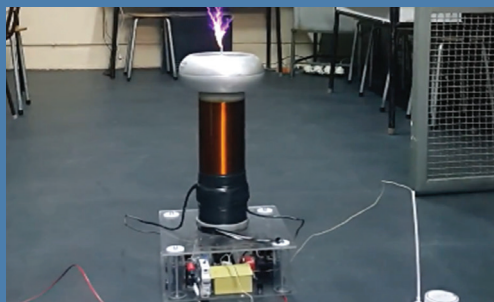


Problemy Współczesnej Inżynierii

Wybrane zagadnienia elektroniki i inżynierii biomedycznej



redakcja
Paweł A. Mazurek
Tomasz N. Kołtunowicz
Andrzej Kociubiński
Piotr Z. Filipek
Sebastian Styła

Politechnika Lubelska
Lublin 2018

Problemy Współczesnej Inżynierii

Wybrane zagadnienia elektroniki i inżynierii
biomedycznej



Politechnika Lubelska
Wydział Elektrotechniki i Informatyki
ul. Nadbystrzycka 38A
20-618 Lublin

Problemy Współczesnej Inżynierii

Wybrane zagadnienia elektroniki i inżynierii
biomedycznej

redakcja:

Paweł A. Mazurek

Tomasz N. Kołtunowicz

Andrzej Kociubiński

Piotr Z. Filipek

Sebastian Styła



Politechnika Lubelska
Lublin 2018

Recenzenci:

członkowie Komitetu Naukowego VII Sympozjum Elektryków i Informatyków

Skład: Tomasz Kołtunowicz, Paweł A. Mazurek

Publikacja wydana za zgodą Rektora Politechniki Lubelskiej

© Copyright by Politechnika Lubelska 2018

ISBN: 978-83-7947-304-5

Wydawca: Politechnika Lubelska

ul. Nadbystrzycka 38D, 20-618 Lublin

Realizacja: Biblioteka Politechniki Lubelskiej

Ośrodek ds. Wydawnictw i Biblioteki Cyfrowej

ul. Nadbystrzycka 36A, 20-618 Lublin

tel. (81) 538-46-59, email: wydawca@pollub.pl

www.biblioteka.pollub.pl

Druk: TOP Agencja Reklamowa Agnieszka Łuczak

www.agencjatom.pl

Elektroniczna wersja książki dostępna w Bibliotece Cyfrowej PL www.bc.pollub.pl

Nakład: 60 egz.

SPIS TREŚCI

	PRZEDMOWA	7
1	AUTONOMICZNY SYSTEM NAMIERZANIA OBIEKTÓW <i>MICHAŁ STYŁA</i>	9
2	PROJEKT I WYKONANIE PÓŁPRZEWODNIKOWEJ CEWKI REZONANSOWEJ <i>KAROL PAWELEC, PAWEŁ OKAL, PRZEMYSŁAW ROGALSKI</i>	22
3	PROJEKT I WYKONANIE STANOWISKA LABORATORYJNEGO DO SYMULACJI LAPAROSKOPOWYCH <i>ERYK PRZERWA</i>	33
4	KONCEPCJA UKŁADU ZAWIESZENIA DO POJAZDU KLASY FORMULA STUDENT W PROGRAMIE VI – MOTORSPORT <i>ARKADIUSZ GITA</i>	45
5	POLIMERY BIODEGRADOWALNE A ZASTOSOWNIA MEDYCZNE <i>MAGDALENA KOZIEŁ</i>	56
6	ZASTOSOWANIE FOLIOWYCH CZUJNIKÓW PIEZOELEKTRYCZNYCH W MONITOROWANIU FALI TĘTNA W WYBRANYCH PUNKTACH UKŁADU TĘTNICZEGO <i>AGATA PAWŁOWSKA, PAWEŁ SURTEL</i>	68
7	ZASTOSOWANIE NANOTECHNOLOGII W MEDYCYNIE <i>BERNADETTA MICHALIK</i>	83
8	ZASTOSOWANIE CZUJNIKA KWARCOWEGO W URZĄDZENIU DO BADANIA KROPLI KRWI W DIAGNOSTYCE CHOROBY ALZHEIMERA <i>MARLENA ZARZECZNA, MICHAŁ SURTEL</i>	94
9	MODEL I SYMULACJE TRANZYSTORA MOSFET Z WĘGLIKA KRZEMU <i>ALEKSANDRA PEDRYC, AGNIESZKA MARTYCHOWIEC</i>	103
10	PROJEKT I TECHNOLOGIA ZŁĄCZA SCHOTTKY’EGO Z WĘGLIKA KRZEMU <i>AGNIESZKA MARTYCHOWIEC, ALEKSANDRA PEDRYC</i>	117
11	PROJEKT I TECHNOLOGIA KONDENSATORÓW GRZEBIENIOWYCH Z MIEDZI <i>MACIEJ SZYPULSKI, DAWID ZARZECZNY, TOMASZ LIZAK, KRZYSZTOF MUZYKA</i>	128

12	<i>BADANIE WYTRZYMAŁOŚCI POŁĄCZEŃ DRUTOWYCH DLA WYBRANYCH METALIZACJI</i>	141
	TOMASZ LIZAK, KRZYSZTOF MUZYKA, DAWID ZARZECZNY, MACIEJ SZYPULSKI	
13	<i>PROJEKT I TECHNOLOGIA KONDENSATORÓW GRZEBIENIOWYCH DO MONITOROWANIA HODOWLI KOMÓREK</i>	155
	DAWID ZARZECZNY, MACIEJ SZYPULSKI, TOMASZ LIZAK, KRZYSZTOF MUZYKA	
14	<i>WYBRANE PROBLEMY W DIAGNOSTYCE OBWODU STEROWANIA SILNIKIEM TDI</i>	169
	KRZYSZTOF KOWALCZYK	
	<i>SPONSORZY I INSTYTUCJE WSPIERAJĄCE VII SYMPOZJUM NAUKOWE ELEKTRYKÓW I INFORMATYKÓW</i>	181
	<i>PATRONI VII SYMPOZJUM NAUKOWEGO ELEKTRYKÓW I INFORMATYKÓW</i>	182

PRZEDMOWA

Szanowni Uczestnicy i Sympatycy Sympozjum, Czytelnicy

Kolejne, siódme Sympozjum Naukowe Elektryków i Informatyków SNEiI 2017 już za nami. Cykliczna studencka inicjatywa Kół Naukowych oraz Samorządu Studenckiego Politechniki Lubelskiej stała się ważnym wydarzeniem w życiu akademickim naszego Wydziału.

Patronat honorowy nad Sympozjum objął Lubelski Oddział Stowarzyszenia Elektryków Polskich. W gronie patronujących instytucji było także Polskie Towarzystwo Informatyczne – Koło w Lublinie oraz Lubelski Oddział Polskiego Towarzystwa Zarządzania Produkcją. Po raz kolejny mogliśmy też liczyć na wsparcie naszych uczelnianych władz – Jego Magnificencji Rektora Politechniki Lubelskiej prof. dr hab. inż. Piotra Kacejko oraz Pani Dziekan Wydziału Elektrotechniki i Informatyki prof. dr hab. inż. Henryki D. Stryczewskiej.

Gorące podziękowania składamy naszym zaproszonym prelegentom. Otwarty panel rozpoczął wykład prof. dr inż. Ryszarda Sikory, dhc Zachodniopomorskiego Uniwersytetu Technologicznego w Szczecinie. Kolejnym prelegentem był dr inż. Waldemar Gawron, koordynator prac badawczych VIGO System S.A. Ostatnią prezentacją tego panelu "Innowacyjne rozwiązania i systemy teletechniczne KZŁ Bydgoszcz" wygłosili Andrzej Weidemann oraz Jerzy Pietrzyk z Kolejowych Zakładów Łączności z Bydgoszczy – Grupa PKP Informatyka.

Tradycyjnie, celem naszego Sympozjum była wymiana wiedzy i doświadczeń praktycznych w gronie młodych członków społeczności akademickiej, przedstawicieli lokalnych władz oraz przemysłu. Tematyka Sympozjum analogicznie jak rok wcześniej obejmowała obszary dziedzin techniki powiązanych z kierunkami studiów realizowanymi na naszym Wydziale – elektrotechniką, mechatroniką, informatyką oraz inżynierią biomedyczną.

Nasze Sympozjum jest wspianą okazją do przeglądu osiągnięć i oceny dorobku inżyniersko-naukowego naszych studentów. Duża liczba głoszonych referatów (około czterdzieści) mobilizują członków kół naukowych i kadrę uczelni do działań wspierających. Oddana w Państwa ręce monografia jest opracowaniem naukowym zawierającym artykuły dotyczące działań jakie studenci naszej uczelni realizują w ramach własnych projektów i prac badawczych.

Jako redaktorzy i współorganizatorzy SNEiI 2017 cieszymy się z dobrego odbioru naszej serii „*Problemy Współczesnej Inżynierii*” i raz jeszcze dziękujemy wszystkim instytucjom za udzielone wsparcie.

Redaktorzy

AUTONOMICZNY SYSTEM NAMIERZANIA OBIEKTÓW

1. WSTĘP

Gwałtowny rozwój technik mikroprocesorowych na początku lat 70. XX. w. spowodował głęboki stopień automatyzacji wielu procesów przemysłowych ograniczając w ten sposób udział człowieka do nadzoru i kontroli.

Projekt miał na celu opracowanie urządzenia w formie robota – wieżyczki, którego działanie w skrócie będzie bazować na systemie mikroprocesorowym sterującym, za pomocą czujek bądź czujników, dwoma silnikami elektrycznymi, nakierowując się w ten sposób na obiekty o żądanych parametrach fizycznych. Jeden z silników odpowiada w takiej konfiguracji za ruch horyzontalny, a drugi za wertykalny ramienia z wyposażeniem. Urządzenie to zostało skonstruowane głównie pod kątem zastosowań paramilitarnych takich jak np. ASG (ang. Air Soft Gun) oraz paintball lub ochrony mienia i monitoringu.

Niewiele jednak osób zdaje sobie sprawę z tego, że pomimo znacznego upowszechnienia się elektroniki i automatyki w obecnych czasach większość technologii tego typu, zanim trafi do przemysłowego czy cywilnego użytku ma swoją genezę w wojsku i jest determinowana przez światowe wyścigi zbrojeń i ciągle zmieniające się warunki pola walki, a nie tylko przez chęć podniesienia komfortu i standardów życia [1].

2. CEL I ZAŁOŻENIA PROJEKTOWE

Właściwie podejście do tematu jest uwarunkowane w pierwszej kolejności od analizy bliźniaczych rozwiązań technicznych dostępnych na rynku oraz od wyeksponowania ich poszczególnych składowych. Znacznym utrudnieniem okazał się poziom dostępu do specyfikacji technicznych urządzeń pokrywających się formą z założonym projektem. Wynika to z faktu, że tego typu konstrukcje są wytwarzane przez zakłady zbrojeniowe i producentów profesjonalnych urządzeń do systemów alarmowych na zamówienie podmiotów różnego rodzaju i w każdym przypadku ilość informacji na ich temat w literaturze i Internecie jest ograniczona i ściśle kontrolowana. Pomimo tego udało się wyszczególnić charakterystyczne rozwiązania wg których można przystąpić do doboru technologii [1]:

- The Super aEgis II
- Phalanx CIWS
- autonomiczna wieżyczka strażnicza z modulem LIDAR i czujką PIR.

¹ Politechnika Lubelska, Wydział Elektrotechniki i Informatyki

Dwie pierwsze pozycje to urządzenia wykonane przez profesjonalne zakłady zbrojeniowe, które doczekały się już wdrożenia (rys. 1, 4). Ostatnia natomiast jest projektem osoby prywatnej do zastosowań komercyjnych. Potwierdziło to fakt, że wszystkie komponenty potrzebne do wykonania projektu są ogólnie dostępne, a wykonanie go – możliwe.

Analiza rozwiązań pozwoliła na podzielenie projektu na sześć zasadniczych części. Składać się więc na niego będą [1]:

- jednostka sterująca
- silniki napędowe wieży
- czujki / czujniki
- system zasilania
- podstawa / konstrukcja wsporcza
- typ wyposażenia jakim urządzenie będzie dysponować.



Rys. 1. The Super aEgis II produkcji Dodaam Systems Ltd. (źródło: [5])

3. JEDNOSTKA STERUJĄCA

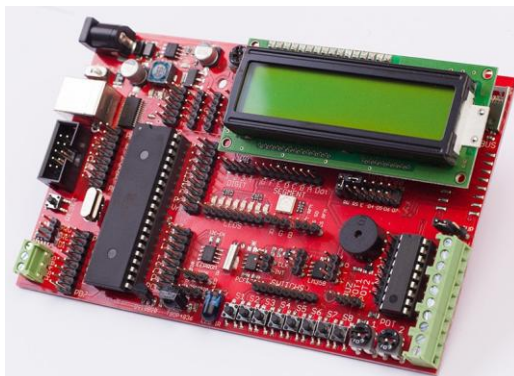
Przez ostatnie kilkanaście lat pojawiło się na rynku wiele platform sprzętowych, które mogłyby posłużyć celom rozważanego projektu. Przychylnym faktem jest odejście od dedykowanych pojedynczym procesorom języków typu Asembler tworzących kod maszynowy dla mikroprocesorów na podstawie kodu źródłowego na rzecz bardziej uniwersalnych platform programistycznych jakimi są na przykład języki wysokiego poziomu takie jak C, C++ czy Python. Zwiększa to elastyczność i możliwości wprowadzania zmian czy nawet zmienienia

platformy na inną mając tylko na uwadze różnicę w architekturach tych systemów mikroprocesorowych [2].

Dobór konkretnej jednostki sterującej był uwarunkowany od wielu niekiedy sprzecznych ze sobą parametrów. Na podstawie najważniejszych z nich zostały wytypowane następujące pozycje:

- rodzina platform komputerowych Raspberry Pi
- platforma Intel Edison
- rodzina mikrokontrolerów AVR ATmega (rys. 2).

Przyjęte założenia sprowadzały się do wybrania jednostki, która pozwoli na aplikację mobilną (energooszczędną) przy czym moc obliczeniowa będzie wystarczająca do odpowiednio szybkiej reakcji urządzenia na bodźce z otoczenia. Istotną kwestią był tu też rodzaj i liczba oferowanych portów. Z punktu widzenia sterowania napędami wieży duży nacisk został także położony na możliwość użycia modulacji PWM.



Rys. 2. Układ zastosowany w projekcie – płytki EvB 5.1 v5 z procesorem ATmega32A [6]

Modulacja w ujęciu ogólnym jest procesem manipulowania danym parametrem sygnału (zazwyczaj jednym) przy zachowaniu stałości pozostałych. W przypadku gdy w PWM steruje się wypełnieniem przebiegu, reszta parametrów takich jak amplituda czy częstotliwość pozostaje niezmienna. Metodę taką można wykorzystywać do sterowania prędkością obrotową silników prądu stałego lub przemieszczeniem silników skokowych i serwonapędów [1].

4. SILNIKI NAPĘDOWE WIEŻY

Z punktu widzenia projektu zastosowane silniki będą musiały spełniać kilka zasadniczych kryteriów. O ile we wszelkiego rodzaju przemyśle nacisk kładziony jest przede wszystkim na możliwości sterowania prędkością obrotową, tutaj będzie on wymaganiem drugorzędym, a w szczególnych przypadkach pomijalnym. Największą wartość będą przedstawiać silniki elektryczne o jak najko-

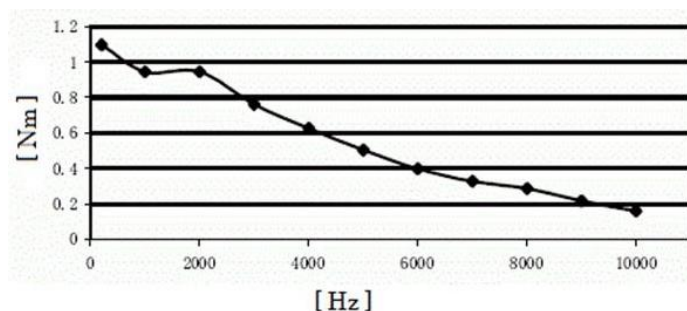
rzystniejszym i precyzyjnym sterowaniu przemieszczeniem np. kątem obrotu wałka.

Jednostki napędowe, które zostały poddane analizie i porównaniom zgodnie z wcześniejszym opisem to [1]:

- silniki dwufazowe
- serwonapędy
- silniki skokowe zwane też krokowymi bądź impulsowymi.

Wykorzystanie silnika dwufazowego wiązałoby się z koniecznością dobrania bądź zaprojektowania falownika (ze względu na zasilanie akumulatorowe). Przy obecnej sprawności tych urządzeń dochodzącej niekiedy do 95% nie stanowiło by to problemu (np. drony) gdyby nie fakt, że silnik ten sam w sobie oferuje dosyć małą precyzję sterowania przemieszczeniem. Dodać do tego można ryzyko utraty kontroli nad maszyną przez zjawisko zwane niesamohamownością czyli sytuację dla, której przy zerowym poziomie sygnału sterującego (sterowanie fazowe) jednostka zaczyna zachowywać się jak silnik jednofazowy. W silnikach tych ponadto nie jest możliwe wykorzystanie modulacji PWM.

Serwonapędy w przeciwieństwie do silników dwufazowych oferują bardzo wysokie precyzje regulacji przemieszczenia. Pomimo zastosowania w tych urządzeniach prostego w zrozumieniu sterowania PWM sprostanie wysokim wymaganiom dokładnościowym wymagałoby jednak dodatkowego użycia nietypowych sterowników (regulatorów) dysponujących modelem matematycznym dla danego serwonapędu. Możliwości, które będzie projektantowi stwarzać serwonapęd będą wzrastać wraz z zapewnianą mu przez sterownik mocą obliczeniową. Posiadają też często ograniczenia w postaci maksymalnego kąta obrotu nie przekraczającego 360° , a nawet 180° .



Rys. 3. Charakterystyka pracy przykładowego, czteroprzewodowego silnika skokowego FL57STH56-2804A o momencie trzymającym 1,2 Nm [7]

Alternatywnym silnikiem dla serwonapędu do tego rodzaju aplikacji jest silnik skokowy. Podobnie jak serwonapęd oferuje bardzo wysokie standardy sterowania z wykorzystaniem zewnętrznego regulatora. Nie posiada jednak ograniczeń przy obracaniu się jak jego „konkurent”, a jedynie najmniejszy kąt obrotu

jaki jest w stanie wykonać (tzw. krok). Często rekompensowane jest to przez odpowiedni sterownik, który jest w stanie zwiększyć jego rozdzielczość. Cechuje się dużą mocą i momentem obrotowym. Jego największymi wadami są stosunkowo niskie prędkości obracania się oraz silna zależność generowanego momentu od częstotliwości podawanego sygnału sterującego co widać na rysunku 3. Jedną z metod zwiększania dynamiki pracy silników tego typu jest ich wstępny rozruch za pomocą sygnałów o niższych częstotliwościach. Zalety silnika skokowego i proste metody kompensacji jego wad sprawiły, że to właśnie on został zastosowany w projekcie [3].

5. CZUJKI / CZUJNIKI

Reakcje urządzenia na bodźce wysyłane ze środowiska, w którym będzie się znajdowało będą ściśle uzależnione od rodzaju zastosowanych czujników. Tak jak jednostka sterująca jest dla niego swojego rodzaju „mózgiem”, silniki „mięśniami” tak czujniki będą pełnić rolę „oczu i uszu” i ich poprawny dobór będzie mocno pożądanym. Z uwagi na charakter wykonywanej pracy muszą być to układy skupiające się przede wszystkim na parametrach fizycznych typowych dla organizmów żywych. Należy więc przede wszystkim skupić się na „wspólnym mianowniku” dla wszystkich istot żywych (np. człowiek, zwierzę) i coraz częściej obiektów nieożywionych (np. drony) czyli ruchu i temperaturze.

Nie biorąc pod uwagę coraz bardziej popularnych czujników dualnych wyszczególnić można następujące sensory / detektory rozpoznające żądane wielkości fizyczne [1]:

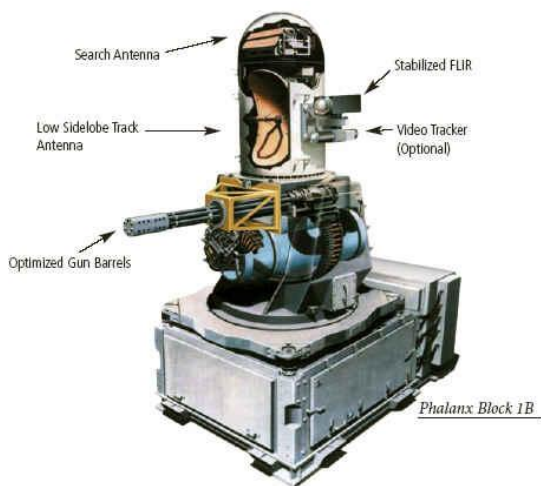
- czujniki ultradźwiękowe
- czujniki radarowe i lidarowe
- czujki PIR (pasywne detektory promieniowania podczerwonego).

Czujniki ultradźwiękowe są bardzo dobrym rozwiązaniem dla założonego urządzenia. W pewnych zastosowaniach są one jednak niewystarczające głównie ze względu na swój krótki zasięg działania w porównaniu z czujnikami wykorzystującymi fale elektromagnetyczne. Między innymi ze względu na tę cechę są stosowane głównie w systemach unikania kolizji między obiektami bądź medycynie (USG). Większość ich szkodliwych zależności od czynników atmosferycznych wynika z natury fali mechanicznej.

Czujniki radarowe i lidarowe są często przedstawiane razem ze względu na ich podobne zastosowanie. Służą do przeczesywania przestrzeni, wykrywania znajdujących się w niej obiektów oraz określania ich położenia i odległości z tym, że robią to za pomocą dwóch różnych zakresów promieniowania elektromagnetycznego. Są to najbardziej precyzyjne i dalekosiężne ale i jednocześnie najbardziej wymagające przedstawione w tym artykule czujniki. Dowodem na to są wszelkie wdrożenia tych technik w systemach antybalistycznych i przeciwcwilotniczych na całym świecie takich jak np.: Phalanx CIWS, MEADS, Patriot,

Hawk czy Nike-Hercules. Nasuwa to wniosek, że zastosowanie ich w tak małej aplikacji może sprawić problemy z zagwarantowaniem zasilania oraz miejsca do instalacji. Względem założonego projektu mogło to wywołać problemy z bezkolizyjnym rozmieszczeniem wszystkich elementów, czyli silników wieży i dodatkowych anten razem z ich ewentualnymi napędami. Zastosowanie mapowania przestrzeni za pomocą triangulacji spotęgowało by ten skutek [4].

Najrozsądniejszym rozwiązaniem okazały się detektory promieniowania podczerwonego. Ich kompaktowa budowa przy większym zasięgu od czujników ultradźwiękowych oraz możliwość prostego rozszerzenia ich kąta widzenia za pomocą tzw. soczewki Fresnela bez konieczności zastosowania dodatkowych obrotowych elementów sprawiła, że to właśnie one zostały zainstalowane w urządzeniu. Pobór prądu dla pojedynczego detektora (PIR HC-SR501) w stanie czuwania wyniósł zaledwie $50\mu\text{A}$ [1].



Rys. 4. Wykorzystujący radarowe techniki namierzania obiektów system przeciwlotniczy Phalanx CIWS produkcji Raytheon Company [8]

6. SYSTEM ZASILANIA

Każdy układ elektryczny i elektroniczny zastosowany w urządzeniu został dobrany na napięcie 12 V. Jest to jeden z najpopularniejszych i najczęściej stosowanych poziomów. System składa się z [1]:

- akumulatora żelowego AKUELL BT-02-105 o napięciu 12V, maksymalnym prądzie 105A i pojemności 7Ah (1 szt.)
- układów Pololu A4988 (2 szt.) służących do sterowania (część logiczna napięcie z zakresu 3 do 5,5V) i zasilania (część zasilająca napięciem 8÷35V) wybranych silników skokowych JK42HS40-0504 (2 szt.)

- przetwornic step-down opartych na układzie XL4011SE1 (2 szt.) – służą one wyłącznie do zasilenia części logicznej sterowników Pololu A4988
- ujednoliconego systemu połączeń elektrycznych: listew kołkowych użytych m.in. jako punkty zbiorcze zasilania i masy urządzenia (1 kpl.), przewodów połączeniowych do listew kołkowych żeńsko-żeńskich (3 kpl.) oraz uzupełniających je przewodów do płytek testowych męsko-męskich (1 kpl.).

Płytki EvB 5.1 v5 (1 szt.) oraz detektory PIR HC-SR501 (13 szt.) posiadają już wbudowane stabilizatory na zadane napięcia i nie jest konieczne stosowanie żadnych układów pośredniczących w ich zasilaniu. W przypadku wprowadzania przez dobrane przetwornice zakłóceń do obwodów sterowników jako element zastępczy zostały wybrane stabilizatory na napięcie 5V i prąd 1,5A L780SCV3 (2 szt.).

7. WYPOSAŻENIE, PODSTAWA / KONSTRUKCJA WSPORCZA

W zgodzie z konstytucją i normami życia społecznego obowiązującymi na terenie danego państwa należy podejść do doboru wyposażenia zgodnie z prawem i z daleko idącym dystansem. Pomijając sam fakt zastosowania broni palnej czy nawet repliki w postaci urządzenia miotającego kulki BB 6mm z tworzywa sztucznego, porcelany lub metalu (AEG ASG) bądź tzw. markera używającego jako amunicji kul gumowych wypełnionych farbą (paintball) należy uwzględnić szereg innych zależności.

W związku z dydaktycznym celem niniejszej pracy i stosunkowo niewielkiej mocy użytych napędów jako urządzenie wskazujące poprawność działania zostanie użyte źródło światła o niewielkiej masie i własnym zasilaniu bądź w analogiczny sposób wskaźnik laserowy. Zniweluje to także pośrednio konieczność tworzenia dodatkowych połączeń stałych pomiędzy korpusem, a elementami obrotowymi [1].

Urządzenie zaopatrzone w elementy obracające się musi posiadać stabilną i solidną podstawę. Musi ona być zdolna udźwignąć masę własną i tą, która ma być na nią nałożona przy czym należy również uwzględnić możliwości przeciążeń wynikające z praw fizyki takie jak, np. bezwładność, występowanie dźwigni, albo trzecia zasada dynamiki Newtona odniesiona do odrzutu powstałego przy strzelaniu z broni palnej. Wzrost liczby i udział tych i innych czynników komplikuje model matematyczny urządzenia i nakłada na konstruktora wymóg użycia coraz to mocniejszych materiałów oraz technik wykonania.

Mechanika jako dziedzina wiedzy przy ogromie zagadnień elektrycznych i informatycznych nie była tematem przewodnim tej pracy, a więc wnikanie w prawa nią rządzące będzie wyjątkowo zdawkowe. Dla zadań prostych, nie stwarzających wielkiego ryzyka znaczenie mają materiały o niewielkich wymia-

rach, lekkie i podatne na kształtowanie. Zagadnienia na, które należało zwrócić uwagę przy wyborze komponentów to [1]:

- jeżeli urządzenie nie będzie przytwierdzone do podłoża na stałe czyli, np. za pomocą śrub, nitów, spawów itd. do podłogi, ściany, sufitu bądź innej konstrukcji o znacznie większej masie to należy zagwarantować mu niezbędną liczbę punktów podparcia tak aby podczas wykonywania ruchu bądź odrzucie nie straciło równowagi. Im więcej punktów podparcia tym podstawa staje się stabilniejsza;
- jeżeli urządzenie dysponuje masą przemieszczającą się u jego szczytu, wytrącającą je z równowagi, należy również dążyć do obniżenia środka ciężkości lub odpowiednio dociążyć urządzenie;
- jeżeli urządzenie ma nierównomierny rozkład masy względem osi obrotu urządzenia (w tym przypadku z powodu wykonania elementów łączących ze sobą silniki skokowe) należy go zrównoważyć. Przy wariantach szybko obrotowych może to wprawiać cały układ w niepożądane wibracje.

Ostatecznie jako konstrukcja wsporcza zostanie użyty statyw na kamerę bądź aparat fotograficzny. Jest on wykonany z lekkich i łatwych w obróbce materiałów (głównie aluminium) co pozwoliło na łatwe dostosowanie go do potrzeb urządzenia i transport. Regulowany trójnóg pozwala na wygodne badanie zachowania się detektorów w zależności od wysokości względem podłoża. Ze względu na fakt, że urządzenia audio-video mocowane na takich statywach są stosunkowo ciężkie stworzyło to możliwość na bezpieczny montaż wszystkich elementów [1].



Rys. 5. Wzmocniony statyw na aparat fotograficzny (źródło: [1])

Pomimo solidnej konstrukcji statyw musiał być dodatkowo wzmocniony za pomocą trzech aluminiowych teowników (po jednym między każdą parę nóg). Usztywniło to konstrukcję i pozwoliło na stworzenie gniazda pod akumulator żelowy. Spowodowało to, że statyw stał się przez to częściowo nieskładany (rys. 5) [1].

8. KONSTRUKCJA UKŁADU MECHANICZNEGO I ELEKTRYCZNEGO URZĄDZENIA

Po wykonaniu wzmocnień statywu nastąpiło zainstalowanie opasującej górny poziom mocowań nóg statywu obręczy z detektorów PIR (rys. 6). Zostało do tego użytych 12 spośród 13 czujek. Utworzona w ten sposób obręcz odpowiada za ruch wieżyczki lewo / prawo. Pozostała jedna sztuka została użyta do pokrycia przestrzeni nad urządzeniem wymuszając ruchy góra / dół (rys. 8). Idea obręczy polega na gęstym wypełnieniu przestrzeni dookoła wieży za pomocą siatki promieni podczerwonych tak, aby mogła ona z odpowiednią co do zasięgu dokładnością określać kierunek, z którego może pochodzić zagrożenie. Przy zastosowaniu odpowiednio dużej ilości czujek tzn. większej od 3 sztuk promienie sąsiadujących ze sobą układów zaczynają wzajemnie na siebie nachodzić co przy wybranej liczbie 12 sztuk spowodowało powstanie 24 stref zadziałania co na dystansie 7m (maks. zasięg detektorów) i kącie widzenia 100° (przy użyciu soczewki Fresnela) jest wystarczające. Powodem wybrania takiej liczby detektorów był także zamiar otrzymania jak najrówniejszego podziału przestrzeni pokrywanej przez pojedynczą czujkę i dwie sąsiadujące ze sobą co zagwarantowałyby równomierny obrót urządzenia za poruszającym się obiektem. Otrzymana proporcja to 1:2, czyli na zamianę: pojedyncza czujka 20°, para 40°, pojedyncza 20° i tak dalej.

Na podstawie kombinacji zer i jedynek wysyłanych przez czujki do mikrokontrolera, silniki będą się ustawiać w ściśle określonych położeniach - bezpośrednich po otrzymaniu sygnału z pojedynczej czujki i pośrednich, przy otrzymaniu sygnału z dwóch lub więcej czujek.

Montaż silników został zrealizowany za pomocą aluminiowych obejm i podstawki statywu w przypadku silnika realizującego ruch horyzontalny, a w przypadku wertykalnego za pomocą aluminiowego hub'a i dwóch mocowań dedykowanych do danego silnika. Jedno mocowanie utrzymuje silnik w żądanej pozycji zaś drugie zostało dodane w celu stworzenia platformy montażowej na czujkę PIR oraz ewentualną masę wywarzającą układ. Sposób mocowania silników obrazuje rysunek 6 [1].



Rys. 6. Obręcz z detektorami i silniki skokowe zainstalowane na statywie [1]

Po wypełnieniu przestrzeni pomiędzy nogami statywu za pomocą podzespołów niezbędnych do działania czyli: akumulatora żelowego, układu mikroprocesorowego, przetwornic i sterowników układ został połączony za pomocą wcześniej już wspomnianego systemu listew kołkowych i zworek. Ułatwi to wstępną konfigurację urządzenia i testy oraz możliwe zmiany. Kompletną już konstrukcję mechaniczną i elektryczną można prześledzić na rysunku 8.

Duża liczba połączeń elektrycznych wmusiła zorganizowanie obwodów w przejrzysty sposób dlatego też, aby utworzyć z pojedynczych połączeń sygnałowych grupy funkcyjne, zastosowano listwy kołkowe oraz zróżnicowaną kolorystykę przewodów. Zastosowany przykład listwy kołkowej z miejscami na 40 połączeń przedstawiono na rysunku 7. Analogicznie zostały zrealizowane połączenia zbiorcze zasilania układów napięciem 12V z akumulatora żelowego oraz masa urządzenia. Listwy są zalutowane i wpinane na końcówki biegunowe akumulatora za pomocą standardowych złączek. Z uwagi na fakt, że konstrukcja jest w dużej mierze złożona z elementów aluminiowych przewodzących prąd elektryczny w miejscach newralgicznych dookoła punktów do, których przykręcone zostały płytki nałożono warstwę izolacji. Przestrzeń wewnątrz wieży tworząca ostrosłup stwarza dogodne warunki do wykonania obudowy okablowania za pomocą nóg statywu [1].



