

System automatycznych pomiarów rynometrycznych (5)

Tomasz Kuśmierczyk

Studenckie Koło Naukowe Cybernetyki, Politechnika Warszawska

Streszczenie: Celem projektowanego systemu jest analiza i rozpoznawanie obrazów trójwymiarowych twarzy. Wykorzystując dostępne metody analizy i narzędzia algorytmiczne, dąży się do pozyskania z danych pochodzących ze skanerów 3D informacji dotyczących wymiarów nosa. Artykuł prezentuje innowacyjny pomysł zastosowania deskryptorów punktów w połączeniu z metodami uczenia maszynowego do detekcji wybranych obszarów w skanach 3D. W pracy opisano algorytmu i omówiono różne jego aspekty, pokazano również praktyczne rezultaty zastosowania algorytmu do analizy twarzy.

Słowa kluczowe: spin image, uczenie maszynowe, rozpoznawanie twarzy, chmura punktów

Ostatnie dwa artykuły poświęcone zostały jednemu ze znanych narzędzi rozpoznawania obrazów trójwymiarowych: deskryptorom punktów. Przypomnijmy, że rozumie się przez nie pewne struktury danych służące charakteryzowaniu i rozróżnieniu punktów obrazu 3D. W części 3. opisano kilka różnych typów deskryptorów, m.in. *spin images* (obrazy obrotów). Część 4. poświęcono uszczegółowieniu różnych ich aspektów: wyboru sąsiedztwa, doboru skali i skwantowania, obliczaniu odległości między deskryptorami czy strategiom wyboru punktów.

Część 5. będzie dopełnieniem poprzednich dwóch – opisana zostanie koncepcja praktycznego zastosowania deskryptorów w identyfikacji regionów obrazu 3D. Odbiega ona od innych pomysłów, które przybliżone zostały w części 4. Przypomnijmy, opisano dwa typy rozwiązań: podejście polegające na selekcji pojedynczych punktów charakterystycznych bazujące na klasyfikacji metodami uczenia maszynowego oraz podejście, w którym dopasowuje się deskryptory zbudowane na podzbiorze punktów analizowanego obiektu do analogicznego podzbioru wzorca. Zaproponowane rozwiązanie przypomina podejście drugie, ale jest od niego istotnie różne:

- celem nie jest stwierdzenie poziomu zgodności regionu ze wzorcem ale wyselekcjonowanie podregionu,
- faza dopasowywania jest jedynie częścią zaproponowanego algorytmu.

Algorytm lokalizacji regionu

Głównym elementem zaproponowanego rozwiązania jest algorytm lokalizacji regionu. Jego zadaniem jest identyfikacja części danych (podzbioru chmury punktów) odpowiadających wzorcowi (np. odszukanie nosa w skanie twarzy). Polega on na zastosowaniu łącznego dopasowywania deskryptorów punktów do deskryptorów wzorca w połączeniu z uczeniem maszynowym.

Schemat ogólny algorytmu:

1. Obliczenie deskryptorów na podzbiorze analizowanej chmury punktów:

Na wejściu algorytmu podawana jest chmura punktów, z której wybierany jest podzbiór punktów. Na podzbiorze punktów wyliczane są deskryptory.

2. Dopasowanie do wzorców (przygotowanie zbioru uczącego):

Za pomocą jednego z możliwych podejść (np. algorytmu węgierskiego opisanego w cz. 4.) deskryptory podzbioru punktów z analizowanej chmury punktów dopasowywane są do wybranych deskryptorów punktów z wzorca. Część zostaje dopasowana, a część nie. W ten sposób uzyskiwana jest informacja, w której części przestrzeni położenia (x, y, z) punktów znajduje się poszukiwany obszar, a w której nie.

Wynikiem kroku nr 2 algorytmu są dwa zbiory deskryptorów (wraz z odpowiadającymi im punktami). Jeden zbiór zawiera deskryptory punktów, które przyporządkowano do regionu poszukiwanego. Powinny się one znaleźć w grupie punktów chmury wejściowej należących do zlokalizowanego regionu. Drugi zbiór składa się z elementów odrzuconych. Powinny one zostać losowo rozmieszczone po powierzchni poza tym regionem.

