

ROZDZIAŁ 4

METODY SZTUCZNEJ INTELIGENCJI DO WSPOMAGANIA DIAGNOSTYKI PATOLOGII TKANEK

**prof. dr hab. inż. Stanisław OSOWSKI, dr hab. inż. Tomasz MARKIEWICZ, dr inż. Michał KRUK, prof.
dr hab. Wojciech KOZŁOWSKI**

4.1. Wprowadzenie

Rozwój elektronicznej technologii obliczeniowej w powiązaniu z metodami sztucznej inteligencji otworzył nowe możliwości prowadzenia badań i rozwiązywania problemów o dużej złożoności obliczeniowej. Powstały nowe narzędzia wydobywania i analizy informacji uzyskanej z pomiarów, dostarczające nowych możliwości poznawczych, dotąd niedostępnych dla człowieka. Nastąpił znaczący postęp w przetwarzaniu i gromadzeniu danych w czasie akceptowalnym przez użytkownika. Jedną z dziedzin, w której komputer znajduje coraz szersze zastosowanie, jest inżynieria biomedyczna, operująca ogromnym potencjałem trudnych, nierozwiązanych dotąd zadań, ważnych dla rozwoju medycyny.

Obecny poziom rozwoju nauk komputerowych pozwala efektywnie wdrażać do praktyki teorię i osiągnięcia sztucznej inteligencji, dziedziny informatyki dotyczącej metod i technik wnioskowania symbolicznego przez komputer oraz symbolicznej reprezentacji wiedzy stosowanej podczas tego wnioskowania.

Zastosowanie technik komputerowych w powiązaniu z metodami sztucznej inteligencji umożliwia z jednej strony automatyczną analizę wielu złożonych problemów biomedycznych, a z drugiej staje się motorem dalszego rozwoju i postępów w rozumieniu zjawisk patologicznych zachodzących w organizmie człowieka i związanych z nimi nowoczesnych metod leczenia.

Istotne znaczenie dla dalszego postępu badań medycznych odgrywa rozwój automatycznych metod analizy i diagnostyki medycznej, możliwy dzięki zastosowaniu komputerów. Jednym z ważnych kierunków badań w tym zakresie jest analiza obrazów medycznych, pełniących ważną rolę w zakresie oceny diagnostycznej stanów patologicznych tkanek i możliwości podjęcia specjalizowanego leczenia. Dotyczy to zarówno różnego rodzaju chorób nowotworowych jak i stanów zapalnych organów człowieka. Istotnym elementem rozpoznania choroby, stanu jej zaawansowania i możliwości podjęcia właściwego leczenia jest analiza obrazu chorej tkanki ukierunkowana na rozpoznanie zmian chorobowych i ich ocenę z punktu widzenia sposobu dalszego leczenia.

Rozdział ten poświęcony będzie przedstawieniu wybranych metod sztucznej inteligencji w zastosowaniu do automatycznej analizy obrazów medycznych wspomagających diagnostykę lekarską. Utworzenie komputerowego systemu wspomagającego diagnostykę wymaga wykonania wielu kolejnych etapów. Do najważniejszych należą: segmentacja obrazu ukierunkowana na wyodrębnienie istotnych obiektów graficznych (np. pojedynczych komórek) podlegających analizie, przyporządkowanie tym obiektom deskryptorów numerycznych charakteryzujących je w sposób jednoznaczny, ocena i selekcja tych deskryptorów mająca na celu wyodrębnienie najważniejszych elementów informacji (cech diagnostycznych) oraz zastosowanie tych cech jako zbioru wejściowego dla systemów klasyfikacyjnych dokonujących końcowego rozpoznania obiektów obrazu.

4.4. Segmentacja obrazu

Celem segmentacji jest wydzielenie obiektów obrazu podlegających rozpoznaniu. Pierwszym krokiem jest zwykle skalowanie obrazu odpowiednio do przyjętych wartości referencyjnych. Można to osiągnąć na wiele sposobów, np. rozciąganie histogramu, przesunięcie czy przeskalowanie histogramu za pomocą jednej lub dwu wartości referencyjnych. Możliwość skutecznego zastosowania określonego sposobu skalowania zależy od właściwości obiektów poddanych skalowaniu (np. duża różnorodność barw obiektów tego samego typu). W przypadku obrazów medycznych najczęściej wybiera się jedną wartość referencyjną, na przykład tła przetransformowanego do barwy białej. Najciemniejszy punkt obrazu pozostaje wówczas niezmienny, natomiast pozostałe są transformowane liniowo. Po wykonaniu skalowania obraz poddawany jest dalszym etapom przetwarzania i segmentacji.

Wyróżnić można wiele metod stosowanych w procesie segmentacji, z których najważniejsze to [3,7,22]

- Segmentacja przez progowanie
- Segmentacja poprzez wykrywanie krawędzi
- Segmentacja przez rozrost obszaru
- Segmentacja oparta na statystycznej ocenie jednorodności obszaru
- Segmentacja oparta na klasyfikacji.

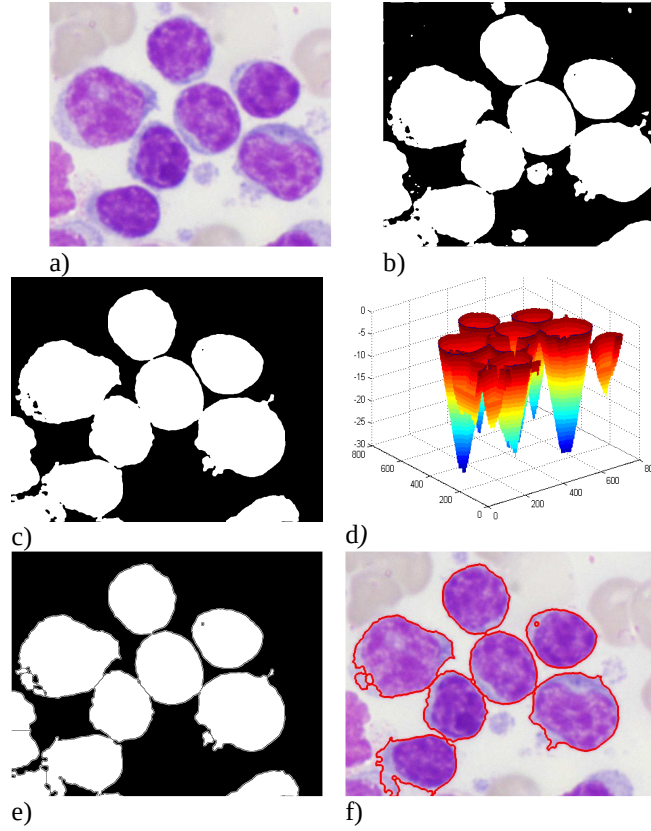
Wszystkie wymienione metody różnią się w sposobie podejścia do ekstrakcji obiektów. W przypadku progowania podstawą są operacje logiczne na pikselach i przypisanie im wartości binarnej 0 lub 1. Wykrywanie krawędzi może odbywać się przy wykorzystaniu filtrów gradientowych (np. Canny, Sobela, Prewitta). W metodzie rozrostu obszaru startuje się z określonej liczby punktów początkowych dołączając do nich następne o podobnym stopniu szarości. Segmentacja związana ze statystyczną oceną jednorodności obszaru wymaga określenia wzorców poszczególnych typów obszaru obiektów i porównaniu aktualnego stopnia szarości ze wzorcem. Segmentacja wykorzystująca klasyfikatory stosuje się bądź samoorganizację (np. K-means, C-means) bądź sieci neuronowe (np. SVM, MLP, RBF).

Dla uzyskania najlepszych wyników segmentacji stosuje się powszechnie operacje morfologiczne, takie jak erozja, dylatacja, otwarcie i zamknięcie czy bardziej zaawansowane, np. transformacja wododziałowa (ang. watershed transformation) przeprowadzane na obrazie binarnym. Przy transformacji do skali szarości a następnie postaci binarnej można za punkt wyjścia wybrać dowolną reprezentację barwną obrazu. Często bardzo dobre wyniki uwidaczniające lepiej brzegi obszarów uzyskuje się przy transformacji obrazu różnicowego w dwu wybranych składowych barw, np. niebieskiej i zielonej.

Skuteczna i powtarzalna dla wielu obrazów jest metoda segmentacji przez progowanie (zwykle dla wartości progu z przedziału (0.8 – 0.98) w skali znormalizowanej (0 - 1)). Tak powstały obraz binarny można poddać w pierwszej kolejności operacjom morfologicznym np. zamykania. Operacja taka jest przeprowadzana na kolorze białym i usuwa wszelkie drobne pozostałości nie należące do obiektu, wygładzając jednocześnie ich kontury. Element strukturujący (SE) może mieć postać kwadratu, linii, dysku, koła lub sześciokąta.

Kolejnym etapem jest wygenerowanie macierzy odległości poszczególnych pikseli komórek od tła i zastosowanie segmentacji metodą działów wodnych. Generacja macierzy odległości odbywa się przy użyciu określonego kształtu elementu strukturującego [22]. Macierz taka jest generowana przy wielokrotnym powtórzeniu operacji erozji wybranym elementem SE. Wszystkie piksele usunięte z powierzchni komórek w kolejnym kroku otrzymują odległość równą temu krokowi. Na tak otrzymanej macierzy odległości przeprowadzana jest procedura segmentacji metodą działów wodnych prowadząca do

wydzielenia poszczególnych komórek. Na rys. 4.1 przedstawiono przykładowo kolejne etapy segmentacji obrazu szpiku kostnego prowadzące do wydzielenia komórek krwiotwórczych.



Rys. 4.1. Ilustracja kolejnych etapów ekstrakcji komórek krwiotwórczych poprzez operacje morfologiczne: a) obraz wejściowy, b) obraz binarny, c) wynik filtracji morfologicznej, d) odwrócona mapa odległości, e) wynik działania operacji działów wodnych, f) wydzielone komórki

Duży wpływ na wynik segmentacji ma wybór progu przy tworzeniu obrazu binarnego. Przy bardziej złożonych obrazach przyjęcie progu ustalanego metodą prób i błędów jest mało skuteczne. Lepsze wyniki uzyskuje się zwykle przy zastosowaniu automatycznego progowania, np. metodą Otsu [21].

Na podstawie histogramu obrazu metoda ta pozwala wyznaczyć próg optymalnie separujący obiekt od tła obrazu. Dla każdej możliwej wartości progu wyznaczane są wariancje klas σ_1^2, σ_2^2 oraz wariancja międzyklasowa σ_B^2 , gdzie klasy 1 i 2 nazywać będziemy odpowiednio obiektem i tłem. Określenie tych wielkości przeprowadza się na podstawie znormalizowanego histogramu poprzez wyznaczenie prawdopodobieństwa przynależności piksela do klasy 1 i 2, oznaczonego odpowiednio jako ω_1 i ω_2 . Następnie oblicza się średnie poziomy szarości obydwu klas i całego obrazu, oznaczone odpowiednio jako μ_1, μ_2, μ_T . Jeśli p_i oznacza i -tą wartość histogramu, to wariancje dla wartości progu k są określone wzorami [21]:

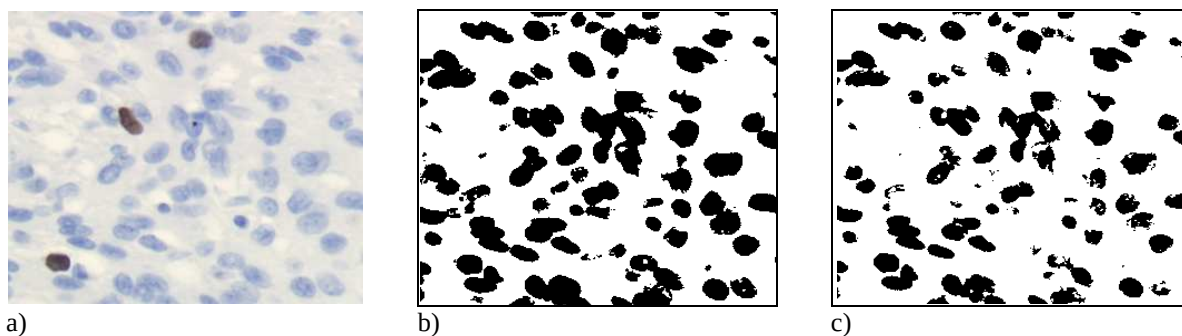
$$\sigma_1^2 = \sum_{i=0}^{k-1} \frac{(i - \mu_1)^2 p_i}{\omega_1}, \quad (4.1)$$

$$\sigma_2^2 = \sum_{i=k}^{t_{\max}} \frac{(i - \mu_2)^2 p_i}{\omega_2}, \quad (4.2)$$

$$\sigma_B^2 = \omega_1(\mu_1 - \mu_T)^2 + \omega_2(\mu_2 - \mu_T)^2 = \omega_1\omega_2(\mu_2 - \mu_1)^2. \quad (4.3)$$

Próg k , dla którego tak wyznaczona wariancja σ_B^2 przyjmuje wartość maksymalną, jest progiem optymalnym.

Na rys. 4.2 zobrazowano efekt operacji progowania obrazu nowotworu meningiomy z różnymi wartościami progu [14]. Rysunek 4.2a przedstawia wystandaryzowany obraz źródłowy, rys. 4.2b – wynik dla progu wyznaczonego metodą Otsu ($t_1 = 199$), a rys. 4.2c – wynik progowania dla wartości progu $t_1 = 215$. Jak widać dla wartości progu wyznaczonej metodą Otsu pola powierzchni odpowiadające profilom większości zabarwionych jąder komórek zostały prawidłowo wydzielone z tła. Tylko niektóre profile jąder komórek nie zostały w całości wydzielone z uwagi na ich niejednorodne jasne zabarwienie. Również obszar tła został prawidłowo odróżniony, nie tworząc żadnych większych artefaktów w otrzymanym obrazie binarnym profili jąder komórek. W przypadku za wysokiego progu (rys. 4.2c) znacząca liczba profili jąder komórek została wydzielona tylko fragmentarycznie.

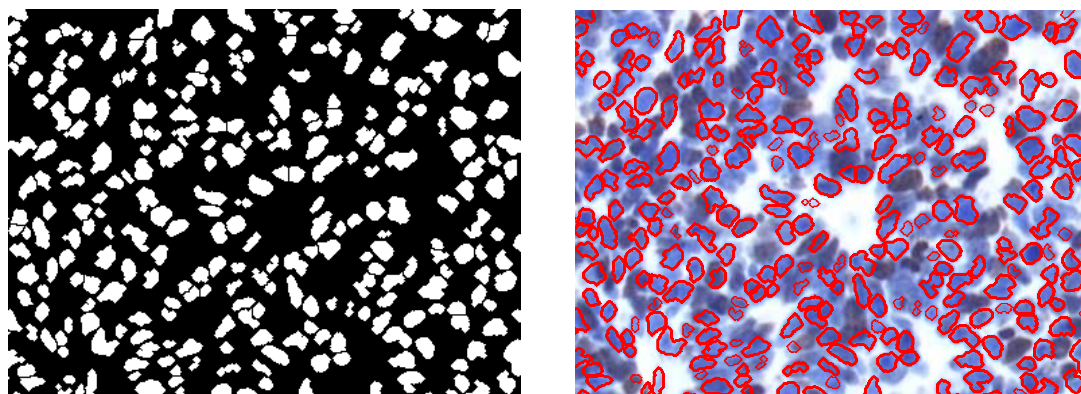


Rys. 4.4. Wyniki operacji progowania obrazu nowotworu meningiomy przy zastosowaniu progowania jednokrotnego: a) obraz oryginalny, b) wynik binarny uzyskany przy progu ustalonym metodą Otsu ($t_1 = 199$), c) wynik binarny uzyskany przy progu $t_1 = 215$

Inną skuteczną metodą segmentacji jest progowanie sekwencyjne. Algorytm ten opiera się na kolejno wykonywanych operacjach progowania na pikselu f obrazu przy krokowym zwiększaniu progu dolnego t_1 i ustalonej wartości progu górnego t_2 zgodnie ze wzorem

$$T_{[t_1, t_2]}(f) = \begin{cases} 1 & \text{dla } t_1 < f < t_2 \\ 0 & \text{else} \end{cases} \quad (4.4)$$

Na starcie przyjmuje się najmniejszą wartość $t_1 = 1$ zwiększając ją sukcesywnie aż do wartości maksymalnej. W każdym kroku algorytmu sprawdzana jest wielkość wydzielonych obiektów pod względem spełnienia górnego i dolnego kryterium pola wydzielonych obiektów. W każdej iteracji odseparowane obiekty spełniające te kryteria są dodawane do obrazu binarnego będącego maską zliczonych profili obiektów. Jeśli profil obiektu został w danym kroku wydzielony, w kolejnych krokach przy zwiększonej wartości progu otrzymuje się również ten sam odseparowany obiekt bez zmiany reprezentacji wydzielonych obiektów.



Rys. 4.3. Wynik segmentacji komórek w obrazie neuroblastomy w postaci: a) maski binarnej profili jąder, b) obrazu oryginalnego z naniesionymi konturami komórek

Przykładowy rezultat zastosowania progowania sekwencyjnego do segmentacji profili jąder komórek w obrazie nowotworu neuroblastomy przedstawiony jest na rys. 4.3 [14].

4.3. Generacja deskryptorów numerycznych obrazu

Automatyczne rozpoznawanie obrazów wymaga stworzenia opisu numerycznego tych obrazów w postaci odpowiednio zdefiniowanych deskryptorów, najlepiej charakteryzujących te obrazy. W przypadku obrazów medycznych można zastosować różne sposoby opisu, na przykład [7,27]

- deskryptory teksturalne – opisujące tekstury, a więc rozkłady statystyczne odcieni szarości lub kolorów, bez dogłębnej analizy szczegółów obiektu,
- deskryptory geometryczne – charakteryzujące właściwości geometryczne obrazu, takie jak pole powierzchni, obwód, wypukłości, symetria i inne parametry opisujące kształt obiektu,
- deskryptory kolorymetryczne – określające w sposób statystyczny zmienność kolorów obiektu przy pomocy momentów statystycznych, np. wartość średnia, wariancja, skośność, kurtoza,
- deskryptory morfologiczne – wartości deskryptorów wyżej wymienionych, określonych po zastosowaniu operacji morfologicznych na badanym obrazie, np. stosunek powierzchni obiektu przed i po wykonaniu określonej operacji morfologicznej. Deskryptory te w sposób bezpośredni opisują zmiany w kształcie obiektu powstałe po przeprowadzeniu operacji morfologicznych, na przykład wygładzeniu konturu.