Rys. 7. Listwa kołkowa odpowiadająca za komunikację detektorów z mikrokontrolerem [1]



Rys. 8. Kompletny układ mechaniczny i elektryczny [1]

9. PODSUMOWANIE

Postępując według analizy dokonanej w pracy udało się opracować i wykonać prototyp robota – wieżyczki wykonującego założone czynności. Zastosowane silniki krokowe okazały się wystarczające go manewrowania lekkimi przedmiotami takimi jak wskaźnik laserowy bądź latarka. Dalsze zwiększanie obciążenia byłoby niewskazane ze względu na ryzyko przeciążenia uszkodzenia. Za instalowanie pełnowymiarowej repliki broni palnej stałoby się realne dopiero po

przeskalowaniu projektu, czyli zainstalowaniu silników o większym momencie obrotowym oraz modyfikacji zasilania układu.

Zmiana rozmiaru kroku na mniejszy od 1/2 w stosunku do liczby zastosowanych czujek okazała się zbędna. Nie przyniosła ona bowiem wystarczających korzyści względem strat wynikających z ograniczenia prędkości obrotowej silników. Zastosowanie materiałów miękkich, łatwych w obróbce przyniosło zarówno korzyści jak i straty.

Wykorzystanie aluminium przyspieszyło i uprościło proces dostosowania konstrukcji do podzespołów oraz ułatwiło nanoszenie korekt. Było to jednak powodem deformacji elementów w procesie montażu. Ze względu na ten fakt, urządzenie stało się wrażliwe na uszkodzenia mechaniczne np. lekkie uderzenie w obręcz z czujkami spowodowało wgięcie obręczy. Wywołało to zmiany w jej kątach widzenia oraz zaburzyło równomierny rozkład siatki promieni podczerwonych co w przełożeniu na praktykę objawiło się w rozbieżnościach pomiędzy zajmowaną pozycją silnika, a realnym położeniem namierzanego obiektu.

Sposób wykonania połączeń elektrycznych silnika krokowego oraz detektora odpowiadającego za ruch wertykalny wieżyczki z resztą podzespołów okazał się wystarczający to testów działania ale nie wystarczający do ciągłej pracy. W momencie wykonania przez silnik krokowy odpowiadający za ruch horyzontalny dwóch pełnych obrotów w jedną stronę przewody uległy nadmiernemu skręceniu. W związku z tym pojawiło się ryzyko rozpięcia w miejscach łączy. Z każdym kolejnym obrotem zwiększałoby się prawdopodobieństwo spowodowania przez nie zwarcia. Rozwiązać to można na dwa sposoby. Pierwszy to zastosowanie pierścieni ślizgowych, które pozwolą na swobodny obrót elementów i nie dopuszczą do skręcenia się przewodów. Drugi polega na wprowadzeniu korekty do programu urządzenia polegającej na ograniczeniu kąta obrotu czyli, np. obrót silnika o kąt przekraczający 180° nakazywałby wieżyczce obracanie się w kierunku przeciwnym. Należy jednak pamiętać, że tego typu zmiany w programie mogą negatywnie wpłynąć na czas naprowadzania się urządzenia na obiekt.

Na podstawie powyższych rozważań można stwierdzić, że cel pracy został osiągnięty [1].

LITERATURA

- [1] Styła M., *Autonomiczny system namierzania obiektów*, Praca inżynierska, Lublin 2017
- [2] Francuz T., *Język C dla mikrokontrolerów AVR od podstaw do zaawansowanych aplikacji*, Wydaw. Helion, 2011
- [3] Kosmol J., *Napędy mechatroniczne*, Wydaw. Politechniki Śląskiej, Gliwice, 2013
- [4] Gajek A., Juda Z., *Czujniki*. Wydaw. Komunikacji i Łączności, Warszawa, 2011

- [5] DoDaam Systems, <http://www.dodaam.com/eng/sub2/menu2.php>, zasoby z dnia 28.04.2017
- [6] Micros Kraków – Hurtownia części elektronicznych, <http://www.micros.com.pl/>, zasoby z dnia 28.04.2017
- [7] Silniki krokowe – Silniki krokowe i sterowniki silników krokowych, <http://silniki-krokowe.com.pl/>, zasoby z dnia 28.04.2017
- [8] NavWeaps, <http://www.navweaps.com/>, zasoby z dnia 28.04.2017

PROJEKT I WYKONANIE PÓŁPRZEWODNIKOWEJ CEWKI REZONANSOWEJ

1. WPROWADZENIE

Tytułarna cewka rezonansowa nazywana jest także transformatorem czy też cewką Tesli, nazwa ta pochodzi od nazwiska jej wynalazcy. Nikola Tesla [1], jeden z największych wynalazców przełomu XIX i XX wieku, człowiek, który w ogromnej skali przyczynił się do rozwoju elektrotechniki. Jego transformator zbudowany został w 1891 r. i miał służyć do badań nad bezprzewodowym przesyłem energii na odległość. Urządzenie było źródłem wysokiego napięcia o dużej częstotliwości, lecz nie pozwalało na osiągnięcie pierwotnego celu działania cewki. Dzięki swoim parametrom transformator znalazł zastosowanie w badaniach związanych z techniką wysokich napięć i materiałami wykorzystywanymi w elektrotechnice.

Pomysł na zbudowanie tego urządzenia powstał podczas zajęć techniki wysokich napięć, gdy zainteresowaliśmy się stojącą w laboratorium cewką Tesli zbudowaną między innymi przez mgr inż. Przemysława Rogalskiego, który wytłumaczył nam jej zasadę działania. Podczas przeszukiwania literatury dotyczącej tego zagadnienia znaleźliśmy wiele rozwiązań które można zastosować przy budowie cewek Tesli. Najciekawszymi projektami były układy które zamiast wykorzystywać źródła wysokiego napięcia wykorzystywały napięcie sieciowe, oparte na mostkach tranzystorowych układy półprzewodnikowe. Konstrukcje te mają wiele zalet, z których największą jest brak konieczności ustalania dokładnych parametrów elektrycznych urządzenia. w klasycznych transformatrach, należy dobrać parametry tak dokładnie, aby obwód pierwotny i wtórny był w stanie rezonansu. Istnieją dwie wersje półprzewodnikowych cewek rezonansowych, jedna z nich zachowuje rezonans w układzie pierwotnym i wtórnym, są to najbardziej wydajne cewki, lecz koszt i złożoność tych projektów odwiodła nas od realizacji tego typu urządzeń. Zdecydowaliśmy się na wykonanie projektu, w którym inaczej niż w wymienionych wyżej rozwiązaniach rezonans występuje tylko w obwodzie wtórnym transformatora. Założeniem naszego urządzenia było stworzenie wydajnego obwodu pierwotnego transformatora, który sterowany jest za pomocą układu elektronicznego. Zadaniem układu elektronicznego jest badanie częstotliwości sygnału z uzwojenia wtórnego i dostosowanie do niego częstotliwości obwodu pierwotnego.

¹ Politechnika Lubelska, Koło Naukowe Materiałoznawstwa Elektrycznego i Techniki Wysokich Napięć „MELJON”

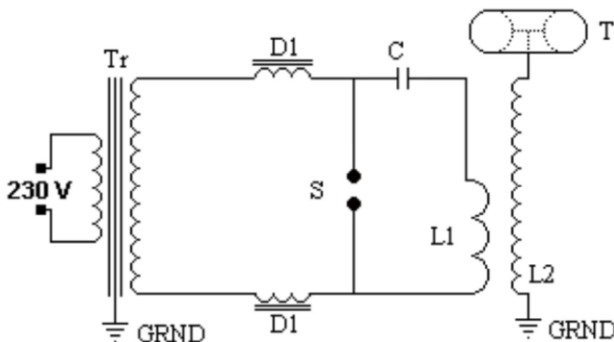
Naszym głównym źródłem informacji były strony, blogi i fora internetowe, część była w języku polskim, jednak najczęściej o tej tematyce pisano w języku angielskim. Najlepiej opisane projekty i zjawiska związane z naszą pracą znaleźliśmy na stronach entuzjastów transformatorów Tesli skupiających się w Stanach Zjednoczonych. Znalezione przez nas rozwiązania zarówno w zakresie parametrów transformatora, zasilania i sterowania układem zasilającym zainspirowały nas do budowy własnego projektu. Zdobyta wiedza w zakresie budowania urządzeń typu SSTC (Solid State Tesla Coil, z ang. półprzewodnikowa cewka Tesli) natchnęły nas do stworzenia własnego projektu. Opracowaliśmy projekt budowy urządzenia, które zbudowaliśmy i zbadaliśmy w ramach prac Koła Naukowego Materiałoznawstwa Elektrycznego i Techniki Wysokich Napięć „Meljon”.

2. ZASADA DZIAŁANIA URZĄDZENIA

Ogólny zarys urządzenia

Transformator Tesli jest urządzeniem elektrycznym które służy do wytwarzania wysokiego napięcia [2]. Składa się z dwóch części (Rys. 1):

- obwodu pierwotnego w który wchodzi kondensatora wysokonapięciowego o dużej pojemności, cewki o niskiej indukcyjności, iskiernika oraz transformatora zasilającego
- obwodu wtórnego w który wchodzi kondensator w postaci toroidu o małej pojemności i cewki o dużej indukcyjności.



Rys. 1. Schemat klasycznej cewki Tesli [3]

W klasycznym transformatorze Tesli zasada działania jest następująca. Oba uzwojenia są sprzężone magnetycznie a ich elementy są dobrane tak, by miały tę samą częstotliwość rezonansową. Cykl pracy urządzenia składa się z etapów:

- kondensator uzwojenia pierwotnego jest ładowany przez transformator.

- po osiągnięciu wysokiej wartości chwilowej napięcia na kondensatorze układ zamyka się przez przerwę iskrową drgając z częstotliwością rezonansową. Jest to stan pracy w którym wypadkowa reaktancja obwodu wynosi zero powodując przepływ bardzo dużego prądu przez uzwojenie pierwotne.
- silne pole magnetyczne wytworzone w uzwojeniu pierwotnym indukuje napięcie w uzwojeniu wtórnym które wraz z toroidem stanowi drugi układ rezonansowy który zostaje zamknięty przez wyładowanie elektryczne z ziemią.

Zjawisko tak dużego przyrostu napięcia spowodowane jest różnicą pojemności między obwodem pierwotnym a wtórnym, które opisane jest przez zasadę zachowania energii dla kondensatorów.

Zasada działania półprzewodnikowej cewki rezonansowej

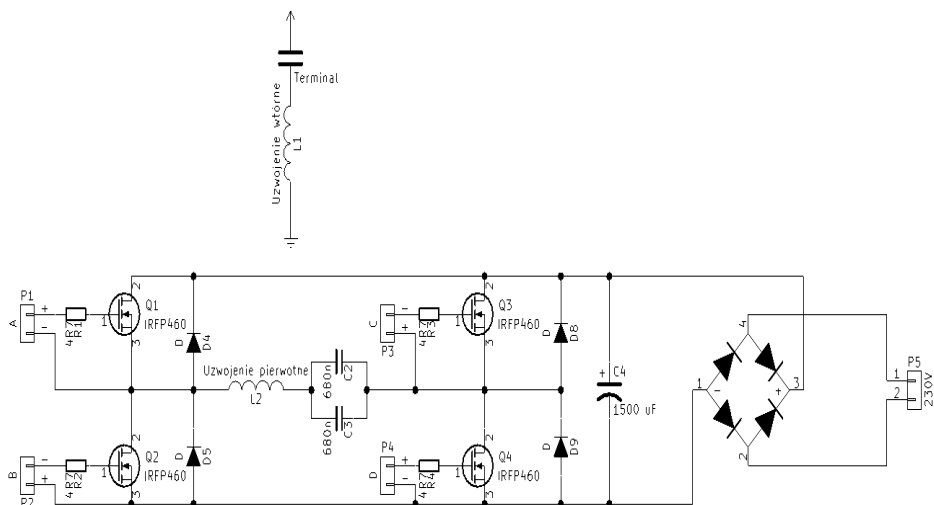
Budowa półprzewodnikowej cewki Tesli odbiega od klasycznego modelu ze względu na zastosowanie nowoczesnych rozwiązań w celu zwiększenia wydajności i składa się z [4]:

- uzwojenia wtórnego
- terminalu
- uzwojenia pierwotnego
- baterii kondensatorów wysokonapięciowych
- układu tranzystorowego mostka zasilającego
- układ elektroniczny sterujący tranzystorami
- przerywacza.

Uzwojenie pierwotne wraz z baterią kondensatorów wysokonapięciowych zasilane jest z układu tranzystorowego mostka. Na układ ten składają się cztery tranzystory MOSFET sterowane układem elektronicznym, mostek prostowniczy oraz kondensator filtrujący. Uzwojenie pierwotne jest sprzężone magnetycznie z uzwojeniem wtórnym, w którym wyprowadzenia z jednej strony są uziemione zaś z drugiej zakończone terminalem. Układ elektroniczny sterujący tranzystorami wraz z układem przerywacza połączony jest z bramkami za pomocą transformatora o przekładni równej 1.

Uzwojenie pierwotne i wtórne sprzężone są ze sobą magnetycznie. Generacja wysokiego napięcia podobnie jak w klasycznym transformatorze Tesli następuje dzięki przekładni zwojowej i zasadzie zachowania energii dla kondensatorów. Półprzewodnikowa cewka tesli posiada rezonans pojedynczy, zachodzi on tylko i wyłącznie w uzwojeniu wtórnym. Układ elektroniczny za pomocą anteny do niego dołączonej bada częstotliwość pola elektrycznego wytwarzanego przez uzwojenie wtórne, która jest częstotliwością rezonansową i dostosowuje do niej pracę tranzystorów uzwojenia pierwotnego.

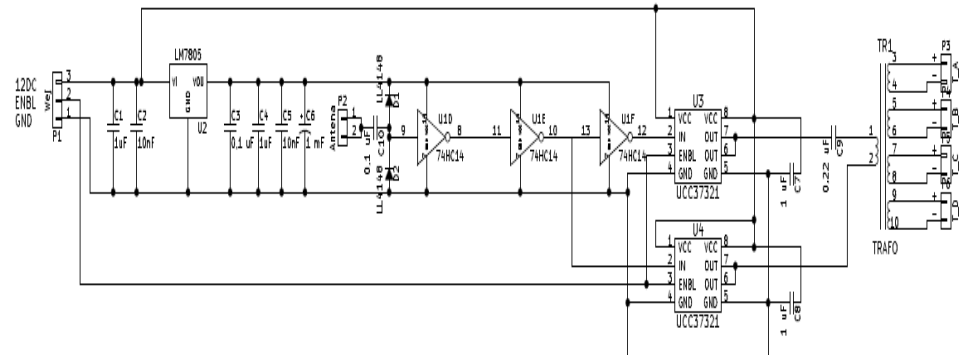
Obwód pierwotny i wtórny



Rys. 2. Schemat obwodu pierwotnego i wtórnego półprzewodnikowej cewki rezonansowej

Zasilanie transformatora składa się z mostka prostowniczego na który podawane jest napięcie poprzez autotransformator, kondensatora filtrującego i mostka tranzystorowego. Mostek wykonany jest z tranzystorów mosfet IRFP460 o dużej wytrzymałości prądowej ciągłej (20A) i impulsowej (80A).

Między dwiema gałęziami układów znajduje się uzwojenie pierwotne wraz z baterią kondensatorów. Dzięki takiemu układowi oraz odpowiedniemu kłucowaniu tranzystorów, mamy kontrolę nad częstotliwością oraz kierunkiem przepływu prądu wyjściowego [5]. Mostek tranzystorowy uzyskuje przewagę nad klasycznymi konstrukcjami tym, że na wyjściu uzyskujemy przebieg bardzo zbliżony do prostokątnego, zamiast przebiegu sinusoidalnego. Różnica ta sprawia, że podczas trwania jednego cyklu przekazywane jest więcej mocy. Straty występują tylko w czasie przełączania się tranzystorów. Do każdego z tranzystorów dołączone zostały diody w celu ochrony przed przepięciami, które może spowodować indukcyjność cewki w trakcie uruchamiania.



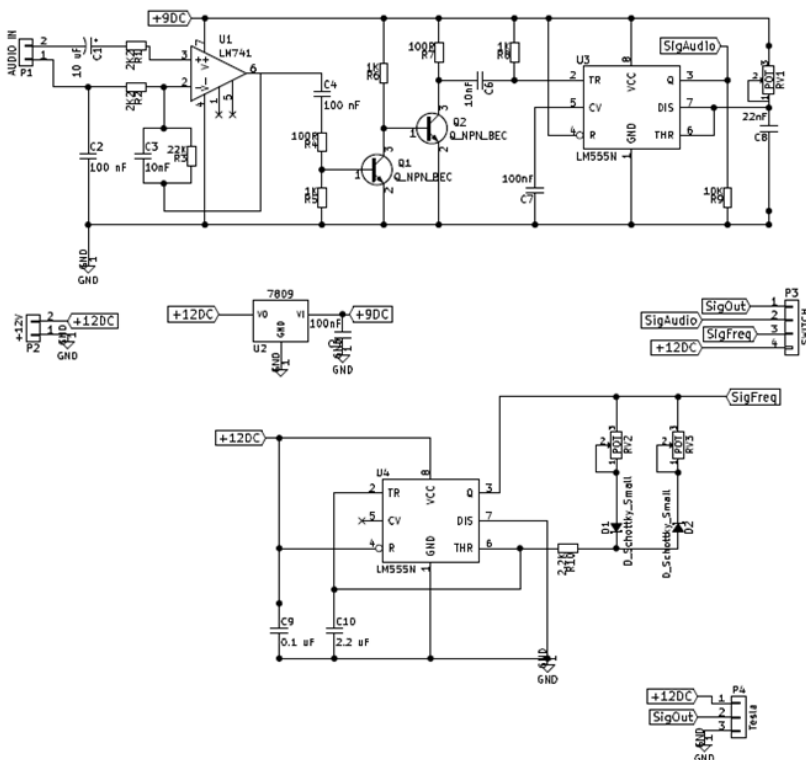
Rys. 3. Schemat elektronicznego układu sterowania tranzystorami

Układ elektroniczny z rysunku 3 jest najbardziej złożonym elementem urządzenia. Jego zadaniem jest odpowiednie taktowanie tranzystorów z częstotliwością rezonansową, poprzez odebranie w antenie sygnału oscylacji układu wtórnego, przetworzenie go na sygnał cyfrowy i odpowiednie przekazanie go na bramki tranzystorów.

Sygnał z anteny jest dyskretyzowany przez inwerter a następnie przez kolejne inwertery odwracany tak, aby podać go w odpowiedniej kolejności na sterowniki tranzystorów. Układ posiada jeden dodatkowy inwerter ze względu na brak dostępności odwracających sterowników tranzystorowych. Dodatkowy przerzutnik odwracający nie stanowi przeszkód w działaniu układu, ponieważ jego opóźnienie jest zbyt małe.

W projekcie zastosowaliśmy dwa rodzaje przerywania sygnału. Jednym z nich jest przerywacz, który sterowany za pomocą timera LM555N pozwala na dobranie częstotliwości pracy transformatora poprzez regulację sygnału prostokątnego dla stanu wysokiego i dla stanu niskiego.

Drugim z układów przerywacza jest układ muzyczny, który ma za zadanie przetwarzanie i filtrowanie podanego sygnału audio. W obwodzie na wejściu zastosowany został filtr dolnoprzepustowy, wzmacniacz i przerywacz w postaci timera LM555N. Filtr pełni tutaj ważną rolę, ponieważ przez częstotliwość transformatora ograniczona przez przerwania LM555N nie moglibyśmy uzyskać sygnału dźwiękowego dla ścieżek o wysokich częstotliwościach. Obwód ten obciążony jest przymusem podawania sygnału audio w którym dominują niskie częstotliwości, aby uzyskać pożądaną efekt.



Rys. 4. Schemat układu przerywacza

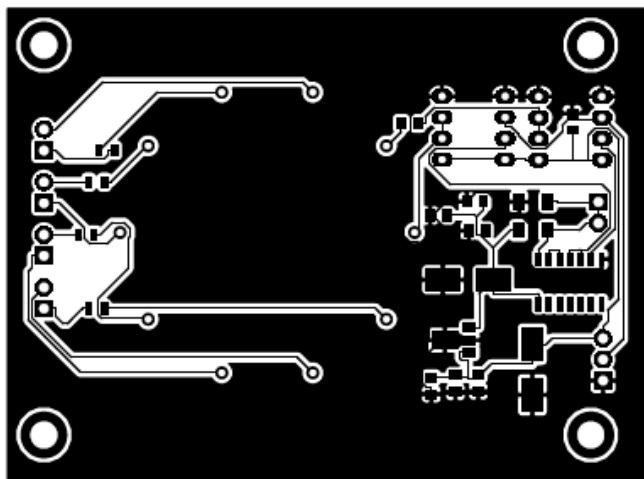
Oba układy mają za zadanie modyfikować pracę transformatora tak, aby była on nieciągly. Pozwala to na podawanie częstotliwości które są słyszalne i których przerywanie może być widzialne dla człowieka. Urządzenie zostało zaprojektowane na częstotliwość ok. 250kHz, częstotliwość ta jest zbyt duża, aby człowiek mógł zarejestrować jej dźwięk czy także zmiany wizualne.

3. WYKONANIE URZĄDZENIA

Wytwarzanie płytek drukowanych

Pierwszym etapem wytwarzania płytek było utworzenie ich schematów jak i spisu elementów jak na Rys. 3 i Rys.4. Następnie należało dobrać wymiary płytki i ułożyć elementy tak aby były rozłożone optymalnie co do swojej roli. Zdecydowaliśmy się utworzyć dwa układy, jeden z nich do sterowania tranzystorami został umieszczony tuż przy układzie zasilania, drugi zaś miał być „pilotem” którym można by było dokonywać regulacji w bezpiecznej odległości od urządzenia. Schematy jak i projekty płytek zdecydowaliśmy się wykonać w pro-

gramie KiCad, który pozwolił także na wykonanie wydruków masek służących do wykonania układów drukowanych.



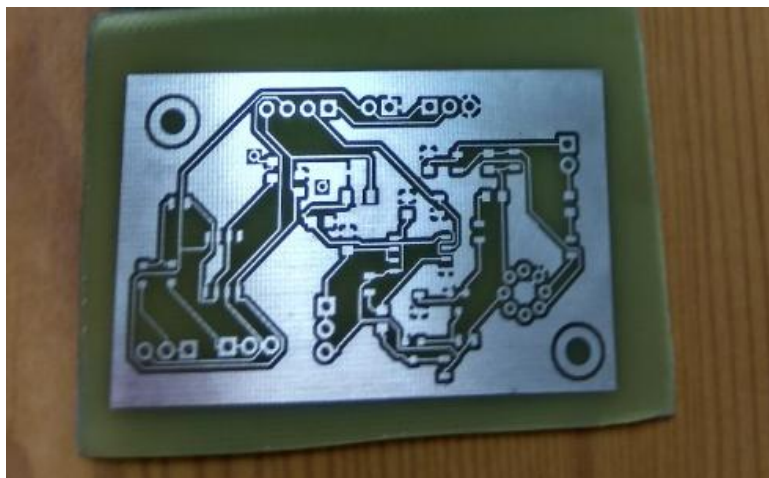
Rys. 5. Maska płytki elektronicznego układu sterowania tranzystorami

Na rysunku 5. Przedstawiony został projekt maski elektronicznego układu sterowania tranzystorami. W podobnej technologii została wykonana także maska płytki elektronicznego układu przerywacza

Układy elektroniczne zostały wykonane za pomocą metody fotochemicznej. Pierwszym etapem było naświetlanie światłem ultrafioletowym laminatu pokrytego substancją fotoczułą masek naszych płytek. Następnie docięte, naświetlone płytki zostały wywołane dzięki 7% roztworowi metakrzemianu sodu. Miał on za zadanie rozpuszczenie emulsji światłoczułej w miejscach gdzie była ona naświetlania.

Aby pozbyć się zbędnej warstwy miedzi z płytek została zastosowana metoda wytrawiania. Miała ona za zadanie za pomocą 30% roztworu nadsiarczanu amonu usunąć zbędną miedź z płytki.

Aby ograniczyć proces utleniania się miedzi na płytce drukowanej postanowiono zastosować metodę cynowania. Proces ten miał za zadanie zapobiec reakcji miedzi z powietrzem, a także ułatwić proces lutowania elementów. W tym celu nałożyliśmy na laminat pastę z cyną do lutowania miękkiego, następnie w procesie podgrzewania stacją z gorącym powietrzem udało nam się nanieść cienką warstwę cyny na płytkę.



Rys. 6. Gotowa płytka przed wycięciem do prawidłowych rozmiarów

Wykonanie obwodów silnoprądowych

Pierwszym etapem wykonania docelowego transformatora Tesli było wykonanie uzwojenia wtórnego, najtrudniejszego z elementów. W tym celu wykorzystaliśmy rurę hydrauliczną PCV o średnicy 11cm i za pomocą drutu nawojowego o średnicy 2,5mm wykonany został obwód wtórny. Aby uniknąć zwarcie między zwojami, uzwojeniami a także przebić od terminalu należało wykonać izolację w postaci lakieru elektroizolacyjnego. Zbudowana została amatorska konstrukcja nawijarki, która ułatwiła zbudowanie cewki. Uzwojenie zgodnie z projektem posiada 1000 zwojów, z jedynym wyprowadzeniem do terminalu a z drugim na terminal.

Uzwojenie pierwotne nawinięte zostało bezpośrednio na uzwojenie wtórne w celu zwiększenia sprawności transformatora, wykonane zostało przewodem miedzianym izolowanym o średnicy 2mm. Ważnym aspektem było dobra izolacja między uzwojeniami, ponieważ lakier elektroizolacyjny i izolacja przewodu nie wydawała nam się wystarczająca zastosowany został papier transformatorowy i taśma izolacyjna PCV. Nawinięte zostało 8 zwojów na które została zastosowana kolejna warstwa papieru transformatorowego i taśmy izolacyjnej.

Ostatnim etapem było stworzenie terminalu, z założenia miał mieć on kształt toroidalny i najczęściej wykonywany jest z elastycznych, karbowanych rur aluminiowych które sprawiają, że występują niewielkie uloty z wypustek które zmniejszają sprawność urządzenia. Terminal został wykonany z dwóch oczyszczonych patelni skręconych za pomocą śruby. Wypustki po rączkach zaklejone zostały aluminiową taśmą, a śruba od strony uzwojenia została wyszlifowana, aby nie posiadała ostrych krawędzi generujących uloty czy też punktów, w których mogło by nastąpić przebicie.



Rys. 7. Uzwojenie wtórne po nawinięciu



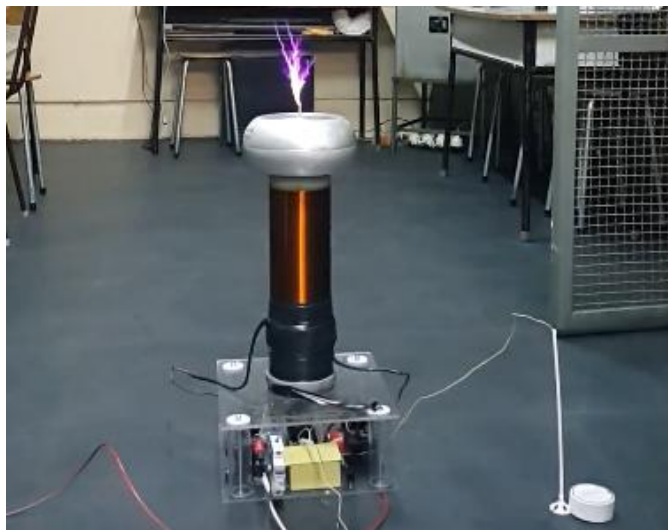
Rys. 8. Zdjęcie karkasu z nawiniętymi uzwojeniami i założonym terminalem

4. PRACA URZĄDZENIA

Gotowy do działania transformator przysporzył nam problemy z uruchomieniem, ponieważ wymagał on autotransformatora o wysokiej wydajności prądowej, a także pomiar prądu mógł odbywać się tylko za pomocą amperomierza cęgowego. Urządzenie wzbudza się w okolicach 20V pobierając prąd około 11A. Wyładowania widoczne są przy około 30V i wraz ze zwiększaniem napięcia prąd rośnie, aczkolwiek nieliniowo.

Cewka Tesli przez nas pracuje w 3 trybach, ciągłym, przerywanym i muzycznym. W trybie ciągłym uzyskuje największą moc, przy napięciu 100V pobiera z sieci moc około 1,6kW. W trybie przerywanym i muzycznym moc jej zależy od wypełnienia sygnału. W każdym z trybów powoduje olbrzymie zakłócenia, radia w otaczających pomieszczeniach przestają działać, w smartfonach w odległości około 4 metrów od urządzenia przestają działać ekrany dotykowe, a także w trybie przerywanym potrafi wymuszać akcje na otaczających je urządzeniach takie jak przyciszenie i zgłaśnianie telefonów.

Z pomocą naszego urządzenia udało nam się uzyskać zjawiska takie jak wiatr elektronowy czy wzbudzenie pobliskich lamp wyładowczych do działania (w odległości około 5 metrów). Nie udało nam się zmierzyć długości wyładowań transformatora, ponieważ rozstaw używanego przez nas iskiernika ostrzowego był zbyt mały, maksymalnym był 50 cm co jest imponującym wynikiem, jest to urządzenie o największej mocy wyładowań na Katedrze Urządzeń Elektrycznych i Techniki Wysokich Napięć.



Rys. 9. Transformator podczas pracy w trybie ciągłym

Transformator pracuje na relatywnie wysokiej częstotliwości, występuje tutaj tak zwany efekt naskórkowy. Pomimo tak wysokiego napięcia dzięki naskórkowości możliwym jest „dotknięcie” łuku elektrycznego przez człowieka. Podczas testów efektu naskórkowego nie był możliwy jednak bezpośredni kontakt z łukiem elektrycznym przez jego temperaturę, która mogła dokonać poparzeń skóry, jednak poprzez metalowe elementy nie są odczuwalne jakiejkolwiek skutki porażenia prądem elektrycznym.

Urządzenie daje ogromne możliwości do badań nad szybkozmiennymi wyładowaniami czy polem elektromagnetycznym. Obecnie używane jest do badań naukowych dotyczących perkolacji. Transformator jest też doskonałym narzędziem do obrazowania wielu zjawisk fizycznych, a także może być stosowany do celów widowiskowych szczególnie dzięki długim wyładowaniom i trybowi muzycznemu.

LITERATURA

- [1] https://pl.wikipedia.org/wiki/Nikola_Tesla (zasoby z dnia 28.05.2017)
- [2] M. Tilbury, *The Ultimate Tesla Coil design and construction guide*, McGraw Hill, 2008
- [3] <http://teslacoil.republika.pl/> (zasoby z dnia 28.05.2017)
- [4] <http://teslacoil.pl/sstc/> (zasoby z dnia 28.05.2017)
- [5] <http://www.forbot.pl/forum/topics20/teoria-mostek-h-h-bridge-kompendium-dla-robotyka-vt111.html> (zasoby z dnia 28.05.2017)

PROJEKT I WYKONANIE STANOWISKA LABORATORYJNEGO DO SYMULACJI LAPAROSKOPOWYCH

1. ZARYS TEORETYCZNY ZABIEGÓW LAPAROSKOPOWYCH

Pojęcie laparoskopii

Laparoskopia jest zabiegiem wziernikowania jamy otrzewnej polegającym na wprowadzeniu specjalistycznego instrumentarium poprzez powłoki brzuszne celem przeprowadzenia zabiegów czy diagnozowania stanów patologicznych [7].

Zaletami tej metody są m.in.: zmniejszenie utraty krwi, skrócenie okresu rekonwalescencji, redukcja powikłań. Wadą jest fakt, iż nie we wszystkich przypadkach laparoscopia jest możliwa, operacje laparoskopowe trwają dłużej niż operacje tradycyjne, użycie kosztownego sprzętu wymaga treningu [1].

Sprzęt laparoskopowy

Do instrumentarium laparoskopowego zaliczane są: system obrazowania, instrumenty oraz urządzenia wytwarzające i podtrzymujące odnię otrzewnową, trokary, laparoskopowe instrumenty operacyjne i inne.

System obrazowania składa się z laparoskopu ze światłowodem i kamerą, źródła światła oraz monitora, na którym wyświetlany jest obraz. Laparoskop to sztywna rurka posiadająca okular wraz z zestawem soczewek, dzięki którym światło nie ulega rozszczepieniu, uzyskując szerokokątny obraz o zadawalającej rozdzielczości. Zwykle laparoskopy mają średnicę 10mm lub 5mm, ale kiedy zachodzi taka potrzeba, stosuje się laparoskopy o średnicy 1,6–2mm. Kąt widzenia mieści się w przedziale 0°–45°[13].

Przesyłanie wiązki światła ze źródła do laparoskopu jest możliwe dzięki kablowi światłowodowemu. Posiadają one wewnątrz włókna szklane. Są one giętkie, co pozwala na swobodne posługiwanie się laparoskopem. Źródło światła zawiera zasilacz o dużej mocy, żarówki, układ chłodzący, układ optyczny skupiający światło oraz układ regulacji natężenia i mocy światła. Wykorzystywane są żarówki halogenowe o mocy 150W lub w przypadku, gdy niezbędne jest mocniejsze światło, żarówki ksenonowe o mocy 300W, pozwalające na około 250 godzin pracy [4]. W ostatnich latach na rynek zostały wprowadzone również źródła światła LED.

