz i Ambuj Tewari. “Stochastic Methods for ℓ_1 regularized Loss Minimization”. W: Journal of Machine Learning Research 12 (lip. 2011), s. 1865–1892. ISSN: 1532-4435. URL: <http://dl.acm.org/citation.cfm?id=1953048.2021059>.
- [70] Jack Sherman i Winifred J. Morrison. “Adjustment of an Inverse Matrix Corresponding to a Change in One Element of a Given Matrix”. W: The Annals of Mathematical Statistics 21.1 (1950), s. 124–127. DOI: 10.1214/aoms/1177729893. URL: <https://doi.org/10.1214/aoms/1177729893>.
- [71] Connor Shorten i Taghi M Khoshgoftaar. “A survey on image data augmentation for deep learning”. W: Journal of Big Data 6.1 (2019), s. 1–48.
- [72] Nitish Srivastava i in. “Dropout: A Simple Way to Prevent Neural Networks from Overfitting”. W: Journal of Machine Learning Research 15.56 (2014), s. 1929–1958. URL: <http://jmlr.org/papers/v15/srivastava14a.html>.
- [73] Jiexiong Tang i Guang-Bin Huang. “Extreme Learning Machine for Multilayer Perceptron”. W: IEEE TRANSACTIONS ON NEURAL NETWORKS AND LEARNING SYSTEMS 27.4 (2016).
- [74] Martin A. Tanner i Wing Hung Wong. “The Calculation of Posterior Distributions by Data Augmentation”. W: Journal of the American Statistical Association 82.398 (1987), s. 528–540. ISSN: 01621459. URL: <http://www.jstor.org/stable/2289457> (term. wiz. 09.09.2024).
- [75] R. Tibshirani. “Regression Shrinkage and Selection via the Lasso”. W: Journal of the Royal Statistical Society (Series B) 58 (1996), s. 267–288.
- [76] Robert Tibshirani i in. “Sparsity and smoothness via the fused lasso”. W: Journal of the Royal Statistical Society Series B (2005), s. 91–108.

- [77] Robert Tibshirani i in. “Strong Rules for Discarding Predictors in Lasso-Type Problems”. W: Journal of the Royal Statistical Society Series B: Statistical Methodology 74.2 (list. 2011), s. 245–266. ISSN: 1369-7412. DOI: 10.1111/j.1467-9868.2011.01004.x. eprint: https://academic.oup.com/jrsssb/article-pdf/74/2/245/49513900/jrsssb_74_2_245.pdf. URL: <https://doi.org/10.1111/j.1467-9868.2011.01004.x>.
- [78] A. N. Tikhonov. “On the solution of ill-posed problems and the method of regularization”. W: Dokl. Akad. Nauk SSSR 151.3 (1963), s. 501–504.
- [79] Paul Tseng i Sangwoon Yun. “A Coordinate Gradient Descent Method for Non-smooth Separable Minimization”. W: Math. Program. 117 (mar. 2009), s. 387–423. DOI: 10.1007/s10107-007-0170-0.
- [80] Vladimir N. Vapnik. The nature of statistical learning theory. Springer-Verlag New York, Inc., 1995. ISBN: 0-387-94559-8.
- [81] Jie Wang i in. “A Safe Screening Rule for Sparse Logistic Regression”. W: Neural Information Processing Systems. 2013. URL: <https://api.semanticscholar.org/CorpusID:2499725>.
- [82] L.P. Wang i C.R. Wan. “Comments on ‘The Extreme Learning Machine’”. W: IEEE Transactions on Neural Networks 19.8 (2008), s. 1494–1495.
- [83] Peter M. Williams. “Bayesian Regularization and Pruning Using a Laplace Prior”. W: Neural Computation 7 (1995), s. 117–143.
- [84] Max A Woodbury. Inverting modified matrices. Statistical Research Group, Memo. Rep. 42. Princeton, N. J: Princeton University, 1950.
- [85] Gui-Bo Ye i Xiaohui Xie. “Split Bregman Method for Large Scale Fused Lasso”. W: Computational Statistics and Data Analysis 55.4 (kw. 2011), 1552–1569. ISSN: 0167-9473.
- [86] Guo-Xun Yuan, Chia-Hua Ho i Chih-Jen Lin. “An Improved GLMNET for L1-regularized Logistic Regression”. W: Journal of Machine Learning Research 13.1 (czer. 2012), s. 1999–2030. ISSN: 1532-4435. URL: <http://dl.acm.org/citation.cfm?id=2503308.2343708>.
- [87] Guo-Xun Yuan i in. “A Comparison of Optimization Methods and Software for Large-Scale L1-Regularized Linear Classification”. W: Journal of Machine Learning Research 11 (grud. 2010), 3183–3234. ISSN: 1532-4435.

- [88] Ciyu Zhu i in. “Algorithm 778: L-BFGS-B: Fortran Subroutines for Large-scale Bound-constrained Optimization”. W: ACM Trans. Math. Softw. 23.4 (grud. 1997), s. 550–560. ISSN: 0098-3500. DOI: 10.1145/279232.279236. URL: <http://doi.acm.org/10.1145/279232.279236>.
- [89] Hui Zou i Trevor Hastie. “Regularization and variable selection via the Elastic Net”. W: Journal of the Royal Statistical Society, Series B 67 (2005), s. 301–320.
- [90] A.M. Zyarah i D. Kudithipudi. “Extreme learning machine as a generalizable classification engine”. W: 2017 International Joint Conference on Neural Networks (IJCNN). IEEE, 2017, s. 3371–3376.

Załącznik A

Do pracy załączono zestaw kodów źródłowych w językach C++ i Python, zawierający wszystkie metody zaimplementowane w ramach pracy oraz większość eksperymentów. Przed uruchomieniem należy zainstalować pakiety zależne (m.in. MKL, Armadillo, Python/scikit-learn) oraz zbudować odpowiednie biblioteki wywołując polecenie `make`.