Tak wygenerowane opisy numeryczne lepiej lub gorzej charakteryzują obraz (obiekt) pod względem zdolności rozróżniania różnych obiektów. Po odpowiedniej selekcji posłużą one stworzeniu wektora cech diagnostycznych dla klasyfikatora. Należy przy tym pamiętać o normalizacji, czyli sprowadzeniu wartości różnych cech do zbliżonych poziomów. Normalizacja taka może być przeprowadzona na różne sposoby, na przykład podzielenie wartości oryginalnych przez odpowiadającą im wartość maksymalną, normalizacja statystyczna, itp [4,19].

Istnieje wiele metod opisu tekstury. Różnią się one liczbą wydzielanych deskryptorów i podstawami matematycznymi ich generacji. Do najważniejszych można zaliczyć algorytm GLCM Haralicka, algorytm Unsera, Gabora, Markowa, algorytm geometryczny Chena, algorytm wykorzystujący opis fraktalny, fourierowski, falkowy itp. Obszerny opis algorytmów tekstualnych można znaleźć w pracy [28].

Deskryptory geometryczne określające kształt i powierzchnię obiektów należą do najczęściej używanych. Są zwykle generowane przy użyciu masek binarnych obiektu. Należy zaznaczyć, że największy wpływ na jakość deskryptorów geometrycznych ma poprawność przeprowadzonej segmentacji obrazu. Spośród tego typu deskryptorów obiektów graficznych do najczęściej wykorzystywanych należą: pole powierzchni mierzone w pikselach, długość najdłuższej i najkrótszej osi obiektu, ekscentryczność, powierzchnia wielokąta wypukłego opisanego na obiekcie, promień koła o takiej samej powierzchni co obiekt, stosunek powierzchni obiektu do powierzchni najmniejszego prostokąta opisanego na nim, obwód, średnia odległość pikselowa od centralnego piksela obiektu, zwartość mierzona jako stosunek kwadratu obwodu obiektu do jego powierzchni, symetria osiowa, symetria pól względem najdłuższej osi czy liczba wklęsłości obiektu.

Cechy kolorymetryczne wykorzystuje się do opisu statystycznego obrazu dla podstawowych kolorów (zwykle RGB). Mogą one dotyczyć całego obrazu, jego gradientu bądź innych sposobów opisu kolorów, np. histogramu. Pod uwagę bierze się zwykle wartości średnie, wariancję, skośność oraz kurtozę [19].

4.4. Selekcja cech diagnostycznych

Deskrytory numeryczne omówione w punkcie poprzednim stanowią maksymalny zbiór potencjalnych cech diagnostycznych, które mogą być wykorzystane w automatycznym rozpoznaniu wzorca reprezentującego badany obiekt. Badania pokazały, że użycie maksymalnego zestawu cech nie prowadzi do najlepszych wyników [8], gdyż nie są one jednakowo ważne w procesie rozpoznania wzorców. Pewne cechy w procesie rozpoznania pełnią funkcję szumu pomiarowego, pogarszając możliwość rozpoznania różnych typów obiektów. Ważnym elementem procesu staje się zatem ocena jakości cech i opracowanie metod ich selekcji przy tworzeniu optymalnego wektora wejściowego $\mathbf{x}$, na podstawie którego przeprowadzone będzie ostateczne rozpoznanie klasy obiektu. Zdolność generalizacji klasyfikatora jest powiązana zarówno z wymiarem wektora wejściowego, ale również ze składem cech i takim ich doбором, który najlepiej różnicuje różne typy obiektów. W wyniku badań stwierdzono, że cechy wysoce skorelowane mają zwykle niekorzystny wpływ na jakość klasyfikacji, dominując nad innymi i tłumiąc w ten sposób ich korzystne działanie.

W badaniu jakości cech można zastosować dwie strategie [8,20]. W pierwszej bada się każdą cechę niezależnie od zastosowanej metody klasyfikacji (tzw. filtrowanie cech), oceniając ich jakość pod kątem różnicowania klas bez uwzględniania końcowej fazy czyli klasyfikacji. Cechy o najgorszym wskaźniku różnicowania są odrzucane, a pozostałe tworzą wektor wejściowy $\mathbf{x}$. W praktyce okazuje się, że włączenie równoległego działania wielu cech na raz może zmienić „jakość” danej cechy. Pewne cechy (nawet te gorsze) współpracując ze sobą wzbogacają się nawzajem, podnosząc wzajemnie swoją wartość diagnostyczną. Stąd w praktyce zalecane jest wartościowanie poszczególnych cech działających jednocześnie.

W drugim podejściu do selekcji cech diagnostycznych (tzw. metoda powłoki) selekcja optymalnego zbioru cech odbywa się przy ścisłej współpracy algorytmu selekcyjnego z klasyfikatorem dokonującym rozpoznania obiektów. Należy podkreślić, że wynik selekcji w dużej mierze zależy od zastosowanej metody, a zbiór wyselekcjonowanych cech może znacznie różnić się między sobą.

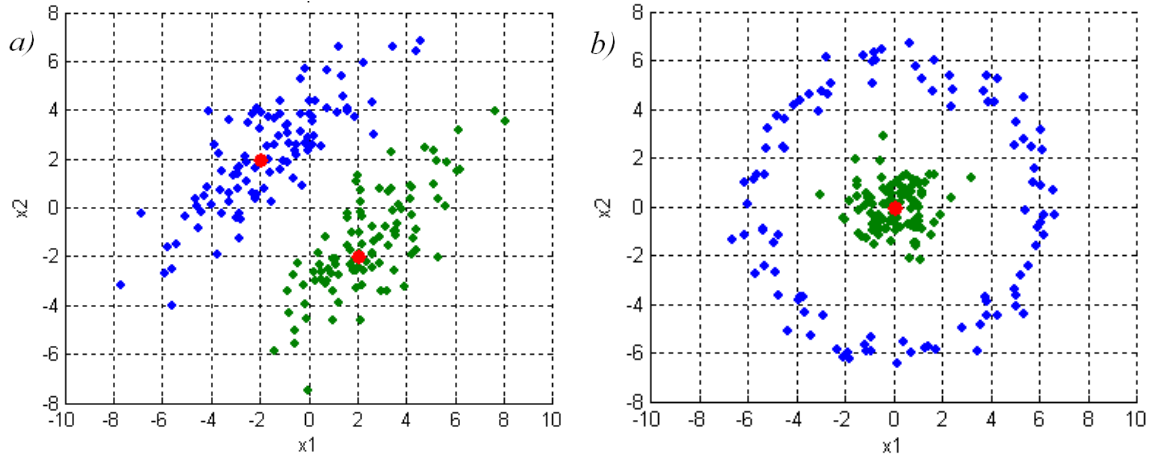
Istnieje bogata literatura dotycząca metod selekcji [8,20]. W tym wykładzie omówimy kilka wybranych metod, które zdaniem autorów dobrze sprawdzają się w praktyce przetwarzania obrazów medycznych.

4.4.1. Metoda Fishera

Selekcja Fishera stanowi jedną z najbardziej intuicyjnych metod oceny pojedynczej cechy pod względem jej wartości dyskryminacyjnej w kontekście realizowanego problemu klasyfikacji binarnej. Cecha ma zdolność rozróżnienia pomiędzy dwoma klasami, jeśli centrum (średnia arytmetyczna) c_1 wartości cechy dla klasy pierwszej jest maksymalnie oddalona od centrum c_2 wartości cechy dla klasy drugiej, oraz gdy odchylenia standardowe cechy należące do poszczególnych klas σ_1 oraz σ_2 są niewielkie, czyli punkty skupione są wokół centrów klas. Oba warunki można ująć w postaci współczynnika istotności Fishera zdefiniowanego w postaci [4]

$$S_{12}(f) = \frac{|c_1 - c_2|}{\sigma_1 + \sigma_2} \quad (4.5)$$

Duża wartość miary $S_{12}(f)$ oznacza dobrą zdolność dyskryminacyjną cechy f pomiędzy klasami, a mała oznacza że dane należące do obu klas są rozproszone i potencjalnie przemieszane ze sobą, co dyskwalifikuje ją jako cechę diagnostyczną. Należy zauważyć, że model oceny cech według kryterium Fishera nie uwzględnia synergii cech w zakresie dyskryminacji, czyli ich współdziałania w zakresie możliwości rozróżnienia pomiędzy wzorcami należącymi do dwóch różnych klas, co może prowadzić do sytuacji przedstawionych na rys. 4.4.



Rys 4.4. Współczynnik istotności Fishera dla cech x_1 oraz x_2 : a) $F_{12}(x_1)=0.98$, $F_{12}(x_2)=0.91$, b) $F_{12}(x_1)=0$, $F_{12}(x_2)=0$

4.4.4. Selekcja według korelacji z klasą

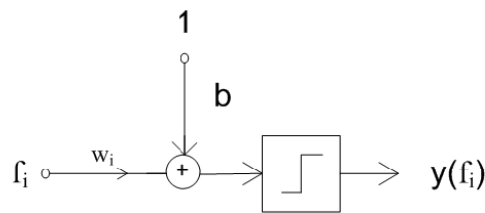
W tej metodzie kolejność cech jest ustawiana według wartości korelacji cechy f z klasą. Korelacja ta jest definiowana wzorem [4,20]

$$S(f) = \frac{\sum_{k=1}^K P_k (m_k - m)^2}{\sigma^2 \sum_{k=1}^K P_k (1 - P_k)} \quad (4.6)$$

We wzorze tym m jest średnią wartością cechy f dla wszystkich danych, m_k średnią wartością cechy dla k -tej klasy, σ^2 jest wariancją cechy f , a P_k prawdopodobieństwem *a priori* wystąpienia k -tej klasy w zbiorze danych. Im większa wartość tej miary tym większa wartość dyskryminacyjna cechy f przy rozróżnianiu klas. Metoda ta ma te same wady co metoda Fishera, gdyż nie uwzględnia kolektywnej współpracy poszczególnych cech na rzecz klasyfikacji obiektów należących do różnych klas.

4.4.3. Zastosowanie jednowarstwowej sieci SVM

W tej metodzie trenuje się sieć jednowarstwową SVM (liniową) dla odwzorowania danych uczących charakteryzujących modelowany proces przy zastosowaniu tylko jednej cechy na raz (rys. 4.5) [20]



Rys. 4.5. Sieć jednowarstwowa SVM do oceny dyskryminacyjnej cechy f_i

Trenuje się tyle sieci SVM ile jest cech, notując za każdym razem błędy odwzorowania analizowanego procesu. Wartość diagnostyczna danej cechy jest mierzona poprzez dokładność uzyskanego odwzorowania matematycznego procesu. Im większa dokładność (mniejszy błąd odwzorowania procesu przy użyciu pojedynczej cechy) tym większa jest zdolność dyskryminacyjna rozważanej cechy diagnostycznej.

4.4.4. Selekcja przy zastosowaniu wielowejsciowej liniowej sieci SVM

W metodzie tej trenowana jest liniowa sieć SVM przy użyciu pełnego zbioru potencjalnych cech diagnostycznych. Każda cecha diagnostyczna stanowi zatem jeden z sygnałów wejściowych dla takiej sieci. Sygnał wyjściowy liniowej sieci SVM określony jest zależnością

$$y(\mathbf{x}) = \text{sgn}(\mathbf{w}^T \mathbf{x} + b) \quad (4.7)$$

w której $\mathbf{w}$ jest wektorem wag opisującym cechy diagnostyczne, natomiast b jest wartością polaryzacji. Cecha odpowiadająca wadze o najniższej wartości bezwzględnej wytrenowanej sieci SVM ma najmniejszy wpływ na wynik klasyfikacji.

4.4.5. Selekcja przy zastosowaniu nieliniowego jądra

W tym podejściu podobieństwo dwu wzorców scharakteryzowanych przez cechy diagnostyczne zawarte w wektorach $\mathbf{x}_i$ oraz $\mathbf{x}_j$ mierzy się przy wykorzystaniu funkcji jądra $K(\mathbf{x}_i, \mathbf{x}_j)$. Przy oznaczeniu miary podobieństwa dwu wektorów poprzez $D(\mathbf{x}_i, \mathbf{x}_j)$ mamy

$$D(\mathbf{x}_i, \mathbf{x}_j) = K(\mathbf{x}_i, \mathbf{x}_j) \quad (4.8)$$

Wybór jądra może być w zasadzie dowolny, choć jednym z najlepszych kandydatów jest jądro gaussowskie, z uwagi na to, że generuje wartości znormalizowane w przedziale $[0, 1]$, jest funkcją monotoniczną a przy tym spełnia warunki niezmienniczości względem przesunięcia wektorów, gdyż $K(\mathbf{x}_i + \mathbf{x}_0, \mathbf{x}_j + \mathbf{x}_0) = K(\mathbf{x}_i, \mathbf{x}_j)$.

Rozważmy problem klasyfikacji binarnej o danych należących do dwu klas. Załóżmy, że klasa pierwsza reprezentowana jest przez n_1 , a klasa druga przez n_2 wzorców. Średnia miara podobieństwa danych należących do k -tej klasy ($k=1, 2$) może być wówczas określona wzorem

$$m_k = \frac{1}{n_k(n_k - 1)} \sum_{i,j=1}^{n_1, n_2} D(x_i^{(k)}, x_j^{(k)}) \quad (4.9)$$

Dobra cecha powinna charakteryzować się jak największym podobieństwem dla wzorców należących do tej samej klasy. Oznacza to, że wartości wskaźników m_1 i m_2 określonych wzorem (4.9) powinny być jak największe.

Z kolei podobieństwo cechy opisującej wzorce należące do dwu różnych klas powinno być jak najmniejsze. Średnie podobieństwo międzyklasowe opisać można wzorem

$$d_{12} = \frac{1}{n_1 n_2} \sum_{i,j=1}^{n_1, n_2} D(x_i^{(1)}, x_j^{(2)}) \quad (4.10)$$

We wzorze tym podobieństwo obliczane jest dla wzorców reprezentujących różne klasy. Na podstawie miar m_1 , m_2 oraz d_{12} można zdefiniować miarę separowalności dwu klas przez cechę f w postaci

$$S_{12}(f) = \frac{m_1 + m_2}{d_{12}} \quad (4.11)$$

Duża wartość $S_{12}(f)$ oznacza dobre własności cechy f w dyskryminacji między dwoma klasami. Mała wartość oznacza mieszanie się wartości cechy dla obu klas i złe zdolności dyskryminacyjne. Sposób zastosowania tej miary w selekcji cech może przybrać różną formę. Najlepiej jest badać jakość każdej cechy w pełnym środowisku wszystkich cech. W związku z powyższym przeprowadza się selekcję w sposób rekurencyjny eliminując kolejno ze zbioru pełnego cechę pojedynczą, za każdym razem zmniejszając wymiar wektora cech o jedną cechę eliminowaną ze zbioru. Dla każdego przypadku obliczana jest wartość miary $S_{12}(f)$ po wyeliminowaniu cechy i -tej

$$S_{12}^{(-i)}(f) = \frac{m_1^{(-i)} + m_2^{(-i)}}{d_{12}^{(-i)}} \quad (4.12)$$

Największe wartości współczynnika dyskryminacyjnego $S_{12}^{(-i)}(f)$ wskazują na cechy, których usunięcie zwiększa wartość dyskryminacyjną pozostałego wektora cech. W ten sposób selekcjonujemy cechy które powinny zostać usunięte ze zbioru. Cechy pozostałe definiują zbiór cech optymalnych.

Istotną sprawą jest dobór hiperparametrów, w szczególności szerokość σ funkcji gaussowskiej zastosowanej w charakterze jądra. Najprostszą stosowaną metodą jest dobór tej wartości według wzoru $\sigma = \sqrt{N/2}$, gdzie N jest wymiarem pełnego wektora cech.

4.4.6. Selekcja przy zastosowaniu transformacji PCA

Zastosowanie transformacji PCA (ang. *Principal Component Analysis*), w selekcji cech wymaga przeprowadzenia pierwszej fazy wstępnego liniowego przetwarzania zbioru deskryptorów według wzoru [19]

$$\mathbf{y} = \mathbf{W}\mathbf{x} \quad (4.13)$$

w którym $\mathbf{W}$ jest macierzą ortogonalną Karhunen-Loewe określaną na podstawie rozkładu macierzy autokowariancji danych według wartości własnych (ang. *eigenvalue decomposition*). Oryginalny zestaw deskryptorów (wektor $\mathbf{x}$) jest transformowany w wektor przekształcony $\mathbf{y}$, w taki sposób, że najważniejsze porcje informacji (cechy diagnostyczne) są skumulowane w kilku pierwszych składnikach wektora $\mathbf{y}$ (szeregowane nierosnąco według ich wartości wariancji). Składowe wektora $\mathbf{y}$ są więc traktowane jako potencjalne cechy diagnostyczne.

Faza ostatecznej selekcji cech w oparciu o analizę składowych głównych polega na redukcji wymiarów wektora $\mathbf{y}$ poprzez odrzucenie jego składników o najmniejszej zawartości energetycznej, a więc o najniższej wariancji. Podejście to jest motywowane hipotezą iż cechy o niskiej wartości wariancji stanowią szum pomiarowy i nie wnoszą istotnych zależności na potrzeby realizowanego modelu klasyfikacji statystycznej a wręcz mogą go zakłócić.