¹ Politechnika Lubelska, Wydział Elektrotechniki i Informatyki

Kamera umieszczona na drugim końcu laparoskopu przekazuje obraz pola operacyjnego do monitora. Współczesne kamery składają się z układów scalonych, posiadają obiektywy ze zmienną ogniskową, a barwy regulowane są automatycznie. Obecnie możliwe jest stosowanie kamer dwu- i trójwymiarowych. Przy używaniu kamery trójwymiarowej operator zakłada specjalne okulary [5].

Igła Veressa to narzędzie, od którego zaczyna się zabieg laparoskopowy. Służy do wytworzenia odmy otrzewnowej. Posiada tępą końcówkę, która po przeprowadzeniu przez powłoki brzuszne osłania ostrze, co zapobiega uszkodzeniu narządów wewnętrznych. Za utrzymanie odmy otrzewnowej odpowiada insuflator. Wprowadza on około 10 litrów gazu (najczęściej dwutlenku węgla) na minutę zapewniając wyznaczone ciśnienie 12–15mm Hg. Po uzyskaniu oczekiwanego ciśnienia insuflator wyłącza się, a jeśli ciśnienie spada zostaje ponownie uruchomiony przez zestaw czujników. Płukania i odsysania płynnej zawartości jamy brzusznej dokonuje akwapurator [1].

Dostęp laparoskopu i instrumentarium laparoskopowego do pola operacyjnego zapewniają trokary. Są to przyrządy posiadające kanały przez które wprowadza się narzędzia. Trokary posiadają ostry grot, aby możliwe było przeprowadzenie ich przez powłoki brzuszne. Dzięki jednokierunkowej zastawce niemożliwe jest uchodzenie gazu z jamy otrzewnej [7].

Laparoskopowe instrumenty operacyjne mają za zadanie umożliwić operatorowi chwytanie, preparowanie, przecinanie lub koagulowanie tkanek. Posiadają one wymienne końcówki pozwalające na wybór najbardziej odpowiedniej w danym momencie.

Szkolenie laparoskopowe

Mimo, iż od pierwszych odkryć i doświadczeń w dziedzinie laparoskopii minęło około dwieście lat, to dopiero w latach 80. XX. wieku rozwinęła się technika operowania z wykorzystaniem tych rozwiązań dzięki ówczesnemu rozwojowi technologicznemu oraz technicznemu [6].

Badania przeprowadzone na Uniwersytecie Medycznym w Łodzi wykazały wyraźną poprawę wyników testów sprawności w posługiwaniu się instrumentarium laparoskopowym po wcześniejszych treningach na symulatorach. Zaobserwowano również współzależność pomiędzy wynikami testów a rozwojem integracji sensoryczno-motorycznej u badanych [3].

W skład szkolenia z zakresu laparoskopii wchodzi cztery etapy. Pierwszym z nich jest wprowadzenie teoretyczne do nowej metody operacyjnej. Uczestnicy zapoznają się ze sprzętem wykorzystywanym podczas zabiegów, sposobem wytwarzania i utrzymywania odmy, dowiadują się jak przygotowywać i używać poszczególnych narzędzi laparoskopowych oraz poznają metody cięcia, koagulacji lub preparowania tkanek. Poznają przeciwwskazania i powikłania pooperacyjne [4].

Drugi etap to laboratoryjne ćwiczenia praktyczne. Na początku tego etapu lekarze nabierają umiejętności posługiwania się aparaturą przy użyciu symulatorów i тренаżerów. Symulacja jest techniką, która zastępuje lub wzmacnia doświadczenia lekarz – pacjent w kontrolowanych i nie zagrażających zdrowiu warunkach. Zastępuje aspekty realnego świata w interaktywny sposób [10]. Na tym etapie ćwiczenia symulacyjne, zamiast na sali operacyjnej, nie wymagają narażania pacjentów. Nie wydłuża się czasu operacji, a także eliminuje się ryzyko wystąpienia zakażenia, jakie mogłoby wystąpić przy wprowadzeniu na salę większej grupy osób. Podczas treningu symulacyjnego nie potrzebna jest nieustanna obecność nauczyciela [2, 8].



Rys. 1. Przykładowy тренаżer laparoskopowy

Następnie uczestnicy kontynuują szkolenie na wypreparowanych narządach zwierząt oraz przeprowadzają zabiegi na zwierzętach. Doskonają wtedy swoją precyzję w posługiwaniu się sprzętem laparoskopowym stwarzając warunki takie, jak podczas operacji człowieka.

Dopiero po takim przygotowaniu rozpoczyna się etap trzeci i szkolenie przenosi się na salę operacyjną. Uczestnicy uczą się jak prawidłowo przygotować sprzęt laparoskopowy, jak ułożyć pacjenta na stole operacyjnym, poznają zadania całego zespołu operacyjnego.

Ostatnim, czwartym etapem szkolenia jest dwutygodniowy staż w klinice, gdzie wykonuje się kilka zabiegów laparoskopowych dziennie. Początkowo jako asysta, później jako operator prostszych zabiegów, aż w końcu wykonując pod okiem doświadczonego chirurga samodzielny zabieg np. cholecystektomię.

Oczywista jest konieczność ciągłego poszerzania swojej wiedzy i umiejętności na sympozjach, dodatkowych szkoleniach lub konferencjach, czy oglądając filmy z przeprowadzonych zabiegów laparoskopowych [4].

2. PROJEKTOWANIE SYMULATORA LAPAROSKOPOWEGO

Analiza dostępnych symulatorów laparoskopowych

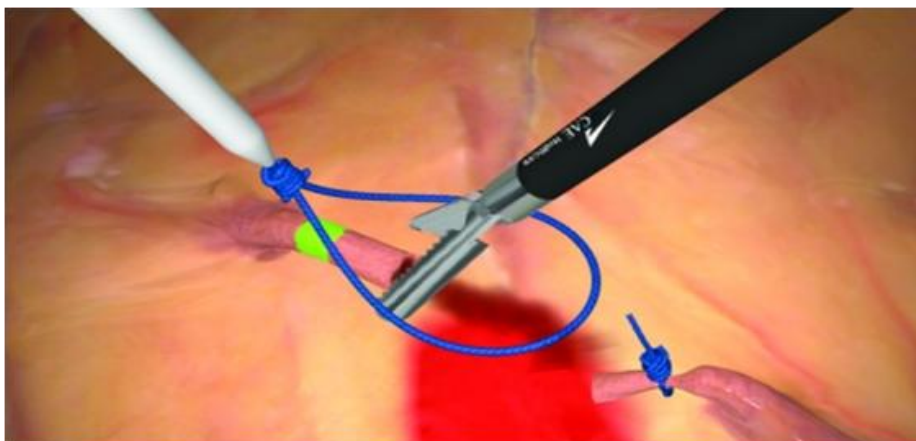
Proces projektowania stanowiska rozpoczęto od zapoznania się z metodą zabiegów laparoskopowych, analizy symulatorów i trenażerów laparoskopowych obecnie dostępnych na rynku.

Modele dostępne na rynku można podzielić na dwie kategorie. Pierwsza z nich to proste trenażery, najczęściej zawierające dwa instrumenty laparoskopowe w postaci kleszczy chwytnych i korpus z planszą do wykonywania podstawowych ćwiczeń z zakresu koordynacji sensoryczno-motorycznej lub najprostszych zabiegów chirurgicznych. Przeważnie nie zawierają kamery i wyświetlacza. Ich ceny wahają się w granicach kilku do kilkunastu tysięcy złotych.



Rys. 2. Przykład trenażera laparoskoowego firmy Grena

Drugim typem symulatorów są te oparte na systemach komputerowych stwarzające wirtualne środowisko, próbujące naśladować realne warunki operacji laparoskopowych. Instrumentarium laparoskopowe posiada wbudowane sensory, pozwalające na odzwierciedlanie ich rzeczywistych ruchów w symulacji. Zakup takiego sprzętu to wydatek od dwudziestu tysięcy złotych wzwyż.



Rys. 3. Obraz z ekranu symulatora Lap VR

Określenie wymagań i założeń projektu

Rozpoczynając projektowanie należy wyznaczyć wymagania jakie muszą być zrealizowane. Określono następujące założenia w projekcie stanowiska symulacji laparoskopowych: stworzenie warunków jak najwierniej odwzorowujących warunki zabiegów laparoskopowych, opracowanie stanowiska ewoluującego, tak, by użytkownik korzystając z niego mógł stale się rozwijać, zmniejszenie ceny symulatora w porównaniu do oferty rynkowej oraz umożliwienie wygodnego transportu urządzenia.

Należy zwrócić uwagę, iż nie zawsze wszystkie założenia będą mogły być w pełni zrealizowane. Czasem możliwości spełnienia poszczególnych założeń wykluczają się nawzajem. Wtedy konieczny jest wybór ważniejszego wymagania i dążenie do jego realizacji.

Koncepcja stanowiska laboratoryjnego do symulacji laparoskopowych

Pierwotna koncepcja symulatora zakładała wykonanie modelu torsu człowieka na podstawie gipsowego odlewu z otworami, przez które możliwe by było wprowadzanie narzędzi laparoskopowych. Wewnątrz niego znajdowałaby się kamera transmitująca obraz z wnętrza fantomu oraz plansza do wykonywania zadań rozwijających koordynację oko – ręka. Jednak wykonanie modelu okazało się czasochłonne i bardzo kosztowne, a takie rozwiązanie spełniałoby głównie funkcje estetyczne, dlatego postanowiono zastąpić go pudełkiem. Wymiary pudełka dobrano tak, by użytkownik mógł swobodnie operować wewnątrz narzędziami laparoskopowymi. Pudełko pozbawione jest dna, aby można go było stawiać na wymiennych planszach.

Do obrazowania wnętrza symulatora planowano wykorzystać kamerę internetową zainstalowaną na stałe w pudełku. Zrezygnowano z ruchomego obrazowania, tak jak ma to miejsce przy zabiegach laparoskopowych, ponieważ operator w trakcie treningu ma zajęte obie ręce instrumentarium laparoskopowym i do kierowania torem wizyjnym niezbędna byłaby druga osoba. Okazało się jednak, że pole widzenia zwykłej kamery internetowej w pudełku jest niewystarczające. Z tego powodu zdecydowano użyć szerokokątnej kamery Genius WideCam 1050, o polu widzenia 120°, którą można podłączyć do każdego komputera za pomocą złącza USB.

Źródło światła stanowi dwanaście diod LED. Umieszczone są na okrągłej płytce, przyklejonej do jednej ze ścian pudełka. Oświetlenie jest zasilane dzięki podłączeniu do komputera złączem USB. Zastosowanie takiego sposobu zasilania nie jest przypadkowe. Dzięki temu uniknięto montowania dodatkowych źródeł energii, a wykorzystano, potrzebny do wyświetlania obrazu komputer.

Narzędzia laparoskopowe udało się zdobyć dzięki uprzejmości Działu Aparatury Medycznej przy Wojewódzkim Szpitalu Specjalistycznym w Lublinie. Są to dwa wysterylizowane kleszcze chwytne nieużywane już do zabiegów. Ważnym założeniem projektu była możliwość ewoluowania stanowiska wraz z rozwojem ćwiczącego. Właśnie dlatego zrezygnowano z jednego, stałego zadania znajdującego się wewnątrz symulatora na rzecz wymiennych plansz, które po postawieniu na nich pudełka stanowić będą jego dno. Takie rozwiązanie pozwoli na dostosowanie poziomu trudności zadań do umiejętności operatora. Ponadto po opanowaniu plansz dołączonych do symulatora możliwe jest stworzenie własnych dostosowanych do indywidualnych potrzeb trenującego. Jest to innowacyjne rozwiązanie pozwalające na kreatywną twórczość użytkownika.



Rys. 4. Wygląd symulatora z zewnątrz

Stanowisko laboratoryjne tworzy pudełko i trzy kompletne plansze. Na początku szkolenia laparoskopowego lekarze nabywają nowych zdolności sensoryczno-motorycznych, dlatego projektując zadania do symulatora inspirowano się zabawkami, którymi bawią się dzieci wykształcające swoją koordynację oko – ręka. Ważne przy tworzeniu plansz było to, aby zadania nie były zbyt proste, a rozwiązanie ich zajmowało więcej niż tylko kilka sekund.



Rys. 5. Zabawki wykształcające koordynację sensoryczno- motoryczną u dzieci

Zadanie na planszy numer 1 polega na przeprowadzaniu sznurka przez obręcze o średnicy od 27mm do 35mm, umiejscowionych na wysokości od 8cm do 10cm ustawionych pod różnym kątem. Dodatkowym utrudnieniem dla operatora jest fakt, iż przez jedną obręcz sznurek musi być przeprowadzony przy użyciu tylko jednego graspera.



Rys. 6. Plansza nr 1

Plansza numer 2 to puzzle. Użytkownik ma za zadanie dopasować kilka brył o podstawach prostych figur geometrycznych do odpowiednich gniazd. Bryły zostały zaprojektowane tak, aby możliwe było wygodne uchwycenie ich narzędziami laparoskopowymi i sprawne umieszczenie na swoich miejscach.

*Rys. 7. Plansza nr 2*

W skład planszy numer 3 wchodzi dwa zadania. Pierwsze z nich to rozciągnięcie gumki na wszystkie bolce. Drugie zadanie wymaga cierpliwości i determinacji. Wewnątrz materiałowego rękawa o długości 20cm znajduje się kulka. Trenujący musi przeprowadzić ją z jednego końca rękawa na drugi. Nie jest to łatwe, gdyż rękaw został zaprojektowany tak, by kulka wewnątrz niego nie poruszała się swobodnie. Operator musi przeciskać ją przez całą długość rękawa kleszczami.

*Rys. 8. Plansza nr 3*

Poniższa tabela przedstawia kosztorys wykonania symulatora. Uwzględniono w niej przedmioty zakupione w celu realizacji projektu.

Tabela 1. Kosztorys symulatora

Lp.	Element symulatora	Cena [zł]
1	Sklejka na pudełko i plansze	30
2	Materiały do obróbki drewna	60
3	Kamera internetowa	100
4	Uchwyty do pudełka	6
5	Filament do druku 3D	60
6	Wkręty, nakładki, nakrętki	20
7	Suma	276

Analizując powyższą tabelę oraz ceny rynkowe dostępnych symulatorów można jednoznacznie stwierdzić, iż założenie obniżenia ceny zostało w pełni zrealizowane.

3. WYKONANIE STANOWISKA SYMULACYJNEGO

Prace wykonywane w drewnie

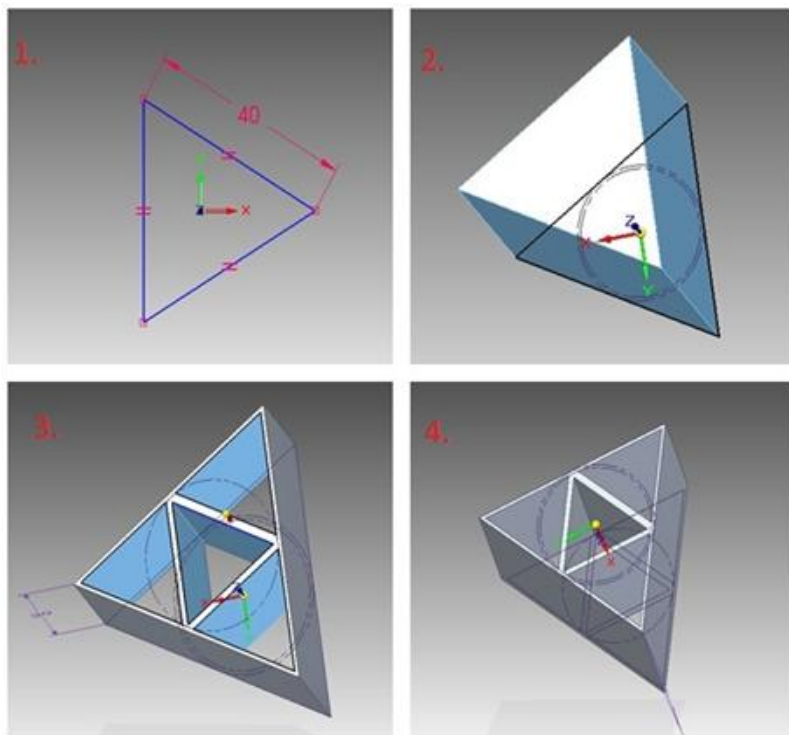
Aby pudełko było stosunkowo lekkie, co stanowi ważną zaletę podczas transportu, ale i wytrzymałe postanowiono użyć sklejki tradycyjnej o grubości 4mm. Jest to płyta złożona z kilku warstw drewna połączonych ze sobą klejem. Warstwy mają naprzemiennie prostopadłe ułożone włókna [11].

Wymiary pudełka to 40cm szerokości, 35cm długości oraz 20cm wysokości. Zostały one dopasowane do zasięgu narzędzi laparoskopowych. Konstrukcje wzmacniają kantówki tworzące szkielet pudełka. Na bocznych ścianach pudełka zostały wywiercone otwory o średnicy 4cm, przez które wprowadza się graspery do wnętrza symulatora. Gabaryty plansz zostały dobrane tak, aby doskonale pasowały do przykrywającego je pudełka.

Aby wygładzić wszystkie nierówności i uzupełnić nieszczelności zastosowano akrylową masę szpachlową. Następnie wszystkie elementy drewniane zostały pokryte dwoma warstwami lakierobejcy. Preparat ten nadał powierzchni delikatny połysk i podkreślił naturalny rysunek drewna, jednocześnie zabezpieczając go przed np. zadrapaniami [12]. Podnoszenie pudełka ułatwiają dwa uchwyty zamontowane na bocznych jego ścianach. Natomiast na wewnętrznych ścianach zainstalowane jest źródło światła oraz kamera internetowa do transmisji obrazu.

Modelowanie przestrzenne w programie Solid Edge

Niemal wszystkie elementy znajdujące się na planszach symulatora zostały wyprodukowane metodą druku trójwymiarowego. Aby było to możliwe, należało uprzednio przygotować w programie typu CAD modele przestrzenne tych elementów. Idealny do takich zadań jest wyżej wspomniany program Solid Edge.



Rys. 9. Etapy powstawania jednego z elementów planszy nr 2: 1. Szkic podstawy modelu, 2. Przeciąganie- nadawanie wysokości bryle, 3. Wycinanie zbędnych fragmentów bryły, 4. Gotowy model przestrzenny elementu

Druk trójwymiarowy

W ostatnich latach na popularności zyskała nowa metoda produkcji – wydruki trójwymiarowe, inaczej 3D. Możliwe było to poprzez dynamiczny rozwój programów komputerowych do modelowania przestrzennego. Dzięki temu, jedyne ograniczenia jakie stoi nam na przeszkodzie przy wykorzystywaniu tych urządzeń to nasza wyobraźnia.

Proces powstawania wydruku trójwymiarowego składa się z kilku etapów. Przedstawia je poniższy diagram.

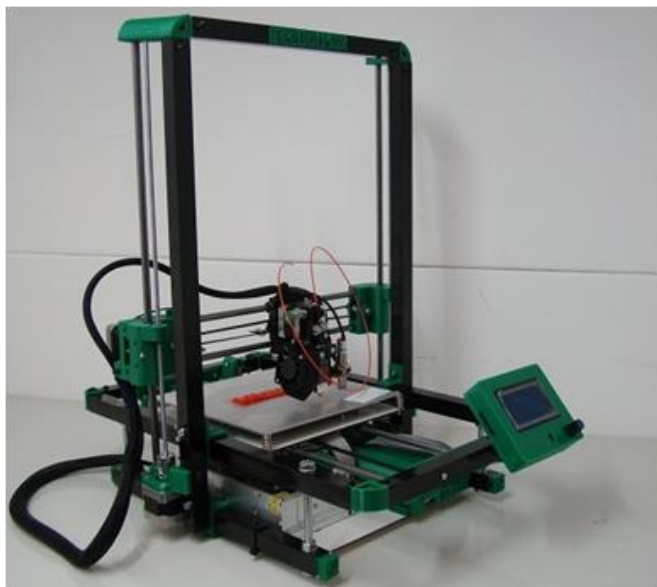


Rys. 10. Etapy druku 3D (źródło własne)

Pierwszy etap polega na przygotowaniu wirtualnego modelu przestrzennego pożądanego bryły w programie typu CAD. Następnie przy wykorzystaniu wspomaganego komputerowego CAM model ten jest dzielony na warstwy oraz programowane są ruchy nakładania materiału. Po tych etapach następuje wytwarzanie obiektu techniką przyrostową w maszynie sterowanej numerycznie. Finalnie otrzymano gotowy obiekt i przystąpiono do jego obróbki [9].

Elementy plansz zostały wydrukowane przy użyciu techniki druku 3D, po uprzednim przygotowaniu ich przestrzennych modeli w programie Solid Edge o formacie plików STL. Pliki te zostały przetworzone przez program typu CAM – Slic3r na G-code, czyli definicję operacji, jakie ma wykonać drukarka 3D podczas produkcji konkretnego elementu.

Do wydruków wykorzystano drukarkę Prusa i3 pracującą w metodzie FDM. Posiada ona port kart pamięci SD, dzięki któremu G-code zostaje wprowadzony do urządzenia.



Rys. 11. Drukarka Prusa i3

Materiałem wykorzystanym przy produkcji elementów był polilaktyd (PLA). Jest to polimer w pełni biodegradowalny, który można uzyskać z odnawialnych surowców naturalnych. Do tego celu stosuje się np. mączkę kukurydzianą. Materiał ten znalazł szerokie zastosowanie w inżynierii biomedycznej, do produkcji resorbowalnych nici chirurgicznych bądź implantów stomatologicznych [66]. Polilaktyd jest ogólnie dostępny na rynku w formie żyłki nawiniętej na szpulkę w różnorodnych kolorach oraz kilku średnicach.

4. PODSUMOWANIE I WNIOSKI

Przedstawione na początku pracy cele zostały zrealizowane. Wykonano w pełni sprawne i dostosowane do umiejętności operatora, realizujące założenia ekonomiczne, mobilne stanowisko laboratoryjne odwzorowujące warunki zabiegów laparoskopowych. Urządzenie testowane było przez konstruktora oraz studentów II roku Inżynierii Biomedycznej na Politechnice Lubelskiej podczas zajęć z Elektronicznej Aparatury Medycznej. Zarówno konstruktor, prowadzący laboratorium, jak i uczestniczący w nich studenci byli zadowoleni z ćwiczeń na symulatorze nie zostały zgłoszone żadne uwagi odnośnie stanowiska.

Projekt nadaje się, aby wprowadzić go do powszechnego stosowania podczas szkolenia laparoskopowego lekarzy. Oczywiście jest, iż istnieje ciągła potrzeba ulepszania symulatora. Rozwiązania wykorzystane podczas projektowania stanowiska pozostawiają otwarte drzwi do jego udoskonalania.

LITERATURA

- [1] Fibak J., *Chirurgia: repetytorium*, Wydawnictwo Lekarskie PZWL, Warszawa 1998.
- [2] Gruca Z., Kobiela J., Stefaniak T., *Przydatność symulatorów w nauczaniu małoinwazyjnej techniki operacyjnej w chirurgii*, „Wideochirurgia i inne zabiegi małoinwazyjne”, nr 3, 2008, s. 30–34
- [3] Jabłoński J., *Search for the diagnostic tools for assessment of laparoscopic skills*, “Family Medicine & Primary Care Review”, nr 15, 2013, s. 112–113
- [4] Kostewicz W., *Chirurgia laparoskopowa*. Wydawnictwo Lekarskie PZWL, Warszawa 2002
- [5] Mishra R., *Textbook of Practical Laparoscopic Surgery*, Jaypee Brothers Medical Publishers Ltd., New Delhi 2013
- [6] Mouret P., *How I developed laparoscopic cholecystectomy*, “Ann Acad. Med.”, Singapore, 1996, nr 25, s. 744
- [7] Noszczyk W., *Chirurgia*, Wydawnictwo Lekarskie PZWL, Warszawa 2007
- [8] Scott D., *Laparoscopic Training on Bench Models: Better and More Cost Effective than Operating Room Experience?* American College of Surgeons 85th Annual Clinical Congress, Surgical Forum, San Francisco, CA, October 13, 1999, s. 272–283
- [9] Siemiński P., Budzik G., *Techniki przyrostowe: druk drukarki 3D*, Oficyna Wydaw. Politechniki Warszawskiej, Warszawa 2015
- [10] Torres K. et al., *Simulation techniques in the anatomy curriculum: review of literature*, “Folia Morphol”, 2014, nr 73, s. 1–6
- [11] Vigue J., *Prace w drewnie*, Wydawnictwo Arkady, Warszawa 2010
- [12] <http://www.altax.pl/zabezpiecz-swoj-dom-przed-wizyta-niechcianych-gosci-czesc-i-drewno-wewnatrz/> (zasoby z dnia 3.12.2016)
- [13] http://www.optec.pl/produkty/ap_med/laparo.html (zasoby z dnia 27.10.2016)

KONCEPCJA UKŁADU ZAWIESZENIA DO POJAZDU KLASY FORMULA STUDENT W PROGRAMIE VI – MOTORSPORT

1. WPROWADZENIE

Zagadnienie projektowania zawiesznień istniało znacznie wcześniej przed pojawieniem się pojazdów silnikowych. Każdy pojazd wymaga układu łączącego nadwozie z kołami. Pierwsze pojazdy konne wyposażone były w sztywne połączenia pomiędzy miejscem przebywania pasażerów, a osiami kół. Takie połączenia były bardzo niekomfortowe dla podróżujących. Dlatego już XVI w. w dworskich karocach zaczęto stosować elementy sprężysto – tłumiące w postaci pasów skórnych. Późniejszy okres przyniósł wynalazek w postaci pierwszych resorów[1]. Rozwiązania te spełniały podstawowe zadanie. W pewnym stopniu tłumiły drgania powstające podczas jazdy i zmniejszały siły dynamiczne. Od obecnych zawiesznień w pojazdach silnikowych wymaga się znacznie więcej [2]:

- dobrej izolacji nadwozia od sił powstających na styku opony i nawierzchni oraz zapewnienie możliwie stałego obciążenia koła
- umożliwienie skrętu kół oraz zagwarantowanie jego odpowiedniego położenia względem powierzchni drogi
- ograniczenie ruchu kół do odpowiedniego zakresu, przeciwdziałanie przechyłom nadwozia przy przyspieszaniu, hamowaniu i pokonywaniu zakrętów
- tłumienie drgań wysokiej częstotliwości, ograniczenie powstawania hałasu, zapewnienie odpowiedniego komfortu pasażerom.

Jednoczesne spełnienie wszystkich powyższych wymagań jest zadaniem niezwykle trudnym. W związku z tym dąży się do pewnych kompromisów. W pojazdach używanych na co dzień ważne jest zapewnienie odpowiedniego komfortu. Wymagania dotyczące wygody mają tu stosunkowo większą wagę niż w pojazdach wyścigowych, gdzie najistotniejsze znaczenie ma zapewnienie odpowiedniej przyczepności do nawierzchni oraz zapewnienie odpowiednich osiągnięć.

W zależności od przeznaczenia pojazdu stosuje się różne typy zawiesznień. Podstawową klasyfikację ze względu na przemieszczenia jednego koła danej osi względem drugiego stanowi podział na zawieszania zależne, niezależne i pół niezależne. Zawieszania zależne stosowane są przeważnie w samochodach ciężarowych gdzie istotna jest duża wytrzymałość oraz możliwość przeniesienia du-

¹ Politechnika Lubelska, Koło Naukowe Podstaw Inżynierii Produkcji

punktów zawieszenia. W menu programu wybiera się typ zawieszenia, natomiast w oknach dialogowych wprowadza koordynaty poszczególnych punktów, sztywności sprężyn, współczynniki tłumienia amortyzatorów itp. Dzięki temu zmian geometrycznych zawieszenia, a co za tym idzie charakterystyki dokonuje się błyskawicznie poprzez wpisanie nowych współrzędnych. Co więcej program pozwala na wizualizację zaprojektowanego zawieszenia oraz pokazuje różnice pomiędzy poprzednimi wersjami projektu. Zrzut ekranu z modułu VI – SuspensionGen zaprezentowano na rysunku 1.

Zawieszenia w bolidach wyścigowych

W wielu seriach wyścigowych takich jak Formula 1, Formula 2, Formula 3 czy wyścigi Formuły Student stosuje się podwójne wahacze poprzeczne jako elementy wodzące w układach zawiesznień bolidów[6]. Przykłady bolidów F2 i F3 zaprezentowano na rysunkach 2 i 3.



Rys. 2. Przykład bolidu Formuły 2 [7]

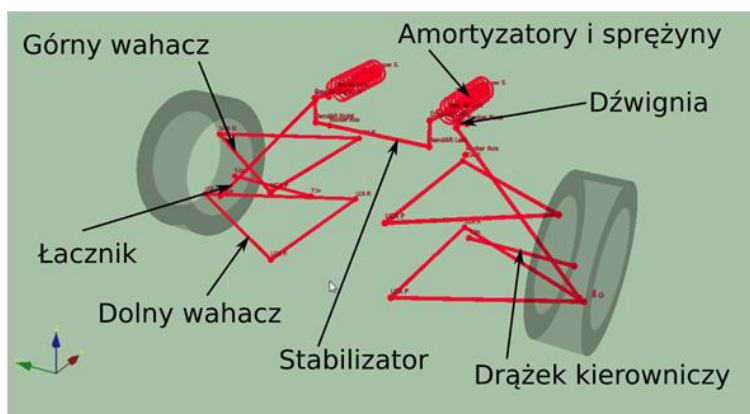
Koncepcja tego typu zawieszenia zwanego z j. angielskiego *Double Wishbone* lub *Four Bar Link Independent Suspension System* została zastosowana po raz pierwszy w przednim zawieszeniu samochodu wyścigowego w latach 30. XX. w. Natomiast w latach 60. XX. w. weszła do powszechnego użycia i jest wykorzystywana do dziś. Powodami stosowania tego typu zawieszenia były: niska masa nieresorowana, nisko położony środek bocznego obrotu zawieszenia oraz szerokie możliwości ustawiania geometrii porównaniu z innymi rozwiązaniami zawiesznień [9]. Ponadto w takim zawieszeniu niezależnym stosuje się dodatkowo stabilizatory, które ograniczają przechyły nadwozia podczas pokonywania zakrętów [10]. Dodatkowo zapewniona jest możliwość dowolnego przestrzennego rozmieszczenia elementów sprężystych i tłumiących w układzie pull lub push.

Dzięki temu konstruktor ma możliwość zdecydowania, czy przy pionowym przemieszczeniu koła są one ściskane czy rozciągane. W celu przenoszenia sił na wyżej wymienione komponenty używa się łączników oraz dźwigni [9]. Wszystkie przedstawione powyżej cechy decydują o szerokim zastosowaniu tego rozwiązania we współczesnych bolidach wyścigowych.



Rys. 3. Przykład bolidu Formuły 3[8]

Przykład zawieszenia Double Wishbone wraz z opisem części składowych zaprezentowano na rysunku 4.



Rys. 4. Przykład zawieszenia Double Wishbone

Konkurs Virtual Formula 2017

Celem przeprowadzenia badań było przygotowanie możliwie najbardziej konkurencyjnego pojazdu do międzynarodowego konkursu VI – Grade Virtual Formula 2017. Zadaniem biorących udział w konkursie zespołów jest optymalizacja osiągnięć elektrycznego, wirtualnego modelu samochodu wyścigowego. Celem jest minimalizacja czasu okrążenia na zadanym torze, uzyskanie największego przyspieszenia na 100m i jak najszybszy przejazd po torze w kształcie ósemki (tzw. Skidpad). Zwycięża zespół, który osiągnie największą ogólną liczbę punktów z poszczególnych zadań. Zgodnie z wymaganiami stawianymi przez organizatorów pojazd musiał spełniać wszystkie wymogi pojazdu klasy Formula Student [11].

2. ZAŁOŻENIA KONCEPCYJNE DO BUDOWY MODELU

Pierwszym etapem projektowania zawieszenia było przyjęcie odpowiednich założeń konstrukcyjnych. Główne cele związane z ulepszeniem modelu dostarczonego przez organizatorów [12, 13, 14]:

- zmniejszenie rozstawu kół obydwu osi
- zmniejszenie skoku zawieszenia
- redukcja zmian kąta zbieżności względem przemieszczenia pionowego koła
- zachowanie korzystnej charakterystyki zmian kąta pochylenia koła
- obniżenie środka bocznego przechyłu zawieszenia.

Założenia do projektu przedniego zawieszenia:

- typ zawieszenia – podwójne wahacze poprzeczne nierównoległe z krótszym wahaczem górnym, połączenie sprężyny i amortyzatora z dolnym wahaczem w układzie push,
- rozstaw osi kół przednich 1250mm wobec 1400mm w modelu bazowym,
- skok zawieszenia ograniczony do 80mm wobec 95mm w modelu bazowym,
- wprowadzenie dodatniego promienia zataczania o wartości 40mm.

Założenia do projektu tylnego zawieszenia:

- typ zawieszenia – podwójne wahacze poprzeczne nierównoległe z krótszym wahaczem górnym, połączenie sprężyny i amortyzatora z górnym wahaczem w układzie push,
- rozstaw osi kół przednich 1200mm wobec 1400mm w modelu bazowym,
- skok zawieszenia ograniczony do 80mm wobec 95mm w modelu bazowym.