Powyższe założenia nie zawsze są spełnione. Jednym ze sposobów poprawiających działanie fazy dopasowywania jest powtórzenie powyższej procedury dla różnych wzorców. Np. generowanych jest 300 deskryptorów punktów, których 100 losowo wybranych dopasowywanych jest do wzorca nr 1, kolejnych 100 do wzorca nr 2 i ostatnich 100 do wzorca nr 3. W ten sposób unika się problemu wyboru właściwego wzorca, który wystarczająco dobrze reprezentowałby wszystkie możliwości.

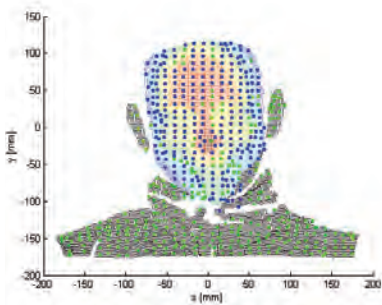
Dla poprawy jakości tak otrzymanych danych uczących, można przefiltrować zbiory dopasowanych i niedopasowanych punktów/deskryptorów. Korekta zbioru opiera się na założeniu o spójności poszukiwanego regionu. Pojedyncze punkty i małe grupy punktów oddalone od większości dopasowanych punktów zostają odrzucone. Szczegóły algorytmu filtracji przybliżone zostaną w następnej części artykułu.

3. Zbudowanie klasyfikatora:

Zbiór punktów o współrzędnych (x, y, z) odpowiadających punktom z poszukiwanego regionu oraz tych, które nie należą do regionu tworzą zbiór treningowy klasyfikatora. Zadaniem klasyfikatora jest nauczenie się rozróżniania, czy dany punkt należy do poszukiwanego regionu, czy nie, na podstawie współrzędnych (x, y, z) . Klasyfikatorem mogą być sieci neuronowe, SVM lub inne typowe narzędzie uczenia maszynowego.

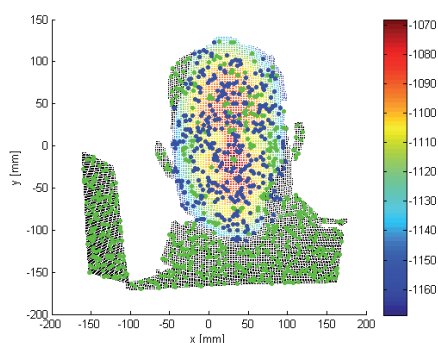
4. Klasyfikacja wszystkich punktów skanu:

Wykorzystując zbudowany wcześniej klasyfikator, rozdziela się wszystkie punkty chmury na dwie grupy. Ilustrację działania algorytmu zawierają rysunki. Wyobraźmy sobie, że pozyskaliśmy skan korpusu człowieka i chcemy zidentyfikować w nim obszar twarzy. Próbkujemy punkty skanu, wybierając ich podzbiór (ze uwagi na wydajność czasową nie można użyć wszystkich – czas dopasowania do punktów wzorca jest co najmniej kwadratowy). Następnie dla tych punktów obliczamy deskryptory i dopasowujemy do deskryptorów wzorca. Część z nich zostaje dopasowana – punkty oznaczone na niebiesko, część nie – punkty oznaczone na zielono. Widać, że część punktów dopasowana została błędnie. Problem rozwiązuje częściowo algorytm filtrowania opisany w kolejnej części, częściowo wykorzystując zdolność generalizacji klasyfikatora. Zbiór punktów umożliwia indukcję klasyfikatora, który dalej wykorzystany jest do sklasyfikowania wszystkich punktów obrazu. Część z nich zostaje odrzucona (kolor czarny), część uznana za należące do poszukiwanego regionu (kolor odpowiadający głębokości – twarz).