4.4.7. Selekcja przy zastosowaniu LDA

Ta metoda realizuje proces wstępnej selekcji cech w oparciu o transformację LDA (ang. *Linear Discriminant Analysis*), zwaną liniową dyskryminacją Fishera. W odróżnieniu od omówionej w poprzednim podrozdziale metody składników głównych, proces liniowej dyskryminacji Fishera uwzględnia etykiety przynależności do klas przekształcanego zbioru danych. Stanowi zatem nadzorowaną metodę ekstrakcji cech, która w odróżnieniu od stosowanego w PCA kryterium maksymalizacji wariancji kolejno ekstrahowanych cech, maksymalizuje współczynnik istotności Fishera zdefiniowany wzorem (4.5). O ile w metodzie selekcji Fishera współczynnik istotności wyznaczany był dla każdej z cech osobno, w fazie ekstrakcji cech selekcji LDA poszukiwana jest taka liniowa transformacja oryginalnego zestawu atrybutów, która zmaksymalizuje współczynnik istotności Fishera wzdłuż kolejnych kierunków przyjętego układu współrzędnych. Wektory transformaty LDA dla kolejnych kierunków ortogonalnych odnajdywane są przez maksymalizację funkcji $J(\mathbf{w})$ względem wektora wag [8,9],

$$J(\mathbf{w}) = \frac{\mathbf{w}^T \mathbf{S}_b \mathbf{w}}{\mathbf{w}^T \mathbf{S}_w \mathbf{w}} \quad (4.14)$$

gdzie $\mathbf{S}_b$ jest macierzą rozproszenia międzyklasowego natomiast $\mathbf{S}_w$ macierzą rozproszenia wewnątrzklasowego. Transformata LDA odnajdywana jest poprzez rozwiązanie uogólnionego problemu własnego (*eigen-decomposition*) dla wymienionych macierzy. Eksploatowana w tym przypadku heurystyka opiera się na założeniu, że cech o dużej zdolności dyskryminacyjnej należy szukać wśród tych cech, które stosunkowo dobrze dyskryminują liniowo analizowany zbiór danych. W kolejnych iteracjach algorytmu bieżący

zbiór cech powiększany jest zatem o pierwszą niewykorzystaną cechę przetransformowanego zbioru danych, a więc cechę o największym współczynniku istotności Fishera.

4.4.8. Selekcja przy zastosowaniu ICA

Selekcja ICA (ang. *Independent Component Analysis*) stanowi algorytm selekcji cech, oparty o analizę składowych niezależnych. ICA, podobnie jak PCA realizuje liniową transformację zbioru cech oryginalnych, dąży jednak do maksymalizacji statystycznej niezależności ekstrahowanego zbioru atrybutów w oparciu o statystyki wyższego rzędu.

U podstaw koncepcji ekstrakcji cech na bazie analizy składowych niezależnych leży hipoteza, iż stan analizowanego procesu jest w pełni determinowany przez szereg niezależnych czynników (*independent components*), których bezpośrednia obserwacja ze względów praktycznych nie jest zazwyczaj możliwa. Przyjmuje się, że wyniki pomiarów stanowią liniową kombinację wektora czynników niezależnych, o nieznanych bliżej wagach. W ramach techniki ICA przyjmowany jest zatem liniowy model opisujący relacje pomiędzy wielkościami mierzonymi a wartościami składowych niezależnych, który w formie macierzowej można wyrazić wzorem $\mathbf{X}(t) = \mathbf{W}\mathbf{Y}(t)$, gdzie $\mathbf{X}$ jest wektorem obserwacji, $\mathbf{Y}$ stanowi szukany wektor składowych niezależnych a $\mathbf{W}$ jest nieznaną macierzą mieszającą, wyrażającą wpływ, jaki poszczególne składowe mają na rezultat prowadzonych obserwacji. Tak sformułowany problem nazywany jest również ślepą separacją sygnałów (*blind signal separation*) a dekompozycja według składowych niezależnych stanowi jedną z bardziej znanych metod stosowanych dla jego rozwiązania[2].

W zastosowaniu do selekcji cech diagnostycznych przyjmuje się, że kolejne wartości indeksu t reprezentują numer przeprowadzonej serii m równoległych pomiarów analizowanego procesu, a więc numer wzorca. Wektory $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_m$ stanowią wówczas kolejne cechy diagnostyczne i nie muszą opisywać tej samej wielkości fizycznej. Mechanizm ekstrakcji cech ICA bazuje na centralnym twierdzeniu granicznym, mówiącym, że rozkład sumy odpowiednio wystandaryzowanych, niezależnych zmiennych losowych zbiega do rozkładu normalnego dla wzrastającej liczby składników. Ważoną sumę dwóch niezależnych zmiennych losowych cechuje zatem rozkład bliższy do normalnego, aniżeli każdy z oryginalnych komponentów tej sumy. Dekompozycja ICA w zastosowaniu do selekcji cech polega zatem na takim doborze macierzy mieszającej $\mathbf{W}$, aby czynniki składowe $\mathbf{Y}(t) = \mathbf{W}^{-1}\mathbf{X}(t)$ miały rozkład jak najmniej zbliżony do normalnego. Monografia [2] jest doskonałym opracowaniem dotyczącym metod ślepej separacji danych. Po przeprowadzeniu tego typu separacji można stwierdzić z dużym prawdopodobieństwem że odnalezione komponenty będą statystycznie niezależne. Do najczęściej eksploatowanych heurystyk wyrażających miarę *normalności* analizowanych zmiennych losowych należą: kurtoza, negentropia, informacja wzajemna oraz estymator największej wiarygodności [2].

4.4.9. Zastosowanie algorytmu genetycznego w selekcji cech

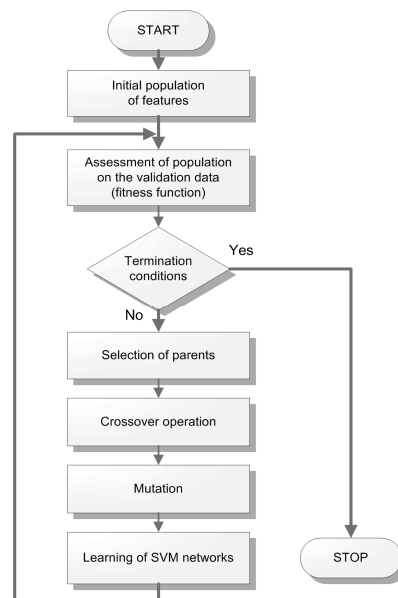
Selekcja optymalnego podzbioru cech diagnostycznych przy zastosowaniu algorytmu genetycznego stanowi alternatywną technikę przeszukiwania przestrzeni cech i zaliczana jest do metod globalnych optymalizacji [1,17], gdyż charakteryzuje się równoległym przetwarzaniem wielu punktów przestrzeni rozwiązań, pozbawionym tendencji do wpadania w lokalne ekstrema optymalizacji. Idea algorytmu genetycznego zaczerpnięta została z biologicznego zjawiska naturalnej ewolucji organizmów wskutek ich oddziaływania z otoczeniem. Na grunt matematyczny przeniesione zostały biologiczne mechanizmy ewolucji prowadzące do sukcesywnego dostosowywania się osobników do warunków zewnętrznych: selekcja, reprodukcja (krzyżowanie) oraz mutacja.

Przy zastosowaniu algorytmu do selekcji cech poszczególnym deskryptorom numerycznym (potencjalnym cechom diagnostycznym) przyporządkowuje się wartości

binarne, stanowiące wektor (chromosom) poddany operatorom genetycznym selekcji, krzyżowania i mutacji [24]. Potencjalne cechy diagnostyczne są więc kodowane przez ciąg binarny o liczbie elementów równej liczbie wszystkich deskryptorów numerycznych wygenerowanych w procesie przetwarzania obrazu. W ciągu tym, logiczna wartość bitu określa czy dana cecha zostanie uwzględniona (wartość bitu 1) czy też nie (wartość zerowa bitu). W realizowanej strategii funkcja przystosowania f_{fit} jest bezpośrednio definiowana przez błąd walidacji krzyżowej osiągnięty przez nieliniowy klasyfikator SVM na wyselekcjonowanym zbiorze cech, z poprawką promującą zbiory o mniejszej liczności. Funkcja ta wyrażona jest zależnością

$$f_{fit} = -\left(error + \alpha \frac{N_c}{N} \right) \quad (4.15)$$

gdzie *error* to błąd walidacji krzyżowej klasyfikatora definiowany jako norma euklidesowa błędów dopasowania wyników aktualnych do wartości zadanych, N_c - liczba cech bieżącego rozwiązania, N - liczebność oryginalnego zbioru potencjalnych cech, natomiast α określa współczynnik kary za dużą liczbę wykorzystanych cech. Zwiększanie współczynnika kary orientuje algorytm w kierunku rozwiązań o małej liczbie cech kosztem osiągniętej dokładności, co zabezpiecza algorytm przed nadmiernym dopasowaniem do zbioru danych trenujących i stanowi narzędzie pozwalające sterować zdolnością generalizacji klasyfikatora.



Rys. 4.6. Schemat blokowy selekcji cech przy użyciu algorytmu genetycznego

Na rys. 4.6 przedstawiono schemat blokowy selekcji cech przy zastosowaniu algorytmu genetycznego. Zakończenie działania procedury następuje w chwili osiągnięcia zakładanej wartości funkcji przystosowania lub jeśli w kilkunastu kolejnych pokoleniach (iteracjach) nie zostanie odnotowana istotna poprawa funkcji przystosowania.

4.5. Klasyfikatory neuronowe

W kolejnym etapie przetwarzania obrazów zwłaszcza przy rozpoznaniu i klasyfikacji istotną rolę pełnią klasyfikatory. W chwili obecnej najbardziej popularne są klasyfikatory neuronowe, stanowiące zwykle nieliniową strukturę zdolną do rozdzielenia wzorców nieseparowanych liniowo. Do najbardziej skutecznych należą sieć perceptronu wielowarstwowego (MLP), sieć radialna (RBF) czy sieć SVM [9,19,23].

4.5.1. Klasyfikator SVM

Spośród wymienionych rozwiązań za najbardziej efektywny uważa się klasyfikator SVM stanowiący uniwersalne rozwiązanie stosujące różne funkcje jądra, w tym liniowe, gaussowskie, wielomianowe czy sigmoidalne. Istotną nowością w stosunku do klasycznych rozwiązań sieci neuronowych jest zastosowanie w uczeniu maksymalizacji marginesu separacji (tolerancji) między dwoma klasami, czyniące tę sieć stosunkowo mało wrażliwą na zaburzenia danych wejściowych. W rozdziale tym ograniczymy się wyłącznie do sieci SVM, odwołując się w przypadku pozostałych sieci do pozycji [9]. Sieć SVM jest w ogólności siecią liniową działającą na danych zrzutowanych nieliniowo przy użyciu funkcji wektorowej $\phi(\mathbf{x})$. Funkcja ta odwzorowuje oryginalną przestrzeni N -wymiarową $\mathbf{x}$ w przestrzeń K -wymiarową zdefiniowaną przez funkcję $\phi(\mathbf{x})$. Równanie hiperpłaszczyzny separującej dwie klasy przyjmuje wówczas postać [23,25]

$$y(\mathbf{x}) = \mathbf{w}^T \phi(\mathbf{x}) + b_0 = \sum_{j=1}^K w_j \phi_j(\mathbf{x}) + b_0 = 0, \quad (4.16)$$

gdzie $\phi(\mathbf{x}) = [\phi_1(\mathbf{x}), \dots, \phi_K(\mathbf{x})]^T$ a $\mathbf{w}$ jest wektorem wag $\mathbf{w} = [w_1, \dots, w_K]^T$. W uczeniu sieci SVM należy tak ukształtować hiperpłaszczyznę separującą (wagi $\mathbf{w}$ oraz b_0), aby uzyskać maksymalizację marginesu separacji dwu klas przy minimalnej liczbie błędów klasyfikatora dla danych uczących. Matematycznie problem pierwotny uczenia sieci SVM sprowadza się do minimalizacji funkcji celu $\phi(\mathbf{w}, \xi)$ [20,23,25]

$$\min \phi(\mathbf{w}, \xi) = \frac{1}{2} \mathbf{w}^T \mathbf{w} + C \sum_{i=1}^p \xi_i \quad (4.17)$$

przy ograniczeniach funkcyjnych dla każdej pary danych uczących ($i = 1, 2, \dots, p$)

$$d_i (\mathbf{w}^T \phi(\mathbf{x}_i) + b_0) \geq 1 - \xi_i, \quad (4.18a)$$

$$\xi_i \geq 0, \quad (4.18b)$$

gdzie $C > 0$ jest stałą regularyzacji dobieraną przez użytkownika, $\xi_i \geq 0$ jest nieujemną zmienną dopełniającą podlegającą optymalizacji, p – liczbą danych par uczących $(\mathbf{x}_i, d_i)$.

Definiując funkcję Lagrange’a, zadanie pierwotne uczenia upraszcza się do programowania kwadratowego względem mnożników Lagrange’a α_i . Ponadto wszystkie operacje na etapie uczenia i testowania sieci wykonuje się przy użyciu funkcji jądra zdefiniowanej w postaci $K(\mathbf{x}_i, \mathbf{x}) = \phi^T(\mathbf{x}_i) \phi(\mathbf{x})$, spełniającej warunki Mercera [25]. Najczęściej stosowane w praktyce funkcje jądra to

- funkcja wielomianowa $K(\mathbf{x}_i, \mathbf{x}) = (\mathbf{x}_i^T \mathbf{x} + \gamma)^p$
- funkcja gaussowska $K(\mathbf{x}_i, \mathbf{x}) = \exp(-\gamma \|\mathbf{x} - \mathbf{x}_i\|^2)$
- funkcja sigmoidalna $K(\mathbf{x}_i, \mathbf{x}) = \tanh(\beta \mathbf{x}_i^T \mathbf{x} + \theta)$

W wyniku zastosowania funkcji Lagrange’a problem pierwotny uczenia sieci SVM opisany równaniami (4.12) i (4.13), jest transformowany do problemu dualnego w postaci zadania programowania kwadratowego, zdefiniowanego w postaci [23]

$$\max Q(\alpha) = \sum_{i=1}^p \alpha_i - \frac{1}{2} \sum_{i=1}^p \sum_{j=1}^p \alpha_i \alpha_j d_i d_j K(\mathbf{x}_i, \mathbf{x}_j) \quad (4.19)$$

przy ograniczeniach

$$\sum_{i=1}^p \alpha_i d_i = 0, \quad (4.20a)$$

$$0 \leq \alpha_i \leq C. \quad (4.20b)$$

Rozwiązanie problemu dualnego pozwala obliczyć optymalne wartości mnożników Lagrange’a, na podstawie których wyznacza się optymalny wektor wag sieci

$\mathbf{w} = \sum_{i=1}^{N_{sv}} \alpha_i d_i \phi(\mathbf{x}_i)$. W równaniu tym sumowanie odbywa się tylko względem wskaźników i dla których mnożniki Lagrange’a przyjmują wartości niezerowe (wektory $\mathbf{x}_i$ odpowiadające im nazywane są wektorami podtrzymującymi – support vectors). W tworzeniu wag sieci SVM biorą więc udział tylko te wektory $\mathbf{x}_i$, dla których odpowiadające im mnożniki Lagrange’a są niezerowe. Sygnał wyjściowy $y(\mathbf{x})$ sieci w trybie testowania jest zdefiniowany w postaci

$$y(\mathbf{x}) = \sum_{i=1}^{N_{sv}} \alpha_i d_i K(\mathbf{x}_i, \mathbf{x}) + b \quad (4.21)$$

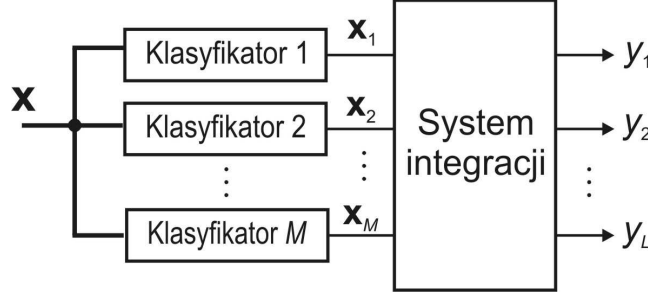
i zależy wyłącznie od wektorów podtrzymujących $\mathbf{x}_i$ oraz przyjętej postaci funkcji jądra $K(\mathbf{x}_i, \mathbf{x})$. Szybkie i skuteczne metody algorytmiczne rozwiązania problemu dualnego są szeroko znane i przedstawione między innymi w pracach [23].

Z uwagi na to, że jedna sieć SVM jest w stanie odróżniać tylko dwie klasy od siebie, dla problemu klasyfikacji wieloklasowej należy skonstruować wiele takich sieci i zastosować odpowiednią metodę wyłaniania zwycięscy. Jednym z najczęściej stosowanych rozwiązań jest metoda „jeden przeciw jednemu” [20,23]. W metodzie tej dla zbioru M rozpoznawanych klas konstruuje się $\frac{M(M-1)}{2}$ klasyfikatorów typu SVM rozróżniających za każdym razem 2 klasy danych ze zbioru uczącego, kolejno parowanych ze sobą. Po wytrenowaniu wszystkich sieci przystępuje się do właściwej klasyfikacji przy założeniu konkretnej wartości wektora wejściowego $\mathbf{x}$. Jeśli sieć SVM rozróżniająca dwie klasy (i -tą od j -tej) wskazuje na i -tą klasę, należy zwiększyć o jeden sumę wskazań do tej klasy. W przeciwnym przypadku należy zwiększyć o jeden sumę wskazań do klasy j -tej. Klasa o największej liczbie zwycięstw wśród $M(M-1)/2$ sieci jest uważana za zwycięską (wektor $\mathbf{x}$ zalicza się ostatecznie do tej klasy).