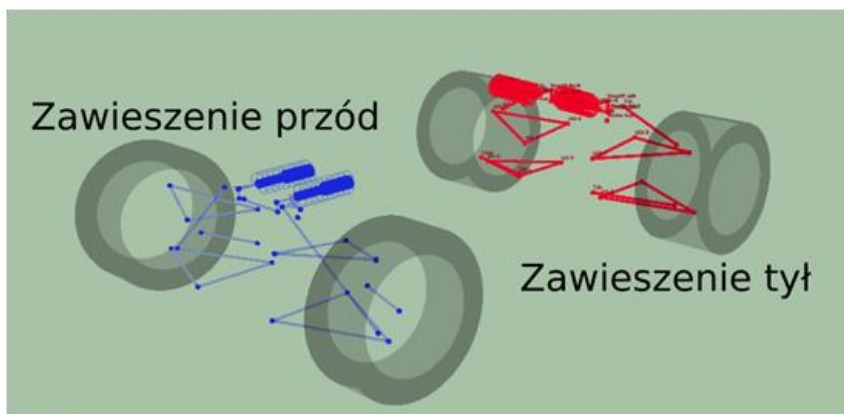
Modyfikacji charakterystyk dokonywano poprzez zmianę położenia głównych punktów mocowania elementów zawieszenia w przestrzeni.

3. UZYSKANE REZULTATY

W rezultacie uzyskano szereg charakterystyk oraz wizualizację gotowego zawieszenia. Na rysunku 5 przedstawiono rzeczoną wizualizację. Rysunki 6–9 przedstawiają kolejno zmianę kąta pochylenia koła, zbieżności, położenia środka bocznego przechyłu oraz przeciwdziałanie obniżaniu przodu pojazdu podczas hamowania. Wszystkie charakterystyki zostały wykreślone w funkcji przemieszczenia pionowego koła.

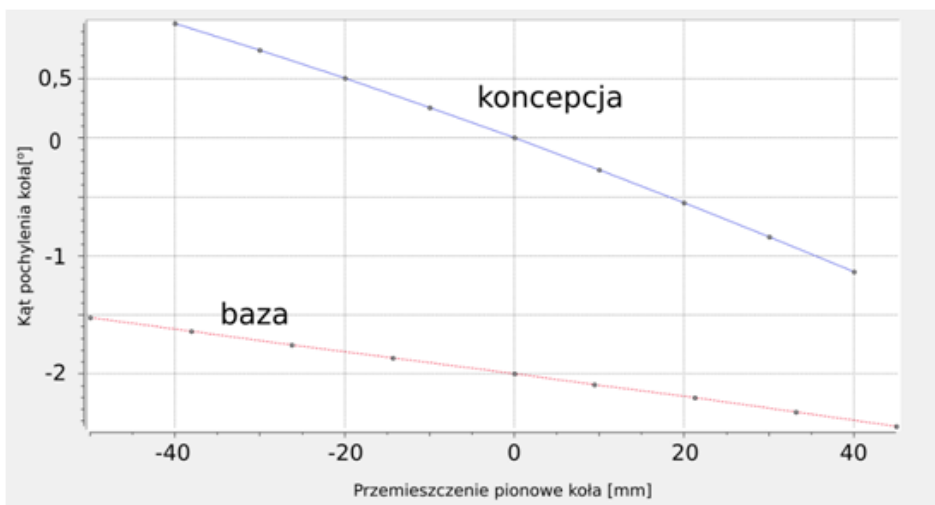
Charakterystyki zawieszenia tylnego znajdują się na rysunkach 10–12. Znajdują się na nich krzywe zmian kąta pochylenia koła, zbieżności, położenia środka bocznego przechyłu w funkcji przemieszczenia pionowego koła.

Różnice w położeniu krzywych zmian kąta pochylenia koła na charakterystykach dla przedniego i tylnego zawieszenia wynikają z przyjęcia innych podstawowych kątów pochylenia koła w zawieszeniu bazowym (-2° dla przedniego, -1° dla tylnego) i koncepcyjnym (0°).



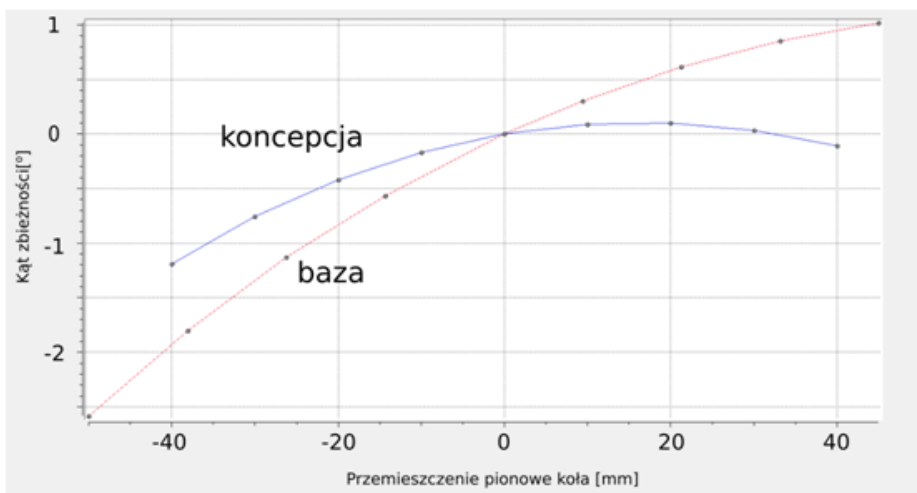
Rys. 5. Wizualizacja zawieszenia pojazdu Formula Student

Wizualizacja prezentuje rozplanowanie przestrzenne komponentów zawieszenia. Umieszczenie amortyzatorów i sprężyn śrubowych w tylnym zawieszeniu ponad elementami wodzącymi zapewniło więcej miejsca dla układu przeniesienia napędu.



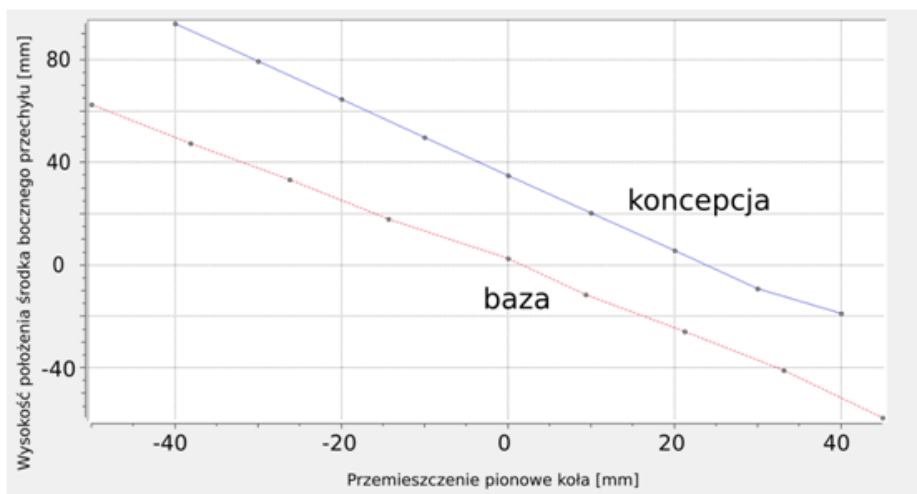
Rys. 6. Charakterystyka zmian kąta pochylenia koła dla zawieszenia przedniego

Modyfikacje zawieszenia przedniego pozwoliły na zachowanie korzystnej charakterystyki zmian kąta pochylenia koła. Na podstawie przebiegu krzywej można stwierdzić, że kąt ten zmniejsza się wraz z większym ugięciem zawieszenia. Takie zachowanie poprawia przyczepność pojazdu w zakrętach.



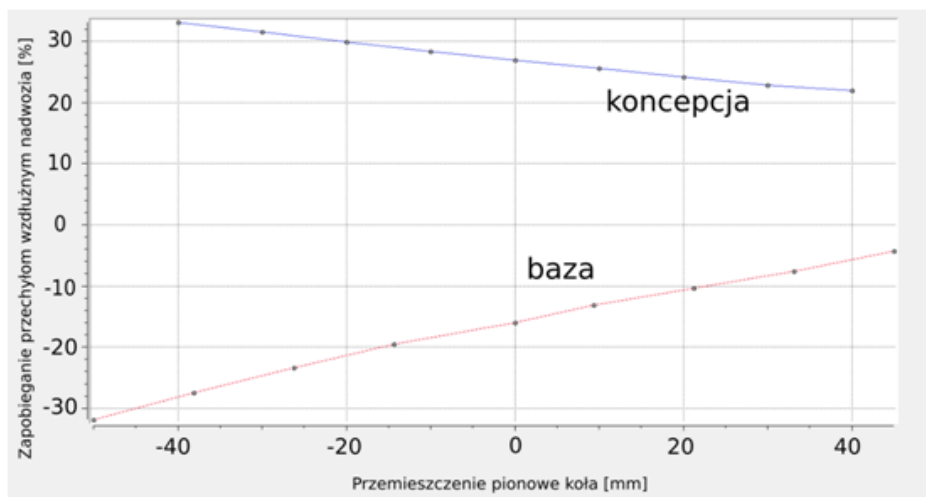
Rys. 7. Charakterystyka zmian zbieżności dla zawieszenia przedniego

Wprowadzone ulepszenia znacznie poprawiły przebieg zmienności kąta zbieżności. Zakres zmian jest znacznie niższy w zawieszeniu koncepcyjnym niż w bazowym.



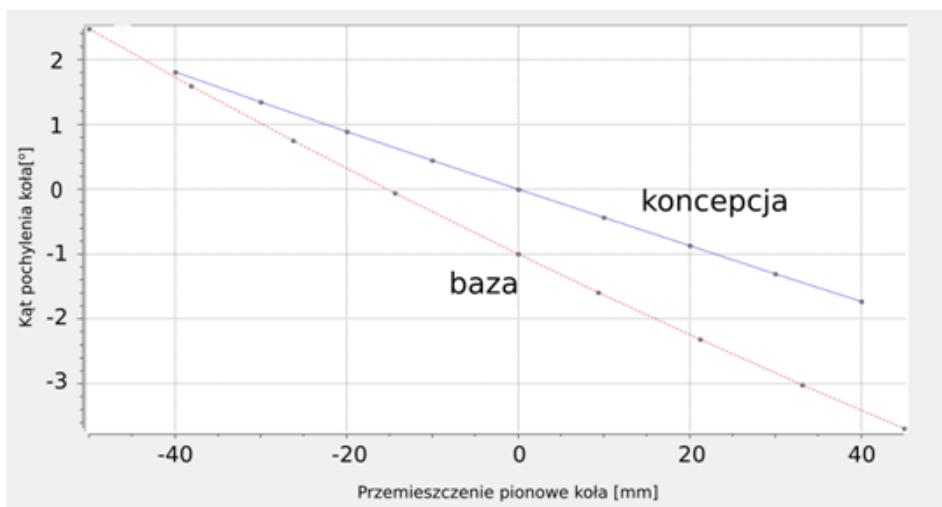
Rys. 8. Charakterystyka zmian położenia środka bocznego przechyłu dla zawieszenia przedniego

Niestety nie powiodła się próba obniżenia położenia środka bocznego przechyłu nadwozia. Podniósł się on ok. 30mm. Jednakże pozostałe pozytywne aspekty modyfikacji pozwoliły na akceptację tej wartości.



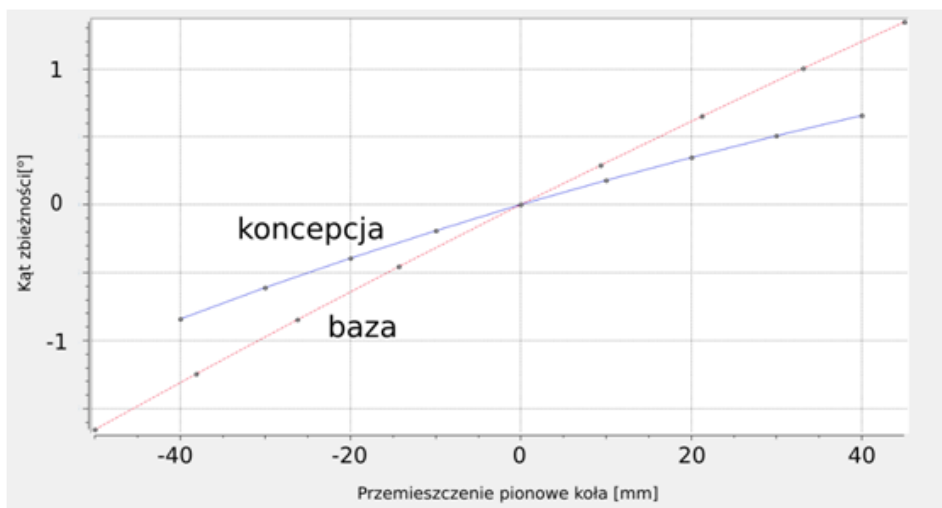
Rys. 9. Charakterystyka przeciwdziałania przechyłom wzdłużnym nadwozia dla zawieszenia przedniego

Zastosowane zmiany kinematyki pozwoliły również na zwiększenie zdolności pojazdu do przeciwdziałania przechyłom wzdłużnym bolidu. Dzięki temu ograniczone zostało „siadanie” przodu pojazdu podczas hamowania.



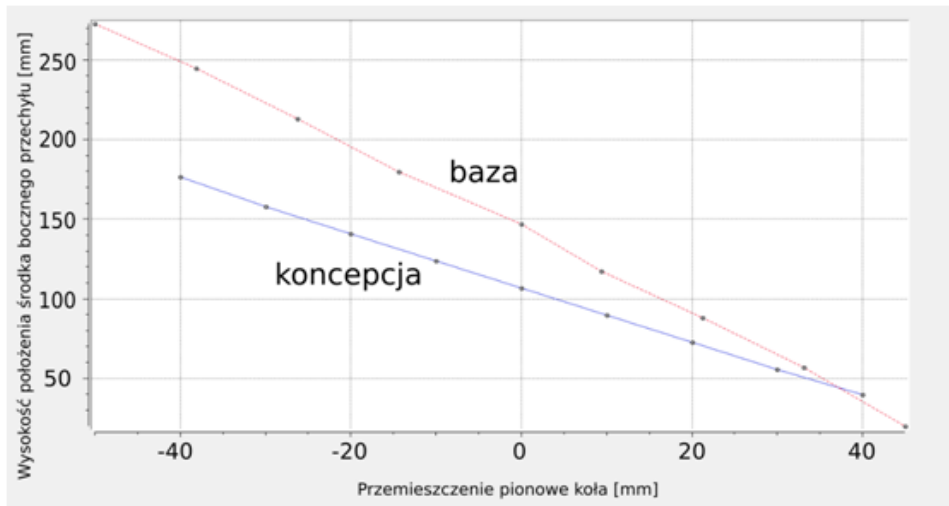
Rys. 10. Charakterystyka zmian kąta pochylenia koła dla zawieszenia tylnego

W zawieszeniu tylnym zmiana charakterystyki kąta pochylenia koła, podobnie jak w zawieszeniu przednim pozostała korzystna. W dalszym ciągu wpływa pozytywnie na sterowność samochodu.



Rys. 11. Charakterystyka zmian zbieżności dla zawieszenia tylnego

Uzyskano również znacznie niższy zakres zmienności kąta zbieżności w stosunku do tylnego zawieszenia bazowego.



Rys. 12. Charakterystyka zmian położenia środka bocznego przechyłu dla zawieszenia tylnego

Dążenia do obniżenia środka bocznego przechyłu w zawieszeniu tylnym powiodły się. Dodatkowo jego położenie pozostało wyżej niż w zawieszeniu przednim. Powyższa charakterystyka wskazuje na poprawę stabilności pojazdu w zakręcie.

4. PODSUMOWANIE

Przeprowadzone badania pozwoliły na przygotowanie modelu kompletnego układu zawieszenia przedniego i tylnego do pojazdu klasy Formula Student. Uzyskane charakterystyki są satysfakcjonujące. Zastosowane modyfikacje w zawieszeniu koncepcyjnym pozwoliły na poprawę przebiegu krzywych względem zawieszenia bazowego. Uzyskane przestrzenne położenia głównych punktów mocowań można wykorzystać przy projekcie rzeczywistego pojazdu.

Implementacja koncepcyjnych modeli zawieszeń w wirtualnym pojeździe pozwoliła na redukcję czasu przejazdu jednego okrążenia o 2 sekundy. Stanowi to zysk na poziomie ok. 2,5%. Wynik ten osiągnięto tylko poprzez zastosowanie nowego układu zawieszenia. Co więcej, zmniejszenie rozstawu osi względem pojazdu bazowego poprawi zwrotność pojazdu, pozwoli lepiej wykorzystać dostępną szerokość toru, przez co możliwe będzie zastosowanie łagodniejszej ścieżki przejazdu i zwiększenie prędkości przejazdu. W związku z tym przewidywane są kolejne redukcje czasu pokonywania okrążenia.

LITERATURA

- [1] *Encyklopedia Powszechna PWN*, Państwowe Wydawnictwo Naukowe, Warszawa, 1984
- [2] Happian-Smith J. i inni, *An Introduction to Modern Vehicle Design*, Reed Educational and Professional Publishing Ltd, Great Britain, 2002
- [3] Gillespie T.D., *Fundamentals of Vehicle Dynamics*, Society of Automotive Engineers, Warrendale, 1992
- [4] Tandel A. i inni, *Modeling, Analysis and PID Controller Implementation on Double Wishbone Suspension Using SimMechanics and Simulink*, 12th Global Congress on Manufacturing and Management, Vellore, 2014, s.1274–1281
- [5] Fudała W., Nowak R., *Badania kinematyki i dynamiki zawieszenia i układu kierowniczego pojazdu Formula Student*, „Inżynierowie nowej ery”, 2010, s. 119–128
- [6] Chepkasov S., *Suspension Kinematics Study of the "Formula SAE" Sports Car*, International Conference on Industrial Engineering ICIE 2016, Yekaterinburg, 2016, s.1280–1286
- [7] Formula 2 Official, <http://www.fiaformula2.com/> (zasoby z dnia 02.06.2017)
- [8] Formula 3 Official, <http://www.fiaf3europe.com/2017-drivers-and-teams/> (zasoby z dnia 02.06.2017)
- [9] Smith C., *Tune to win*, AERO PUBLISHERS INC., United States, 1978
- [10] Bielecki B., *Fizyka samochodów wyścigowych*, “Scientific Bulletin of Chełm Section of Mathematics and Computer Science”, nr 1, 2009, s. 15–21
- [11] Konkurs Virtual Formula 2017, <http://www.vi-grade.com> (zasoby z dnia 02.06.2017)
- [12] Saurabh S. Y., *Design of Suspension System for Formula Student Race Car*, 12th International Conference on Vibration Problems ICOVP, 2015, Guwahati, 2016, s.1138–1149
- [13] Formula Student Rules 2017, <https://www.formulastudent.de/fsg/rules/> (zasoby z dnia 02.06.2017)
- [14] Sayers M., *Standard Terminology for Vehicle Dynamics Simulations*, The University of Michigan Transportation Research Institute (UMTRI), Michigan, 1996

POLIMERY BIODEGRADOWALNE A ZASTOSOWANIA MEDYCZNE

1. WSTĘP

W dzisiejszych czasach tworzywa polimerowe są masowo używane we wszystkich obszarach życia, między innymi w rolnictwie, motoryzacji, przemyśle spożywczym i opakowaniowym oraz w medycynie. W związku z tak szerokim zastosowaniem polimerów obserwuje się znaczący wzrost ilości odpadów z tworzyw sztucznych.

Kompostowalność i biodegradowalność stają się pożądaną właściwością produktów jednorazowego użycia i opakowań, na które wciąż rośnie zapotrzebowanie. Chęć zmniejszenia emisji dwutlenku węgla oraz potrzeba uniezależnienia od paliw kopalnych wpływa na poszukiwanie biopochodnych materiałów jako alternatyw dla tradycyjnych tworzyw i daje możliwość ochrony środowiska.

Cieężko jest sobie wyobrazić funkcjonowanie współczesnej medycyny bez tworzyw sztucznych, które są wykorzystywane między innymi jako rękawiczki ochronne, strzykawki, implanty. Tworzywa sztuczne napędzają postęp w medycynie, a co za tym idzie polepszają jakość życia a nawet są w stanie je wydłużyć.

Tworzywa polimerowe w medycynie znalazły zastosowanie jako przyrządy jednorazowego użytku oraz przedmioty wysoce rozwinięte technologicznie. Do drugiej grupy zastosowań możemy zaliczyć między innymi implanty ortopedyczne, stenty oraz protezy naczyń krwionośnych wykorzystywane w kardiologii, sztuczne rogówki, aparaty słuchowe a także nośniki leków.

Podstawowymi i najważniejszymi właściwościami jakie muszą spełniać polimery stosowane w medycynie jest nietoksyczność, możliwość sterylizacji bez wywołania zmian fizykochemicznych oraz biokompatybilność czyli możliwość działania z żywym systemem jakim jest organizm. Ponadto biopolimery w zależności od zastosowania powinny charakteryzować się między innymi odpowiednią wytrzymałością mechaniczną.

2. CHARAKTERYSTYKA POLIMERÓW STOSOWANYCH W MEDYCYNIE

Polimery stosowane w medycynie należą do grupy polimerów specjalnych zastosowań. Stawiany jest im szereg wymagań, które muszą spełnić, aby mogły funkcjonować w tak specyficznej dziedzinie życia jaką jest medycyna.

¹ Politechnika Lubelska, Wydział Elektrotechniki i Informatyki

Polimery można klasyfikować w różny sposób, w stosunku do przyjętego kryterium podziału. Trzeba jednak pamiętać, że nie ma wystarczająco precyzyjnego kryterium podziału, który zawierałby wszystkie ujęcia tak zróżnicowanej grupy, jaką są polimery.

Polimery biomedyczne ze względu na zastosowanie można podzielić na:

- polimery o krótkim czasie stykania się materiału polimerowego z zewnętrzną częścią organizmu, np. podczas rehabilitacji czy diagnozowania.
- polimery o długim czasie styku materiału polimerowego z zewnętrzną częścią organizmu, np. soczewki kontaktowe, protezy kończyn.
- polimery wykorzystane do budowy urządzeń, narzędzi oraz części aparatury, np. dreny, przewody do hemodializy.
- polimery stosowane jako środki farmakologiczne do wprowadzania leków do organizmu.
- polimery wykorzystywane do budowy części, które są na stałe wszczepiane do wnętrza organizmu, np. sztuczne zastawki serca, nici chirurgiczne.

Na początku oceny przydatności tworzywa polimerowego w medycynie, uwzględnia się wytrzymałość mechaniczną, kształt oraz inne własności fizyczne polimerów. Na podstawie badań fizycznych można sprawdzić reakcję na strukturę cząsteczek tworzywa lub nieodpowiednią wielkość. Ponadto można określić nierówności powierzchni tworzywa, mogące uszkadzać tkanki na skutek tarcia.

Wyklucza się polimery o niewłaściwych wymiarach kryształitów oraz cząsteczek wypełniaczy za pomocą badania, które dokonuje pomiaru dyfrakcji promieni rentgenowskich.

Wymagania są stawiane materiałom w zależności od tego jaką funkcję mają pełnić w medycynie oraz przez jaki czas mają być bezpiecznie użytkowane, a określa się to osobno dla każdego materiału polimerowego na podstawie właściwości użytkowych i funkcji jakie ma spełniać. Czy będą z nich wykonane implanty, czy przedmioty o krótkotrwałym czasie kontaktu z organizmem żywym.

Wszystkie polimery o zastosowaniu medycznym spełniają następujące warunki [6]:

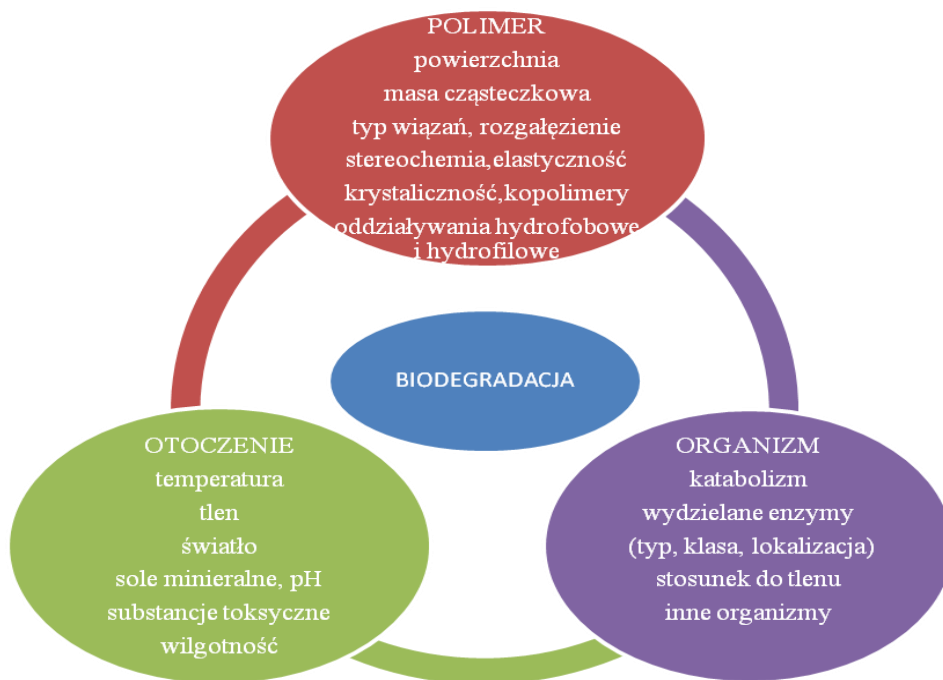
- muszą być łatwe do utrzymania w czystości –czyszczenie i wyjaławianie
- są odporne na działanie: środków chemicznych, odkażających, myjących, a także działanie wysokiej temperatury, czynników fizjologicznych i promieniowania rentgenowskiego
- nie mogą być toksyczne, ani działać rakotwórczo
- muszą odznaczać się biogodnością, to znaczy, że nie mogą powodować stanów zapalnych, alergii a także reakcji immunologicznych i mutagennych
- są łatwe do formowania, co pozwala na nadanie odpowiedniej formy użytkowej bez degradacji polimeru
- muszą być biokompatybilne, inaczej mówiąc biotolerancyjne.

3. BIODEGRADACJA MATERIAŁÓW POLIMEROWYCH

Pod pojęciem degradacji znajdują się zarówno procesy kończące się zmniejszeniem masy molowej, a także procesy, które prowadzą do sieciowania i powstania struktur rozgałęzionych. Degradacja może następować pod wpływem wielu, różnych czynników: chemicznych takich jak powietrze; fizycznych na przykład ultradźwięki, promienie jonizujące, światło słoneczne; biologicznych czyli enzymy, grzyby, bakterie; a także pod wpływem podwyższonej temperatury [4].

Degradacja jest procesem nieodwracalnym i wieloetapowym, który powoduje wyraźne zmiany w strukturze chemicznej polimeru. W trakcie degradacji w łańcuchu głównym polimeru pękają wiązania kowalencyjne, co skutkuje utratą właściwości użytkowych polimeru. Ponadto zmienia się właściwości: fizyczne i chemiczne na przykład struktura chemiczna, ciężar cząsteczkowy; a także właściwości mechaniczne, połysk, przezroczystość itp. [7].

Na rysunku 1 zostały przedstawione czynniki wpływające na szybkość biodegradacji.



Rys. 1. Czynniki wpływające na szybkość biodegradacji

W związku z czynnikami powodującymi biodegradację wyróżniamy [7]:

- degradację abiotyczną – następuje pod wpływem ciepła, promieniowania elektromagnetycznego, jak również aktywnych związków chemicznych
- degradację biotyczną – jest wywołana działaniem czynników biologicznych czyli organizmów żywych, zwłaszcza enzymów produkowanych przez rozmaite mikroorganizmy, takie jak grzyby czy bakterie. Degradacja biotyczna nazywana jest również degradacją biologiczną.

Biodegradacja tworzyw polimerowych następuje w dwóch etapach. Pierwszy polega na bezpośrednim enzymatycznym rozszczepianiu makrocząsteczek wraz z metabolizmem produktów rozpadu lub sukcesywnej asymilacji enzymatycznej makrocząsteczek zaczynających się od końców łańcucha. Drugi etap to rozpad oksydacyjny makrocząsteczki razem z metabolizmem fragmentów (jest to biodegradacja pośrednia) [7].

Badania biodegradacji wykonuje się w celu oszacowania podatności danego materiału na rozpad mikrobiologiczny. Polegają one na określaniu podatności pojedynczych związków organicznych oraz ich mieszanin na rozkład mikrobiologiczny.

4. POLIMERY BIODEGRADOWALNE

Polimery biodegradowalne zaliczają się do klasycznych termoplastów, przejawiają dobre własności fizykochemiczne, fizykomechaniczne, a także biodegradowalne. Modyfikując strukturę łańcucha polimerowego można wpłynąć na czas życia tego typu polimerów. Przebieg biodegradacji zależy również od warunków środowiska, rodzaju mikroorganizmów czynnych, wiązań sieciujących, ciężaru cząsteczkowego, wielkości i kształtu wyrobu itp. Proces ten prowadzi do powstania małych cząsteczkowych produktów nie będących zagrożeniem dla środowiska [1].

Polilaktyd (PLA)

Powszechnie PLA stanowi blisko 40% wszystkich wytwarzanych z surowców odnawialnych tworzyw biodegradowalnych.

Polilaktyd charakteryzuje duża polarność, co przyczynia się do jego mniejszej adhezji. Należy do grupy termoplastycznych poliestrów i jest wytwarzany z surowców odnawialnych. Może być poddawany procesom przetwórstwa tworzyw na typowych urządzeniach bez potrzeby ulepszania sprzętu.

PLA można wytworzyć dwoma sposobami. Pierwszy jest tańszy, sprowadza się do bezpośredniej polikondensacji kwasu mlekowego w wyniku czego produkt ma dość mały ciężar cząsteczkowy, co skutkuje gorszymi właściwościami mechanicznymi. Następny sposób to metoda dwuetapowa, która polega na otrzymaniu laktydy z kwasu mlekowego a dalej na jego polimeryzacji, czyli tak zwana polimeryzacja z otwarciem pierścienia laktydu. Wyprodukowanie poli-

laktydu tą metodą ułatwia otrzymanie polimeru o dużym ciężarze cząsteczkowym (powyżej 50 000) i dobrych właściwościach mechanicznych. W zależności od miejsca gdzie pęka wiązanie pierścienia, polimeryzacja polilaktydu może zachodzić według dwóch mechanizmów jonowych. Anionowym jeżeli rozerwaniu ulegnie wiązanie pomiędzy tlenem a węglem grupy acylowej; natomiast kationowym, gdy rozerwanie będzie miało miejsce pomiędzy tlenem a węglem grupy alkilowej [2, 8].

PLA ulega całkowitej biodegradacji (produkty jego rozkładu są naturalnymi metabolitami), następnie są przyswajany przez żywe organizmy między innymi bakterie. Można również poli(kwas mlekowy) poddawać recyklingowi [5].

Pierwszą fazą degradacji polilaktydu jest hydroliza, w rezultacie czego powstaje kwas mlekowy oraz oligomery – chiralność atomów węgla nie zmienia się. Drugą fazą degradacji jest mineralizacja i rozpad. Degradacja w zależności od stopnia krystaliczności może trwać od kilku tygodni do nawet kilku lat [3].

Poprzez proces kopolimeryzacji można modyfikować mechaniczne właściwości polilaktydu. Najczęściej stosowanym do tego kopolimerem jest glikolid. W ten sposób powstaje PLGA czyli poli(laktyd-ko-glikolid, który jest materiałem amorficznym, całkiem biodegradowalnym, a także nadającym się do kompostowania. Przez dodanie kaprolaktonu do PLA otrzymujemy elastomeryczny polimer, który ma budowę amorficzną i jest całkowicie biodegradowalny [2].

Polilaktyd znajduje zastosowanie w dwóch głównych nurtach; pierwszy to masowe wykorzystanie jako tworzywo termoplastyczne wytwarzane z surowców odnawialnych i potrafiące samorzutnie degradować hydrolitycznie i biologicznie po założonym z góry czasie eksploatacji; drugi nurt to wykorzystanie tego polimeru w sposób specjalistyczny – biomedyczny [1].

Zastosowanie polilaktydu w medycynie

Polilaktyd ze względu na swoją termoplastyczność, biodegradowalność, dobre właściwości mechaniczne, a także nietoksyczność znalazł swoje zastosowanie między innymi w medycynie i farmacji. PLA główne zastosowania znalazł w systemach kontrolowanego dostarczania i uwalniania leku, do produkcji śrub ortopedycznych, produkcji hydrożeli oraz w inżynierii tkankowej. Naprawa uszkodzonej tkanki kostnej jest dużo bardziej efektywna poprzez wykorzystanie rusztowań polimerowych z polilaktydu.

Możliwość wykorzystania takich polimerów w procesach inkapsulacji leków odnosi się do obszarów, w których tradycyjne dostarczenie leków jest mało skuteczne, ze względu na częstotliwość i miejsce podawania. Poprzez zmianę właściwości polimerowego nośnika leku w środowiskach o różnych wartościach pH, możemy wpłynąć na miejsce i ilość uwalnianego leku. Stosowane jako nośniki leków nanowłókna wytwarzane są metodą elektroprzędzenia z mieszaniny roztworu polimeru z lekiem.