Rys. 1. Ilustracja działania algorytmu lokalizacji regionów [1]

Fig. 1. Results presentation for region localization algorithm [1]



Rys. 2. Ilustracja działania algorytmu lokalizacji regionów dla 399 deskryptorów w zbiorze uczącym i perceptronu dwuwarstwowego o 5 neuronach ukrytych [1]

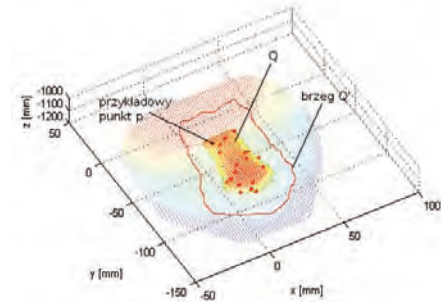
Fig. 2. Results presentation for region localization algorithm (399 spin images; multilayered perceptron) [1]

Przygotowanie wzorca

Opisany algorytm wymaga przygotowania wzorca, wyizolowanego fragmentu skanu, który będzie reprezentacją pewnego regionu, np. wzorec nosa. Dla wzorca obliczane

są deskryptory, które dopasowywane są do deskryptorów wyliczonych na analizowanym skanie.

Fazę przygotowania danych pokazuje rys. 3. Przez Q oznaczono obszar wzorca – nos. Na jego powierzchni wybrano T_w punktów p , dla których wyliczone zostaną deskryptory. Należy tu uwzględnić nie tylko obszar Q , ale również pewne jego otoczenie Q' . Dzięki temu w deskrypcji uwzględniona zostanie informacja nie tylko na temat kształtu wzorca, ale również „kontekstu” jego występowania.



Rys. 3. Sposób przygotowania deskryptorów wzorca [1]

Fig. 3. Spin images calculation in pattern image [1]

Opisana metoda zawiera kilka niedopowiedzeń. Sposób wyboru punktów, dla których obliczane są deskryptory opisany zostanie w ostatniej części. Odpowiedź na pytanie, jaki obszar uwzględnić we wzorcu, by uzyskać najlepszą skuteczność dopasowania, zależy w dużym stopniu od parametrów deskryptorów. Podobnie rzecz się ma z wyborem parametru T_w .

Wyliczenie deskryptorów

Podstawowym problemem algorytmu analizy skanu jest zdecydowanie, dla których punktów zbudowane zostaną deskryptory i ile ich należy obliczyć. Przygotowane deskryptory dopasowywane będą w kolejnych fazach do deskryptorów wzorca, należy więc zoptymalizować proces wyboru punktów względem jakości dopasowywania.

Wiadomo, ile deskryptorów wyliczono wcześniej dla wzorca. Wiadomo też, w jaki sposób zostały one rozmieszczone na powierzchni. Nie wiadomo jednak, ile i jak je wyliczyć dla analizowanego skanu. Na skanie, oprócz obszaru, który powinien zostać dopasowany X , znajduje się duży obszar, który nie powinien zostać dopasowany Y . Dąży się do tego, żeby w obszarze X znalazło się mniej więcej tyle samo deskryptorów, ile jest ich we wzorcu i żeby zostały one rozmieszczone w sposób najbardziej zbliżony. W przypadku optymalnym wszystkie deskryptory zostaną dopasowane w kolejnej fazie do wzorca – obszaru Q , a deskryptory z obszaru Y – odrzucone. Można założyć, że deskryptory rozmieszczane są równomiernie na powierzchni obiektów. Należy zadbać jedynie, by gęstość rozmieszczenia na powierzchni (punkt/mm²) deskryptorów we wzorcu i analizowanym skanie była taka sama. Problemem jest jednak to, że wzorec i analizowany obiekt mogą mieć różną rozdzielczość, np. wybór co setnego punktu może w przypadku skanu o wyższej rozdzielczości dać 3 punkty/mm² i np. 1 punkt/mm² dla niższej rozdzielczości. Naturalnym wydaje się powiązanie liczby deskryp-

torów, które powinny zostać wyliczone z rozdzielczością skanu – s (liczba punktów na jednostkę powierzchni):