Należy zaznaczyć, że w procesie projektowania klasyfikatora SVM ważny jest prawidłowy dobór stałej regularyzacji C . Jej rolą jest kontrola relacji pomiędzy złożonością sieci (liczbą wektorów podtrzymujących) i nieseparowalnością danych użytych w uczeniu. Niska wartość C oznacza akceptację większej liczby nieseparowalnych danych uczących. Duża wartość C skutkuje niższą liczbą błędów klasyfikacji danych uczących, ale też bardziej złożoną strukturą sieci (większa liczba wektorów podtrzymujących) co prowadzić może do pogorszenia generalizacji. Optymalna wartość C może być ustalana oddzielnie dla każdej rozpoznawanej pary klas, na przykład poprzez wykonanie serii eksperymentów z udziałem danych uczących i weryfikujących. W przypadku zastosowania nieliniowej funkcji jądra, na przykład gaussowskiej, optymalizacja parametrów C i σ powinna przebiegać równocześnie. Za optymalną kombinację tych parametrów uważa się taką, dla której błąd testowania na danych weryfikujących jest najmniejszy.

4.5.4. Zespół klasyfikatorów

W przypadku trudnych zadań klasyfikacyjnych, do których z reguły należą zadania przetwarzania obrazów medycznych, wynik klasyfikacji końcowej można poprawić poprzez zastosowanie równoległe wielu klasyfikatorów na raz. Każdy z nich wykonuje to samo zadanie, generując własny wynik, zwykle różniący się dla każdego zastosowanego klasyfikatora. Uwzględnienie ich wszystkich na raz poprzez odpowiednią integrację daje możliwość poprawy dokładności rozpoznania [13]. Warunkiem jest niezależność działania poszczególnych klasyfikatorów jak również porównywalność dokładności każdego z nich. W wyniku integracji spodziewamy się dokładności lepszej od najlepszego z nich. Na rys. 4.7 przedstawiony jest schemat ogólny zespołu klasyfikatorów.



Rys. 4.7. Schemat zespołu klasyfikatorów

Zespół złożony jest z wielu klasyfikatorów. Każdy z nich niezależnie od siebie dokonuje rozpoznania i klasyfikacji wzorców reprezentowanych przez wektor $\mathbf{x}$ podając swój własny werdykt do integratora. Zadaniem integratora jest wypracowanie końcowej decyzji klasyfikacyjnej, przypisującej wzorec do określonej klasy. Zastosować można różne strategie integracji wyników. Najprostszą jest głosowanie większościowe, w którym za zwycięską klasę uważa się tę na którą wskazało najwięcej klasyfikatorów. Tego typu rozwiązanie nie gwarantuje wyników lepszych od najlepszego klasyfikatora indywidualnego, gdyż klasyfikator najgorszy, znacznie odbiegający jakością od najlepszego, może pogorszyć wynik końcowy. Można tego uniknąć stosując głosowanie większościowe ważone, w którym głos każdego klasyfikatora brany jest z wagą uwzględniającą dokładność działania poszczególnych klasyfikatorów. Klasyfikatorowi najlepszemu przypisuje się wówczas wagę najwyższą, a najgorszemu najniższą. Określanie wag może mieć różne rozwiązanie z których najprostsze może być zapisane w postaci

$$w_{ij} = \frac{\eta_{ij}^m}{\sum_{k=1}^M \eta_{ik}^m} \quad (4.22)$$

gdzie η_{ik} oznacza dokładność k -tego klasyfikatora przy rozpoznaniu i -tej klasy (na danych uczących), m – wykładnik uwypuklający wpływ różnych klasyfikatorów na wynik rozpoznania, a M liczbę klasyfikatorów w zespole.

Innym rozwiązaniem integracji może być zastosowanie dodatkowego klasyfikatora, którego sygnałami wejściowymi są wyniki klasyfikatorów indywidualnych. Dobre wyniki integracji klasyfikatorów uzyskuje się również przy zastosowaniu transformacji PCA lub zmodyfikowanej metody Bayesa [13]

Zgodnie z metodą Bayesa dla zespołu klasyfikatorów określa się prawdopodobieństwo $\mu_i(\mathbf{x})$ wskazania na i -tą klasę na podstawie aktualnych wskazań pojedynczych klasyfikatorów oraz ich znanych rezultatów klasyfikacji na zbiorze danych uczących. Jest ono określone wzorem

$$\mu_i(\mathbf{x}) = \prod_{j=1}^M \frac{cm_{i,s_j}^{(j)} + 1/N}{n_i + 1} \quad (4.23)$$

w którym n_i jest liczbą danych uczących należących do klasy i natomiast $cm_{i,s_j}^{(j)}$ oznacza wartość elementu (i,s_j) -tego macierzy niezgodności j -tego klasyfikatora wskazującego na klasę s_j zamiast na i -tą.

W przypadku zastosowania transformacji PCA w roli integratora dokonuje się w pierwszej kolejności redukcji wymiaru danych [13]. Rozpatrzmy dane uczące zespół w postaci macierzy $\mathbf{A}$ o wymiarze $p \times (c \cdot M)$, gdzie p jest liczbą danych uczących, M – liczbą klasyfikatorów w zespole, c – liczbą klas podlegających rozpoznaniu. Każdy klasyfikator generuje na wyjściu wektor binarny $\mathbf{z}$ zawierający 1 na pozycji rozpoznanej klasy i zero poza

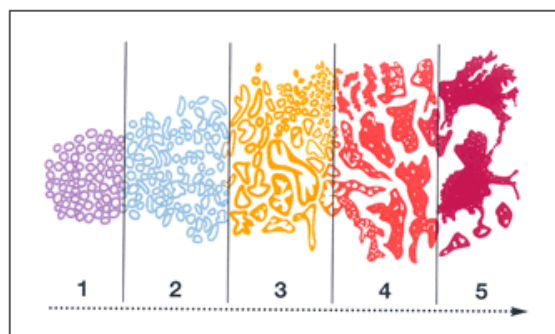
nią. Każdy wiersz macierzy $\mathbf{A}$ powstaje jako złączenie sygnałów wyjściowych M klasyfikatorów dla aktualnego wektora wejściowego poddanego analizie. W pierwszym kroku transformuje się wszystkie wiersze macierzy $\mathbf{A}$ w przestrzeń o zredukowanym wymiarze K (w stosunku do wartości cM) przy zastosowaniu liniowej transformacji PCA. W ten sposób każdy wiersz (wektor) $\mathbf{z}$ macierzy $\mathbf{A}$ jest transformowany w wektor $\mathbf{y}=\mathbf{Wz}$ o wymiarze K zredukowanym w sposób optymalny, przy zachowaniu maksimum informacji oryginalnej. Zbiór wektorów $\mathbf{y}$ jest w drugim kroku algorytmu użyty jako zbiór uczący klasyfikatora drugiego stopnia, pracującego w roli układu decydującego o końcowym wyniku rozpoznania.

4.6 Przykłady przetwarzania obrazów medycznych

4.6.1. Segmentacja i parametryzacja struktur histologicznych dla oceny skali Gleasona w prostaty

Rak stercza (prostaty) zajmuje u mężczyzn trzecie miejsce w zestawieniu zapadalności i czwarte miejsce w zestawieniu zgonów z powodu nowotworu (6-7%). Do oceny stopnia zaawansowania raka prostaty stosuje się powszechnie skalę Gleasona [5,6]. W skali tej mikroskopowemu obrazowi tkanki prostaty przypisuje się liczby od 1 do 5. Liczba 1 odpowiada komórkom nowotworowym najlepiej zróżnicowanym przypominającym budową tkanki zdrowe gruczołu krokowego. Stopień 5 odpowiada komórkom nisko zróżnicowanym (najbardziej złośliwym), a więc tkance chorej wykazującej typowe cechy komórek rakowych, tworzącej z sąsiednimi tkankami tzw. nacieki nowotworowe. Pozostałe cyfry: 2, 3 i 4 oznaczają odpowiednie zaawansowanie stadiów pośrednich między tymi stanami skrajnymi.

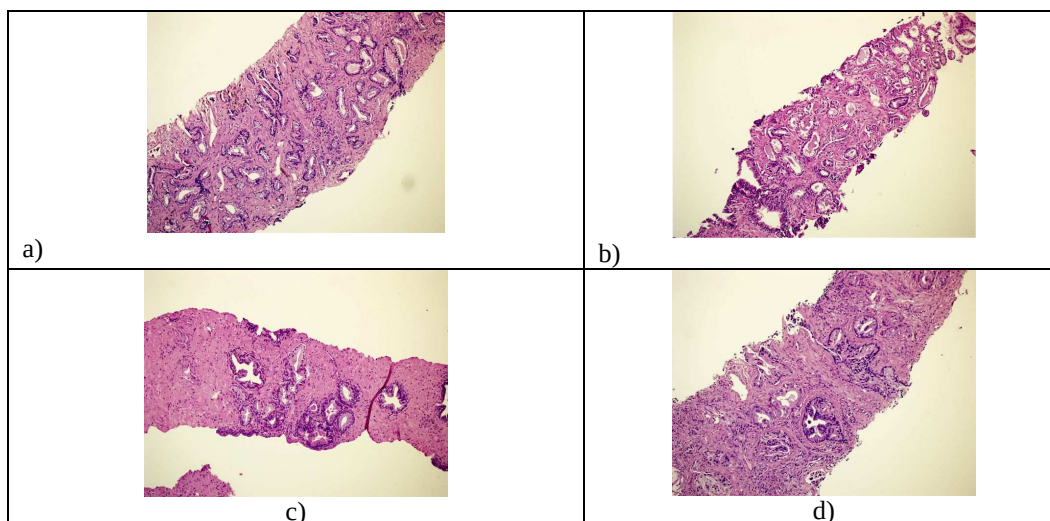
Na rys. 4.8 przedstawiono poglądowo interpretację skali Gleasona od 1 do 5, powiązaną z wyglądem komórek i struktur tworzących tkankę prostaty. W zależności od skali widać wyraźne zróżnicowanie struktur histologicznych tworzących tkankę.



Rys.4.8. Powiązanie wyglądu tkanki prostaty z liczbami od 1 do 5 stosowanymi w skali Gleasona

W raku prostaty mogą wystąpić obszary o różnych stopniach zaawansowania zmian chorobowych. Dla uzyskania bardziej obiektywnego wyniku ocenę stopniową stosuje się do dwu typów histoarchitektonicznych przeważających pod względem objętości obszaru zmienionego. Sumując wyniki z obu obszarów otrzymuje się wynik sumacyjny w skali Gleasona. Skala Gleasona zawiera się wówczas między 2 a 10. Wyższy wynik oznacza większe ryzyko szybszego wzrostu i rozprzestrzeniania się raka prostaty. W praktyce stopnie 2-4 odpowiadają komórkom dobrze zróżnicowanym, (najmniej złośliwym), 5-7 średnio zróżnicowanym, 8-10 komórkom źle zróżnicowanym (najbardziej złośliwym). Wartości poniżej 7 interpretowane są jako nowotwory mniej agresywne [5,6,].

Z biopsji tkanki tworzony jest preparat mikroskopowy służący patomorfologowi do diagnozy schorzenia. Na rys. 4.9 przedstawiono typowe obrazy mikroskopowe tkanek prostaty odpowiadające różnym wartościom skali Gleasona.



Rys.4.9. Obrazy wycinków stercza z różną oceną w skali Gleasona: a) Gleason 2+2, b) Gleason 3+2, c) Gleason 4+4, d) Gleason 5+5

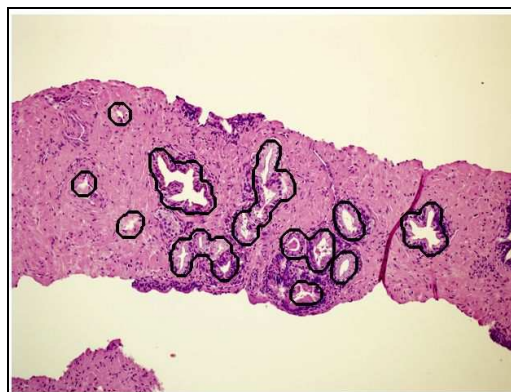
Rys. 4.9a odpowiada wartościom 2+2, rys. 4.9b – wartościom 3+2, rys. 4.9c – wartościom 4+4 i rys. 4.9d wartościom 5+5. Widać wyraźne zróżnicowanie struktur (utkania) komórek w tkance w zależności od stopnia zaawansowania choroby nowotworowej (wartości sumacyjnej skali Gleasona).

Dla rozwiązania problemu automatycznej oceny skali Gleasona na podstawie obrazu należy dokonać ekstrakcji cew gruczołowych tworzących struktury histologiczne z obrazów biopsji, przyporządkować im deskryptory numeryczne, które przy zastosowaniu klasyfikatora umożliwią automatyczne powiązanie ich z wartościami skali Gleasona przypisanymi badanej tkance.

Obraz wycinka stercza zapisany był w powiększeniu 100x i rozdzielczości 621x464 pikseli. Ważnym etapem przetwarzania obrazu jest algorytm segmentacji cew gruczołowych (struktur histologicznych) z obrazów tkanki. Zastosowano rozwiązanie bazujące na wykrywaniu jasnych przestrzeni wewnątrz obrazu tkanki. Wykrycie jasnego pola o odpowiednio dużej powierzchni będzie utożsamiane w algorytmie z wykryciem cewy. Do ekstrakcji zaproponowano algorytm wykorzystujący funkcje morfologii matematycznej [3,22]. Algorytm ten można przedstawić w następujących punktach:

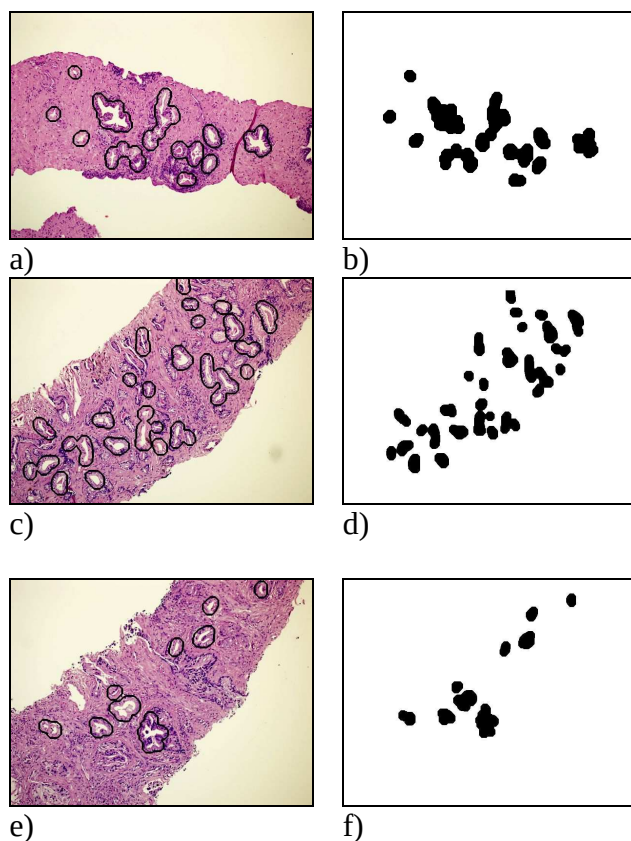
- Wczytanie obrazu wejściowego.
- Rozciągnięcie histogramu dla poprawy kontrastu i sprowadzenie tła obrazu wyjściowego do barwy białej.
- Określenie obrysu tkanki, w tym
 - o binaryzacja obrazu
 - o otwarcie – usunięcie najmniejszych elementów
 - o uzyskanie konturu całej tkanki przy zastosowaniu procesu „kontur przez erozję”.
- Powtórna binaryzacja obrazu przy białym tle.
- Otwarcie – usunięcie najmniejszych obiektów w obrazie.
- Dylatacja – powiększenie obiektów pozostałych w obrazie po operacji otwarcia.
- Zastosowanie gradientu morfologicznego do powstałego obrazu dla uzyskania obrysów poszczególnych cew gruczołowych i konturów całej tkanki.
- Dylatacja obrazu zawierającego obrysy cew dla poszerzenia konturu obrysów.
- Odjęcie obrysu tkanki od obrazu zawierającego obrysy cew dla wyeliminowania obrysu samej tkanki.
- Wypełnienie otworów w obiektach zamkniętych. W wyniku uzyskuje się wypełnione powierzchnie cew gruczołowych.

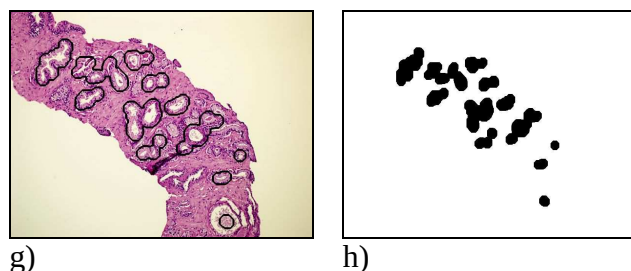
- Wynik podlega następnie etykietowaniu przypisującemu numery kolejnym cewom. Wynikiem segmentacji jest obraz binarny zawierający wydzielone z obrazu kształty cew gruczołowych. Dla oceny jakości segmentacji cew przeprowadza się jeszcze wizualizację porównawczą, polegającą na nałożeniu obrysów cew na obraz wyjściowy. Pozwala to na wzrokową ocenę poprawności przeprowadzonej segmentacji. Wynik takiego porównania dla obrazu z rys. 4.9a przedstawiono na rys. 4.10.



Rys.4.10. Obraz wyjściowy tkanki z rys. 4.9a z nałożonymi wynikowymi obrysami cew gruczołowych

Na rys. 4.11 przedstawiono z lewej strony obrazy oryginalne a z prawej odpowiadające im lokalizacje cew gruczołowych. Jak widać na ilustracji algorytm dobrze radzi sobie z cewami odseparowanymi od siebie (nie stykającymi się). Można zauważyć pewną tendencję algorytmu do łączenia cew położonych bardzo blisko siebie, łącząc je w skupiska. Ten aspekt może mieć pewien wpływ na wynik oceny w skali Gleasona oraz na parametryzację obrazu.