5. METODY BIODEGRADACJI – BADANIA WŁASNE

W badaniach został wykorzystany polilaktyd PLA L175, który jest dedykowany do procesów wytłaczania i termoformowania. W tabeli 1 zostały przedstawione wybrane właściwości PLA L175.

Tab. 1. Właściwości PLA L175 według danych producenta

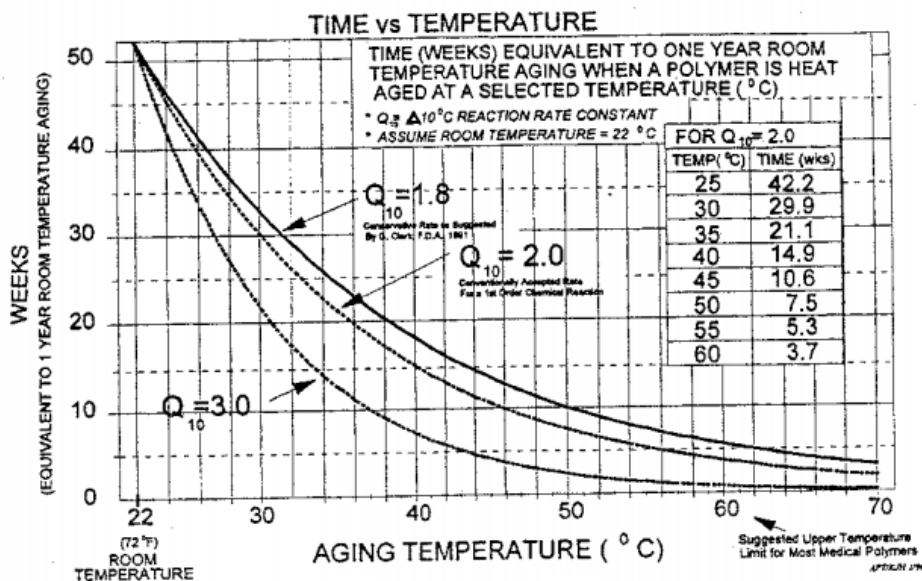
Własność		Wartość	Jednostki
Aplikacja	Termoformowanie, włókna, wytłaczanie	–	–
Właściwości fizyczne	Gęstość	1,24	g/cm ³
	Czystość optyczna	>99	% L-izomeru
Przetwarzanie	Wskaźnik szybkości płynięcia MFI	6	g/10min
	Temperatura topnienia T _m	175	°C
	Temperatura zeszklenia T _g	57	°C
	Podsuszanie przed tworzeniem	Tak	–
Właściwości mechaniczne	Moduł sprężystości	3500	MPa
	Wytrzymałość na rozciąganie	50	Mpa
	Wydłużenie przy zerwaniu	<5	%

Próbki zostały przygotowane z folii powstałej w wyniku procesu wytłaczania z rozdmuchiowaniem folii na linii technologicznej W25/25. Próbki wykonane z biodegradowalnej folii z polilaktyd PLA L175 zostały poddane procesom biodegradacji abiotycznej pod wpływem temperatury oraz biodegradacji biotycznej czyli biologicznej.

Degradacja abiotyczna (termiczna)

Zgodnie z normą ASTM F1980-07(2011) przyjęto temperaturę starzenia w komorze cieplnej 60°C ze względu na to, że jest to górna granica temperatury dla większości polimerów stosowanych w medycynie. Przyjęto, że Q_{10} ma wartość 2, ponieważ ta wartość jest powszechnie stosowana do obliczeń współczynnika starzenia (na podstawie równań Arrheniusa). Bardziej agresywne współczynniki są przyjmowane, gdy badany system jest wystarczająco dobrze scharakteryzowany w literaturze.

Na podstawie wykresu znajdującego się na rysunku 2 przyjęto, że próbki będą degradować przez okres 4 tygodni w temperaturze 60°C co odpowiada degradacji w temperaturze pokojowej przez jeden rok.



Rys. 2. Wykres przedstawiający zależność czasu starzenia od temperatury [ASTM F1980-07 (2011)]

W komorze termicznej umieszczono 12 próbek wyciętych z folii PLA 175. Co tydzień były wyciągane po trzy próbki folii i były one poddawane rozciąganiu na maszynie wytrzymałościowej. Badania były powtarzane co tydzień przez cztery tygodnie w ciągu których zbadano łącznie 12 próbek folii wykonanej z polilaktyd PLA L175.

Degradacja biotyczna (biologiczna)

W celu przeprowadzenia badań dotyczących wpływu degradacji biotycznej (biologicznej) na właściwości mechaniczne, próbki PLA zostały umieszczone w ziemi kompostowej na 5 tygodni zgodnie z normą PN-EN 13432:2002.

Co tydzień dokonywano odczytu wilgotności powietrza, która wahała się w granicy od 69 do 92%. Co tydzień również wyciągano po 3 próbki PLA w celu zbadania ich właściwości mechanicznych na maszynie wytrzymałościowej. Czynności te były powtarzane przez 5 tygodni.

Badanie właściwości mechanicznych

Podczas badań wykorzystywano maszynę wytrzymałością firmy Zwick/Roell model Z010 wyposażoną w uchwyty śrubowo-klinowe typu 8306 przedstawione na rysunku 3.



Rys. 3. Próbkki umieszczone w uchwytach maszyny Zwick/Roell Z010

Podczas badań siła wstępna wynosiła 0,1 MPa, prędkość badania 50 mm/min, a odległość między uchwytami przy pozycji startowej była równa 110 mm. Badania były wykonane zgodnie z normami: PN-EN ISO 527-1, PN-EN ISO 527-2 oraz PN-EN ISO 527-3.

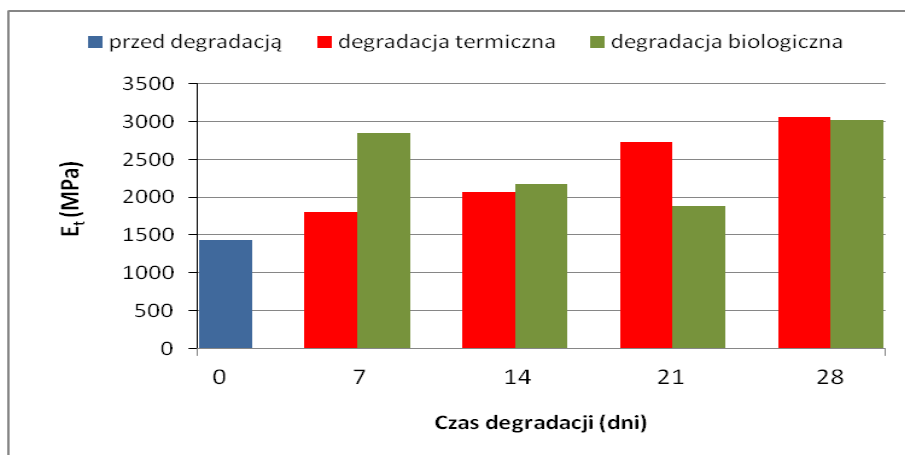
W tabeli 2 przedstawiono porównanie właściwości wytrzymałościowych próbek z polilaktydu L175 poddanych uprzednio procesom degradacji termicznej oraz degradacji biologicznej.

Tab. 2. Porównanie wartości wybranych właściwości podczas degradacji termicznej i biologicznej

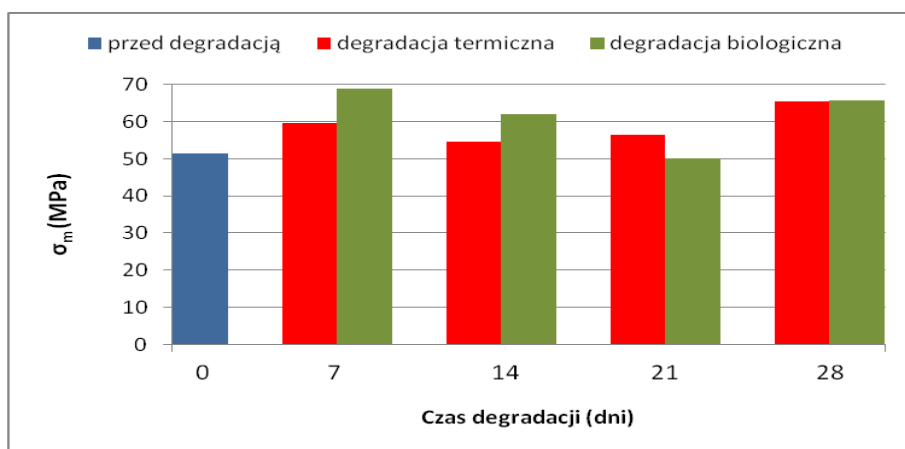
Lp.	Czas degradacji (dni)	Degradacja termiczna			Degradacja biologiczna		
		E_t (MPa)	σ_m (MPa)	ϵ_m (%)	E_t (MPa)	σ_m (MPa)	ϵ_m (%)
1	7	1806	60	2,6	2857	69	2,1
2	14	2071	55	2,5	2173	62	2,0
3	21	1819	62	3,0	1889	50	1,7
4	28	3063	65	2,4	3029	66	2,1

6. ANALIZA WYNIKÓW

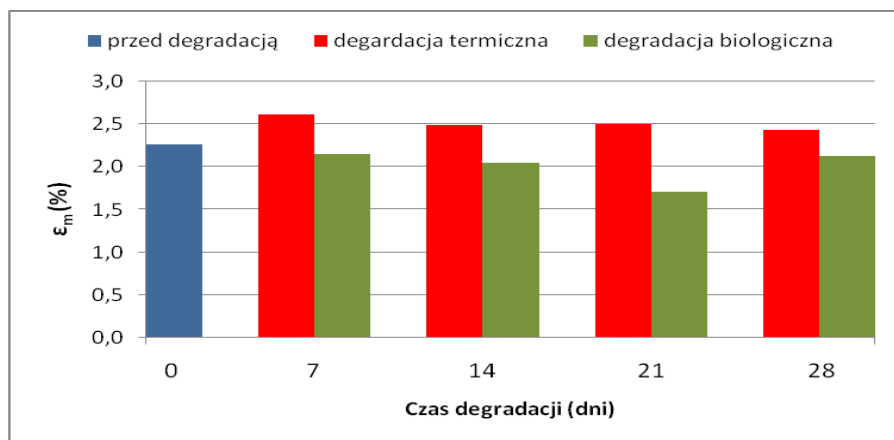
Poniżej dokonano porównania wybranych właściwości tj. modułu Younga, naprężeń i względnego odkształcenia liniowego, próbek poddanych degradacji termicznej i biologicznej. Zobrazowane jest to odpowiednio na rysunkach 4–6.



Rys. 4. Modułu Younga próbek w czasie degradacji termicznej i biologicznej



Rys. 5. Wartości naprężeń próbek w czasie degradacji termicznej i biologicznej



Rys 6. Względne odkształcenie liniowe próbek w czasie degradacji termicznej i biologicznej

Analizując wykres przedstawiony na rysunku 4 można zauważyć wzrost wartości modułu Younga dla próbek poddanych degradacji termicznej w trakcie trwania starzenia. Przeciwnieństwem tego są wartości modułu odkształcalności liniowej dla próbek poddanych degradacji biologicznej, których to wartość modułu Younga maleje.

Na rysunku 5 można zaobserwować, że wartość naprężeń próbek z folii PLA w trakcie degradacji termicznej jak i biologicznej maleje w sposób liniowy.

Dokonując analizy rysunku 6 można zauważyć nieznaczny liniowy spadek wartości względnego odkształcenia liniowego zarówno próbek poddanych degradacji termicznej jak i degradacji biologicznej.

Podsumowując wartość wytrzymałości na rozciąganie próbek wyciętych z folii PLA podczas degradacji termicznej w stosunku do wytrzymałości na rozciąganie próbek nie degradowanych:

- po 7 dniach wzrosła o 15,9%,
- po 14 dniach wzrosła o 6,2%,
- po 21 dniach wzrosła o 19,6%,
- po 28 dniach wzrosła o 27%.

Analizując wartości modułu sprężystości próbek PLA poddanych degradacji termicznej, można zauważyć, że wartość ta wzrosła w stosunku do wartości modułu Younga mierzonej przed degradacją próbek o:

- 25% po 7 dniach,
- 43,3% po 14 dniach,
- 25,9% po 21 dniach,
- 111,9% po 28 dniach.

Wytrzymałość na rozciąganie próbek poddanych degradacji biologicznej zmieniała się w następujący sposób:

- po 7 dniach wzrosła o 33,8% w stosunku do wartości wytrzymałości na rozciąganie próbek nie poddanych degradacji biologicznej
- po 14 dniach wzrosła o 20,4% w stosunku do wartości wytrzymałości na rozciąganie próbek nie poddanych degradacji biologicznej
- po 21 dniach zmalała o 2,5% w stosunku do wartości wytrzymałości na rozciąganie próbek nie poddanych degradacji biologicznej
- po 28 dniach wzrosła o 27,8% w stosunku do wartości wytrzymałości na rozciąganie próbek nie poddanych degradacji biologicznej
- po 35 dniach wzrosła o 24% w stosunku do wartości wytrzymałości na rozciąganie próbek nie poddanych degradacji biologicznej.

W przypadku próbek PLA poddanych degradacji biologicznej, wartość modułu sprężystości w stosunku do wartości modułu sprężystości próbek nie będących pod wpływem działania procesu degradacji biologicznej wzrosła o następujące wartości:

- po 7 dniach o 97,7%
- po 14 dniach o 50,4%
- po 21 dniach o 30,7%
- po 28 dniach o 109,5%
- po 35 dniach o 112,7%.

7. PODSUMOWANIE

Można stwierdzić, że zarówno degradacja termiczna jak i biologiczna mają wpływ na wartości wybranych właściwości polimerów biodegradowalnych. Na podstawie przeprowadzonych badań na maszynie wytrzymałościowej firmy Zwick/Roell Z010 dokonano obserwacji zmian wartości właściwości mechanicznych takich jak moduł odkształcalności liniowej, naprężenia rozciągające, względne odkształcenie liniowe. Badania przeprowadzono przed degradacją oraz trakcie trwania degradacji termicznej i biologicznej folii PLA L175.

Moduł odkształcalności liniowej próbki biodegradowanej termicznie przez 28 dni wzrasta nawet o 111,9%, a naprężenia rozciągające osiągają wartość o 27% większą w porównaniu do próbki nie poddanej procesowi degradacji termicznej.

W przypadku degradacji biologicznej naprężenia oraz moduł odkształcalności liniowej wzrastają podobnie jak w sytuacji degradacji termicznej. Moduł odkształcalności liniowej dla próbki poddanej degradacji biologicznej przez 28 dni wzrósł o 109,5 %, a po degradacji trwającej 35 dni o 112,7% w stosunku do próbki nie degradowanej. Naprężenia osiągają maksymalnie 33,8% większą wartość podczas degradacji niż naprężenia w próbce nie poddanej takiemu procesowi degradacji.

W trakcie badań udowodniono, że degradacja biologiczna i termiczna wpływają na zmianę wartości właściwości mechanicznych w podobny sposób. Jednak większe zmiany w krótszym czasie można zaobserwować w przypadku degradacji termicznej.

LITERATURA

- [1] Duda A., Penczek S., *Polilaktyd [Poli(kwas mlekowy)]: synteza, właściwości i zastosowania*, „Polimery”, 2003, 48, nr 1, s. 16–27
- [2] Foltynowicz Z., Jakubiak P., *Poli(kwas mlekowy)- biodegradowalny polimer otrzymywany z surowców roślinnych*. „Polimery”, 2002, 47, nr 11–12, s. 769–774
- [3] Gołębiewski J., Gibas E., Malinowski R., *Wybrane polimery biodegradowalne - otrzymywanie, właściwości, zastosowanie*. „Polimery”, 2008, 53, nr 11–12, s. 799–807
- [4] Grabowska B., *Biodegradacja tworzyw polimerowych*. „Archives of Foundry Engineering”, 2010, vol. 10, s. 57–60
- [5] Mucha M., *Polimery a ekologia*. Politechnika Łódzka, 2002
- [6] Praca zbiorowa. Kalińska D., Kuś H., Zwinogrodzki J., *Tworzywa sztuczne w medycynie*. WNT, 1970
- [7] Stachurek I., *Problemy z biodegradacją tworzyw sztucznych w środowisku*. „Zeszyty Naukowe Wyższej Szkoły Zarządzania Ochroną Pracy w Katowicach”, 2012, nr 1(8), s. 74–108
- [8] Szlezyngier W., Brzozoski K.Z., *Tworzywa sztuczne. Polimery specjalne i inżynierijne*. T. 2. Wydawnictwo Oświatowe FOSZE, 2012.

ZASTOSOWANIE FOLIOWYCH CZUJNIKÓW PIEZOELEKTRYCZNYCH W MONITOROWANIU FALI TĘTNA W WYBRANYCH PUNKTACH UKŁADU TĘTNICZEGO

1. WSTĘP

Choroby układu krążenia są nadal jednym z najczęstszych przyczyn zgonów. Pomimo istotnego postępu, jaki poczyniono w dziedzinie przeciwdziałania umieralności z tej przyczyny, choroby kardiologiczne nadal są powodem 46% wszystkich zgonów w Polsce. Jednym z najważniejszych czynników powodujących wzrost ciśnienia tętna oraz wzrost ciśnienia skurczowego są: sztywność tętnic i zjawisko odbicia fali tętna. W związku z tym faktem badanie sztywności tętnic oraz wartości centralnego ciśnienia tętna jest jak najbardziej zasadne, zmiany zachodzące w naczyniach krwionośnych wyprzedzają w czasie rozwój chorób układu krążenia. Z tego względu nieinwazyjne, niskokosztowe oraz powtarzalne techniki pomiarów tych parametrów stały się w ostatnich latach przedmiotem zainteresowania.

W niniejszym artykule przedstawiono analizę możliwości zastosowania czujników piezoelektrycznych do rejestrowania biosygnalów na przykładzie fali tętna w monitorowaniu chorób układu krążenia. Przetestowanie powszechnie dostępnych, łatwych do adaptacji w układzie pomiarowym oraz przystępnych cenowo foliowych czujników piezoelektrycznych pod tym kątem przydatności jest zadaniem jak najbardziej wartym uwagi.

Zjawisko piezoelektryczności odkryli bracia Jacques i Pierre Curie w 1880 roku podczas badania kryształów kwarcu. Stwierdzili, że ściskanie lub rozciąganie kwarcowych płytek powoduje powstawanie na ich powierzchni ładunków elektrycznych. Rok po odkryciu piezoelektryczności bracia Curie zweryfikowali postulowane przez Lippmanna zjawisko odkształcenia mechanicznego kryształu kwarcu umieszczonego w polu elektrycznym. Zjawisko piezoelektryczne jest odwracalne, a zależność między polem elektrycznym a odkształceniem jest liniowa. Występuje ono w kryształach, które nie mają środka symetrii [7]. Kryształy te pod wpływem odkształceń mechanicznych polaryzują się w określonych kierunkach.

¹ Politechnika Lubelska, Koło „MicroChip” Wydział Elektrotechniki i Informatyki

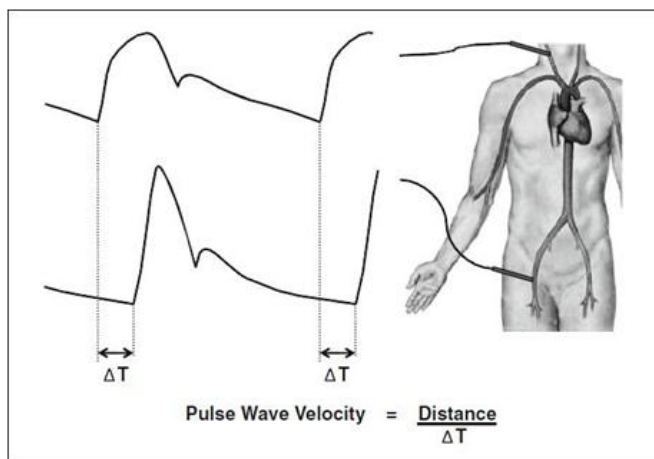
Występowanie fali tętna to zjawisko naturalnego odkształcenia ścian naczyń krwionośnych spowodowane wyrzutem krwi z serca pod dużym ciśnieniem. Aorta jest miejscem, w którym rozpoczyna się fala. Odkształcenia naczynia występują cyklicznie i są związane z pracą serca a podatność na nie sprawia, że aorta pełni rolę zbiornika krwi i sprawuje tym samym ważną funkcję w układzie krwionośnym. Według prawa Poiseuille'a ciecz przepływa tylko przy różnicy ciśnień na końcach układu rur, w których znajduje się ciecz. W chwili, gdy pompa przestaje tłoczyć ciecz i ciśnienie spada do zera, przepływ cieczy ustaje. Zgodnie z tym prawem, podczas rozkurczu serca krew powinna przestać płynąć. Nie dzieje się tak za sprawą elastyczności ścian, które pozwalają na zgromadzenie sporej ilości krwi i wtłoczenie jej do układu krwionośnego w czasie rozkurczu serca. Aorta przy niezmienionej chorobowo elastyczności jest w stanie zmagazynować ponad 60% objętości frakcji wyrzutowej lewej komory serca, przepuszczając bezpośrednio na obwód tylko niecałe 40% objętości. Zmagazynowana objętość jest wtłaczana podczas rozkurczu, by zapewnić ciągły przepływ krwi. Opisywane procesy możliwe są do monitorowania min. poprzez rejestrację fali Korotkowa.

2. MONITOROWANIE FALI TĘTNA W UKŁADZIE TĘTNICZYM CZŁOWIEKA

Diagnostyka w chorobach układu krążenia wykorzystuje informacje zawarte w sygnale tętna. Z punktu widzenia lekarza prowadzącego badania pacjenta ważna jest informacja zarówno o szybkości rozprzestrzeniania się fali tętna tzw. fali Korotkowa jak również jej zmiana kształtu w różnych punktach układu tętniczego.

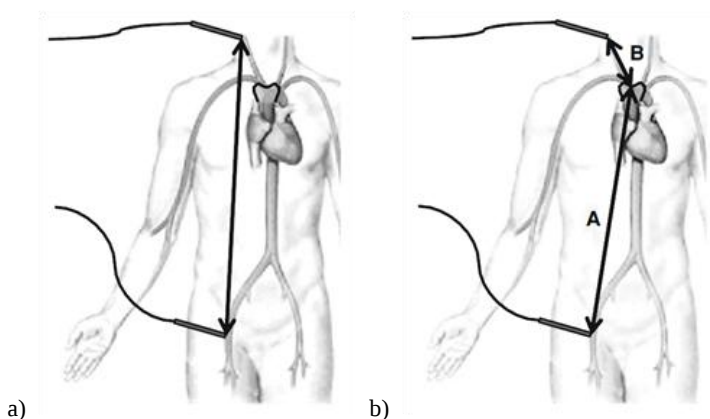
2.1 Pomiar regionalnej sztywności tętnic

Jednoczesny zapis przebiegu fal tętna w różnych punktach pozwala wyznaczyć jej prędkość. Tętno wyczuwane palpacyjnie (np. tętnica szyjna, promieniowa, podobojczykowa, udowa). Należy określić długość między miejscami rejestracji fal. Prędkość oblicza się dzieląc dystans, jaki pokonała fala na odcinku między punktami rejestracji przez opóźnienie między ramionami wstępującymi przebiegów fal tętna (odcinek Δt na rysunku 1.) – droga[m]/czas[s]. Ograniczeniem metody jest to, że otyłość brzuszna nie pozwala na dokładne określenie dystansu między punktami rejestracji wyczuwalnego tętna [1]. Mimo to, ocena prędkości fali tętna w relacji między tętnicą szyjną i udową została uznana za złoty standard pomiaru sztywności naczyń. Pomiar wykonany między tymi tętnicami pozwala na dokonanie oceny sztywności aorty.



Rys. 1. Sposób pomiaru relacji czasowej między falami tętna [3]

Należy zwrócić uwagę na sposób określania odległości między tętnicą szyjną wspólną i udową. Przez lata odległość między tymi tętnicami określano w sposób bezpośredni, w prostej linii tak, jak jest to zobrazowane na rysunku 2. a).



Rys. 2. a) bezpośredni sposób pomiaru odległości tętnic
b) metoda subtraktywna odległość = A-B [3]

Innym sposobem pomiaru jest metoda subtraktywna, w której punktem odniesienia jest wcięcie szyjne mostka (Rys. 2. b). W pierwszym etapie należy określić odległość tętnicy szyjnej wspólnej i udowej a następnie odjąć je od siebie, by finalnie otrzymać odpowiednią odległość pokonywaną przez falę tętna. Ostatnie badania pokazują, że pośredni sposób pomiaru odległości tętnic jest bardziej skorelowany z faktycznym dystansem anatomicznym pomiędzy nimi, zwłaszcza, że ten sposób pomiaru bierze pod uwagę fakt przemieszczania się fali

w przeciwnych kierunkach [2]. Wiąże się to z tym, że w momencie dotarcia fali tętna do tętnicy szyjnej, ten sam dystans zostanie pokonany przez falę propagującą się w przeciwnym kierunku. Odpowiedni dobór sposobu określenia dystansu pomiędzy tętnicami jest o tyle ważny, że w bardzo znaczący sposób wpływa na ostateczną wartość prędkości fali tętna. Przykładowe obliczenia zostały wykonane dla relacji czasowej między falami tętna równej 100ms, bezpośredniej odległości między tętnicami równej 650 mm oraz odległości tętnicy szyjnej od wcięcia mostka i tętnicy udowej do wcięcia mostka równych odpowiednio 100mm i 560mm. Dla metody bezpośredniej wartość PWV wyniosła 6,5m/s ($PWV = 650 \text{ mm}/100\text{ms}$), natomiast dla metody pośredniej wartość 4,6m/s ($PWV = (560 \text{ mm} - 100\text{mm})/100\text{ms}$). Dla tego przypadku wartość obliczana według metody konwencjonalnej była większa aż o 32% od wyniku uzyskanego dzięki innej metodzie. To dowodzi, że odpowiedni sposób oszacowania wzajemnej odległości miejsc pomiaru fali jest niezbędny do uzyskania poprawnej wartości PWV . Kolejnym z zalecanych sposobów szacowania odległości, dającym wyniki bardziej zbliżone do rzeczywistych niż w metodzie bezpośredniej jest używanie do obliczeń 80% wartości dystansu zmierzonego bezpośrednio pomiędzy tymi tętnicami [3].

2.2 Warunki pomiaru prędkości fali tętna

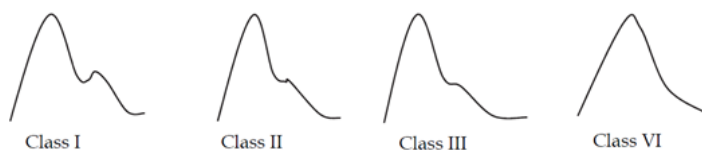
Spełnienie rekomendowanych zaleceń co do przeprowadzania pomiaru i należytego przygotowania osoby badanej jest gwarancją otrzymania powtarzalnych wyników. Do standardowych wymagań należą:

- minimalny okres przygotowawczy pacjenta: pacjent powinien odpocząć w cichym pomieszczeniu przynajmniej 10, a najlepiej około 30 minut
- pacjent powinien powstrzymać się od palenia, spożywania posiłków i napojów zawierających kofeinę przynajmniej na 3 godziny przed badaniem. Palenie ostro zwiększa sztywność naczyń, spożycie posiłku natomiast zmniejsza opór naczyń przy jednoczesnym zwiększeniu częstości skurczów i objętości frakcji wyrzutowej
- powstrzymanie od spożywania alkoholu przynajmniej na 10 godzin przed badaniem
- pacjent podczas badania nie powinien mówić. To pozwoli wyeliminować krótkoterminowe wpływy na sztywność naczyń. Pacjent nie powinien również być zbyt senny. Podczas snu również wzrasta sztywność tętnic.
- Ważnym aspektem jest wybór odpowiedniej pozycji pacjenta. Badanie zazwyczaj przeprowadza się w pozycji leżącej lecz można je przeprowadzić również w pozycji siedzącej. Należy pamiętać, że ciśnienie krwi różni się w zależności od pozycji pacjenta
- jeżeli konieczne jest powtórzenie badania, powinna je wykonywać ta sama osoba w możliwie identyczny sposób [4]

2.3 Analiza kształtu fali tętna

W przypadku wykrywania konkretnego ryzyka obliczenia oparte na punktach o specjalnym znaczeniu są czułe. Biorąc pod uwagę kilka czynników ryzyka ciężko za ich pomocą, oszacować ogólną kondycję układu sercowo-naczyniowego. Konkretny kształt fali tętna jest wynikiem między innymi mikrokrośnięcia, czynności fizjologicznej serca, naczyń krwionośnych,. Im więcej załączonych informacji, tym dokładniejsza klasyfikacja kształtu fali tętna.

Względna stabilność fali tętna jest, że osoba badana jest odprężona, a pomiar jest przeprowadzany w cichym pomieszczeniu. W takich warunkach analiza fali tętna daje powtarzalne wyniki. Podobieństwo fali tętna nie zmienia się znacząco u ludzi o podobnym stanie zdrowotnym układu krwionośnego, nawet przy zmianach częstości skurczów serca i siły, z jaką jest wyrzucana krew z lewej komory. Dzięki temu przebiegi fali tętna można podzielić na kilka klas charakterystycznych dla różnych stanów zdrowotnych układu krwionośnego:



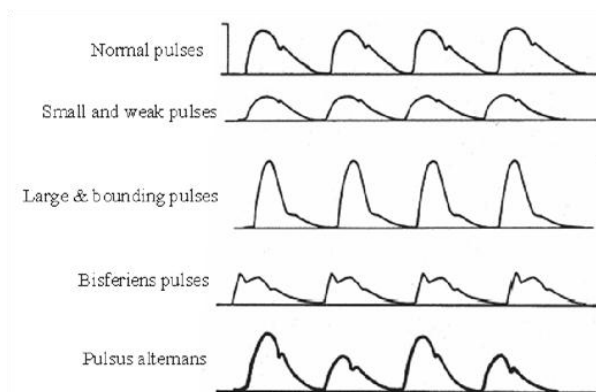
Rys. 3. Klasy tętna charakteryzujące stan zdrowotny [5]

Klasa I – wyraźne wcięcie dykrotyczne na ramieniu zstępującym

Klasa II – wcięcie dykrotyczne nie tworzy się, a zbocze pozostaje poziome

Klasa III – brak obecności wcięcia oraz widoczna zmiana kąta opadania ramienia zstępującego

Klasa IV – brak wcięcia dykrotycznego



Rys. 4. Klasy tętna charakteryzujące stan zdrowotny [6]

Ta klasyfikacja skupia się na wcięciu dykrotycznym, który został uznany za wskaźnik sztywności tętnic [5]. Inny badacz wyróżnił pięć typów fali tętna. Skupił się na bardziej szczegółowym opisie fali tętna i przyczynach danego kształtu fali tętna. Jego wyniki są przedstawione na poniższym rysunku.

Każde odstępstwo od normalnego pulsu odzwierciedla patologiczne zjawiska zaistniałe w układzie krwionośnym. Kształt fali tętna obrazuje pracę i fizjologię układu. Analiza fali tętna pod względem wizualnym pozwala na rozpoznanie możliwych schorzeń, które doprowadziły do patologicznego funkcjonowania układu. Tętno małe i słabe w porównaniu do tętna normalnego, z wolniej narastającym ramieniem wstępującym wskazuje na zmniejszoną objętość wyrzutową. Tego rodzaju patologia występuje przy niewydolności serca, hipowolemii (deficyt krwi w układzie krążenia), zwężeniu zastawki aortalnej. Tętno o dużej amplitudzie i szybko narastającym ramieniu wstępującym jest wynikiem zwiększonej frakcji wyrzutowej, zmniejszonego oporu naczyń obwodowych lub obu tych zjawisk na raz. Takie tętno może być objawem gorączki, nadczynności tarczycy, anemii, niedomykalności zastawki aorty, zmniejszonej podatności ścian aorty na odkształcenie w wyniku zeszywnienia. Zwiększona objętość wyrzutowa może być następstwem zbyt wolnego rytmu serca, np. przy bradykardii. Tętno dwubitne (pulsus bisferiens) zawiera dwie dodatnie fale ciśnienia w czasie trwania tej samej fazy skurczu serca. Przyczyną może być niedomykalność zastawki aorty lub jednocześnie zwężenie i niedomykalność tej zastawki. Tętno naprzemienne (pulsus alternans) zmienia amplitudę ze skurczu na skurcz, chociaż cykliczność tych zmian jest zasadniczo zachowana. Występuje przy nadciśnieniu tętniczym, niewydolności lewej komory i chorobie wieńcowej [6].

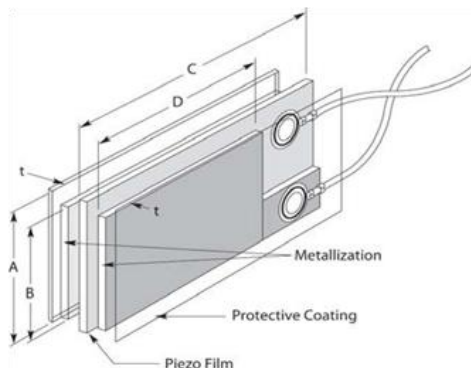
3. ZASTOSOWANIE PIEZOELEKTRYCZNYCH CZUJNIKÓW FOLIOWYCH

Czujniki piezoelektryczne wykorzystywane do rejestracji fali tętna zbudowane są jako inteligentne, wyposażone w układ mikroprocesorowy mikrofony rejestrujące biosygnaly. Najważniejszym elementem czujnika jest przetwornik piezoelektryczny. Właściwości przetworników piezoelektrycznych prezentowane w niniejszym rozdziale jednoznacznie wskazują na ich przydatność w monitorowaniu fali tętna.