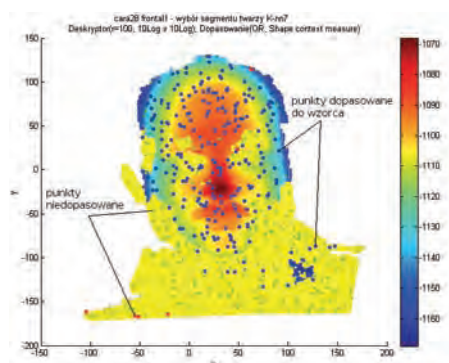
$$T = T_w \frac{N}{s} \frac{s_w}{N_w}$$

gdzie: s – rozdzielczość analizowanego skanu [punkt/mm²], s_w – rozdzielczość wzorca [punkt/mm²], T – liczba deskryptorów na analizowanym skanie, T_w – liczba deskryptorów na wzorcu, N – liczba punktów analizowanego skanu, N_w – liczba punktów wzorca.

Parametryzacja algorytmu

Działanie algorytmu zależy od zestawu parametrów. Sterują one jego zachowaniem w kolejnych fazach. Są to:

- liczba i wybór wzorców – powinny dobrze definiować klasę poszukiwanych danych. Jednocześnie nie powinno być ich

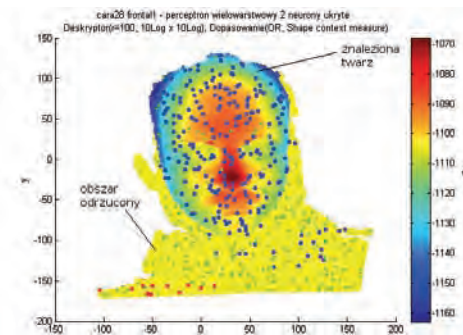


Rys. 4. Zachowanie algorytmu lokalizacji twarzy dla klasyfikatora k-nn dla $n = 7$ [1]

Fig. 4. Face localization algorithm results for k-nn classifier ($n = 7$) [1]

zbyt dużo, by nie mnożyć kosztów obliczeniowych;

- parametry charakteryzujące deskryptory (promień otoczenia, rozdzielczość, typ skali) i sposób ich dopasowywania (typ filtru, stosowana odległość) – powinny być dostosowane do spodziewanej zmienności danych;
- sposób wyboru punktów, dla których zbudowane zostaną deskryptory;
- liczba deskryptorów we wzorcach. Jest naturalnym, że zwiększając ich rozmiar, można poprawić dopasowanie do wzorca i w efekcie lepsze wyszukiwanie obszarów. Jednocześnie zwiększa się czas przetwarzania. Parametr ten powinien zostać uzależniony od oczekiwanej dokładności wyniku;
- parametry filtrowania punktów;
- typ klasyfikatora użytego do klasyfikacji punktów wejściowego obrazu. Rozważane były dwa typy klasyfikatora: k-nn i sieć neuronowa (perceptron wielowarstwowy uczony metodami wstecznej propagacji błędów i Levenberga-Marquardta). Wadą pierwszego jest ryzyko nadmiernego (lokalnego) dopasowania do



Rys. 5. Zachowanie algorytmu lokalizacji twarzy dla perceptronu dwuwarstwowego z 2 neuronami w warstwie ukrytej [1]

Fig. 5. Face localization algorithm results for neural network classifier (2 hidden neurons) [1]

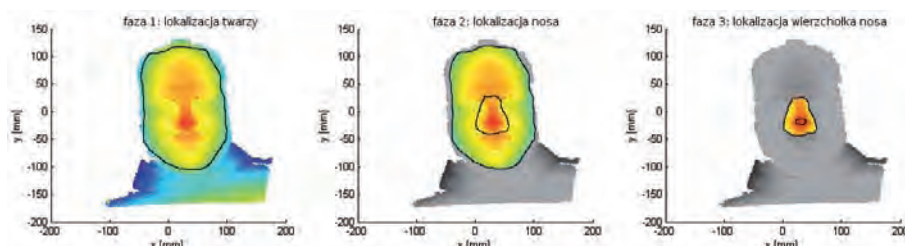
danych. Powoduje to, że wynikowy zbiór punktów jest niespójny i może zawierać nieciągłości. Ryzykiem związanym z sieciami neuronowymi jest możliwość zatrzymania procesu uczenia w minimum lokalnym, co daje złą jakość klasyfikacji. Zaletą jest dobra zdolność uogólniania, na którą można wpływać, manipulując liczbą neuronów ukrytych. Przykład zachowania obu klasyfikatorów pokazano na rys. 4 i 5.