Rys.4.11. Przykładowe wyniki działania algorytmu wydzielającego cewy gruczołowe. Lewa kolumna przedstawia obrazy wyjściowe z obrysem cew dokonany przez działanie algorytmu, prawa kolumna – wyniki działania algorytmu

Istotnymi cechami rozróżniającymi poszczególne oceny w skali Gleasona są parametry geometryczne obrazu związane z rozmieszczeniem, kształtem oraz ilością cew występujących na obrazie tkanki. Spośród najważniejszych parametrów, które należy wziąć pod uwagę w ocenie można wymienić następujące: powierzchnia cew, powierzchnia wypukła, odległość między sąsiednimi cewami, stosunek łącznej powierzchni cew do całej powierzchni obrazu, obwód rzeczywisty oraz obwód wieloboku wypukłego opisanego na obrazie cewy. Na bazie tych parametrów pierwotnych można zdefiniować parametry pochodne mogące mieć wpływ na rozróżnienie poszczególnych stadiów choroby nowotworowej. Przyjmując następujące oznaczenia: d_l – średnica długa, d_s – średnica krótka, A – pole powierzchni cewy, A_c – pole powierzchni wieloboku wypukłego opisanego na cewie, P – obwód rzeczywisty cewy, P_c – obwód wieloboku wypukłego opisanego na cewie zdefiniowano dodatkowo następujące parametry [27].

- Współczynnik eliptyczności

$$F_e = \frac{d_s}{d_l} \quad (4.24)$$

- Współczynnik kolistości

$$F_c = \frac{4\pi A}{P} \quad (4.25)$$

- Współczynnik kolistości wypukłej

$$F_{cc} = \frac{4\pi A_c}{P_c} \quad (4.26)$$

- Współczynnik zwartości

$$F_h = \frac{P^2}{A} \quad (4.27)$$

- Współczynnik zwartości wypukłej

$$F_{ch} = \frac{P_c^2}{A_c} \quad (4.28)$$

- Współczynnik pofałdowania

$$F_{ce} = \frac{P_c}{P} \quad (4.29)$$

- Współczynnik postrzępienia

$$F_r = \frac{A}{A_c} \quad (4.30)$$

Istotną informację diagnostyczną stanowi również odległość między najbliższymi cewami. Do obliczenia tej odległości opracowano następujący algorytm: [11]

1. Określenie największych odległości d_l między końcami cewy (tzw. duża średnica). Linia ta pozwala to na określenie orientacji przestrzennej cewy. Cewy mają zwykle eliptyczny,

- podłużny kształt, dzięki temu przebieg dużej średnicy pozwala zorientować się, jak względem siebie są one położone.
2. Dla każdej średnicy dużej konstruowane są proste prostopadłe, poczynając od początku średnicy, co 20 pikseli, do jej końca. Największa odległość między krańcowymi punktami cewy w tym kierunku określać będziemy jako średnicę małą d_s .
 3. Jeżeli prosta prostopadła przecina dużą średnicę innej cewy, to linię od obrysu jednej do drugiej traktuje się jako linię wchodzącą w skład wyznaczania odległości d_n między nimi. Zazwyczaj takich linii jest wiele, więc wyznacza się wartość średnią i odchylenie standardowe. Odchylenie standardowe pozwala stwierdzić, czy średnia odległość wyznaczona tą metodą jest wartością miarodajną.
 4. Przy pomiarze odległości bierze się pod uwagę tylko najbliższe sobie cewy, przy czym nie może między nimi leżeć żadna inna.

Wykorzystując oprogramowanie zaimplementowane w Matlabie [16] przeanalizowano łącznie 32 obrazy pochodzące od pacjentów z różną oceną skali Gleasona. W analizie uczestniczyło 9 obrazów z oceną 2, 7 obrazów z oceną 3, 12 obrazów z oceną 4 oraz 4 obrazy z oceną 5. Wśród zbioru obrazów poddanych analizie znalazł się tylko jeden z oceną 1. Ilość ta jest niemiernodajna dla wyciągnięcia wniosków, stąd przypadek ten nie został wzięty pod uwagę przy analizie wyników.

Przy prezentacji wyników liczbowych dotyczących parametryzacji zastosujemy dodatkowo następujące oznaczenia:

N – liczba cew gruczołowych w obrazie tkanki

A – powierzchnia cewy

A_s – sumacyjna powierzchnia cew w obrazie tkanki

A_{tk} – powierzchnia całej tkanki na obrazie

d_n – odległość między najbliższymi cewami

d_l – długość dużej średnicy

d_s – długość małej średnicy

Dla uzyskania dokładnej informacji opisującej numerycznie właściwości dyskryminacyjne poszczególnych parametrów charakteryzujących klasy chorobowe zastosowano miarę istotności Fishera określoną wzorem (4.5): Jak podkreślono to przy omawianiu tej metody duża wartość $S_{AB}(f)$ wskazuje na dobre własności dyskryminacyjne danej cechy. Z kolei mała wartość świadczy o tym, że dana cecha słabo rozróżnia daną parę klas, co czyni ją bezużyteczną w rozróżnianiu tych klas ze sobą. Wyniki numeryczne miary Fishera dla poszczególnych deskryptorów, podane w kolejności narastających wartości sumacyjnych dla rozpatrywanych par klas przedstawiono w tabeli 4.1.

Tabela 4.1. Wartości miary Fishera dla poszczególnych deskryptorów

Deskryptor	S23	S24	S25	S34	S35	S45	SUMA
F_{ce}	0,08	0,09	0,10	0,00	0,02	0,02	0,30
F_r	0,00	0,13	0,15	0,11	0,13	0,00	0,52
F_{cc}	0,54	0,04	0,03	0,51	0,49	0,00	1,61
d_l	0,32	0,13	0,33	0,36	0,51	0,17	1,81
P	0,37	0,47	0,59	0,14	0,32	0,18	2,07
A_{tk}	0,24	0,41	0,23	0,65	0,43	0,12	2,08
d_n	0,11	0,46	0,59	0,34	0,48	0,17	2,14
d_s	0,80	0,52	0,63	0,15	0,07	0,07	2,24
F_c	0,91	0,06	0,00	0,64	0,59	0,04	2,24
F_h	0,75	0,45	0,64	0,23	0,09	0,15	2,30
P_c	0,53	0,54	0,75	0,12	0,37	0,21	2,52

A	0,16	0,44	0,68	0,55	0,78	0,22	2,83
A_c	0,16	0,62	0,84	0,72	0,92	0,23	3,48
F_{ch}	1,29	0,49	0,81	0,65	0,37	0,27	3,88
F_e	0,64	0,82	1,31	0,24	0,61	0,29	3,91
N	0,07	0,62	1,27	0,44	0,94	0,60	3,94

Wyniki te w sposób ścisły pokazują, że cechami najsilniej rozróżniającymi wszystkie pary klas są: liczba cew gruczołowych na obrazie, współczynnik eliptyczności, współczynnik zwartości wypukłej oraz powierzchnia wypukła cew. Z kolei najsłabszymi cechami diagnostycznymi pod tym względem są: współczynnik pofałdowania i postrzępienia.

Niezależnie od sumacyjnej wartości dyskryminacyjnej poszczególnych cech można zauważyć, że niektóre cechy mają bardzo dobre własności rozróżnienia jedynie wybranych klas ze sobą. Przykładem takiej cechy jest współczynnik F_c , który pomimo nie najwyższej noty sumacyjnej wykazuje się bardzo wysokim współczynnikiem rozróżnienia klasy drugiej od trzeciej. Oznacza to, że budując automatyczny system klasyfikacyjny warto pokusić się o zróżnicowanie sygnałów wejściowych dla poszczególnych klasyfikatorów pracujących w zespole (np. klasyfikatory SVM pracujące w trybie „jeden przeciw jednemu”).

Dla uzyskania dobrych wyników klasyfikacji kluczowy jest wybór odpowiedniej liczby najbardziej znaczących cech diagnostycznych. Najlepszym sposobem doboru jest tu zastosowanie analizy kowariancyjnej zbioru cech. Na rys. 4.12 przedstawiono rozkład wartości własnych macierzy kowariancji $\mathbf{R}_{xx}$ przy uwzględnieniu pełnego zbioru cech. Jak widać na podstawie tego wykresu trudno wyrokować o rozmiarze optymalnego zbioru, gdyż poszczególne wartości własne maleją monotonicznie.

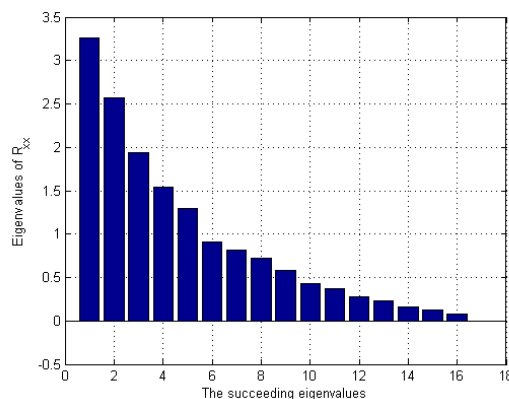
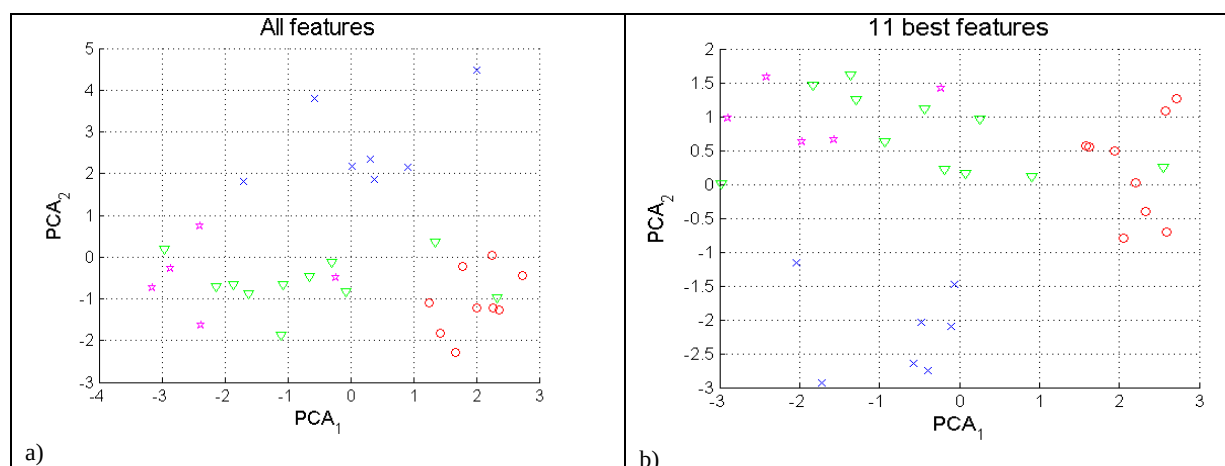


Fig. 4.14. Rozkład wartości własnych macierzy kowariancji dla zbioru wszystkich deskryptorów.

W tym przypadku lepszym rozwiązaniem jest analiza rozkładu danych zrzutowana na dwa największe składniki główne w analizie PCA przy uwzględnieniu różnej liczby najbardziej znaczących deskryptorów. Na rys. 4.13 przedstawiono 2 przykłady takiego rozkładu. Rys. 4.13a dotyczy pełnego zbioru deskryptorów, natomiast rys. 4.13b zbioru ograniczonego do 11 najlepszych deskryptorów. Na wykresach tych znak $\times$ reprezentuje dane dotyczące skali Gleasona 2, kółko – skali 3, znak trójkąta – skali 4 a pięciokąta – skali 5.



Rys. 4.13. Rozkład danych zrzutowanych na 2 najważniejsze składniki główne; a) pełny zbiór deskryptorów, b) zbiór 11 najważniejszych deskryptorów

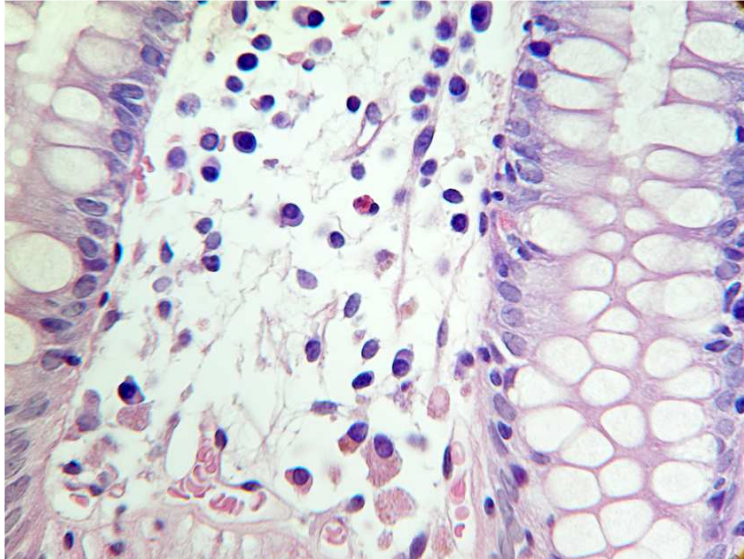
Finalnym krokiem jest przeprowadzenie końcowej klasyfikacji danych przy zastosowaniu klasyfikatora neuronowego. W tej części zadania wykorzystano sieć SVM o jądrze gaussowskim i wartościach hiperparametrów $C=1000$ oraz $\sigma=0.05$. Przy małym dostępnym zbiorze danych (32 próbki) konieczne było zastosowanie metody „leave-one-out”. W implementacji tej metody 31 próbek użyto do uczenia i jedną do testowania, powtarzając eksperyment 32 razy przy różnym doborze próbki testującej. Przy zastosowaniu 17 deskryptorów jako cech diagnostycznych system klasyfikacyjny popełnił 5 błędów (15,6% błędu względnego). Redukcja zbioru cech do 11 najlepszych deskryptorów pozwoliła zmniejszyć liczbę błędnych rozpoznań do jednego (3% błędu względnego).

4.6.4. Rozpoznawanie komórek obronnych w stanach zapalnych jelita

Przedstawienie problemu

Typowym stanem zapalnym jelita jest choroba Leśniowskiego-Crohna zaliczana do grupy nieswoistych zapaleń jelit (IBD). Jest to przewlekły, nieswoisty proces zapalny ściany przewodu pokarmowego powodujący przewlekłą biegunkę, bóle brzucha, niedrożność jelit; guzowaty opór w jamie brzusznej, zmiany okołodbytnicze, gorączkę, spadek masy ciała, zahamowanie wzrostu i osłabienie. Może dotyczyć każdego odcinka jelita, lecz najczęściej lokalizuje się w końcowej części jelita cienkiego oraz początkowej grubego [26].

Ewidentną oznaką choroby jest pojawienie się komórek obronnych w podścielisku komórek. Należą do nich limfocyty, plazmocyty oraz granulocyty kwasochłonne i obojętnochłonne. Dla potwierdzania diagnozy tej choroby i oceny stopnia jej rozwoju konieczne jest rozpoznanie tych komórek i ocena ich liczebności. Ocena taka dokonywana jest na obrazie biopsji komórek barwionych hematoxyliną i eozyną (HE) przy powiększeniu 600x. Przykład obrazu poddanego przetworzeniu przedstawiony jest na rys. 4.14 [11,12]. Przedstawia on typowy wygląd przekroju jelita grubego. Z lewej i prawej strony widoczne są rozległe struktury cew gruczołowych, które na tym etapie przetwarzania powinny być usunięte. Komórki obronne podlegające rozpoznaniu występują w środku obrazu jako pojedyncze ciemne punkty na jasnym tle podścieliska.



Rys. 4.14. Przykład obrazu biopsji okrężnicy w chorobie Leśniowskiego-Crohna

Automatyczne rozpoznanie tych komórek sprowadza się do przeprowadzenia następujących etapów przetwarzania [12]:

- Wydzielenie poszczególnych komórek z obrazu i zapisanie ich w bazie danych.
- Przetworzenie obrazu każdej komórki w taki sposób, aby opisać ją za pomocą zbioru deskryptorów numerycznych.
- Selekcja deskryptorów (cech) najbardziej znaczących, podkreślających różnice między reprezentantami różnych klas przy podobnych wartościach dla tej samej klasy.
- Przeprowadzenie procesu automatycznego rozpoznania (a następnie zliczenia) komórek na podstawie wyselekcjonowanych cech. Ten etap jest przeprowadzany przez wybrane typy klasyfikatorów neuronowych połączonych w zespół ekspertowy wytrenowany na utworzonym zbiorze danych opisanych przez eksperta medycznego.
- Tak wytrenowany klasyfikator może być zastosowany w trybie on-line, dokonując rozpoznania świeżo wydzielonej komórki na podstawie przypisanych jej cech diagnostycznych.

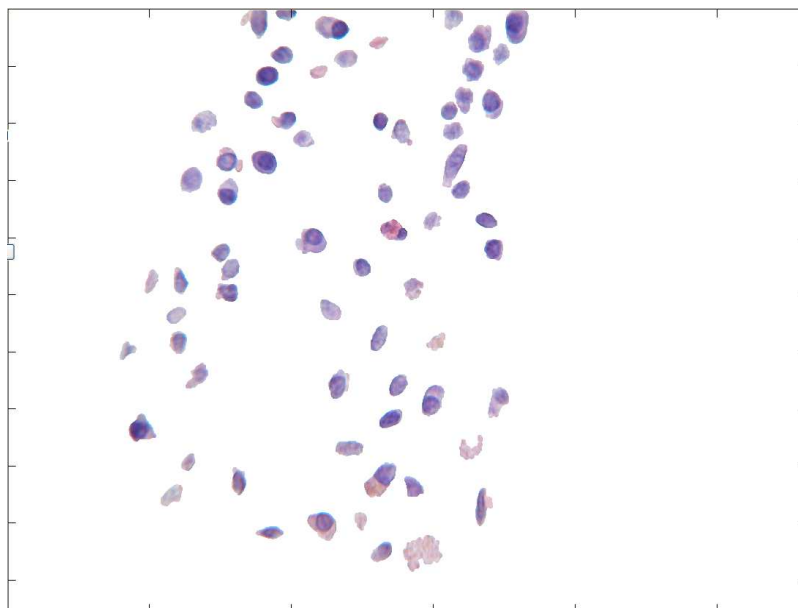
Segmentacja komórek

Pierwszym etapem przetwarzania obrazu jest segmentacja komórek. Zadanie wydzielenia poszczególnych komórek wymaga przeprowadzenia wielu etapów filtracji i segmentacji przy wykorzystaniu operacji morfologicznych. Można je przedstawić jako ciąg operacji opisanych w następujący sposób [11].