Przetworniki piezoelektryczne należą do grupy przetworników generacyjnych. Oznacza to, że wytwarzają wyjściowy sygnał pomiarowy bezpośrednio pod wpływem działania wielkości mierzonej, bez konieczności zasilania ich z zewnątrz. Popularnymi materiałami piezoelektrycznymi używanymi o budowy przetworników są: kwarc, tytanian baru, sól Seigneta. Od czasu, gdy w latach sześćdziesiątych ubiegłego wieku odkryto słaby efekt piezoelektryczny w kości i ścięgnach wieloryba, rozpoczęły się poszukiwania materiałów organicznych, które również mogłyby wykazywać taki efekt. W roku 1969 wykryto silną ak-

tywność piezoelektryczną w polifluorku winylidenu (PVDF). Na bazie tego polimeru wytwarzane są coraz bardziej popularne czujniki foliowe. Jedne z najważniejszych zalet tych czujników to:

- szeroki zakres częstotliwości pracy – $0,001\text{--}10^9\text{Hz}$
- duży zakres mierzonego ciśnienia – $10^{-8}\text{--}10^6\text{psi}$
- wysoka wytrzymałość elektryczna (wytrzymuje natężenie pola $75\text{V}/\mu\text{m}$, przy którym większość czujników ceramicznych ulega depolaryzacji)
- niska impedancja akustyczna, zbliżona do wody, ludzkiej tkanki oraz
- wysoka elastyczność
- duże napięcie wyjściowe – dziesięć razy wyższe niż napięcie generowane przez przetworniki ceramiczne przy takiej samej sile obciążającej
- duża wytrzymałość mechaniczna, udarność (moduł Younga $10^9\text{--}10^{10}\text{Pa}$)
- odporność na wilgoć (absorpcja wilgoci $< 0,02\%$), większość chemikałów, utleniaczy, intensywne promieniowanie ultrafioletowe, promieniowanie jądrowe
- możliwość wytwarzania czujników w najróżniejszych kształtach
- możliwość przymocowania przetwornika używając typowych lepiszczy [8].



Rys. 4. Budowa przetwornika foliowego PVDF [10]

Budowę cienkowarstwowego czujnika piezoelektrycznego reprezentuje trójwymiarowy model z rysunku 4. Film piezoelektryczny jest pokryty z obu stron metalizacjami ze związków srebra z odprowadzeniami na przewody, a całość jest zabezpieczona powłoką ochronną z laminatu. Czujniki te są lekkie, elastyczne, wytrzymałe i dostępne w różnych wariantach grubości i wielkości.

Energia wyjściowa jest proporcjonalna do objętości filmu poddanego siłom wymuszającym. Grubość filmu można dobierać pod kątem optymalizacji wielkości sygnału elektrycznego lub ze względu na spodziewane obciążenia

mechaniczne. Im grubszy film, tym większe napięcie wyjściowe i mniejsza pojemność przetwornika.

Cienkowarstwowe czujniki PVDF mają jednak swoje ograniczenia. Przykładem jest względnie mały współczynnik sprzężenia elektromechanicznego. Współczynnik ten jest miarą stopnia przetwarzania w piezoelektryku energii mechanicznej w elektryczną i odwrotnie. PVDF wypada w porównaniu do materiałów ceramicznych gorzej pod względem tego parametru, szczególnie przy częstotliwości rezonansowej piezoelektryka i przy częstotliwościach niskich. Innym ograniczeniem tych przetworników jest dopuszczalna temperatura pracy, nieprzekraczająca 135°C. Kolejnym ważnym czynnikiem, który wynika z pojemnościowej natury przetwornika jest wrażliwość na zakłócenia elektromagnetyczne, które indukują się na nieosłoniętych elektrodach. Z tego względu zaprojektowano czujniki w wersji ekranowanej, które mogą być używane w miejscach, gdzie występują duże zakłócenia elektromagnetyczne.

Wygenerowaną ilość ładunku lub napięcie wyjściowe powstałe w wyniku konwersji energii mechanicznej na elektryczną może być obliczone z następujących wzorów:

$$D = \frac{Q}{A} = d_{3n}X_n$$

$$V_o = g_{3n}X_n t$$

Osie mechaniczne (n), wzdłuż których działa siła: 1 – oś wzdłużna, 2 – oś poprzeczna, 3- oś pionowa, gdzie:

$n = (1,2,3)$ oś działania siły

D = gęstość ładunku

Q = ładunek

A = powierzchnia czynna elektrody

d_{3n} = współczynnik piezoelektryczny dla osi działania siły [pC/m^2]

X_n = siła działająca w danym kierunku

g_{3n} = współczynnik piezoelektryczny dla osi działania siły [V/m]

t = grubość filmu.

Właściwości piroelektryczne przetworników

Folie PVDF należy do grupy ferroelektryków i tak jak one wykazuje efekt piroelektryczny – wytwarza ładunek elektryczny w odpowiedzi na zmiany temperatury. Szczególnie silnie pochłania fale promieniowania podczerwonego o długości mieszczącej się w przedziale od 7 do 20 μm . Absorbowane przez PVDF widmo promieniowania podczerwonego odpowiada ściśle widmu emitowanemu przez ludzkie ciało. To powoduje, że taki czujnik jest czułym detektorem promieniowania podczerwonego ludzkiego ciała. Efekt piroelektryczny w tych czujnikach jest silny i należy brać go pod rozwagę w szczególności w aplikacjach niskoczęstotliwościowych (od <0,01 do 1 Hz), w celu uniknięcia nakładania się na siebie sygnałów pochodzących zarówno od odkształceń mechanicznych, jak i zmian otaczającej temperatury. Dla wyższych częstotliwości zmiany

temperatury będą mniejsze niż zmiany sygnału użytecznego, co pozwala na odfiltrowanie niepożądanego sygnału piroelektrycznego.

Ilość ładunku i napięcie wygenerowane w wyniku zmian temperatury są dane wzorami:

$$Q = p\Delta T A$$

$$V = \frac{pt\Delta T}{\varepsilon}$$

gdzie:

Q = ładunek

V = napięcie

p = piroelektryczny współczynnik ładunku

ΔT = zmiana temperatury

A = powierzchnia filmu piezoelektrycznego

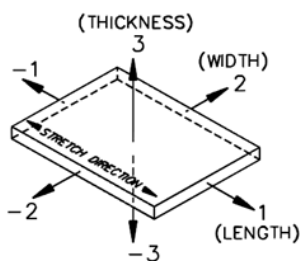
t = grubość filmu piezoelektrycznego

ε = przenikalność elektryczna

$$V = \frac{(30\mu\text{C}/\text{m}^2\text{C}) \cdot (9 \cdot 10^{-6}\text{m}) \cdot (1^\circ\text{C})}{(106 \cdot 10^{(-12)}\text{C}/\text{Vm})} = 2,55\text{V}$$

Tryby pracy foliowych przetworników piezoelektrycznych

Przetworniki piezoelektryczne mają charakter anizotropowy. Ich odpowiedź na odkształcenia, naprężenia, przyłożoną siłę zależy od kierunku działania siły wymuszającej. Elektrody są umieszczone zarówno na górnej jak i dolnej stronie czujnika, tak więc osią elektryczną jest zawsze kierunek „3” (Rys. 5). Oś mechaniczną natomiast może być każda z osi przetwornika, gdyż w każdym z tych kierunków można przyłożyć siłę wymuszającą. Czujnik może być poddawany kompresji w osi „3” lub też rozciągany w pozostałych osiach. Typowo piezolaminaty foliowe pracują w kierunku wzdłużnym osi „1” przy odbiorze częstotliwości niższych od 100kHz, natomiast w trybie kompresji są używane w aplikacjach ultradźwiękowych powyżej częstotliwości 100kHz.



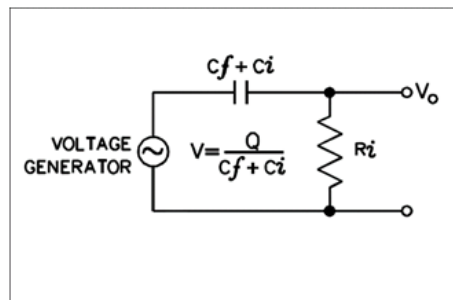
Rys. 5. Klasyfikacja osi przetwornika [9]

Przetwornik zamocowany jak belka obustronnie podparta, pracujący w trybie na zginanie jest w stanie wygenerować kilkaset razy wyższe napięcie niż podczas pracy na ściskanie (tryb kompresji) [9].

Układ współpracujący z przetwornikiem

Budowę docelowej aplikacji należy rozpocząć od doboru obwodu współpracującego. Projektowanie takiego układu należy rozpatrywać przy uwzględnieniu następujących czynników:

- zakres mierzonych częstotliwości i amplitudy sygnału w całym zakresie mierzonych sił wymuszających,
- odpowiednia rezystancja obciążenia, która ogranicza dolną częstotliwość pracy przetwornika,
- dla małych wartości sygnału wyjściowego należy zastosować bufor, aby nie obciążać przetwornika.



Rys. 6. Schemat równoważny przetwornika i rezystancji wejściowej

Jedną z najbardziej krytycznych części interfejsu takiej aplikacji jest rezystancja wejściowa, oznaczona na rysunku 6. jako R_i . W połączeniu z pojemnością przetwornika determinuje dolną częstotliwość graniczną jako filtr górno-przepustowy zgodnie ze wzorem:

$$f_0 = \frac{1}{2\pi RC}$$

gdzie:

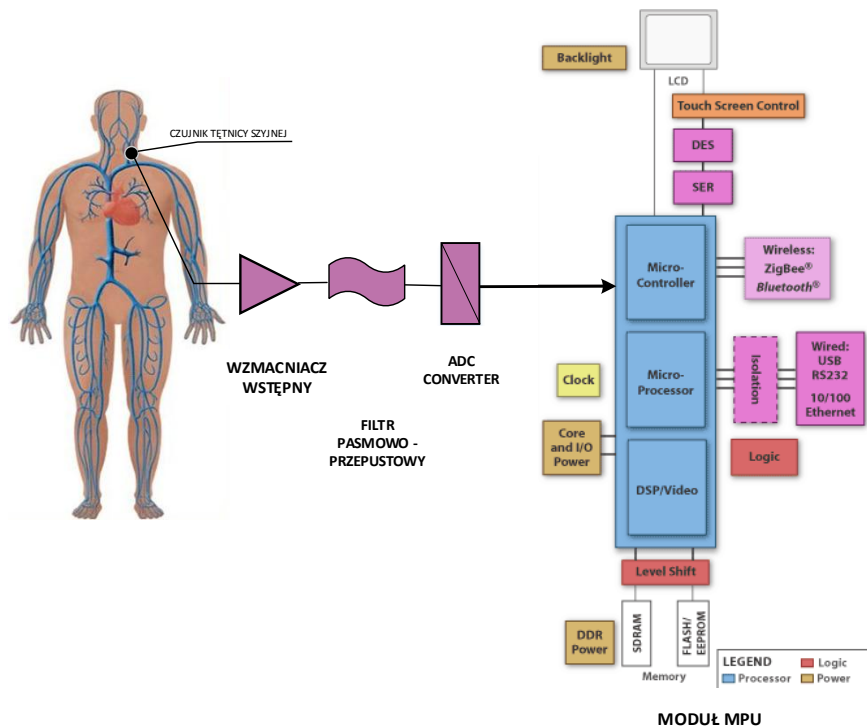
R – rezystancja obciążająca przetwornika,

C – pojemność przetwornika.

W związku z tym, że najniższe częstotliwości składowe fali tętna mogą zaczynać się już w okolicy 1Hz należy zadbać o odpowiednią rezystancję wejściową. Wykorzystany w pracy przetwornik piezoelektryczny ma pojemność równą 500pF. Częstotliwość graniczna w połączeniu z rezystorem o wartości 100M Ω wyniesie 3,2Hz. W związku z powyższym rezystor obciążający przetwornik musi mieć wartość co najmniej 100M Ω .

Model urządzenia do pomiaru fali tętna z wykorzystaniem czujników piezoelektrycznych

Na rys 2.4 przedstawiono opracowany w trakcie badań model przyrządu laboratoryjnego do monitorowania fali tętna w wybranych miejscach układu tętniczego. Przyrząd pozwala zarejestrować przebieg sygnału tętna przy użyciu pojedynczego czujnika.



Rys. 7. Schemat blokowy przyrządu do monitorowania fali tętna

Najważniejszym elementem urządzenia do akwizycji sygnału tętna i jego cyfrowej analizy jest:

- czujnik piezoelektryczny umiejscowiony na szyi, wzmacniacz wstępny z układem precyzyjnego bloku ADC (16–24 bitów) oraz moduł MPU mikrokontrolera z cyfrowym blokiem cyfrowej obróbki sygnału.
- wzmacniacz wstępny – wzmacniacz o dużym wzmocnieniu i dużej impedancji wejściowej. Zadaniem wzmacniacza jest dopasowanie sygnału wyjściowego czujnika do pola napięć przetwornika ADC. Zastosowane wzmacniacze charakteryzują się niskim poziomem szumu, który umożliwia osiągnięcie poziomu szumu sygnału wejściowego (RMS) w przybli-

zeniu 2 mikrowoltów, wymagane do rejestracji sygnałów fizjologicznych o małej mocy.

- filtr pasmowo-przepustowy – sygnały biofizyczne AC, do których należy sygnał tętna rejestrowane poprzez czujniki piezoelektryczne zazwyczaj zawierają stosunkowo dużych offset DC, który zmienia się w funkcji czasu (dryft składowej stałej). Filtry górnoprzepustowe służą do blokowania składowej stałej DC oraz niepożądanego dryfu niskiej częstotliwości lub innych artefaktów. Częstotliwość odcięcia tych filtrów jest konfigurowana w szerokim zakresie.
- szybki multiplexer analogowy – w celu miniaturyzacji urządzenia sygnały z poszczególnych wzmacniaczy są multipleksowane do pojedynczego wyjścia, które dołączone jest do pojedynczego przetwornika ADC.
- przetwornik ADC – (16–24)-bitowy przetwornik ADC zapewnia digitalizację analogowych sygnałów pomiarowych. Przetwornik kontrolowany jest przez standardowy interfejs szeregowy (SPI).
- moduł mikrokontrolera – moduł ten stanowi główną jednostkę kontrolno – pomiarową rejestratora. Poprzez porty oraz interfejsy szeregowy mikrokontrolera dołączone są: przetwornik ADC wzmacniaczy wstępnych, pamięć danych pełniąca funkcję bufora danych, moduł komunikacyjny w standardzie wi-fi/BLE/ZigBee. Moduł mikrokontrolera współpracuje z wbudowanym blokiem DSP, którego zadaniem jest obróbka sygnałów pomiarowych w czasie rzeczywistym, w tym wyznaczenia wskaźników mierzonego tętna fali. Moduł mikrokontrolera posiada program wbudowany zarządzający architekturą urządzenia pomiarowego, odpowiedzialny za realizację funkcji pomiarowych układu.

Akwizycja sygnału

Akwizycja sygnału oraz jego przetwarzanie zostało wykonane w oparciu o aparaturę laboratoryjną wraz z oprogramowaniem firmy Agilent. Akwizycja sygnału z wykorzystaniem wspomnianej aparatury jest prosta i pozwala na szybkie testowanie. Na rysunku poniżej przedstawiono prototyp czujnika do pomiaru fali tętna.

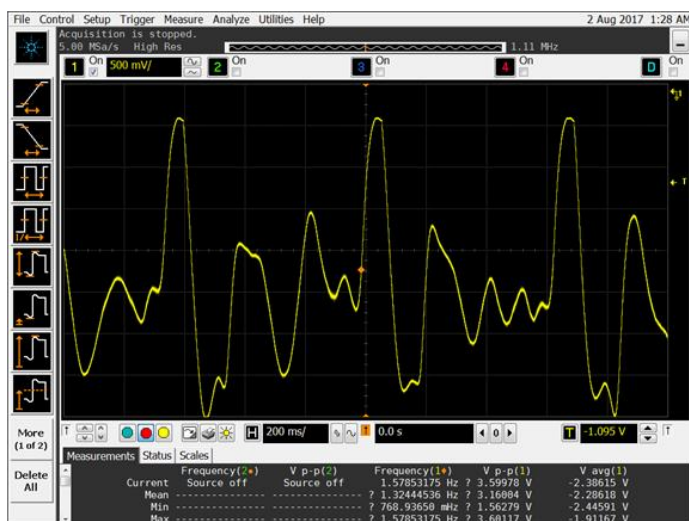


Rys. 8. Wykonany model czujnika do pomiaru fali tętna

Na rys. 9, 10, 11 przedstawiono próbne sygnały fali tętna zarejestrowane na szyi pacjenta – osoby zdrowej.



Rys. 9. Poprawny wynik pomiaru fali tętna na szyi pacjenta



Rys. 10. Poprawny wynik pomiaru fali tętna na szyi pacjenta w powiększeniu



Rys. 11. Wynik pomiaru fali tętna gdy czujnik został umieszczony za blisko rozwidlenia tętnicy szyjnej pacjenta

4. PODSUMOWANIE

Po przeprowadzeniu badań z wykorzystaniem czujników oraz programu nawiązują się pewne wnioski dotyczące dalszego rozwoju projektu. Pierwszym z nich są odpowiednie kwalifikacje, jeżeli chodzi o przeprowadzanie pomiaru.

Krytycznym aspektem jest odpowiednie umiejscowienie czujników nad badanymi tętnicami, co w dużej mierze pozwala na uniknięcie obecności artefaktów pochodzących od naczyń leżących w pobliżu badanej tętnicy. Ten aspekt w szczególności dotyczy tętnicy szyjnej wspólnej. Czujnik w tym przypadku koniecznie musi zostać umieszczony przed rozwidleniem tętnicy, aby uniknąć rejestracji artefaktów i zakłóceń związanych z odbijaniem się fali ciśnieniowej w miejscu rozwidlenia. Na rysunku 2.6 c można zaobserwować przebiegi które pokazują zniekształcenia spowodowane zbyt bliskim umieszczeniem czujnika przy rozwidleniu tętnicy. Oprócz samego umieszczenia czujnika ważny jest też jego docisk do ciała osoby badanej. Czujnik nie może być zamocowany zbyt luźno, gdyż amplituda rejestrowanego sygnału jest wtedy znacznie zaniżona. Opisanie wyżej aspekty są ważne z punktu widzenia wyeliminowania wpływu czynnika ludzkiego na pomiar. To pozwoli na bardziej precyzyjne oszacowanie przydatności samego czujnika do wykonywania tego typu pomiarów.

LITERATURA

- [1] Laurent S., Safar M.E. *Uszkodzenie dużych tętnic: pomiary i znaczenie kliniczne*. W: Mancina G., Grassi G., Kjeldsen S.E. (red.). *Nadciśnienie tętnicze podręcznik European Society of Hypertension*. Gdańsk, Via Medica 2009; 192–202
- [2] Sugawara J., Hayashi K., Yokoi T., Tanaka H., *Carotid-Femoral Pulse Wave Velocity: Impact of Different Arterial Path Length Measurements*. "Artery Res.", 2010 Mar 1; 4(1), p. 27–31
- [3] Salvi P., *Pulse waves. How Vascular hemodynamics affects blood pressure*. Springer 2012
- [4] Van Bortel L.M., Duprez D., Safar M.E. i wsp., *Clinical Applications of Arterial Stiffness*, Task Force III: Recommendations for Users Procedures. *AJH* 2002; p. 445–452
- [5] Gaetano D. Gargiulo, McEwan A., *Advanced biomedical engineering*, 2011
- [6] Bates B., Bickley L.S., Szilagyi P.G., *Bate's Guide to Physical Examination and history taking*. Tenth edition 2009.
- [7] Siergiejew M., *Wstęp do fizyki kryształów*. Uniwersytet Szczeciński 2001
- [8] Chwaleba A., Moeschke B., Ploszajski G., *Eletronika*. WSiP 2014
- [9] Measurement Specialties: *Piezo Film Sensors Technical Manual*
- [10] <http://www.imagesco.com/piezoelectric/> (zasoby z dnia 21.04.2017)

ZASTOSOWANIE NANOTECHNOLOGII W MEDYCYNIE

1. NANOTECHNOLOGIA

Nanotechnologia jest dosyć szybko rozwijającą się dziedziną wiedzy, która łączy w sobie zarówno medycynę jak i przemysł, a także powiązane z nimi obszary życia. Zdefiniować ją można jako nauka obejmująca projektowanie, tworzenie i użytkowanie struktur, których co najmniej jeden z wymiarów jest w skali nano [1]. Dany wytwór może zostać nanomateriałem, gdy po zabiegu miniaturyzacji do rozmiarów nano pojawiają się zmiany właściwości fizycznych i chemicznych w odniesieniu do wyrobu w skali makro, materiału litego. Nanotechnologia obejmuje cztery główne obszary, jakimi są:

- nanoelektronika,
- nanotechnologia molekularna,
- nanomateriały,
- mikroskopy w skali nano.

Obecnie jest to temat przyszłości i ma szerokie zastosowanie zaczynając od elektroniki, transportu, przemysłu spożywczego, wyrobu ubrań, a kończąc na medycynie i farmacji. Okazuje się być nieoceniona podczas produkcji implantów oraz sensorów jak również przy dostarczaniu leków. Naukowcy wiążą z tą techniką spore nadzieje w leczeniu nowotworów.

Nanotechnologia zapoczątkowała w 1959 r. przemówieniem Richarda Feynmana. Jego głównym zamysłem było umieszczenie Encyklopedii Britannica na nośniku danych, nie większym od łepka szpilki. Przełomowym wynalazkiem w rozwoju nanotechnologii był skaningowy mikroskop tunelowy, który służył do wytwarzania nanostruktur rzędu 40–70nm. Za pomocą tego urządzenia powstał napis IBM, a jego odczyt jest możliwy jedynie przy użyciu mikroskopu elektronowego [1].

Materiały takich wielkości określa się mianem „zero-wymiarowe”, ponieważ każdy wymiar zawarty jest w nanoskali. Dzięki temu można je odpowiednio modelować w zależności od przeznaczenia tak, aby właściwości fizyczne, chemiczne i biologiczne dostosować do przyszłej roli danej cząsteczki w systemie. Mają możliwość wnikania do organizmu drogami oddechowymi, przez skórę oraz z pokarmem. Krążące w organizmie cząsteczki mogą odkładać się w narządach [12].

¹ Politechnika Lubelska, Wydział Elektrotechniki i Informatyki

Każdy obiekt możemy rozpatrywać w odniesieniu do jednego bądź wielu wymiarów, analizując jego układ współrzędnych. Nanostruktury również możemy analizować pod tym kątem, co prowadzi do podziału na układy:

- quasi-zerowymiarowe (półprzewodnikowe kropki kwantowe)
- quasi-jednowymiarowe (nanorurki węglowe, nanodrut)
- quasi-dwuwymiarowe.

2. KROPKI KWANTOWE

W ostatnich latach gwałtownie wzrosło zastosowanie luminescencyjnych kropek kwantowych w badaniach biologicznych. Wpłynął na to unikalny rozmiar zależny od optycznych właściwości i coraz to nowsze osiągnięcia naukowe.

Kropki kwantowe to sztucznie wytwarzane kropelki, które mogą zawierać wszystko, od pojedynczego elektronu do zbioru kilku tysięcy. Są strukturą, której każdy wymiar ograniczony jest barierą potencjału. Skutkiem tego jest całkowita kwantyzacja energii. Ich typowe wymiary wahają się od nanometrów do kilku mikrometrów (2–10nm), a ich wielkość, kształt i interakcje mogą być precyzyjnie kontrolowane dzięki wykorzystaniu zaawansowanej technologii nanowytwarzania [1].

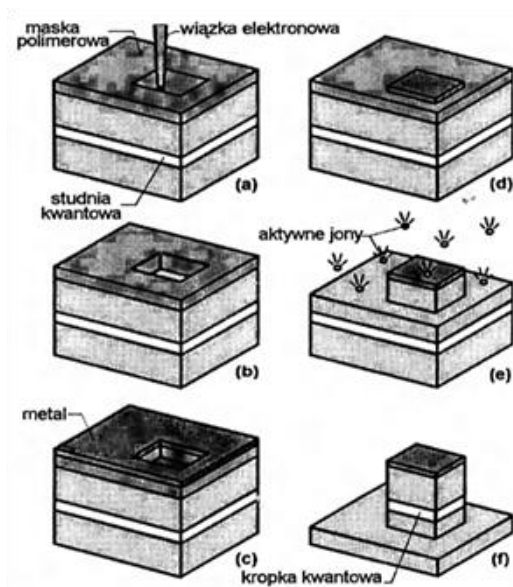
Wytwarzanie kropek kwantowych

Specyficzne właściwości kropek kwantowych zależą od sposobu ich wytwarzania i zastosowanych technologii. Z uwagi na łatwe wytworzenie dobrej jakości układów quasi-wymiarowych, często stosuje się metodę polegającą na stworzeniu potencjału bocznego. Inne zaś metody bazują na powstawaniu mikrostruktur, których zadaniem jest wiązanie elektronów podczas krystalizacji.

Jest jedną z najbardziej rozpowszechnionych metod jest metoda samoorganizacji. Termin ten wskazuje na formowanie cząsteczek w określone struktury przy udziale oddziaływań fizycznych bądź chemicznych, zachodzących pomiędzy atomami. Sposób ten działa na zasadzie osadzania jednego materiału półprzewodnikowego na inny półprzewodnik. Ważna jest wartość stałej siatki krystalicznej. Pierwszy materiał ma wartość dużą, natomiast ten, który stanowi podłoże ma tą wartość małą. Wynikiem tego jest utworzenie, na powierzchni, miejsc o takiej samej stałej jak podłoże, ale nie przekraczając grubości krytycznej. W kolejnym etapie warstwę nieregularnych wysp pokrywa się takim samym materiałem, który wcześniej służył jako podłoże. Metoda ta daje możliwość wytworzenia kropek o zróżnicowanych rozmiarach i kształtach. Kropki kwantowe utworzone tą metodą charakteryzują się niezłymi parametrami optycznymi, jednorodnością kształtu, niewielkimi rozmiarami jak również nie występują defekty brzegowe [11].

Technologia wytrawiania to najwcześniej poznany sposób tworzenia kropek kwantowych. Polega na wytrawianiu w studni kwantowej wypełnionej gazem

elektronowym o charakterze quasi-dwuwymiarowym. Na przygotowane podłoże, zawierające studnie kwantowe nanoszona zostaje warstwa polimeru, po czym całość jest naświetlana wiązką elektronową bądź jonową. Naświetlanie próbki służy wyznaczeniu wzoru, który będzie odpowiadać kształtowi kropki. Następnie maska, czyli warstwa polimeru, zostaje zdjęta w miejscach poddanych naświetlaniu, a na całość powierzchni podłoża nanosi się warstwę metalu. Za pomocą odczynników chemicznych usuwa się nanoszone warstwy z próbki, dzięki czemu miejsca nie chronione przez maskę i naświetlane stanowią czystą powierzchnię próbki z cienką powłoką metalu [6].



Rys 2. Powstawanie kropek kwantowych w procesie wytrawiania. [6]

W 1982 roku po raz pierwszy Ekimow opisał sposób wytwarzania kropek kwantowych z użyciem szklanych matryc o charakterze dielektrycznym. Koloidalne zawiesiny kropek kwantowych powszechnie syntetyzuje się poprzez wprowadzenie półprzewodnikowych prekursorów, w warunkach sprzyjających termodynamicznemu wzrostowi kryształów w obecności półprzewodnikowych środków wiążących, które działają na kinetykę wzrostu kryształów i kontrolują utrzymanie ich rozmiaru w granicach wielkości kwantowych [7].

Medyczne zastosowanie kropek kwantowych

Postępy w dziedzinie nanotechnologii pozwoliły na projektowanie narzędzi opartych na nanocząstkach w celu poprawy diagnozy i zindywidualizowaniu leczenia wielu skomplikowanych chorób. Nową szansą na dokładną diagnozę

wielu chorób są półprzewodnikowe kropki kwantowe, stanowiące nową platformę dla wysoko-przepustowej charakterystyki ilościowej biomarkerów tkanek i innych próbek klinicznych.

Ze względu na swoje unikalne właściwości fizyczne i optyczne, koloidalne kropki kwantowe zostały wykorzystane do opracowania bioczuJNIKÓW. Przyłączenie biocząsteczek na powierzchni kropek, umożliwia generowanie złożonych biokoniugatów (sensorów), które łączą specyficzność biologiczną i funkcjonalną pożądaných cech optycznych. Nano wymiary kropek kwantowych pozwalają stać się im centralnym elementem konstrukcyjnym, który może pomieścić liczne kopie biocząsteczki (np. białko lub DNA) lub kilka biocząsteczek jednocześnie.

Fluorescencyjny rezonansowy transfer energii (FRET) jest szczególnym sposobem transdukcji sygnału ze względu na jego wrażliwość na poziomie interakcji molekularnych. Zastosowano tu zjawisko, w którym między akceptorem i donorem zachodzi wymiana energii. Akceptor pobiera energię od wzbudzonego donora i emituje ją. Przeniesienie energii fluorescencji z cząstki donora do cząstki akceptora zachodzi, gdy odległość pomiędzy dawcą a biorcą jest mniejsza niż krytyczny promień znany jako promień Forstera. Prowadzi to do zmniejszenia emisji dawcy i żywotności stanu wzbudzonego, oraz zwiększenie intensywności emisji akceptora. FRET nadaje się do pomiaru zmian odległości zamiast bezwzględnej odległości, dzięki czemu jest odpowiedni dla pomiarów zmian konformacyjnych białek. Można również monitorować interakcje białek i oznaczać aktywności enzymatyczne. Na dzień dzisiejszy istnieją już nanosensory, których zadaniem jest wykrywanie obecności maltozy, konkretnych sekwencji DNA, aktywności proteaz [8].

Kropki kwantowe umożliwiają obrazowanie *in vivo*. Możliwość wizualizacji rodzimych procesów zachodzących w żywych organizmach jest bezcenne dla zastosowań klinicznych i diagnostycznych. Jest to wciąż dosyć trudne do zastosowania praktycznego ze względu na ograniczenia konwencjonalnego obrazowania i dostępność odpowiednich markerów fluorescencyjnych. W odniesieniu do tego ostatniego, wiele barwników ogranicza krótki czas życia (~1 ns), są podatne na degradację fotochemiczną oraz wykazują niedostateczną jasność fluorescencji. Dodatkowo autofluorescencja tkanek może wykazywać podobne właściwości spektroskopowe, utrudniające rozwiązanie pożądanego sygnału z niechcianego tła. Ze względu na swoje unikalne właściwości fotofizyczne, kropki kwantowe są obiecującą metodą obrazowania fluorescencyjnego w *in vivo* i mogą pokonać wiele ograniczeń charakterystycznych dla typowych barwników.

Obrazowanie *in vivo* chorych komórek i tkanek zapewnia wiele korzyści dla medycyny spersonalizowanej. Zapewnia możliwość diagnozy we wczesnych stadiach choroby, uzyskanie specyficznych informacji dla pacjenta o lokalizacji i rozmiarze rdzenia chorobowego, ocenę niekorzystnych skutków dla zdrowych tkanek, a także monitorowanie postępu choroby oraz odpowiedź na terapię. Konwencjonalne techniki obrazowania medycznego, takie jak USG, tomografia

obrazowania rezonansu magnetycznego (MRI), i tomografia pozytronowa (PET), w większości przypadków nie oferują czułości i rozdzielczości jednocześnie do postawienia diagnozy we wczesnym stadium (np. MRI zapewnia wysoką rozdzielczość, a zarazem niską czułość; natomiast PET – wysoką czułość przy niskiej rozdzielczości). Mikroskopia fluorescencyjna pozostaje najskuteczniejszą techniką profilowania cząsteczkowego zmienionych chorobowo komórek. Jednak obecność barier tkankowych między punktami chorobowymi i sprzętem obrazującym komplikuje wykorzystywanie mikroskopii fluorescencyjnej do obrazowania *in vivo* tkanek biologicznych, które skutecznie pochłaniają i rozpraszają światło widzialne wraz z intensywną produkcją autofluorescencji o szerokim spektrum działania. W przeciwieństwie do fluoroforów organicznych kropki kwantowe posiadają wysoką jasność oraz możliwość multipleksowania wraz z dużymi przesunięciami Stokes'a. Są więc odpowiednim narzędziem do obrazowania *in vivo*, w szczególności szczelina widmowa między wzbudzeniem i emisją kropek kwantowych, która jest znacznie większa niż w przypadku fluoroforów organicznych i może osiągać 300–400nm, w zależności od długości fali światła wzbudzającego. W ten sposób dochodzi do przesunięcia sygnału kropek kwantowych do obszaru o ograniczonej autofluorescencji tkankowej. Wykorzystanie czujników obrazowania bliskiej podczerwieni może jeszcze bardziej zmniejszyć zakłócenia autofluorescencji tkanki i umożliwia obrazowanie *in vivo* z głębszą penetracją i z lepszą rozdzielczością. W ostatnich badaniach *in vivo* morfologii guza naukowcy wykorzystali mikroskopię dwufotonową, aby jednocześnie obrazować naczynia guza (barwione niebieskimi kropkami) i komórki około naczyniowe. Technika ta wykorzystuje nisko energetyczne fotony (z czerwonych i podczerwonych regionów) do wzbudzania kropek kwantowych emitujących w zakresie widzialnym, osiągając znacznie zmniejszone tłumienie światła wzbudzającego przez tkanki wraz ze zmniejszeniem autofluorescencji, jednocześnie umożliwiając wykorzystanie kropek emitujących nad pełnym spektrum światła widzialnego. Niestety metody te są jeszcze w fazie badań na zwierzętach. Obecnie dąży się do zastosowania obrazowania *in vivo* na ludziach, do mapowania węzłów chłonnych i obrazowania guzów [8].