Hierarchiczna lokalizacja regionów

Opisany algorytm pozwala zidentyfikować region odpowiadający wzorcowi w skanie. Ma jednak kilka wad. Pierwszą wadą jest koszt obliczeń. Im większy analizowany skan w stosunku do rozmiaru wzorca, tym więcej należy wygenerować deskryptorów. Dalej, w fazie dopasowywania, znacznie większa jest też liczba możliwych dopasowań. Wydłuża to czas obliczeń co najmniej (w zależności od metody dopasowywania) kwadratowo.

Drugim problemem jest skuteczność dopasowywania. Im większy skan, tym więcej deskryptorów, ale stała liczba odpowiadających wzorcowi. Zwiększa się prawdopodobieństwo błędów w dopasowaniu. Np. możliwa jest sytuacja, że w analizowanym skanie pojawią się obszary o bardzo podobnej strukturze: jeden w obszarze poszukiwanym, a drugi poza nim. Istnieje ryzyko, że dopasowany do deskryptora wzorca zostanie deskryptor z niewłaściwego obszaru, a to obniża jakość powstałego zbioru uczącego dla klasyfikatora.

Ulepszeniem, które redukuje wpływ tych czynników jest zhierarchizowanie procesu lokalizacji regionu. Algorytm lokalizacji regionu wykonywany jest wielokrotnie, a w kolejnych fazach identyfikowane są coraz mniejsze otoczenia poszukiwanego obszaru (rys. 6). Na wejściu podawany



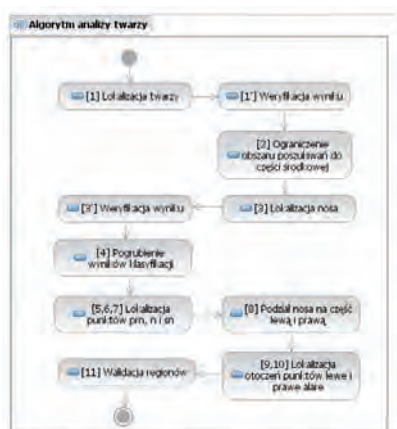
Rys. 6. Ilustracja koncepcji hierarchicznej lokalizacji regionów

Fig. 6. Hierarchical localization of regions. Concept

jest skan korpusu, a celem jest zlokalizowanie wierzchołka nosa. W pierwszej fazie identyfikowany jest obszar twarzy. Następnie lokalizowany jest nos, a dopiero w ostatniej fazie na obszarze nosa wyszukiwany jest wierzchołek. Tak zastosowany algorytm powoduje, że w każdej fazie stosunek powierzchni obszaru odrzucanego do obszaru wybranego jest niewielki. Problemem otwartym pozostaje decyzja, jak dobierać obszary w kolejnych fazach.

Pełny schemat systemu

Przypomnijmy, że postawionym celem było stworzenie systemu pozwalającego na zlokalizowanie punktów charakterystycznych nosa. Schemat centralnego modułu systemu mogącego realizować takie zadanie (rys. 7) zawiera algorytm hierarchicznej lokalizacji regionów zastosowanej do kolejnych obszarów twarzy. Uzupełniony został o dodatkowe fazy, których celem jest weryfikacja wyników częściowych oraz o fazy przyspieszające i poprawiające działanie algorytmu, korzystające z wiedzy dziedzinowej.