- Wczytanie obrazu barwnego $I(x,y,r,g,b)$ biopsji podlegającego przetworzeniu. W oznaczeniu obrazu x,y są współrzędnymi piksela, natomiast r, g, b oznaczają jego intensywność dla kolejnych składowych RGB.
- Zastosowanie algorytmu grupowania K-means dla podziału pikseli na trzy klasy odpowiednio do ich jasności. Piksele najjaśniejsze są przetwarzane na kolor biały i eliminowane z obrazu.
- Usunięcie największych (cewy gruczołowe) i najmniejszych (drobny szum) obiektów utworzonych przez piksele o tej samej skali jasności obrazu. Ten etap przetwarzania został wykonany po przetworzeniu obrazu rzeczywistego na binarny. Eliminuje się obiekty większe niż 3000 pikseli oraz mniejsze niż 40 pikseli.
- Transformacja tak powstałego obrazu poprzez skalę szarości do postaci binarnej.
- Filtracja obrazu przy użyciu operacji morfologicznych (element strukturalny w kształcie

dysku o promieniu 2).

- Generacja mapy odległości pikseli czarnych od białych przy użyciu operacji morfologicznych.
- Zastosowanie algorytmu działów wodnych na mapie odległości. Jako wynik uzyskuje się zwarte obszary reprezentujące poszczególne komórki.
- Wydzielenie obszarów poszczególnych komórek i powrót do naturalnych barw tych obszarów.
- Dodanie konturów do wydzielonych obszarów i zapisanie każdej komórki do bazy danych. Typ każdej komórki przewidzianej w procesie uczenia jest opisany przez lekarza eksperta.

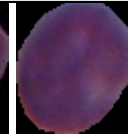
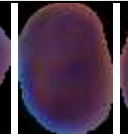
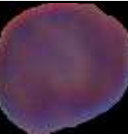
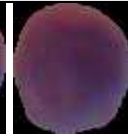
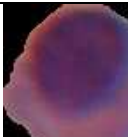
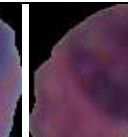
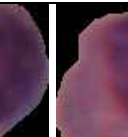
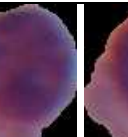
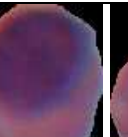


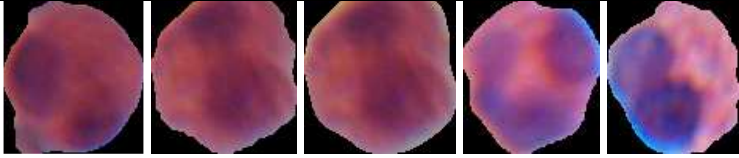
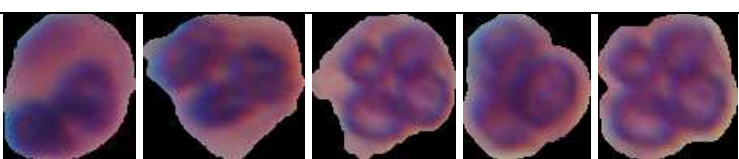
Rys. 4.15. Wygląd komórek wydzielonych przez program z obrazu okrężnicy z rys. 4.14

Na rys. 4.15 przedstawiono komórki wydzielone przez algorytm z obrazu z rys. 4.14 [12]. W tabeli 4.2 przedstawiono przykłady komórek czterech typów poddanych rozpoznaniu. Widoczne jest znaczne zróżnicowanie wyglądu komórek w obrębie tej samej rodziny. Szczególnie widoczne różnice dotyczą granulocytów, dla których może wystąpić różna liczba jąder.

Tabela 4.4.

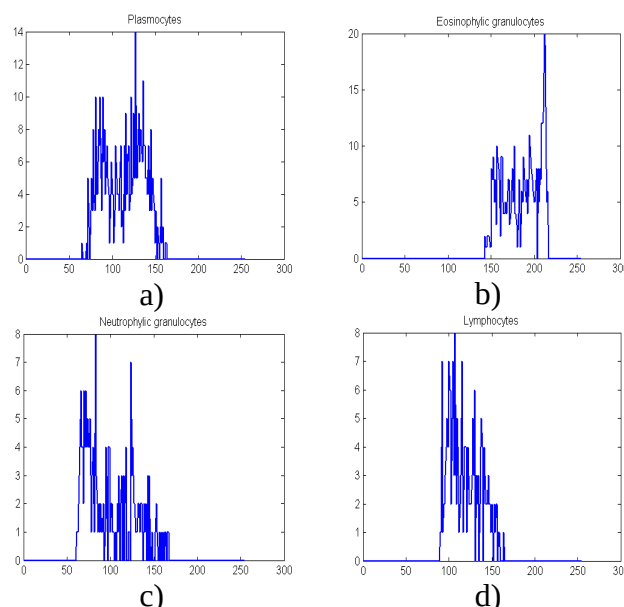
Przykłady komórek reprezentujących limfocyty, plazmocyty oraz granulocyty kwasochłonne i obojętnochłonne

Typ komórek	Przykłady komórek danego typu					
Limfocyty						
Plazmocyty						

Granulocyty kwasochłonne	
Granulocyty obojętnochłonne	

Generacja deskryptorów numerycznych obrazu

Dla potrzeb automatycznego rozpoznania komórek konieczne jest przetworzenie ich obrazów w deskryptory numeryczne, które po odpowiedniej selekcji utworzą cechy diagnostyczne, stanowiące sygnały wejściowe dla klasyfikatorów. W generacji cech wykorzystane zostały parametry opisujące histogramy, geometrię komórek, cechy tekstualne oraz kolorymetryczne.



Rys. 4.16. Histogramy rozkładu jasności pikseli koloru czerwonego dla a) plazmocyty b) granulocyty kwasochłonne, c) granulocyty obojętnochłonne, d) limfocyty

Histogram rozkładu intensywności dla każdego koloru RGB niesie istotną porcję informacji różnicujących różne typy komórek. Przykładowo na rys. 4.16 przedstawiono histogramy dla koloru czerwonego 4 rozważanych typów komórek.

Różnią się one rozpiętością oraz pozycją i wartością maksimum. Dla każdego koloru są to różne wartości. Jako deskryptory numeryczne zastosowano statystyki dotyczące wartości średniej, odchylenia standardowego, skośności i kurtozy dotyczące tych trzech parametrów (dla każdego koloru oddzielnie). Dotyczyły one całej komórki, jądra oraz cytoplazmy (w sumie ponad 40 deskryptorów).

Deskryptory geometryczne opisują różnice w kształcie i wielkości komórek różnych typów. Wśród tych deskryptorów wyróżnione zostały: średnica, obwód, pole powierzchni całej komórki, jądra i cytoplazmy, zwartość, symetria, długość małej i dużej osi, średnia odległość wszystkich pikseli obrazu od środka, a także różnego rodzaju współczynniki utworzone jako stosunek odpowiednich wielkości (w sumie 23 deskryptory). Tabela 4.3

prezentuje przykładowe wartości średnie i odchylenia standardowe różnych typów komórek dla kilku wybranych parametrów geometrycznych. Można zauważyć znaczne różnice w wartościach parametrów dla tych komórek.

Tabela 4.3. Wartości średnie i odchylenia standardowe wybranych parametrów geometrycznych dla 4 rodzajów komórek

	Limfocyty	Plazmocyty	Granulocyty kwasochłonne	Granulocyty obojętnochłonne
Obwód	46±6	70±8	78±9	52±4
Powierzchnia jądra	116±19	180±31	97±27	129±20
Powierzchnia komórki	220±50	461±91	596±120	258±48
Stosunek powierzchni jądra do całej komórki	0.52±0.2	0.39±0.12	0.16±0.1	0.5±0.12

Deskryptory teksturalne opisują w sposób matematyczny wzory charakterystyczne dla danej komórki w sensie stopnia jasności pikseli i relacje zachodzące między nimi. W tym zastosowaniu wykorzystane zostały cechy Haralicka [28] zdefiniowane dla kątów 0, 45° i 90°. Dotyczyły one kontrastu, energii, korelacji, bezwładności, zwartości, entropii oraz średniej sumy wariancji zdefiniowanych dla obrazu reprezentowanego w skali szarości, oddzielnie dla jądra i cytoplazmy (14 deskryptorów).

Ostatnia grupa deskryptorów została utworzona na bazie intensywności pikseli trzech kolorów reprezentowanych w postaci RGB. Są one równe wartościom średnim poszczególnych kolorów dla całej komórki, cytoplazmy oraz jądra. Dodatkowe parametry uzyskano w postaci relacji odpowiednich wartości średnich. W ten sposób wygenerowanych zostało 24 deskryptorów kolorymetrycznych. Wykorzystując wszystkie metody generacji deskryptorów utworzono w sumie 101 potencjalnych cech, które poddano następnie selekcji.

Selekcja cech diagnostycznych

Mając wygenerowany zbiór deskryptorów numerycznych należało dokonać ich selekcji pod względem zdolności różnicowania poszczególnych klas komórek. W rozwiązaniu tego problemu zastosowano miarę istotności Fishera, ustawiając wygenerowane wcześniej deskryptory odpowiednio do wartości tej miary. Jako cechy wybrano najbardziej znaczące deskryptory. Liczba ich dla każdego rodzaju komórki została dobrana eksperymentalnie na bazie danych weryfikujących (25% danych uczących). W eksperymencie użyto danych uzyskanych w Zakładzie Patomorfologii Wojskowego Instytutu Medycznego w Warszawie. Akwizycji obrazów dokonano przy powiększeniu 600x. W wyniku wstępnego przetworzenia utworzono bazę danych komórek pochodzących od 53 pacjentów o różnym stopniu zaawansowania choroby (najmniejsza liczba dotyczy granulocytów obojętnochłonnych pojawiających się w ostatnim stadium choroby). Każda komórka została rozpoznana przez eksperta ludzkiego i jej typ zapisany w bazie danych. Dane liczbowe komórek uczestniczących w eksperymentach numerycznych są przedstawione tabeli 4.4.

Tabela 4.4. Baza danych komórek uczestniczących w eksperymentach numerycznych

	Limfocyty	Plazmocyty	Granulocyty kwasochłonne	Granulocyty obojętnochłonne
Liczba	511	420	253	187

Ze 101 deskryptorów numerycznych w wyniku selekcji metodą Fishera wyodrębniono minimalną liczbę cech, gwarantującą najlepsze wyniki w trybie odtwarzania (na danych weryfikujących). W wyniku tego uzyskano 50 cech diagnostycznych dla limfocytów, 79 dla plazmocytów, 53 dla granulocytów kwasochłonnych oraz 48 dla granulocytów obojętnochłonnych.

Wyniki klasyfikacji danych

Aby uzyskać najbardziej wiarygodne wyniki rozpoznania zastosowano procedurę krosswalidacji. Dostępny zbiór danych podzielono na 5 równych części. Cztery z nich zostały użyte w uczeniu klasyfikatora a jedna część jedynie w odtwarzaniu. Eksperyment został powtórzony 5 razy, za każdym razem zmieniając skład grupy uczącej i testującej. Błędne rozpoznania komórek obliczano jedynie dla danych testujących nie uczestniczących w uczeniu (tryb odtworzeniowy klasyfikatora) jako średnią ze wszystkich 5 prób.

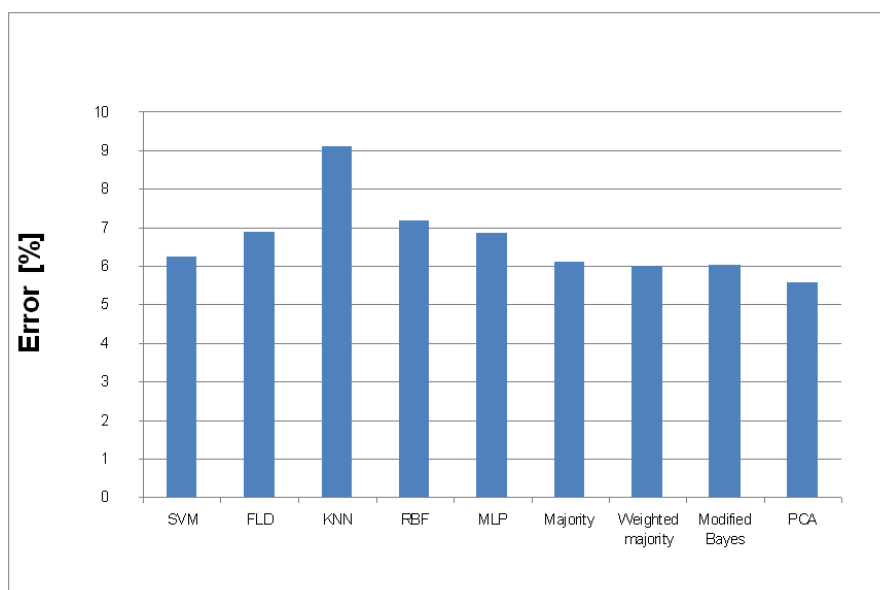
Jako klasyfikatory zostały zastosowane trzy rodzaje sieci neuronowych: SVM o jądrze gaussowskim, RBF oraz MLP) oraz dla porównania klasyfikator liniowy FLD oraz klasyfikator KNN (ang. *K Nearest neighbours*) przy $K=5$. Wyniki statystyczne testów tych klasyfikatorów na danych nie uczestniczących w uczeniu prezentowane są w tabeli 4.5.

Tabela 4.5 Wyniki średnie błędów rozpoznania komórek przy użyciu poszczególnych typów klasyfikatorów określone w stosunku do wyników eksperta ludzkiego

Rodzaj komórek	RBF	MLP	SVM	FLD	KNN
Limfocyty	5.6%	5.3%	3.9%	3.5%	7.6%
Plazmocyty	8.3%	8.3%	6.4%	7.6%	9.3%
Granulocyty kwasochłonne	5.9%	5.1%	5.9%	6.7%	7.5%
Granulocyty obojętnochłonne	10.1%	10.1%	14.3%	14.9%	15.0%
Średnio	7.17%	6.86%	6.24%	6.89%	9.11%

Wyniki pokazują, że najsprawniejszym klasyfikatorem jest sieć SVM, choć niewiele mu ustępuje klasyfikator liniowy FLD. Największe błędy rozpoznania dotyczą granulocytów obojętnochłonnych. Wynika to z ich małej liczby w bazie, zdecydowanie niewystarczającej do właściwego ukształtowania klasyfikatora.

Dalszą poprawę dokładności rozpoznania można uzyskać przy zastosowaniu zespołu klasyfikatorów połączonych w jeden finalny system. Porównano różne metody integracji zespołu: zwykłe głosowanie większościowe, głosowanie większościowe ważone (przy $m=4$), metodę Bayesa oraz PCA (redukcja z oryginalnego wymiaru 20 do 8). Na rys. 4.17 zestawiono wyniki prezentujące średni błąd rozpoznania poszczególnych komórek, uzyskany w procedurze krosswalidacji na danych nie uczestniczących w uczeniu. Najlepszy wynik odpowiadał zespołowi klasyfikatorów integrowanych przy użyci metody PCA, której odpowiadał błąd średni równy 5.58% (poprawa względna o 10.5% w stosunku do najlepszego klasyfikatora indywidualnego – SVM, dla którego średni błąd wynosił 6.24%).



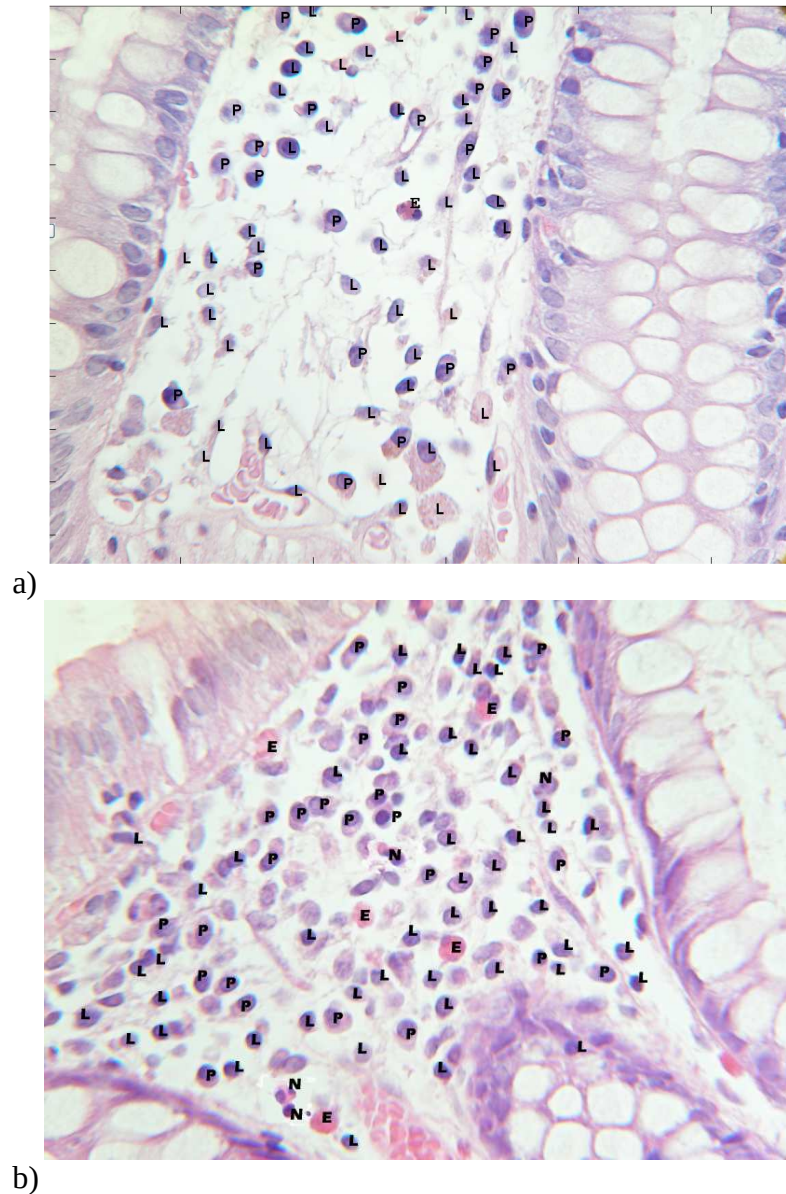
Rys. 4.17. Porównanie średniego błędu rozpoznania komórek przez klasyfikatory indywidualne i zespół klasyfikatorów integrowanych przy zastosowaniu różnych metod integracji

W tabeli 4.6 przedstawiono wyniki szczegółowe klasyfikacji w postaci tak zwanej tabeli niezgodności odpowiadającej najlepszemu wynikowi (integracja PCA zespołu klasyfikatorów). Poszczególne klasy komórek zakodowano przy użyciu symboli: L – limfocyty, P – plazmocyty, E – granulocyty kwasochłonne, N – granulocyty obojętnochłonne. Wielkości diagonalne odpowiadają idealnemu rozpoznaniu klasy, a poza diagonalne rozpoznaniu błędnemu. Wiersze macierzy reprezentują rzeczywistą klasę a kolumny – klasę przypisaną przez układ klasyfikatorów.