Znakowanie immunologiczne można zdefiniować jako wykorzystanie specyficznych przeciwciał do wiązania się z białkiem lub biomolekułą i wizualizację tego zdarzenia za pomocą znakowania kropkami kwantowymi. Dzięki zastosowaniu tej metody można wykrywać różne markery nowotworowe (np. P-glikoproteiny, cytokreatyny, specyficzny antygen błonowy prostaty). Znaczniki te są związane z transformacją lub przerzutami, zwykle wskazują raka i jego postępy. Inne zastosowania obejmują monitorowanie wirusa wewnątrzkomórkowego, typowanie antygenów komórek krwi, wizualizację efektów terapii lekowej na metabolizm komórkowy, a także monitorowanie markerów komórkowych [8].

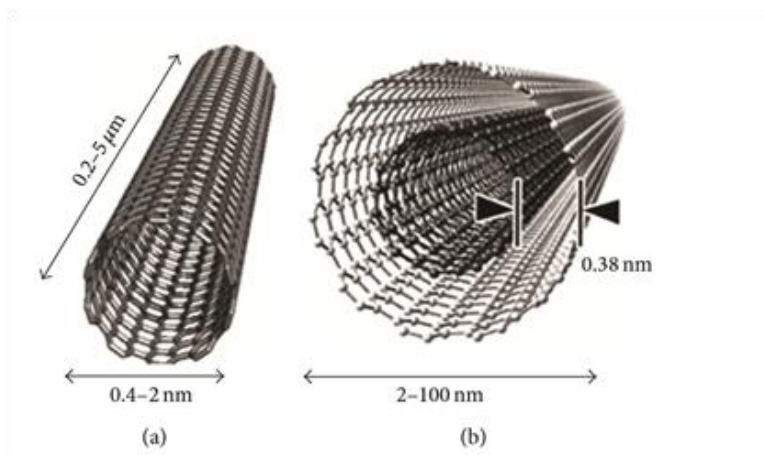
Dokładna identyfikacja głównych celów molekularnych odróżniających chore komórki od zdrowych umożliwiające celowane dostarczanie leku z minimalny-

mi efektami ubocznymi. Nanocząsteczki na bazie nośników leków pokazują ogromny potencjał dla aplikacji z ukierunkowaną dostawą medykamentów, ponieważ chronią ładunek przed degradacją i posiadają profil kontrolowanego uwalniania, jak również integrują wiele ukierunkowanych ligandów na ich powierzchni. Przykładem jest biodystrybucja nośników i ich wychwyt wewnątrzkomórkowy, które mogą być monitorowane za pomocą fluorescencji. Obecnie prowadzone są badania nad dystrybucją doksorubicyny do miejsc zajętych przez komórki rakowe w gruczole krokowym [3].

3. NANORURKI

Nanorurki węglowe (CNT) to alotropy węgla, wykonane z grafenu, mające kształt cylindryczny o średnicy rzędu nanometrów i długości rzędu kilku milimetrów. Ich imponujące strukturalne, mechaniczne i elektroniczne właściwości wynikają z nie dużych rozmiarów i masy, silnej siły mechanicznej i ich wysokiego przewodnictwa elektrycznego i cieplnego.

Dzięki ich wysokiej powierzchni, doskonałej stabilności chemicznej i bogatej elektronicznej strukturze, nanorurki mogą adsorbować lub koniugować z różnymi terapeutycznymi cząsteczkami (tj. leki, białka przeciwciała, DNA, enzymy, itp.). Okazały się być doskonałym sposobem do dostarczania leków wnikać bezpośrednio w komórki i utrzymując lek w stanie nienaruszonym, bez przemiany podczas transportu w organizmie.



Rys. 4. Schematy nanorurek węglowych: (a) o pojedynczych ściankach (SWCNT), (b) wielościenne (MWCNT) [2]

Struktura

W zależności od liczby warstw grafenu, z których zbudowana jest pojedyncza nanorurka, możemy dokonać podziału na pojedyncze nanorurki węglowe (SWCNT) lub wielościenne nanorurki węglowe (MWCNT). SWCNT składają się z jednego cylindra z grafenu o średnicy, która waha się między 0,4–2 nm. W skład MWCNT wchodzi od dwóch do kilku cylindrów współosiowych, każdy wykonany z pojedynczego arkusza grafenu, otaczającego pusty rdzeń. Zewnętrzna średnica utrzymuje się w zakresie od 2 do 100nm, natomiast średnica wewnętrzna mieści się w zakresie 1–3nm, długość jest to 0,2 do kilku m. [5]

Otrzymywanie

W metodzie ablacji laserem czynnikiem odpowiadającym za przemiany fizyczne grafitu z tarczy jest laser. Rurka kwarcowa zawierająca blok grafitu jest ogrzewana w piecu do temperatury około 1200°C. Podczas reakcji utrzymywany jest przepływ argonu. Laser jest używany do odparowania grafitu z kwarcu. Metoda laserowa daje możliwość wyprodukowania nanorurek zarówno jednościennych, jak i wielościennych. Nanorurki wytwarzane z tym sposobem nie muszą być jednakowo proste, a nawet zawierają pewne rozgałęzienia [9].

Uzyskanie nanorurek metodą chemicznego osadzania na różnych materiałach uzyskuje się dzięki pirolizie katalitycznej (CVD). Proces ten polega na chemicznym podziale węglowodoru na podłożu. Mogą tu być wykorzystane węglowodory takie jak benzen, acetylen, metan. Wyładowanie łukowe wzbudza atomy węgla, przez co jest możliwy wzrost nanorurek węglowych. Cząstki metaliczne (żelazo, kobalt, nikiel, miedź) zostają osadzone na nośnikach (tlenek krzemu, tlenek magnezu, tlenek wapnia) i umieszczone różnymi sposobami. Zasadniczo, otwory są wiercone w nośniku, a następnie wszczepia się nanocząstki metaliczne. Dalej, węglowodory, ogrzewa się i rozkłada na podłożu. Węgiel wchodzi w kontakt z cząstkami metali osadzonymi w otworach i zaczynają tworzyć nanorurki w kształcie tunelu [9].

Piroliza wysokotemperaturowa to opracowana przez naukowców metoda syntetyzowania dużych, długich wiązek jednościennych nanorurek w pionie pieca przez pirolizę cząsteczeki heksanu. Te n-heksanowe cząsteczki miesza się z pewnymi substancjami chemicznymi, które stymulują niezależnie wzrost nanorurek. Są one poddane procesowi pirolizy w bardzo wysokiej temperaturze (około 1200–1900°C) w strumieniu wodoru i opcjonalnie innych gazów [9].

Zastosowanie medyczne

Funkcjonalizowane CNT mogą działać jako nośniki dla leków przeciwbakteryjnych, takich jak przeciwgrzybicza amfoterycyna B. CNT dołączają się kowalencyjnie do amfoterycyny B i transportują ją do komórek. Takie połączenie wpływa na zmniejszenie toksyczności przeciwgrzybiczej o prawie połowę

w stosunku do wolnego leku. Funkcjonalizowane CNT mogą również służyć jako dostawca szczepionek.

Wiązanie bakteryjnego lub wirusowego antygeny z CNT pozwala na utrzymanie w stanie nienaruszonym konformacji antygenowej, tym samym indukując odpowiedź przeciwciał. Połączenie funkcjonalizowanych CNT z komórkami B i T epitopów peptydowych może generować system zdolny do wywołania silnej odpowiedzi immunologicznej. Ponadto, nanorurki wykazują aktywność przeciwbakteryjną, ponieważ bakterie mogą być adsorbowane na powierzchni CNT, tak jak w przypadku bakterii *e. coli*. Efekt przeciwbakteryjny nanorurek doprowadza do utleniania wewnątrzkomórkowego glutationu, powodując zwiększenie stresu oksydacyjnego w komórkach bakteryjnych i ewentualnej śmierci komórkowej [2].

Postęp wiedzy o przeszczepie komórek i narządów oraz chemii CNT w ostatnich latach przyczynił się do zrównoważonego rozwoju inżynierii tkanki opartej na nanorurkach i medycyny regeneracyjnej. Nanorurki to materiał zgodny biologicznie, odporny na degradację biologiczną i może być funkcjonalizowany z biocząsteczką do wzmacniania regeneracji narządów. W tej dziedzinie, nanorurki są stosowane jako dodatki w celu wzmocnienia wytrzymałości mechanicznej i przewodności rusztowania tkankowego poprzez włączenie w organizm gospodarza. Połączono karboksylowane SWCNTs z polimerem lub kolagenem, tworząc kompozytowy nanomateriał używany jako rusztowanie regeneracji tkanki [2].

Z powodu swoich małych rozmiarów i dostępnych zewnętrznych modyfikacji, nanorurki są w stanie przekroczyć barierę krew-mózg za pomocą różnych mechanizmów i mogą być przeznaczone do dostawy skutecznych nośników do mózgu. Naukowcy zaobserwowali, że SWCNTs były z powodzeniem wykorzystywane do dostarczenia acetylocholiny do mózgu myszy, dotkniętych chorobą Alzheimera, z wysokim zakresem bezpieczeństwa. Wiele innych funkcjonalizowanych SWCNTs lub MWCNTs z powodzeniem są stosowane jako odpowiednie systemy dostarczania dla leczenia chorób neurodegeneracyjnych albo guzów mózgu. Koniugaty nanorurek z cząsteczkami terapeutycznymi mają lepszy wpływ na wzrost neuronów niż leki stosowane oddzielnie [2].

5. NANODRUTY

Nanodrut to nanostruktura, o średnicy rzędu nanometrów. Można je również określać jako konfiguracja, której stosunek długości do szerokości jest znacząco duży. Alternatywnie nanodruty to nanomateriały, których grubość lub średnica jest ograniczona do kilkudziesięciu nanometrów lub mniej, zaś długość jest dowolna. Istnieje wiele różnych typów w tym nadprzewodzące, metaliczne (np. Ni, Pt, Au), półprzewodnikowe i izolujące (np. SiO₂, TiO₂).

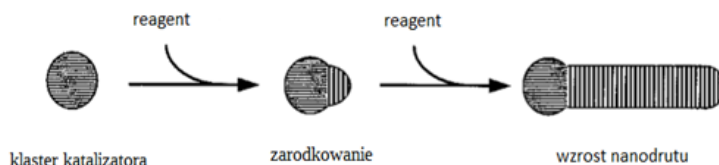
Budowa

Nanodruty wykazują typowe zależności proporcji, w których stosunek długości do szerokości wynosi 1000:1 lub więcej. Często określane są jako jednowymiarowe (1-D) materiały. Posiadają wiele ciekawych właściwości, które nie są widoczne w pojedynczych materiałach lub trójwymiarowych. Jest tak, ponieważ elektrony nanodrutów posiadają kwantowe ograniczenie boczne, a tym samym zajmują poziomy energetyczne, które różnią się od tradycyjnych ciągłych (uporządkowanych) poziomów energetycznych lub zespołów występujących w materiałach pojedynczych [4].

Procesy tworzenia

Metoda wzrostu kryształów VLS jest bez wątpienia najbardziej powszechnie przyjętym podejściem do wzrostu nanodrutów półprzewodnikowych z powodu jego dużej elastyczności. Wzrost nanodrutu techniką VLS obejmuje trzy odrębne etapy (zilustrowane na rys. 10):

- stapianie, polegające na wytworzeniu kropelek cieczy ze stopu, na którym przewód ma być tworzony
- zarodkowanie, którego celem jest umieszczenie substancji na podłożu w postaci pary, która absorbuje się na powierzchni cieczy i dyfunduje do kropli
- wzrost, w którym przesycenie i zarodkowanie na cieczy / ciele stałym interfejsu prowadzi do osiowego wzrostu kryształów.



Rys. 10. Schemat ilustrujący syntezę katalityczną nanodrutów [4]

Metodę samoorganizacji charakteryzuje wzrost bezpośrednio na powierzchni, i równolegle do niej. Powstaje tym sposobem szereg nanodrutów na powierzchni, rozmiar średnicy nie przekracza kilku nanometrów, a odległość między poszczególnymi drutami wynosi maksymalnie dziesięć nanometrów.

Zastosowanie medyczne

Interfejs między nanosystemami i biosystemami staje się jedną z najszerzych i najbardziej dynamicznych dziedzin nauki i techniki, łącząc biologię,

chemię, fizykę, medycynę i biotechnologię. Połączenie różnych obszarów badań obiecuje przynieść rewolucyjne postępy w ramach opieki zdrowotnej, medycyny i nauk przyrodniczych.

Białko w roztworze można wykryć za pomocą urządzeń opartych na krzemowych nanodrutach typu p, w których cząsteczka biotyny, wiążąca się z wysoką selektywnością w stosunku do białka streptawidyny, jest połączona z powierzchnią tlenkową nanodrutów. Nanodruty krzemowe wykorzystuje się do badań w postaci czujników do wykrywania jednoniciowego DNA, gdzie wiązanie tej naładowanej ujemnie polianionowej makrocząsteczki z powierzchnią nanodrutu prowadzi do wzrostu konduktancji. Rozpoznanie docelowych cząsteczek DNA przeprowadza się z sekwencji komplementarnej jednoniciowego materiału, który jest komplementarny do peptydowego kwasu nukleinowego (PNA). PNA jest receptorem dla detekcji DNA od momentu, gdy nienaładowana cząsteczka PNA ma większe powinowactwo i trwałość w porównaniu z odpowiednimi sekwencjami rozpoznawania DNA. Dane detekcji DNA uzyskane z niezależnych urządzeń wykazują bardzo podobne zmiany w przewodności ze wzrostem stężenia DNA. Powtarzalność jest ważnym uzasadnieniem potencjału nanodrutów krzemowych dla rozwoju zintegrowanych czujników, które mogłyby umożliwić wysoką wydajność oraz bardzo czułe wykrywanie DNA w badaniach genetycznych [10].

Cząsteczki organiczne, które wiążą się z białkami, są kluczowe do odkrywania oraz rozwoju produktów leczniczych i tym samym stanowią istotny cel dla czujników. Reprezentatywnym przykładem w tej dziedzinie jest identyfikacja inhibitorów molekularnych kinaz tyrozynowych, które są białkami i pośredniczą w transdukcji sygnałowej w komórkach ssaków poprzez fosforylację reszty tyrozyny substratu białka, wykorzystując ATP. Wiązanie lub hamowanie wiązania ujemnie naładowanego ATP Abl związanego na powierzchni nanodrutu krzemowego jest wykrywane jedynie przy zwiększeniu lub zmniejszeniu przewodności urządzenia. Dane zależne od czasu zarejestrowane z Abl modyfikowanego nanodrutem krzemowym wykazują odwracalny, zależny od stężenia wzrost przewodności na wprowadzanych roztworach zawierających ATP. Przewodność maleje przy stałej koncentracji małych cząsteczek. Zaletą urządzeń są szybkie i bezpośrednie odczyty o wysokiej czułości [10].

Pierwszym przykładem wykazującym zdolność urządzeń opartych na nanodrutach do wykrywania gatunków w ciekłych roztworach przedstawiono w roku 2001 w przypadku stężenia jonów wodorowych i wykrywaniu pH. Podstawowe urządzenie z nanodrutu krzemowego typu p przekształcono do takiego czujnika przez modyfikację powierzchni tlenku, które dają grupy aminowe na powierzchni nanodrutu wraz z naturalnie występującą silanolową grupą tlenową. Aminowe i silanolowe cząstki działają jak receptory dla jonów wodorowych, które ulegają reakcjom protonowania/deprotonowania, zmieniając w ten sposób ładunek powierzchniowy nanodrutu. Urządzenia zmodyfikowane w ten sposób wykazują stopniowy wzrost przewodności pH roztworu, który jest dostarczany za pośred-

nictwem urządzenia mikrostrumieniowego. Wzrost prawie liniowy przewodności jest atrakcyjny z punktu widzenia czujnika, a wynika z obecności dwóch odrębnych grup receptorów ulegających protonowaniu / deprotonacji na różnych zakresach pH. Z mechanistycznego punktu widzenia, zwiększenie przewodności ze wzrostem pH jest zgodne ze zmniejszeniem (zwiększeniem) bieguna dodatniego powierzchni ładunku [10].

LITERATURA

- [1] Szymański P., Markowicz M., Mikiciuk-Olasik E., *Zastosowanie nanotechnologii w medycynie i farmacji*, „Lab”, R. 17, nr 1
- [2] He H., Pham-Huy L. A., Dramou P., Xiao D., Zuo P., Pham-Huy Ch., *Carbon Nanotubes: Applications in Pharmacy and Medicine*, „BioMed Research International” July 2013, vol. 2013
- [3] Salata OV, *Applications of nanoparticles in biology and medicine*, “Journal of Nanobiotechnology” April 2004
- [4] Hu J., Odom T. W., Lieber Ch. M., *Chemistry and Physics in One Dimension: Synthesis and Properties of Nanowires and Nanotubes*, „Accounts of chemical research” 1999, vol. 32, No. 5
- [5] Bhushan Bharat, *Springer Handbook of Nanotechnology*, Wyd. 3, Wydaw. Springer
- [6] Jacak L., Wójs A., *Kropki kwantowe*, „Postępy fizyki” 1998, Tom 49, Zeszyt 1
- [7] Smith A. M., Duan H., Mohs A. M., Nie S., *Bioconjugated Quantum Dots for In Vivo Molecular and Cellular Imaging*, „Advanced Drug Delivery Reviews” August 2008, vol. 60, Issue 11, p. 1226–1240
- [8] Medintz I. L., Mattoussi H., Clapp A. R., *Potential clinical applications of quantum dots*, „International Journal of Nanomedicine” 2008, vol. 3, Issue 2, s. 151–167
- [9] Bachmatiuk A., *Badania nad technologią otrzymywania i własnościami nanorurek węglowych*, Praca doktorska, Politechnika Szczecińska, 2008
- [10] Patolsky F., Lieber Charles M., *Nanowire nanosensors*, „Materials Today” April 2005, vol. 8, Issue 4
- [11] Tytus M., *Własności optyczne kropek kwantowych drugiego rodzaju*, Praca doktorska, Politechnika Wrocławska, 2011
- [12] Rzeszutek J., Matysiak M., Czajka M., Sawicki K., Rachubik P., Kruszewski M., Kapka-Skrzypczak L., *Zastosowanie nanocząstek i nanomateriałów w medycynie*, „Hygeia Public Health” 2014, vol. 49, Issue 3, s. 449–457

ZASTOSOWANIE CZUJNIKA KWARCOWEGO W URZĄDZENIU DO BADANIA KROPLI KRWI W DIAGNOSTYCE CHOROBY ALZHEIMERA

1. WSTĘP

Ciągły rozwój obszarów elektroniki, mechaniki, automatyzacji oraz powiększająca się wiedza z zakresów współczesnej medycyny daje nam coraz to nowe możliwości w zakresie szybszej i dokładniejszej diagnozy różnych chorób. Jedną z możliwości jest pomiar małych ilości mas, przy zastosowaniu m.in. czujników piezoelektrycznych.

Dzisiejsza technologia oferuje nam szeroki wachlarz różnorodnych czujników i elementów elektronicznych stosowanych w elektronicznej aparaturze medycznej. Z całego zakresu możemy wyróżnić czujniki kwarcowe, które połączone w odpowiedniej konfiguracji mogą spełniać różne role np. rolę mikrowagi kwarcowej. Mikrowagi są powszechnie stosowane w środowiskach laboratoryjnych, ponieważ są w stanie wychwycić bardzo małe zmiany masy.

W niniejszym artykule przedstawiono analizę możliwości zastosowania czujnika kwarcowego do pomiar małych mas związanych z badaniem próbki krwi o bardzo małej objętości substancji tj. 10 μ l.

2. CZUJNIKI KWARCOWE I MIKROWAGA KWARCOWA

Sensory kwarcowe są powszechnie dostępnymi czujnikami stosowanym w aparaturze pomiarowej np. mikrowagach elektronicznych. Są często wykorzystywane ze względu na swoje charakterystyki metrologiczne, stosunkowo niską cenę oraz łatwość ich adaptacji w układach elektronicznych.

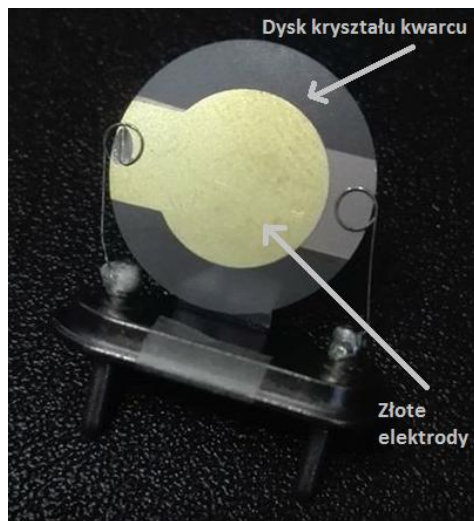
Czujnik kwarcowy jest elementem piezoelektrycznym wykorzystującym efekt piezoelektryczny odwrotny, który został potwierdzony eksperymentalnie przez braci Curie w 1881 roku [3]. Efekt ten powoduje powstanie odkształceń kryształu oraz zmianę jego wymiarów pod wpływem przyłożenia napięcia do elektrod [2, 5].

Zastosowany do przeprowadzenia badań czujnik kwarcowy składa się z płytki kwarcu, nazywanej dyskiem, w kształcie koła o wymiarach:

- grubość – 100 μ m
- średnica – 14 mm.

¹ Politechnika Lubelska, Wydział Elektrotechniki i Informatyki

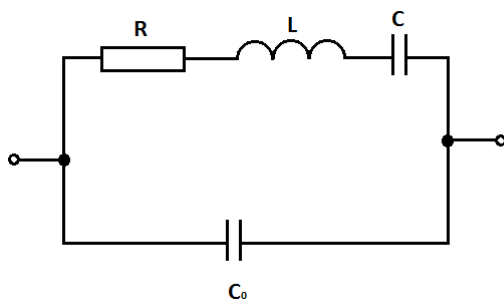
Na powierzchni płytki napyłone są elektrody metaliczne, które zwiększają czułość. Elektrody zapewniają możliwość wzbudzenia drgań czujnika [2].



Rys. 1. Czujnik kwarcowy

Model zastępczy elektryczny – układ RLC

Analizowany czujnik kwarcowy pod względem elektrycznym może być traktowany jako obwód rezonansowy o schemacie zastępczym jak na rysunku 2.



Rys. 2. Schemat zastępczy elektryczny

Wielkość elementów przedstawionych na schemacie zależą od:

- wielkość C oraz L – parametrów mechanicznych płytki czujnika,
- wielkość R – tłumienia drgań obwodu.

Parametrem C_0 określa się pojemność elektrod oraz przewodów łączących [6].

Tab. 1. Przykładowe parametry rezonatorów

f Hz	100 k	500 k	1 M	4 M	10 M	20 M	60 M	120 M
R Ω	400	500	250	20	20	10	30	50
L H	93,8	20,3	3,62	0,100	0,0169	0,0042	0,0035	0,00293
C pF	0,0027	0,005	0,007	0,015	0,015	0,015	0,002	0,0006
C_0 pF	6	6	5	5	3,5	3	5	4

Rezystancja R ma małą wartość. W opisach matematycznych zachodzących przemian jest często pomijalna i przyjmuje wartość 0Ω [9]. Impedancja Z (dla $R=0\Omega$) przedstawia się następująco:

$$Z = \frac{j}{\omega} * \frac{\omega^2 LC - 1}{C_0 + C - \omega^2 LC C_0} \quad (1)$$

gdzie:

L – indukcyjność,

C – pojemność,

C_0 – pojemność elektrod.

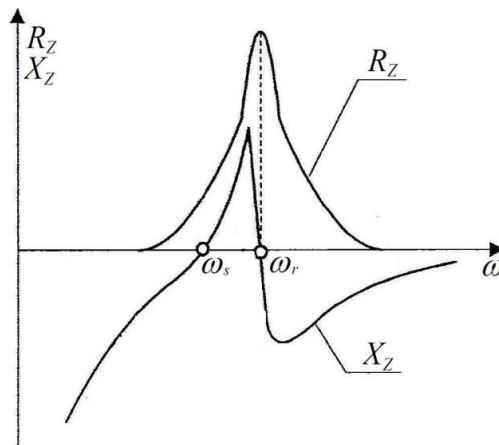
Wyróżniamy dwa rodzaje rezonansu:

- szeregowy ($Z=0$), wyrażony wzorem:

$$f_S = \frac{1}{2\pi\sqrt{LC}} \quad (2)$$

- równoległy ($Z=\infty$)

$$f_R = \frac{1}{2\pi\sqrt{LC}} \sqrt{1 + \frac{C}{C_0}} = f_S \sqrt{1 + \frac{C}{C_0}} \quad (3)$$



Rys. 3 Zmiany impedancji w funkcji częstotliwości [8]

Duża dobroć (stroma charakterystyka fazowa w obrębie f_r) pozwala na uzyskanie dużej stabilności f generowanego sygnału [8]. Dobroć wyrażona jest wzorem:

$$Q = \frac{1}{R} \sqrt{\frac{L}{C}} \quad (4)$$

Innym wzorem na częstotliwość jest wzór związany z wymiarami kwarcu. Zależności te przedstawia wzór 5:

$$f = \frac{nv_q}{2t_q} \quad (5)$$

gdzie:

n – kolejne nieparzyste nadtony (1,3,5...),

t_q – grubość płytki,

v_q – prędkość rozchodzenia się okształcenia ściskającego w kwarcu.

Mikrowaga kwarcowa jest bardzo czułym urządzeniem przeznaczonym do pomiarów bardzo małych zmian masy. Głównym elementem mikrowagi jest czujnik kwarcowy. Występują różne układy pomiaru masy, najważniejsze wymieniono poniżej.

W układzie generatora, stosowana jest jako sensor służący do obserwowania zmian ilości materiału zaaplikowanych na powierzchni płytki kwarcowej poprzez zmierzenie zmiany jej częstotliwości [7]. Częstotliwość rezonansowa jest zależna od ilości cząsteczek zgromadzonych na powierzchni czujnika. Równanie Sauerbrea przedstawia związek pomiędzy zmianą masy a zmianą częstotliwości [4]:

$$\Delta m = -C \frac{\Delta f}{n} \quad (6)$$

gdzie:

Δm – zmiana masy,

Δf – zmiana częstotliwości,

C – stała dla kryształu kwarcu, wynosi $17,7 \text{ ng}/(\text{Hz} \cdot \text{cm}^2)$,

n – numer nadtonu.

Ze wzoru 6 wynika zależność, że zmiana częstotliwość rezonansowej jest proporcjonalna do zmiany masy. Jeśli jest konieczne dostrojenie kwarcu w niewielkim zakresie, podcina się szeregowo do niego dodatkowy kondensator C_s , którego wartość jest większa od wartości C .

Korzystając ze wzorów z rozdziału 1.1, wypadkowa impedancja połączenia wyraża się wzorem:

$$Z' = \frac{1}{j\omega C_s} * \frac{C + C_0 + C_s - \omega^2 LC(C_0 + C_s)}{C_0 + C - \omega^2 LCC_0} \quad (7)$$

gdzie:

L – indukcyjność,

C – pojemność,

C_0 – pojemność elektrod,

C_s – pojemność kondensatora dodatkowego.

przy uwzględnieniu, że $C \ll C_0 + C_s$

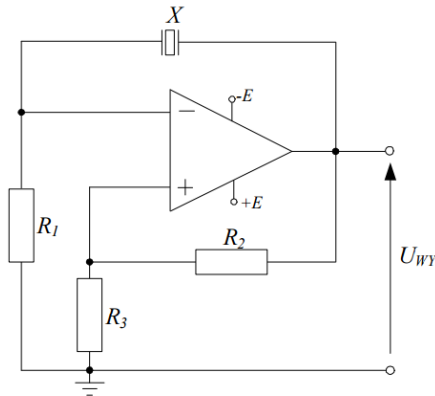
$$f'_S = f_s \left[1 + \frac{C}{2(C_0 + C_s)} \right] \quad (8)$$

z czego wynika względna zmiana częstotliwości:

$$\frac{\Delta f}{f} = \frac{C}{2(C_0 + C_s)} \quad (9)$$

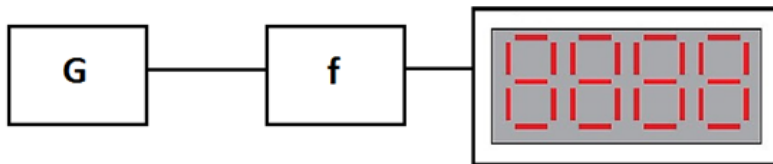
Wartość f_R pozostaje taka sama, dołączenie C_s nie zmienia jej wartości.

Przykładowe zastosowanie rezonatora kwarcowego wraz z miernikiem częstotliwości przedstawione jest jako układ w rezonansie równoległym (rys. 4).



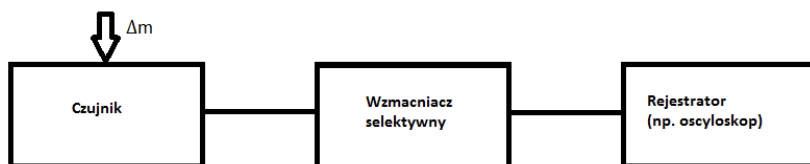
Rys. 4. Schemat generatora kwarcowego z równoległym układem rezonansowym [8]

Schemat wagi z układem częstotliwościomierza przedstawia rysunek 5.



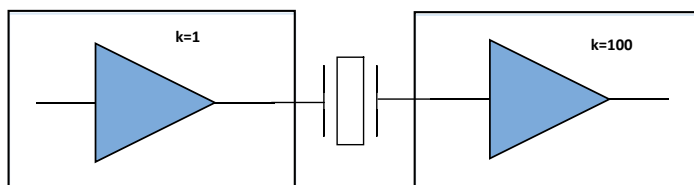
Rys. 5. Schemat wagi z układem częstotliwościomierza

Schemat (rys. 5) składa się z generatora, częstotliwościomierza oraz panelu odczytowego (np. oscyloskop). Częstotliwościomierz jest wyskalowany w jednostkach masy. W układzie wzmacniacza selektywnego, wzmacniacz pobudzany sygnałem sinusoidalnym lub impulsowym – pomiar charakterystyki częstotliwościowej sensora.



Rys. 6. Schemat blokowy mikrowagi z użyciem wzmacniacza selektywnego

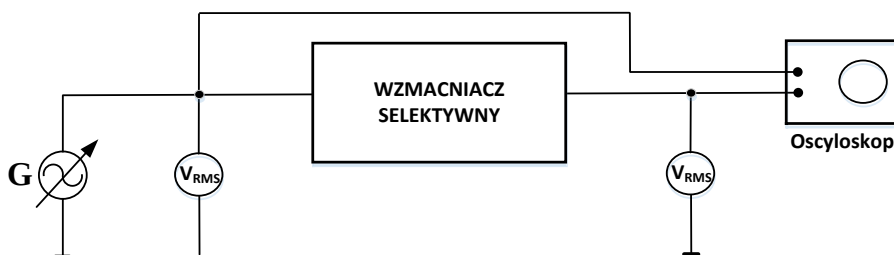
Schemat blokowy wzmacniacza selektywnego prezentuje poniższy rysunek.



Rys. 7. Schemat blokowy wzmacniacza selektywnego

W schemacie blokowym znajdują się dwa wzmacniacze selektywne oraz czujnik kwarcowy. Najważniejszym zadaniem układu jest wzmocnienie sygnału.

Układ pomiarowy do badania charakterystyki częstotliwościowej prezentuje schemat na rysunku 8.