Rys. 7. Przykładowy schemat systemu analizy twarzy [1]

Fig. 7. Sample flow-chart of face analysis system [1]

Na wejście algorytmu podawany jest skan górnej części korpusu. Skan taki może być wstępnie przefiltrowany. Z takiego skanu w kroku 1 wyłuskiwany jest region twarzy. Wybór twarzy (wyznaczanej przez linię włosów i szczęki) na region identyfikowany (lokalizowany) w pierwszej kolejności wynika z założenia, że twarz jest największym obszarem, którego istnienie zakłada się w analizowanej scenie (obrazie 3D). Wybór regionu twarzy ze spójnego segmentu punktów jest specyficzny o tyle, że nie jest możliwe zdeterminowanie we wzorcu jej otoczenia (punktów sąsiadujących z poszukiwanym regionem). W związku z tym deskryptory budowane są wyłącznie w oparciu o punkty twarzy. Otoczenie pozostaje puste. W analizowanych obrazach nie jest ono puste. Wokół twarzy mogą wystąpić fragmenty ubrania czy inne części ciała. W związku z tym odpowiadające sobie punkty wzorca i obrazu mogą różnić się tym, że w tym pierwszym część „kubelków” będzie miała wartość 0. Aby zminimalizować ich wpływ na zachowanie algorytmu dopasowującego, zastosowano filtr AND, który wymusza, by pod uwagę były brane jedynie te elementy histogramu, które w obu deskryptorach są niezerowe.

Faza 1' służy weryfikacji rezultatu uzyskanego w kroku 1. Możliwe jest zastosowanie wielu kryteriów – najprostszym są wymiary geometryczne. Lepszym pomysłem jest zastosowanie sumarycznego kosztu dopasowania deskryptorów. Można przyjąć, że powyżej pewnego kosztu zgłaszany będzie błąd. We wcześniej istniejących rozwiązaniach koszt ten stosowany był już do identyfikacji osób. W związku z tym powinien być wystarczającym kryterium do odróżnienia twarzy od innych obiektów.

W kroku 3 identyfikowany jest nos. Wykorzystując wiedzę na temat położenia nosa, można zredukować obszar poszukiwań do części centralnej twarzy, co wykonywane jest w kroku 2. Służy on przyspieszeniu działania algorytmu.

Kroki 3' i 4 korygują działanie fazy 3. W 3' można wykonać podobną walidację, jak w kroku 1', tu jednak zastosowaną w stosunku do nosa. Krok 4 służy poprawie zachowania algorytmu. Eksperymenty pokazały, że w wyniku nieprecyzyjnego działania klasyfikatora częstym problemem było usuwanie skrajnych obszarów nosa, co powodowało, że skrzydełka nosa znajdowały się poza wyselekcjonowanym obszarem. Dla uniknięcia tego problemu wykonywane jest pogrubienie, tj. do punktów wyselekcjonowanych w fazie 3 dołączane są punkty, które znajdują się od nich nie dalej niż d (np. $d = 3$ mm).

W krokach 5, 6, 7 wykonywana jest lokalizacja regionów otaczających punkty na nosie: n , sn , prn . Średnica tych otoczeń to około 1–2 mm. W celu zidentyfikowania dokładnych położenia punktów wykorzystać można metody specjalizowane (opisane w części 3). Możliwe jest też wyliczenie punktów jako środków ciężkości otoczeń.

Powodem wprowadzenia kroku 8 jest uwzględnienie właściwości obrazów obrotów (ang. *spin images*) polegającej na tym, że nie zależą one od położenia bezwzględnego, a jedynie od względnego rozkładu punktów w otoczeniu. Właściwość ta powoduje, że nie są rozróżniane obszary o symetrycznych charakterystykach (np. deskryptory dla lewej części twarzy, ze względu na symetrię, będą identyczne jak deskryptory wyliczone dla strony prawej). Dlatego nie jest możliwe bezpośrednio zlokalizowanie skrzydełek nosa, gdyż deskryptory wyliczone dla lewego skrzydełka będą się mieszać z tymi dla prawego. Problem można obejść, stosując wyniki punktów 5, 6, 7 do wyznaczenia płaszczyzny rozdzielającej lewą część twarzy od prawej. W tak rozdzielonych obszarach wyszukanie regionów jest już możliwe (krok 9).

Ostatnim punktem procedury jest walidacja uzyskanych współrzędnych punktów. Można to zrobić wykorzystując informację o odległościach między punktami charakterystycznymi.