Tabela 4.6 Macierz niezgodności zespołu klasyfikatorów na danych testujących

	L	P	E	N
L	490	13	4	4
P	11	390	9	10
E	2	4	243	4
N	4	6	5	172

Innym rezultatem działania programu jest przedstawienie graficzne obrazu z naniesionymi automatycznie przez program typami rozpoznanych komórek.



Rys. 4.18. Wyniki graficzne rozpoznania komórek dla pacjenta w pierwszym stadium choroby (a) oraz dalsze stadium (b). Oznaczenia komórek: L – limfocyty, P – plazmocyty, E – granulocyty kwasochłonne, N – granulocyty obojętnochłonne.

Na rys. 4.18 pokazano przykładowe wyniki rozpoznania komórek dla dwu różnych obrazów. Obraz pierwszy reprezentuje pierwszy stopień rozwoju zapalenia, obraz drugi stopień dalszy. Cechą charakterystyczną obrazu drugiego jest pojawienie się komórek obojętnochłonnych (N), charakterystycznych dla stanu głębszego choroby.

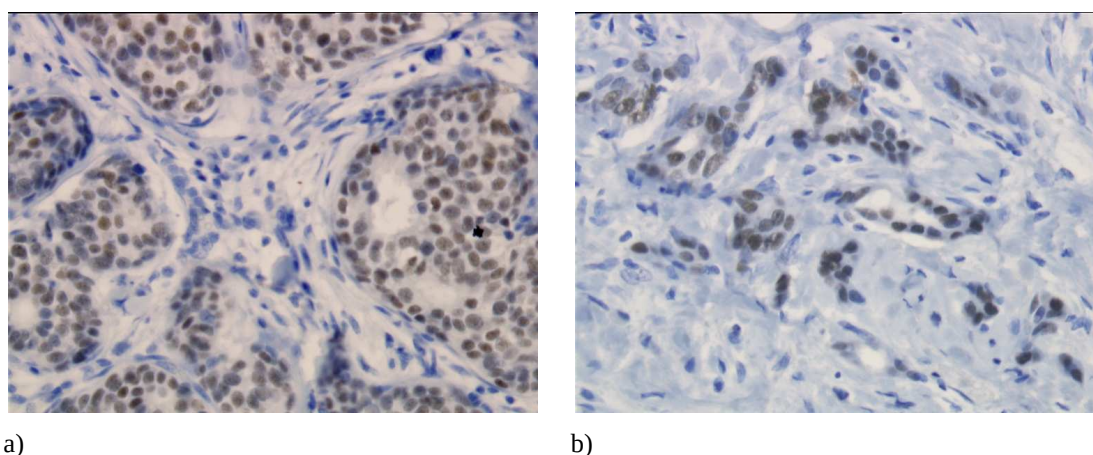
Obraz wyjściowy może podlegać sprawdzeniu i ewentualnej korekcji przez eksperta, w czasie której lekarz może poprawiać wyniki rozpoznania poprzez kliknięcie myszką i korektą ewentualnego błędnego rozpoznania komórki. Wynik poprawiony jest automatycznie uwzględniany w pliku numerycznym zawierającym stan populacji poszczególnych typów komórek.

4.6.3. Analiza ilościowa obrazów raków sutka

Nowotwór sutka jest najczęstszym nowotworem złośliwym występującym u kobiet [10,15]. W podejmowaniu decyzji o sposobie leczenia bierze się pod uwagę wskaźniki predykcyjne. Są to cechy kliniczne, patologiczne i biologiczne używane w celu estymacji

wiarygodności odpowiedzi na określony rodzaj planowanej terapii. Dla przykładu dla predykcji odpowiedzi na leczenie hormonoterapią używa się biochemicznych wskaźników w postaci receptora estrogenu (ER) i receptora progesteronu (PR) [15]. Obydwa są określane na podstawie obrazu tkanki poddanej odpowiedniej reakcji immunohistochemicznej. W aktualnie obowiązujących kryteriach decydujących o zakwalifikowaniu do tego typu leczenia wymagane jest określenie immunoreaktywności, czyli stosunku liczby komórek immunododatnich do liczby wszystkich komórek nowotworowych. Do podania odpowiedniej hormonoterapii kwalifikują się przypadki, dla których wskaźnik ten przyjmuje wartość co najmniej 10%.

Prawidłowa ocena immunoreaktywności jest bardzo trudna i złożona. Pierwszym problemem jest kwestia wybrania do analizy obszaru preparatu odpowiadającego naciekaniu zdrowej tkanki nowotworem (postać inwazyjna) lub wnętrza przewodu mlekowego wypełnionego komórkami nowotworowymi (postać wewnątrzprzewodowa). Na rys. 4.19 zobrazowano przykłady obrazów obu rodzajów raka sutka.



Rys. 4.19. Przykłady rozmieszczenia komórek nowotworowych raka sutka: a) postać wewnątrzprzewodowa, b) postać inwazyjna. Komórkami nowotworowymi są wyłącznie obiekty posiadające duże jądra, od kształtu elipsoidalnego do okrągłego (barwienie ER, $\times 400$)

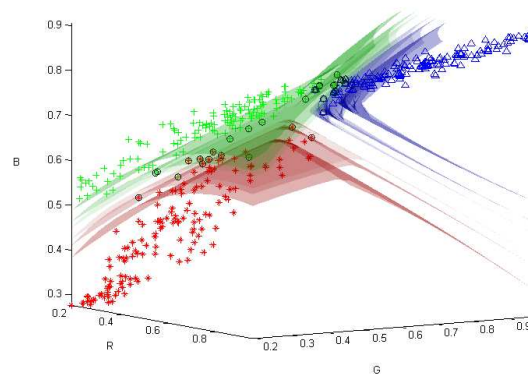
Ważnym problemem jest znaczna zmienność ekspresji barwnika brązowego w jądrach komórek. Stanowi ona istotne utrudnienie w odróżnieniu komórek immunododatnich od immunoujemnych. Dla uzyskania najwyższej dokładności rozpoznania profili komórek należy uwzględnić możliwość punktowego barwienia się podścieliska (tła) preparatu, a jednocześnie rozróżnić profile jąder wybarwionych komórek podścieliska, które nie podlegają zliczeniu. W efekcie obraz poddany standaryzacji może znacząco różnić się od oryginalnego pod względem odcienia profili jąder komórek.

Podstawowym problemem staje się ekstrakcja profili jąder wyłącznie komórek nowotworowych, a więc odróżnienie ich od pozostałych, takich jak komórki podścieliska czy ścian naczyń krwionośnych. Następnym zadaniem jest podział profili jąder komórek nowotworowych na dwie grupy: komórki immunododatnie o jądrach brązowych i immunoujemne o jądrach fioletowych.

Do rozwiązania tego zadania został opracowany specjalizowany algorytm w którym stosuje się najpierw sieć neuronową typu SVM z nieliniową funkcją jądra. Jej zadaniem jest możliwie optymalny podział obrazu przypisujący poszczególne piksele do jednej z trzech klas: klasy 1 zawierającej profile jąder komórek nowotworowych immunododatnich zabarwionych na brązowo, klasy 2 z profilami jąder komórek immunoujemnych, głównie nowotworowych, zabarwionych na fioletowo, oraz klasy 3 zawierającej tło obrazu. Dane wejściowe dla sieci tworzą trzy składowe RGB pikseli obrazu. Z uwagi na rozkład i znaczne

nakładanie się danych odpowiadających różnym klasom (wynikające z podbarwień podścieliska) konieczne jest zastosowanie nieliniowej funkcji jądra (w rozwiązaniu gaussowskiej), umożliwiającej rozwiązanie nieseparowanego problemu klasyfikacyjnego. Przy trzech klasach zastosowano metodę jeden-przeciw-jednemu, konstruując trzy klasyfikatory SVM rozpoznające każdą kombinację dwu klas.

Dane uczące tych sieci zostały wybrane na podstawie trzech przykładowych obrazów i zawierały po 150 pikseli typowych dla każdej klasy (razem 450 danych uczących). Ich rozkład w przestrzeni trójwymiarowej przedstawia rys. 4.20. Widoczne jest, że punkty danych układają się w regiony reprezentujące poszczególne klasy. Punkty oznaczone czerwonym symbolem gwiazdki oznaczają klasę immunododatnią (piksele brązowe), zielonym znakiem plus – immunoujemną (piksele fioletowe) a niebieskim trójkątem – tło. Punkty oznaczone dodatkowo kółkiem są wektorami podtrzymującymi, wyłoniłymi w procesie uczenia sieci SVM. Nałożone na dane płaszczyzny reprezentują kolejne poziomy przynależności do odpowiednich klas (wyznaczone jako minimalna wartość zwracana przez dwie sieci odróżniające daną klasę od pozostałych dwóch klas). Wyraźny rozkład danych na trzy regiony, w niewielkim stopniu zachodzące na siebie, sugeruje relatywnie dobrą przestrzenną separację klas, co jest dobrym prognostykiem do uzyskania satysfakcjonującej generalizacji sieci w procesie testowania (dla danych odpowiadających pikselom, nie biorącym udziału w uczeniu). Należy jednak zauważyć, że klasy nie są całkowicie liniowo separowalne.



Rys. 4.20. Rozkład 3 klas danych uczących sieci SVM

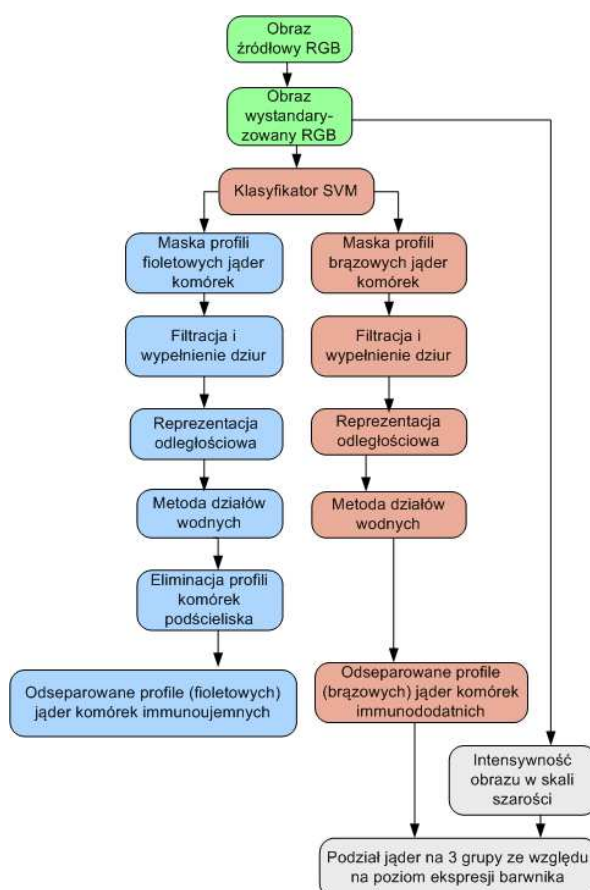
W wyniku zastosowania sieci SVM otrzymuje się maski obszarów należących do trzech klas, które muszą być poddane dalszemu przetwarzaniu. Jest ono realizowane przy użyciu operacji morfologicznych. Dwie z tych masek (jedna odpowiadająca pikselom profili jąder komórek immunododatnich, druga – immunoujemnych) są użyteczne w dalszym przetwarzaniu, mającym na celu rozpoznanie i wyłonienie z nich obiektów o zwartym kształcie tworzącym profil jądra komórki. Maską trzecią nie odgrywa roli w dalszym przetwarzaniu obrazu.

Obie wymienione maski profili jąder komórek poddane są filtracji przy użyciu erozji, dylatacji, otwarcia i zamknięcia w celu wygładzenia brzegów, usunięcia pojedynczych pikseli nie tworzących rozpoznawanych (zwartych) obiektów oraz wypełnienia dziur w potencjalnych profilach. Dla rozdzielenia profili (jąder komórek) stykających się lub częściowo się nakładających, zastosowano reprezentację odległościową i metodę działów wodnych.

Na tym etapie przetwarzania z maski odpowiadającej klasie 1 otrzymuje się obiekty odpowiadające w większości rozdzielonym profilom jąder komórek immunododatnich. W przypadku maski odpowiadającej klasie 2 problemem do rozwiązania pozostaje nadal eliminacja tych, które odpowiadają profilom jąder innych komórek niż nowotworowe (głównie komórkom podścieliska), również koloru fioletowego. Jedynym potencjalnym

źródłem informacji pozwalającej na rozróżnienie komórek pozostaje rozmiar i kształt profili ich jąder. Wydzielenie profili jąder komórek podścieliska na podstawie ich rozmiaru jest bardzo problematyczne, gdyż ich wielkość (choć najczęściej mniejsze od profili jąder komórek nowotworowych) nie są jednoznacznie określona, a przy tym pole przypisane im w procesie segmentacji może być nadmiarowe co czyni je podobnymi do profili jąder komórek nowotworowych. Ponadto należy zauważyć, że komórki nowotworowe cechuje znaczna zmienność biologiczna, również pod względem wielkości, co stanowi dodatkową trudność rozróżniania obu typów komórek.

Ostatecznie jedyną użyteczną cechą mogącą różnicować komórki nowotworowe od pozostałych jest kształt profili ich jąder. Dotychczasowe obserwacje pozwalają przyjąć, że profile jąder komórek nowotworu sutka mają kształt zbliżony do okrągłego podczas gdy profile jąder komórek podścieliska mają kształt wrzecionowaty. Zaprojektowana procedura wykorzystuje ten fakt, stosując kilkukrotną erozję obrazu przy użyciu elementu *SE* w kształcie dysku. Na rys. 4.21 przedstawiono schemat blokowy rozwiązania zastosowany w rozpoznaniu obu klas komórek.

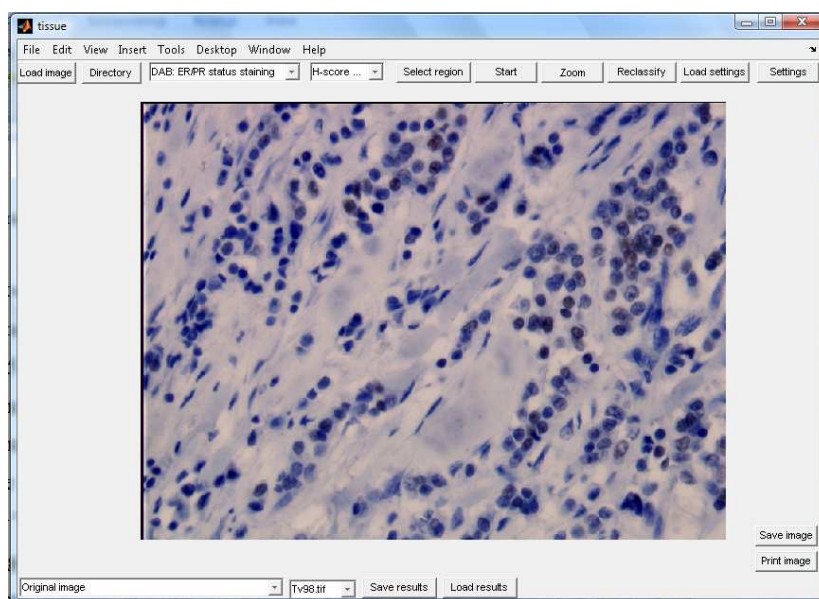


Rys. 4.21. Schemat algorytmu rozpoznania komórek immunododatnich i immunoujemnych na obrazie raka sutka

Ostatnim etapem oceny ilościowej preparatu jest zliczenie komórek immunododatnich i immunoujemnych, reprezentowanych przez profile ich jąder, w obu przetwarzanych maskach, wykonywane w Matlabie przez funkcję *bwlabel*. Z uwagi na różne stosowane na świecie metody oceny preparatów nowotworów sutka z reakcją na ER/PR, konieczne jest zróżnicowanie ostatniego etapu procedury. Możliwy jest wybór oceny standardowej preparatu w postaci liczby komórek immunododatnich i immunoujemnych lub oceny uwzględniającej poziom ekspresji barwnika w komórkach immunododatnich.

W trzystopniowej skali ekspresji, konieczne jest podzielenie komórek immunododatnich na trzy grupy wybarwienia: silne, średnie i słabe. Do rozwiązywania tego zadania wykorzystano informację o intensywności zabarwienia mierzonej jako średnia odległość od koloru białego pikseli tworzących profil jądra danej komórki (po transformacji do skali szarości). Progi podziałów pomiędzy poziomami wybarwienia dzielącymi 3 klasy zostały uzgodnione przez kilku ekspertów na podstawie analizy szeregu wykalibrowanych obrazów. Na rys. 4.21 przedstawiono schemat algorytmu segmentacji implementujący omówione wcześniej kroki, uwzględniający podział komórek immunododatnich na trzy grupy.

Opracowany system automatycznej analizy obrazu umożliwia wykonanie oceny ilościowej serii obrazów mikroskopowych. Ocena wykonywana jest zgodnie z wybraną metodą (na przykład H-score lub wskaźnika R). Graficzny interfejs (rys.4.22) umożliwia wczytanie z pliku pojedynczego zdjęcia lub wczytuje automatycznie wszystkie zdjęcia z wybranego katalogu. Następnie użytkownik systemu ma dodatkową możliwość graficznego zdefiniowania obszaru obrazu podlegającego analizie (na przykład zawierającego tylko postać inwazyjną nowotworu lub tylko obszar wybranego przewodu mlekowego z nowotworem). Jeżeli wczytano zestaw zdjęć, taki obszar można zdefiniować do każdego z nich niezależnie. W przypadku nie zdefiniowania go, domyślnym obszarem analizy ilościowej jest cały obraz. Po przeprowadzeniu wszystkich analiz system daje możliwość odczytania wyniku oddzielnie dla każdego obrazu.

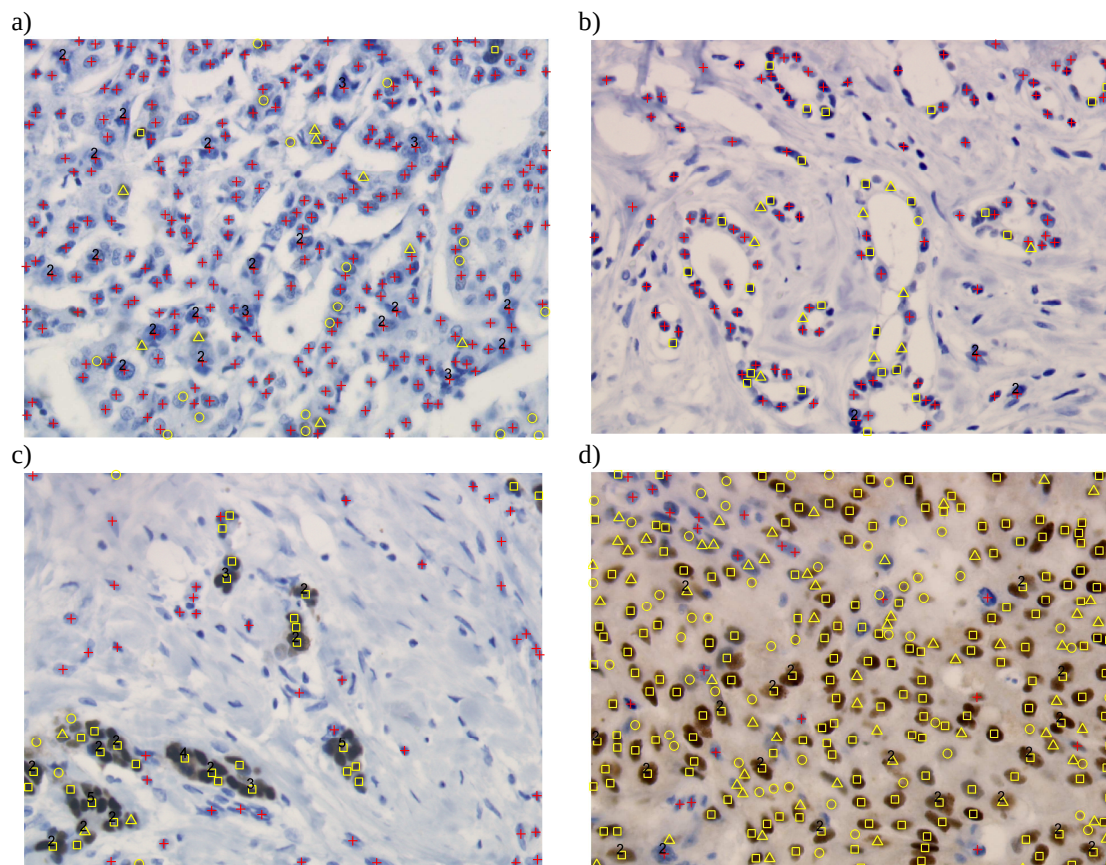


Rys. 4.24. Obraz graficznego interfejsu systemu

Dla oceny dokładności zaprojektowanego systemu użyto obrazów wyselekcjonowanych przez lekarzy zawierających postać inwazyjną nowotworu sutka. Na rys. 4.23 przedstawiono cztery przykładowe wyniki wygenerowane przez system automatyczny prezentujące graficznie rozpoznane profile jąder komórek z podziałem na trzy poziomy intensywności odczynu (wybarwienia jąder komórek immunododatnich). Rozpoznane profile jąder komórek immunoujemnych są oznaczone czerwonym znakiem „+”. Profile jąder komórek immunododatnich są oznaczone w zależności od rozpoznanego poziomu siły odczynu: poziom 1 – kołko koloru żółtego, poziom 2 – trójkąt koloru żółtego, poziom 3 – kwadrat koloru żółtego.

Ponieważ w niektórych przypadkach zdarza się, że system nie jest w stanie rozdzielić nakładających się profili jąder komórek, przeprowadza się analizę rozkładu wielkości pól rozpoznanych obiektów. Te obiekty, które mają wyraźnie większe pole od wartości średniej

dla danego obrazu, traktowane są jako kilka nierozdzielonych profili jąder komórek. Ich liczba jest estymowana proporcją ich pola do średniej wielkości profilu jądra komórki rozpoznanej na tym obrazie. W prezentacji graficznej wyniku dla takich profili jąder komórek zastosowano dodatkowe oznaczenie na czarno w postaci numerycznej (cyfra oznacza estymowaną liczbę profili jąder komórek w klastrze).



Rys. 4.23. Przykładowe wyniki rozpoznania profili jąder komórek raków sutka w odczynach ER/PR z rozróżnieniem poziomu intensywności odczynu (DAB, $\times 400$)

Przeprowadzono również szczegółowe badania ilościowe wielu preparatów przy użyciu opracowanego systemu na bazie obrazów pochodzących od 30 pacjentek (po około 10 zdjęć dla każdej pacjentki). Dokładnie te same obrazy były niezależnie ocenione przez trzech ekspertów, a ich średni wynik (oznaczony symbolem EXP) był brany jako punkt odniesienia przy analizie dokładności zaprojektowanego systemu automatycznego. Podczas przeprowadzania oceny ekspert nie miał wglądu ani do wyników systemu, ani do wyników innych ekspertów. Szczegółowe porównanie otrzymanych wyników otrzymanych dla obrazów z rys. 4.23 zostało zamieszczone w tabeli 4.7. Wyniki zbiorcze dla wszystkich 50 preparatów zawarte są w tabeli 4.8.

W tabeli 4.7 zamieszczone są zarówno wyniki zliczenia profili jąder poszczególnych typów komórek, jak i otrzymane wskaźniki zgodnie ze stosowaną skalą oceny. Wskaźniki te dotyczą procentowej zawartości komórek immunododatnich (oznaczenie R) oraz wskaźnika H w skali H-score zdefiniowując następująco [10,15]:

$$R = \frac{n_{1+} + n_{2+} + n_{3+}}{n_{1+} + n_{2+} + n_{3+} + n_{-}}, \quad (4.31)$$

$$H = \frac{n_{1+}}{n_{1+} + n_{2+} + n_{3+} + n_{-}} + 2 \cdot \frac{n_{2+}}{n_{1+} + n_{2+} + n_{3+} + n_{-}} + 3 \cdot \frac{n_{3+}}{n_{1+} + n_{2+} + n_{3+} + n_{-}}, \quad (4.32)$$

gdzie n_{1+} , n_{2+} , n_{3+} oznaczają odpowiednio liczbę komórek immunododatnich z odczynem na poziomie 1, 2 i 3, a n oznacza liczbę komórek immunoujemnych, wszystkie estymowane liczbą rozpoznanych profili. Średnia różnica (bezwzględna) wartości wskaźnika procentowego R pomiędzy wynikami systemu automatycznego a wynikiem średnim ekspertów, obliczona dla badanych 30 pacjentów wyniosła 0.99% z odchyleniem standardowym 0.52%. W wartościach względnych wynik ten odpowiadał różnicy procentowej równej 3.49% z odchyleniem standardowym 3.09%. W przypadku skali H-score (punktowej) średnia różnica między wynikiem systemu automatycznego a uśrednionymi wynikami ekspertów wyniosła 6.7 punktów co odpowiada 6.75% w wartościach względnych przy odchyleniu standardowym równym 9.42%.

Tabela 4.7. Wyniki zliczenia komórek dla 4 przykładowych obrazów raków sutka z rys. 4.23 w odczynach EP/PR, na podstawie profili ich jąder

Obraz	Ekspert (EXP) lub system (AS)	Liczba komórek immunoujemnych	Liczba komórek immunododatnich z siłą odczynu 1	Liczba komórek immunododatnich z siłą odczynu 2	Liczba komórek immunododatnich z siłą odczynu 3	Wskaźnik R	Wskaźnik H
Rys. 4.11a	AS	278	19	9	2	9.74%	14
	EXP	279	12	12	3	8.82%	15
Rys. 4.11b	AS	102	1	10	26	26.62%	71
	EXP	109	8	7	27	27.81%	68
Rys. 4.11c	AS	43	5	4	54	59.43%	165
	EXP	47	6	12	46	57.66%	151
Rys. 4.11d	AS	24	54	63	185	94.64%	245
	EXP	26	81	44	178	94.10%	214

W tabeli 4.8 przedstawiono wyniki zbiorcze dla obrazów raka sutka wszystkich 30 pacjentek (50 badanych preparatów). Jak widać otrzymane wyniki systemu automatycznego są wystarczająco dokładne w stosunku do eksperta. Należy podkreślić, że między wynikami ekspertów występowały również znaczne różnice (nawet do 20%). Wiążą się one głównie z indywidualną interpretacją eksperta które profile jąder komórek odpowiadają jego zdaniem komórkom nowotworowym (postać inwazyjna). W odróżnieniu od nich wyniki systemu automatycznego są powtarzalne (program uruchomiony na tym samym obrazie generuje zawsze ten sam wynik)..

Tabela 4.8. Wyniki zbiorcze (dla 30 pacjentek) zliczenia komórek dla obrazów raka sutka w odczynach EP/PR, na podstawie profili ich jąder

Ekspert (EXP) lub system (AS)	Średnia liczba komórek immunoujemnych	Średnia liczba komórek immunododatnich z siłą odczynu 1	Średnia liczba komórek immunododatnich z siłą odczynu 2	Średnia liczba komórek immunododatnich z siłą odczynu 3	Średnia wartość wskaźnika R	Średni błąd względny oznaczenia R
AS	118	14	18	71	45.1%	3.5%
EXP	120	17	17	68	44.7%	

4.7. Wnioski końcowe

Tematyka rozdziału dotyczyła zastosowania wybranych metod sztucznej inteligencji (sieci neuronowe, algorytmy genetyczne, zespoły eksperckie) w automatycznej analizie obrazów biomedycznych do wspomagania diagnostyki medycznej w stanach patologicznych tkanek (nowotwory, stany zapalne). Jest to niewielki wycinek ogromnego frontu badawczego związanego z medycyną, który w ostatnich latach jest intensywnie rozwijany na świecie.

Etap analizy obrazów tkanek patologicznych jest szczególnie ważny w badaniach nad nowotworami przy podejmowaniu właściwej decyzji dotyczącej rodzaju i stopnia schorzenia oraz związanego z tym rodzaju leczenia. Dla eksperta medycznego oceniającego i analizującego obraz takiej tkanki jest to zadanie trudne i wymagające żmudnego (często wielogodzinnego) ślęczenia nad mikroskopem lub przed monitorem komputera. Wynik pracy wielu ekspertów nad tymi samymi obrazami jest zwykle zróżnicowany i zależny nie tylko od wiedzy i doświadczenia eksperta, ale również od stopnia jego zmęczenia lub aktualnego stanu psychiczno-fizycznego.

Opracowanie automatycznego systemu komputerowego pozwala wyeliminować pracę ekspertów i zobiektywizować wynik analizy. System komputerowy dobrze przemyślany i rozwiązany generuje wyniki powtarzalne, nie podlega zmęczeniu, może pracować non-stop i pozwala znakomicie przyspieszyć badania nad nowotworami. Zastosowanie zaawansowanych metod przetwarzania obrazu pozwala uzyskać również wyniki ilościowe dotyczące geometrii komórek (np. pole powierzchni, obwód, charakterystykę kształtu, itp.) niemożliwe do uzyskania w manualnej ludzkiej analizie obrazów.

Literatura

- [1] Ashlock D., *Evolutionary computation for modeling and optimization*. Springer 2006, Berlin.
- [2] Cichocki, A., Amari S.I., *Adaptive blind signal and image processing*. Wiley 2003, New York.
- [3] Cytowski J., Gielecki J., Gola A., *Cyfrowe przetwarzanie obrazów medycznych. Algorytmy. Technologie. Zastosowania*, Wydawnictwo Exit 2008, Warszawa.
- [4] Duda R. O, Hart P. E., Stork P., *Pattern classification and scene analysis*. Wiley 2003, New York.
- [5] Duncan W., *Prostate Cancer*. Springer 1981, Berlin, pp. 102-113.
- [6] Gleason D. F., Mellinger G. T., *Prediction of prognosis for prostatic adenocarcinoma by combined histological grading and clinical staging*. Journal of Urology, vol. 111, 1974, pp. 58–64.
- [7] Gonzalez R. C., Woods R. E., *Digital Image Processing*. Addison-Wesley 1992, Reading, Massachusetts.
- [8] Guyon I., Elisseeff A., *An introduction to variable and feature selection*. Journal of Machine Learning Res., vol. 3, 2003, pp. 1158 – 1182
- [9] Haykin S., *Neural Networks, Comprehensive Foundation*. Prentice Hall 1999, New Jersey.
- [10] Kayser G, Radziszowski D, Bzdyl P, Werner M, Kayser K: *Eamus – internet based platform for automated quantitative measurements in immunohistochemistry*. Inf. Conf. Soc. for Cellul. Oncol.(ISCO), Belfast, 2005.
- [11] Kruk M., *Automatyczny system rozpoznawania komórek na podstawie obrazu mikroskopowego wybranej tkanki ludzkiej dla potrzeb diagnostyki medycznej*. rozpr. dokt. Politechniki Warszawskiej, 2008.
- [12] Kruk M., Osowski S., Koktysz R., *Recognition and Classification of Colon Cells Applying the Ensemble of Classifiers*. Computers in Biology and Medicine, vol. 39, 2009, pp. 156-165.
- [13] Kuncheva L., *Combining pattern classifiers: methods and algorithms*. Wiley 2004, New York.

- [14] Markiewicz T., *Wybrane narzędzia wydobywania informacji z obrazów histologicznych w zastosowaniu do wspomagania diagnostyki patomorfologicznej*. rozprawa habilitacyjna Politechniki Warszawskiej, 2010.
- [15] T. Markiewicz, P. Wiśniewski, S. Osowski, J. Patera, W. Kozłowski, R. Koktysz, *Comparative analysis of the methods for accurate recognition of cells in the nuclei staining of the Ki-67 in neuroblastoma and ER/PR status staining in breast cancer*. Analytical and Quantitative Cytology and Histology Journal, vol. 31, 2009, pp. 49-63.
- [16] Matlab: *Image Processing Toolbox - user's guide*. MathWorks 2010, Natick.
- [17] Michalewicz Z., *Genetic Algorithms + Data Structures = Evolution Programs*. Springer 1996, Berlin.
- [18] Nikias C., Petropulu A., *Higher order spectral analysis*. Prentice Hall 1993, New Jersey.
- [19] Osowski S., *Sieci neuronowe do przetwarzania informacji*. OW PW 2006, Warszawa.
- [20] S. Osowski, T. Markiewicz, *Support vector machine for recognition of white blood cells in leukemia* (in G. Camps-Valls, J. L. Rojo-Alvarez, M. Martinez-Ramon, *Kernel methods in bioengineering, signal and image processing*. Idea Group Publishing, London, 2007), pp. 93-123.
- [21] Otsu, N: *A threshold selection method from grey-level histograms*. IEEE Trans. on Systems, Man, and Cybernetics , vol. 9, 1979, pp. 62-66.
- [22] Soile P., *Morphological image analysis, principles and applications*. Springer 2003, Berlin.
- [23] Schölkopf B, Smola A: *Learning with Kernels*. Cambridge, MA: MIT Press 2002, Massachusetts.
- [24] Siroić R., Osowski S., Markiewicz T., Siwek K., *Application of Support Vector Machine and Genetic Algorithm for Improved Blood Cell Recognition*. IEEE Trans. Meas. and Instrum, vol. 58, 2009, pp. 2159-2168.
- [25] Vapnik V: *Statistical Learning Theory*. Wiley 1998, New York.
- [26] Thompson E. M., Price A. B., Altman D. G., Sowter C., Slavin G., *Quantitation in inflammatory bowel disease using computerized interactive image analysis*. J. Clin. Pathol., vol. 38, 1985, pp. 631-638.
- [27] Tadeusiewicz R. and Korohoda P., *Komputerowa analiza i przetwarzanie obrazów*. Wydawnictwo Fundacji Postępu Telekomunikacji 1997, Kraków.
- [28] T. Wagner, *Texture analysis*. (in Jahne, B., Haussecker, H., & Geisser, P., Eds. *Handbook of Computer Vision and Application*. Academic Press 1999, Boston), pp. 275-309.