Wyniki praktyczne

W celu przetestowania systemu przygotowano testowy zbiór 25 obrazów trzech typów:

- 15 obrazów z bazy GavabDB zawierających skany twarzy z przodu – frontal1, frontal2, arriba, abajo, risa – dane dobrej jakości pochodzące z tego samego zbioru, na podstawie którego zbudowano bazę wzorców,
- 5 obrazów z bazy Mechatroniki – obrazy zawierają istotne nieciągłości w okolicach oczu i nosa,

3. 5 obrazów z bazy GavabDB zawierających twarze z profilu – derecha, izquierda.

Każdy z obrazów poddano wstępnej filtracji oraz wyborowi segmentu zawierającego twarz. Wynikowe dane poddano analizie algorytmem hierarchicznej lokalizacji regionów. Aby przetestować jakość algorytmu lokalizacji, w obrazach testowych oznaczono punkty należące i nie należące do poszukiwanego regionu. Tak zaklasyfikowane dane porównano z wynikami automatycznej klasyfikacji. Każdy z punktów analizowanej chmury przyporządkowany został do jednej z grup:

- A (true positive) – punkty w obu przypadkach zaklasyfikowane do poszukiwanego regionu,
- B (false negative) – punkty oznaczone jako należące do regionu, ale odrzucone w automatycznej klasyfikacji,

Tab. 1. Skuteczność lokalizacji nosa [1]

Tab. 1. Nose localization performance [1]

Grupa	Dokładność (acc)		Miara F ₁	
	Średnia	Odchylenie standardowe	Średnia	Odchylenie standardowe
1	96,56 %	0,71 %	79,32 %	3,88 %
2	94,58 %	1,15 %	63,01 %	10,26 %
3	94,24 %	2,32 %	60,97 %	18,21 %

Tab. 2. Skuteczność lokalizacji twarzy [1]

Tab. 2. Face localization performance [1]

Grupa	Dokładność (acc)		Miara F ₁	
	Średnia	Odchylenie standardowe	Średnia	Odchylenie standardowe
1	86,23 %	15,54 %	90,21 %	13,88 %
2	79,79 %	17,18 %	87,30 %	12,08 %
3	69,88 %	20,55 %	78,81 %	15,82 %

Tab. 3. Precyzja identyfikacji punktów charakterystycznych dla obrazów z grupy 1 [1]

Tab. 3. Characteristic points identification precision (test group 1) [1]

Punkt	Liczba obrazów	Przedział odległości [mm]					
		0-3	3-6	6-10	10-20	20-40	40-80
prn	14	71,43 %	21,43 %	7,14 %	0,00 %	0,00 %	0,00 %
n	14	42,86 %	42,86 %	7,14 %	0,00 %	0,00 %	7,14 %
sn	12	25,00 %	33,33 %	16,67 %	16,67 %	0,00 %	8,33 %
al _{lewe}	14	7,14 %	28,57 %	35,71 %	21,43 %	7,14 %	0,00 %
al _{prawe}	14	0,00 %	35,71	28,57 %	28,57 %	7,14 %	0,00 %

Tab. 4. Precyzja identyfikacji punktów charakterystycznych dla obrazów z grupy 2 [1]

Tab. 4. Characteristic points identification precision (test group 2) [1]

Punkt	Liczba obrazów	Przedział odległości [mm]					
		0-3	3-6	6-10	10-20	20-40	40-80
prn	4	2	1	1	0	0	0
n	4	2	0	0	0	0	2
sn	2	1	0	0	0	0	1
al _{lewe}	4	0	0	0	1	3	0
al _{prawe}	4	0	1	0	1	2	0

- C (false positive) – punkty zaklasyfikowane jako nie należące do regionu, ale uznane za twarz w wyniku działania programu,
- D (true negative) – punkty w obu przypadkach odrzucone.

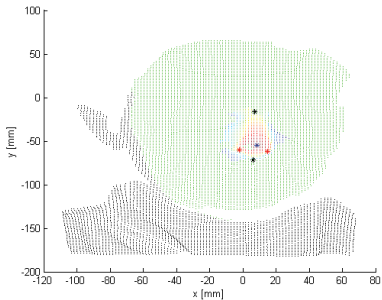
W oparciu o liczności powyższych grup dla każdego z obrazów obliczone zostały miary opisujące jakość klasyfikacji:

- dokładność: $acc = (|A|+|D|)/(|A|+|B|+|C|+|D|)$
- miara F1: $F1 = 2 (precision * recall)/(precision + recall)$, gdzie:
 - o $precision = |A|/(|A|+|C|)$
 - o $recall = |A|/(|A|+|B|)$

Wyniki lokalizacji dla twarzy i nosa przedstawiają tab. 1 i 2.

Tabele 3 i 4 pozwalają ocenić precyzję działania procedury identyfikacji punktów charakterystycznych nosa. Dla każdego z punktów pokazano, w jakim procencie obrazów wynik identyfikacji różni się o daną odległość od wyniku manualnego oznaczenia.

W komórkach zawarto ułamek obrazów należący do przedziału. Przedziały lewostronnie domknięte.



Rys. 8. Wynik działania procedury identyfikacji punktów charakterystycznych

Fig. 8. Sample results of characteristic points identification procedure

W komórkach zawarto liczbę obrazów przypisanych do przedziału. Przedziały lewostronnie domknięte.

Przykład działania pełnej procedury identyfikacji pokazuje rys. 8. Na zielono oznaczono obszar twarzy. Kolorami odpowiadającymi głębi oznaczono region nosa. Czarnymi punktami oznaczono punkty sn i n. Kolorem niebieskim zaznaczono prn, a czerwonymi al.

Podsumowanie

W artykule opisano procedurę lokalizacji regionów. Pokazuje ona praktyczne zastosowanie deskryptorów punktów w połączeniu z uczeniem maszynowym dla potrzeb analizy obrazów trójwymiarowych. Omówiono szczegółowo jej najważniejsze aspekty. Część problemów pozostawiono do doprecyzowania w kolejnej części artykułu. Zademonstrowano również schemat procedury identyfikacji punktów charakterystycznych nosa w oparciu o powyższy algorytm. Pokazano wyniki praktycznych eksperymentów.

Bibliografia

1. Kuśmierczyk T.: *Deskryptory punktów w analizie morfologicznej obrazów trójwymiarowych twarzy*. Praca inżynierska, Politechnika Warszawska, Warszawa 2010.

Automatic nose measurement system, part 5

Abstract: The purpose of designed system is to analyze and recognize three-dimensional face images. Using known techniques, algorithms and tools I am aiming to retrieve nose parameters directly from 3D scans. Current part is devoted to describe innovative method for region identification in 3D scans. Spin images and machine learning are applied. Overall design is shown. Some of the most important aspects are described in details.

Keywords: spin image, machine learning, face recognition, cloud of points

Tomasz Kuśmierczyk

Ukończył studia pierwszego stopnia na Wydziale Elektroniki i Technik Informacyjnych PW na kierunku Informatyka. Kontynuuje studia na poziomie magisterskim na Wydziale Matematyki, Informatyki i Mechaniki UW. Jednocześnie studiuje na Wydziale Fizyki UW na kierunku Fizyka. Zainteresowania naukowe: sztuczna inteligencja i uczenie maszynowe w zastosowaniu do analizy i rozpoznawania obrazów.

e-mail: tk290810@okwf.fuw.edu.pl



REKLAMA



STEROWNIKI.PL

Sterowanie w automatyce portal branżowy



- Aktualności z branży • Pliki • Giełda
- Katalog firm • Baza wiedzy • Praca
- Kalendarz imprez • Kursy • Forum

Wyślij zapytanie ofertowe



i wygraj
pendrive

Reklama Twojej firmy od



490 zł.
netto za rok

ponad
2500 klientów
czekających na
Twoją ofertę

sterowniki.pl Sp. z o.o.
tel. (0-22) 499-88-39
fax (0-22) 205-09-11
www.sterowniki.pl