Rys. 8. Schemat układu pomiarowego

Na podstawie powyższego schematu możliwe jest wyliczenie wzmocnienia k , które przedstawione jest wzorem:

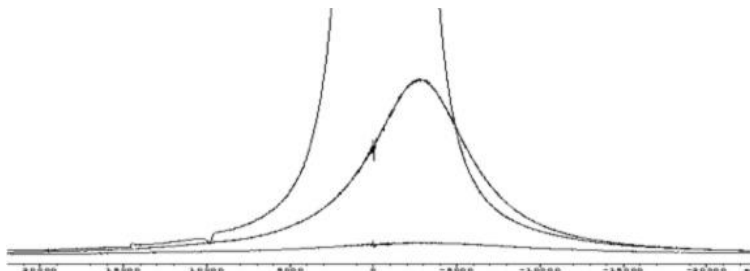
$$k = \frac{U_2}{U_1} \Big|_{f=var} \quad (10)$$

gdzie:

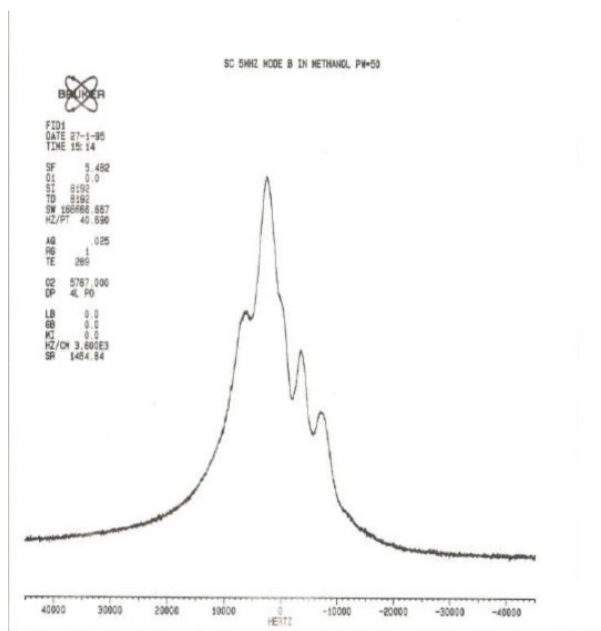
U_1 – napięcie wskazane przez woltomierz 1,

U_2 – napięcie wskazane przez woltomierz 2.

Wzmacniacz selektywny pobudzany jest sygnałem sinusoidalnym lub impulsowym. Dzięki czemu możliwe jest otrzymanie charakterystyki częstotliwościowej. W otrzymanej charakterystyce bardzo ważne jest położenie środka amplitudy, stosunek amplitudy do szerokości połówkowej oraz kształt charakterystyki. Pożądany kształt przedstawiony jest poniżej (rys. 9 i 10).



Rys.9 Charakterystyka [1]

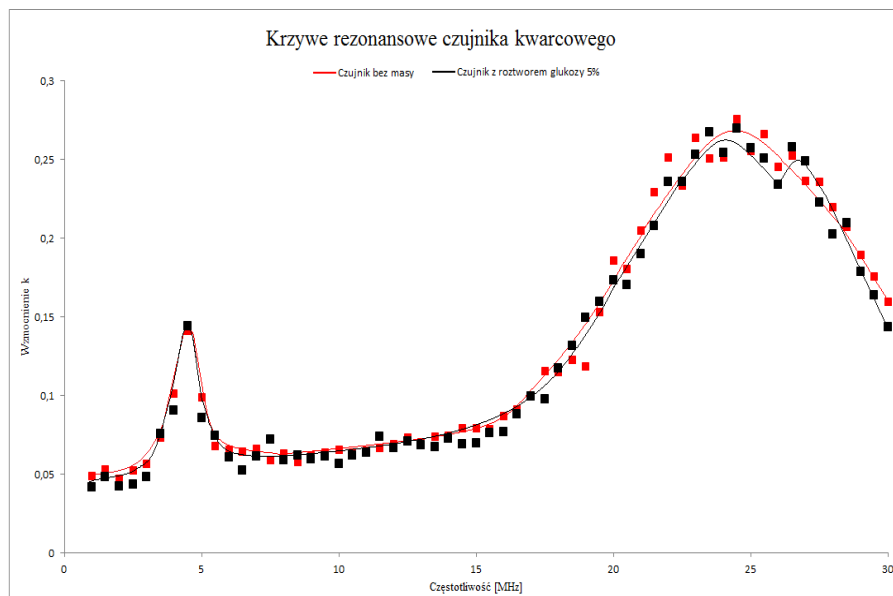


Rys. 10. Charakterystyka czujnika uzyskana za pomocą konsoli BRUKER-a [1]

Powyższe charakterystyki zostały wykonane za pomocą konsoli BRUKER-a dla częstotliwości 5 MHz. Do badań użyto czujnika kwarcowego w układzie ze wzmacniaczem selektywnym pobudzany sygnałem sinusoidalnym lub impulsowym.

2. PODSUMOWANIE

Otrzymane charakterystyki tj. czujnika nieobciążonego masą oraz czujnika z masą (w tym przypadku jest to roztwór glukozy o stężeniu 5%) zostały przedstawione na poniższym wykresie:



Rys. 11. Krzywe rezonansowe czujnika kwarcowego

Z powyższego wykresu można wywnioskować, że pierwsza harmoniczna czujnika kwarcowego znajduje się w granicach ok. 4,5MHz. Kolejna harmoniczna występuje przy częstotliwości 24,5MHz i jest to piąta harmoniczna. Dla tej wartości wzmocnienie jest największe. Widoczne jest również przesunięcie maksymalnej wartości dla czujnika z masą.

Różnice pomiędzy dwoma analizowanymi krzywymi są widoczne, zwłaszcza przy 5 harmonicznej. Krzywa rezonansowa dla zakroplonego czujnika jest nieco węższa oraz lekko obsunięta względem drugiej krzywej.

Uzyskane charakterystyki częstotliwościowe wyraźnie pokazują różnicę przy porównaniu krzywej rezonansowej czujnika nieobciążonego masą i czujnika z założoną ilością masy. Krzywe rezonansowe zostały uśrednione, pozwala na zauważenie przesunięcia wartości szczytowej dla czujnika obciążonego masą. Użyty w badaniach generator sygnału posiadał zakres częstotliwości w zakresie do 30MHz. Laboratorium nie posiada na wyposażeniu generatora o większym zakresie. W celu uzyskania pełnej krzywej rezonansowej należy przeprowadzić badania z użyciem generatora, który posiada zakres do co najmniej 50MHz. Aby

otrzymane wyniki pomiarów były dokładniejsze należy przeprowadzić większą ilość pomiarów (nawet co 10 kHz) w obrębie występowania harmonicznych.

LITERATURA

- [1] *Memorandum clinical neuroengineering*
<http://clinicalneuroengineering.eu/memorandum>, dostęp z dnia 24.10.2016
- [2] Gniewińska B., Klimek C., *Rezonatory i generatory kwarcowe*. Wydawnictwa Komunikacji i Łączności, Warszawa 1980
- [3] Praca zbiorowa pod red. W. Solucha, *Wstęp do piezoelektroniki*. Wydawnictwa Komunikacji i Łączności, Warszawa 1980
- [4] Ćwięka M., *Adsorpcja lizozymu- charakterystyka struktur białkowych z zastosowaniem mikrowagi kwarcowej z monitorowaniem dyssipacji energii oraz powierzchniowego rezonansu plazmonów*. „Zeszyty Naukowe Towarzystwa Doktorantów UJ Nauki Ścisłe”, nr 9 (2/2014), Kraków 2014
- [5] http://www.zstio-elektronika.pl/pliki_t_elektronik/TE_Z3-01.pdf, dostęp z dnia 29.12.2016
- [6] Generatory, <http://elektronik.tl.krakow.pl>, dostęp z dnia 08.01.2017
- [7] Elektronika <http://www.pe.ifd.uni.wroc.pl/Elektronika-W10.pdf>, dostęp z dnia 10.01.2017
- [8] Generatory http://ue.pwr.wroc.pl/wyklad_w10/W14_Generatory.pdf, dostęp z dnia 10.01.2017
- [9] *Układy elektroniczne i technika pomiarowa*. Moduł 6 – generatory sinusoidalnych i niesinusoidalnych, http://wazniak.mimuw.edu.pl_6, dostęp z dnia 02.01.2017

MODEL I SYMULACJE TRANZYSTORA MOSFET Z WĘGLIKA KRZEMU

1. WPROWADZENIE

Rozwój elementów półprzewodnikowych od zawsze związany był z poszukiwaniem materiałów posiadających coraz większą wartość przerwy zabronionej (E_g). Pierwszym stosowanym materiałem w latach 50. był german ($E_g = 0,78$ eV), około 10 lat później pojawił się krzem ($E_g = 1,21$ eV), który zdominował rynek mikroelektroniki na wiele lat. Od lat 90. konstruktorzy elementów półprzewodnikowych skupili się na materiałach o dużej wartości przerwy zabronionej (umownie $E_g > 2$ eV), a najwięcej nadziei zaczęto pokładać w węgliku krzemu ($E_g = 3$ eV) [1]. Pojawienie się węgliku krzemu pozwoliło na znaczny rozwój elektroniki, szczególnie w zakresie urządzeń pracujących przy wysokich napięciach i temperaturach. Czynniki te były jednymi z głównych powodów poszukiwań materiału zastępczego dla krzemu, który okazał się niewystarczający, pomimo doskonale rozwiniętej technologii.

W ostatnich latach dokonywano szeregu badań i procesów technologicznych związanych z obróbką węgliku krzemu. Technologia związana z tym materiałem jest znacznie trudniejsza i droższa niż w przypadku krzemu, głównie spowodowane jest to różnicą w budowie obu materiałów. Wszelkie urządzenia elektroniczne oparte na węgliku krzemu wymagają dużego nakładu pieniężnego, więc niewielkie pomyłki w projektowaniu mają spore odbicie w kwestii finansów. W tym zakresie z pomocą przychodzą programy pozwalające na wykonanie symulacji technologicznych projektowanych urządzeń – zalicza się do nich program Silvaco TCAD. Symulacje w dużym stopniu umożliwiają uniknięcie kosztownych błędów i zaoszczędzenie czasu. pozwalają na obserwowanie wpływu zmian parametrów technologicznych oraz na ich odpowiedni dobór, w zależności od oczekiwanych zachowań urządzenia.

Dotychczas najlepiej dopracowanym urządzeniem wykonanym w technologii węgliku krzemu jest dioda Schottky'ego. Tranzystory są układami znacznie bardziej skomplikowanymi niż złącza Schottky'ego, dlatego prace nad nimi rozpoczęły się później – pierwsze tranzystory MOSFET pojawiły się na rynku dopiero w 2011 roku. Na dzień dzisiejszy parametry tranzystorów krzemowych niemal osiągnają nieprzekraczalne granice związane z samym materiałem, czego nie można powiedzieć o urządzeniach z węgliku krzemu. Potrzeba jeszcze wielu lat i badań nad tym materiałem, by osiągnąć wysoki poziom technologii, jak

¹ Politechnika Lubelska, Koło Naukowe „SEMICON”

w przypadku krzemu. Szybkie tranzystory MOSFET z SiC pozwolą na znaczną miniaturyzację urządzeń takich jak przetwornice czy zasilacze impulsowe, a także będą pożądaną we wzmacniaczach klasy D.

2. WĘGLIK KRZEMU

Węglik krzemu jest od dawna znany jako materiał z dużym potencjałem w zakresie urządzeń pracujących w wysokich temperaturach, przy dużych mocach i wysokich częstotliwościach. Charakteryzuje się on również dużą odpornością na promieniowanie. Jego najważniejszą cechą jest szeroka przerwa zabroniona, wynosząca 2,39–3,3 eV (w zależności od politypu).

Większość tradycyjnych układów scalonych bazujących na krzemie nie jest w stanie pracować w temperaturach powyżej 250°C, w szczególności gdy oprócz wysokiej temperatury pojawiają się również wysokie moce i częstotliwości, a w środowisku pojawia się promieniowanie. Ma to związek z małą przerwą zabronioną, niską przewodnością cieplną oraz niewielką wartością pola krytycznego krzemu. W tym zakresie najlepszym materiałem zastępczym dla krzemu jest węglik krzemu. Znane są również inne półprzewodniki z dużą przerwą zabronioną, takie jak arsenek galu (GaAs), azotek aluminium (AlN), azotek galu (GaN), azotek boru (BN) oraz diament. Jednak węglik krzemu posiada kilka zalet, które czynią go lepszym materiałem niż pozostałe. Do zalet tych należy ogólna dostępność podłoża, znajomość procesów technologicznych oraz możliwość wykorzystania tlenku do stworzenia maski dla procesów, do pasywacji złączy oraz jako dielektryka bramkowego [2].

Tabela 1. Zestawienie podstawowych parametrów półprzewodników [3]

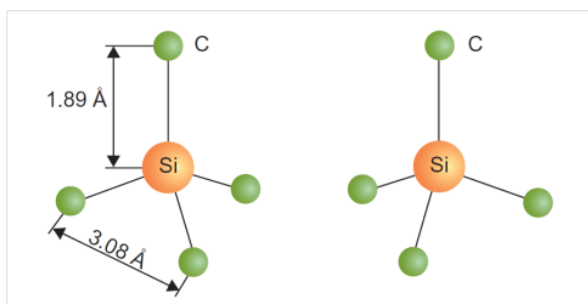
Parametry	Si	3C-SiC	6H-SiC	4H-SiC	Diament
Szerokość przerwy zabronionej, eV	1,12	2,29	3,00	3,20	5,45
Ruchliwość elektronów, cm^2/Vs	1450	800	600	1000	2200
Ruchliwość dziur, cm^2/Vs	470	40	100	115	1600
Prędkość unoszenia elektronów, $\times 10^7 \text{ cm/s}$	1,0	2,5	2,0	2,0	1,5
Krytyczne pole przebiecia, MV/cm	0,25	2,12	2,5	2,2	1–10
Przewodność cieplna, $\text{W/K}\cdot\text{cm}$	1,49	3,6	3,6	3,7	6–20

W tabeli 1 przedstawiono zestawienie fizycznych właściwości trzech głównych politypów węglika krzemu oraz krzemu i diamentu. Duża przerwa energe-

tyczna, pole przebicia, prędkość unoszenia elektronów oraz przewodność cieplna węgla krzemu (zblizona do Cu) umożliwiają stosowanie go przy produkcji urządzeń wysokonapięciowych, wysokoczęstotliwościowych, przewodzących duże prądy i pracujących przy wysokich temperaturach. Przewodność elektryczna warstw SiC może być łatwo kontrolowana w zakresie (10^{14} – 10^{19} cm⁻³) dla typu p i typu n poprzez domieszkowanie. W celu uzyskania domieszki donorowej używa się azotu lub fosforu, dla akceptorowej boru, aluminium lub galu. Dodatkowo, powierzchnia węgla krzemu może być pokryta wysokiej jakości warstwą tlenku uzyskaną poprzez utlenianie termiczne, co jest istotnym czynnikiem przy produkcji urządzeń.

Ponieważ węgiel krzemu posiada pole przebicia dziesięć razy większe od krzemu, urządzenia mocy bazujące na tym materiale mogą być dziesięciokrotnie cieńsze niż w przypadku krzemu. Około trzy razy większa przewodność cieplna pozwala na stosowanie prostych systemów chłodzenia w układach zbudowanych na bazie węgla krzemu. Dzięki tym cechom, przyszłość tego materiału jest obiecująca ze względu na to, że determinują one mniejsze rozmiary urządzeń, mniejsze straty, większą efektywność i łatwiejsze rozpraszanie się ciepła [4–6].

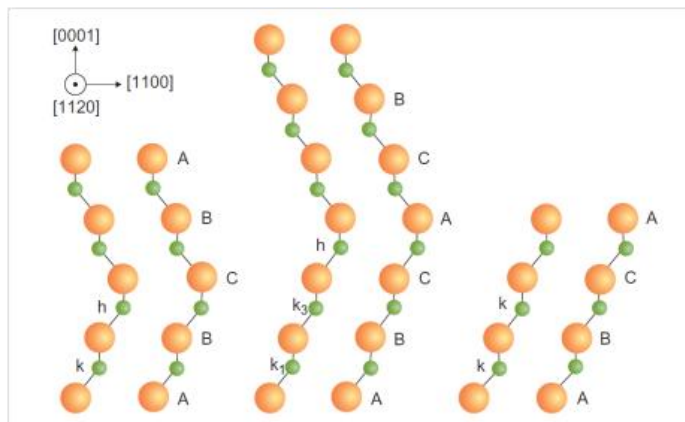
Węgiel i krzem są pierwiastkami IV grupy układu okresowego. Każdy z nich posiada po 4 elektrony walencyjne na zewnętrznej powłoce, które ulegają hybrydyzacji tetraedrycznej sp³, dzięki czemu otrzymywane są bardzo silne wiązania kowalencyjne między atomem krzemu i czterema atomami węgla znajdującymi się wokół. Oba pierwiastki posiadają również zbliżoną elektroujemność (C=2,5, Si=1,8 w skali Paulinga). Węgiel jako pierwiastek bardziej elektroujemny, przyciąga elektrony krzemu, dając wiązanie o charakterze rezonansowym i zyskując składową jonową. Zależności te tworzą tetraedryczną koordynację atomów (rys. 1). Typowymi strukturami sieci krystalicznej dla tego rodzaju wiązań są: typ wurcytu, sfalerytu i romboedryczna.



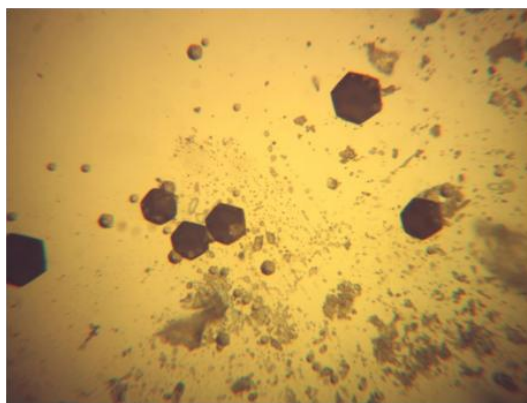
Rys. 1. Tetraedryczna koordynacja atomów [3]

Zaburzenia w strukturach sieci krystalicznej, takie jak dodanie lub usunięcie jednej podwójnej warstwy, powodują utworzenie struktury innego typu – zjawisko to nazywane jest politypizmem [8]. Istnieje ponad 250 różnych struktur poli-

typowych węglik krzemu. Do ich określania stosuje się specjalną metodę opisu stworzoną przed Ramsdela. Złożona jest ona z cyfry, która określa nam liczbę warstw gęstego ułożenia w okresie identyczności danego politypu oraz z litery, która informuje o symetrii komórki elementarnej. Występują 3 odmiany symetrii – regularna (C), heksagonalna (H) i romboedryczna (R) [7]. Najważniejszymi politypami są: 3C, 2H, 4H, 6H, 8H, 9R, 10H, 15R, 19R, 20H, 21H i 24R [2], ale w zastosowaniu praktycznym wykorzystuje 3C, 4H oraz 6H [9] (Rys. 2).



Rys. 2. Podstawowe politypy węglik krzemu (od prawej 4H, 6H, 3C) [3]



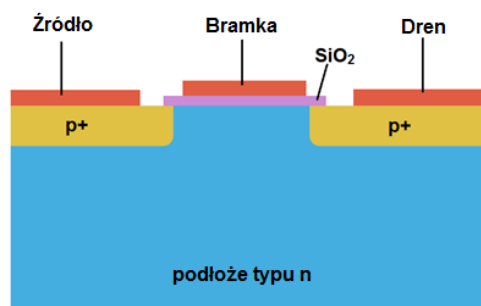
Rys. 3. Defekty w monokryształe węglik krzemu

Czynnikiem, który ogranicza stosowanie węglik krzemu w elektronice jest wysoka gęstość defektów. Najczęstszymi defektami monokryształów (rys. 3) są mikrokanaly oraz dyslokacje śrubowe i krawędziowe. Dzięki procesowi trawienia chemicznego można odkryć dyslokacje oraz policzyć ich gęstości [7].

3. TRANZYSTOR MOSFET – BUDOWA I DZIAŁANIE

Tranzystor MOSFET (ang. *Metal Oxide Semiconductor Field Effect Transistor*) jest to typ tranzystora unipolarnego o strukturze metal-tlenek-półprzewodnik. Istnieją dwie podstawowe kategorie tranzystorów MOSFET – tranzystory z kanałem zubażanym i tranzystory z kanałem wzbogacanym. W ramach pracy zaprojektowany został tranzystor z kanałem wzbogacanym.

Najważniejszą cechą jest odizolowanie elektrody bramki od podłoża za pomocą warstwy dielektryka (zwykle dwutlenku krzemu). Warstwa jest bardzo cienka, rzędu nm, lecz posiada niezwykle dużą rezystancję, dzięki czemu wejściowa rezystancja bramki jest bardzo duża.

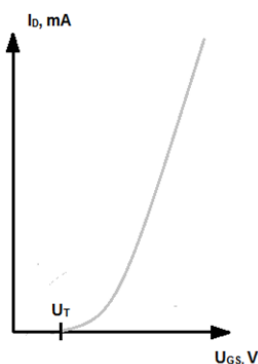


Rys. 4. Tranzystor MOSFET z kanałem wzbogacanym

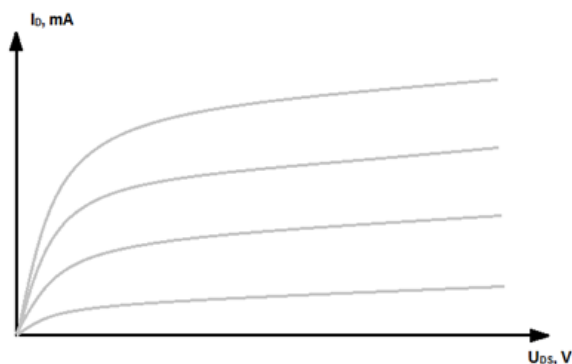
Na rysunku 4 przedstawiono tranzystor o podłożu typu n, w które wprowadzono dwa silnie domieszkowane obszary typu p stanowiące źródło i dren. W przypadku, gdy nie doprowadzone jest żadne napięcie na bramkę, jedno ze złączy działa jako dioda spolaryzowana wstecznie i odcina przepływ prądu. Jeśli do bramki doprowadzi się napięcie (w tym przypadku ujemne), to dziury z obszaru p zostaną przyciągnięte do obszaru tuż pod elektrodą bramki, tworząc chwilowo wąski obszar zwany kanałem. W tej sytuacji, złącze spolaryzowane wstecznie jest bocznikowane i przez indukowany kanał typu p mogą przepływać elektrony między źródłem i drenem. Można też wykonać tranzystor z kanałem typu n, czyli z obszarami typu n wdyfundowanymi w płytkę typu p. W tym przypadku nośnikami ładunku są elektrony i to one są przyciągane pod obszar bramki po wprowadzeniu napięcia dodatniego na elektrodę bramki.

Tranzystor MOSFET posiada dwie podstawowe charakterystyki – wyjściową i przejściową. Charakterystyka przejściowa (rys. 5) jest to zależność prądu drenu (I_D) od napięcia bramka–źródło (U_{GS}) przy stałym napięciu dren–źródło (U_{DS}). Z tej charakterystyki możliwe jest odczytanie napięcia progowego U_T .

Charakterystyka wyjściowa (rys. 6) jest to zależność prądu drenu (I_D) od napięcia dren–źródło (U_{DS}) przy stałym napięciu bramka–źródło (U_{GS}). Obszar charakterystyki wyjściowej można podzielić na dwie części: obszar nasycenia i obszar nienasycenia (liniowy). W zakresie nienasycenia tranzystor MOSFET zachowuje się jak rezystor półprzewodnikowy. Prąd drenu wzrasta liniowo wraz ze wzrostem napięcia dren–źródło. W obszarze nasycenia napięcie U_{DS} nieznacznie wpływa na wartość prądu I_D , a bramka zachowuje właściwości sterujące.



Rys. 5. Przykładowa charakterystyka przejściowa tranzystora MOSFET

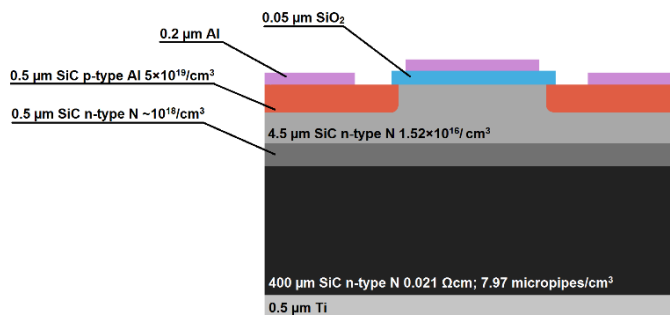


Rys. 6. Przykładowe charakterystyki wyjściowe tranzystora MOSFET

4. MODEL TRANZYSTORA MOSFET

Model tranzystora MOSFET jest oparty na technologii złącza p-n zaprojektowanego w Instytucie Elektroniki i Technik Informacyjnych Politechniki Lubelskiej. Diody te zostały wykonane na płycie o politypii 4H-SiC wyproduk-

wanej przed firmę Cree [9–12]. Płytką składa się z podłoża typu n o rezysytywności $0,021\Omega\text{cm}$, buforowej epitaksjalnej warstwy o grubości $0,5\text{ }\mu\text{m}$ z koncentracją domieszki $1\times 10^{18}\text{cm}^{-3}$ oraz zewnętrznej warstwy epitaksjalnej o grubości $4,5\text{ }\mu\text{m}$ z koncentracją domieszki $1,52\times 10^{16}\text{cm}^{-3}$. Rysunek 7 przedstawia przekrój tranzystora MOSFET zaprojektowanego na opisanym podłożu.



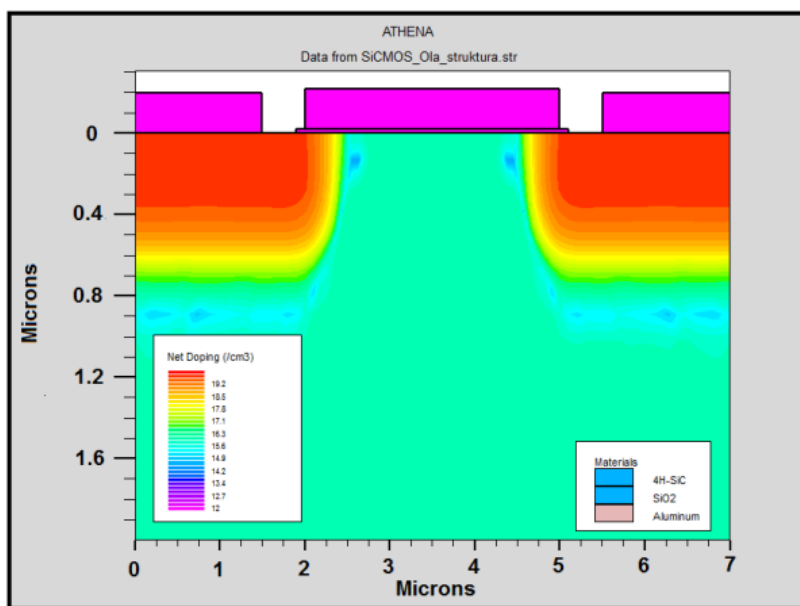
Rys. 7. Schematyczny przekrój projektowanego tranzystora MOSFET

Zadowalające wyniki wykonanych diod p-n pozwoliły na zaprojektowanie tranzystora MOSFET. Ważnym procesem w tym projekcie jest implantacja jonów, czyli wprowadzenie domieszek do materiału w celu utworzenia obszarów o przeciwnym typie domieszkowania niż podłoże. Domieszkowanie może odbywać się również w procesie epitaksji i metodą dyfuzji termicznej. Implantacja pozwala na dokładne określenie lokalizacji i koncentracji domieszki, co jest ciężko wykonalne w przypadku procesu dyfuzji. Dyfuzja termiczna wymaga również stosowania wysokich temperatur powyżej 1500°C , co powoduje znaczną degradację powierzchni. Obszary typu n otrzymuje się najczęściej poprzez implantowanie atomów azotu lub fosforu, a typu p poprzez wprowadzenie akceptorów w postaci boru, glinu oraz galu.

W przypadku implantacji domieszki akceptorowej, użycie boru stwarza bardzo duże problemy, szczególnie na etapie wygrzewania z powodu niskiego stopnia aktywacji ($0,1\%$ dla 1700°C) oraz dużej redyfuzji w trakcie wygrzewania. Korzystniej jest zastosować domieszkę akceptorową w postaci glinu. Spośród trzech możliwych domieszek, glin ma najmniejszą energię jonizacji (około 200 meV). W związku z tym wygodniejsze, choć technicznie trudniejsze jest implantowanie węgla krzemu jonami glinu. Wiąże się to z koniecznością zastosowania źródła jonowego wytwarzającego wiązkę jonów Al^{+} o natężeniu prądu odpowiednio wysokim i stabilnym w czasie [4, 13]. Główną wadą implantacji jest degradacja struktury krystalicznej materiału półprzewodnika. Uszkodzenia mogą sięgać od punktowych defektów spowodowanych pojedynczymi kolizjami przy niskich dozach jonów do całkowitej amorfizacji struktury przy wysokich dozach. W celu odbudowania struktury krystalicznej i aktywacji

zaimplantowanych domieszek poprzez umieszczenie ich w sieci stosuje się po-implantacyjne wygrzewanie.

W projekcie tranzystora zastosowano 4-etapową implantację jonami glinu wykorzystaną przy wytworzeniu diody p-n. Energie implantacji wynosiły kolejno 250keV, 160keV, 100keV i 55keV. Zastosowano wygrzewanie poimplantacyjne w temperaturze 1600°C przez 20 minut w celu aktywowania domieszki. Elektrody drenu, bramki oraz źródła wykonano z aluminium. Podane parametry zostały użyte w programie Athena oprogramowania Silvaco TCAD. Rysunek 8 przedstawia model zaprojektowanego w Silvaco TCAD tranzystora o grubości tlenku bramkowego 20nm i długości kanału 3μm z zaznaczonymi zaimplantowanymi obszarami.



Rys. 8. Przekrój modelu tranzystora MOSFET z zaznaczonymi zaimplantowanymi obszarami

5. WYNIKI SYMULACJI ELEKTRYCZNYCH

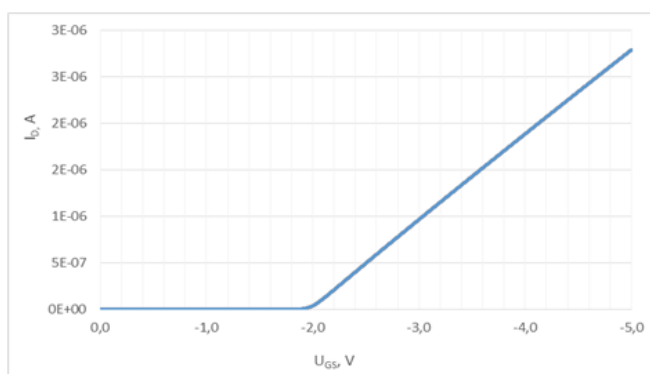
Symulacje elektryczne zostały wykonane w części ATLAS programu Silvaco. ATLAS daje możliwość obserwacji zachowań zaprojektowanego urządzenia w określonych warunkach. Struktura wykonana w części ATHENA została użyta jako struktura wejściowa w części ATLAS.

Charakterystyki prądowo-napięciowe

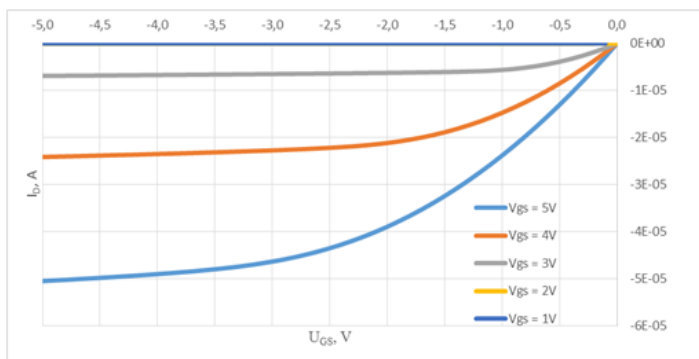
W celu wykonania charakterystyk prądowo-napięciowych określono stałe napięcia dla odpowiednich elektrod. Ustalono napięcie na elektrodzie drenu

o wartości $-0,1\text{V}$, na elektrodzie bramki ustalono kolejno napięcia -1V , -2V , -3V , -4V i -5V .

Następnym etapem było wygenerowanie charakterystyki przejściowej przy stałym napięciu drenu $-0,1\text{V}$ i przy zmianie napięcia na bramce od 0V do -5V ze skokiem $-0,1\text{V}$. Na rysunku 9 przedstawiono charakterystykę przejściową tranzystora o grubości tlenku brankowego 20nm i długości kanału $3\mu\text{m}$. Z podanej charakterystyki można odczytać napięcie progowe, które w podanym przykładzie wynosi około $1,9\text{V}$. Wygenerowano charakterystyki wyjściowe dla omawianego tranzystora (rys. 10) dla stałych wartości napięcia na bramce (-1V , -2V , -3V , -4V , -5V) i przy zmianie prądu drenu w zakresie 0 do -5V z krokiem $0,1\text{V}$.



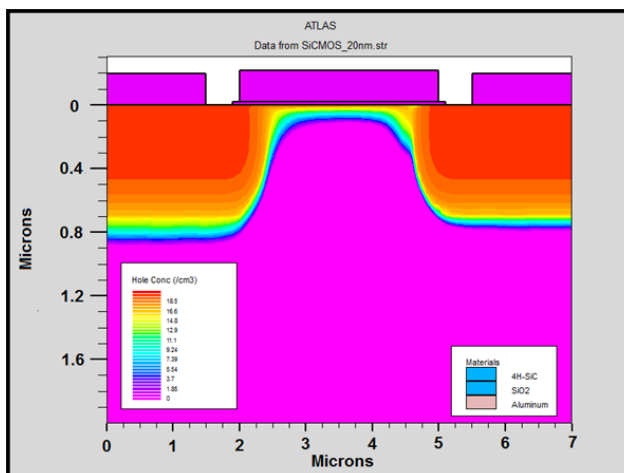
Rys. 9. Charakterystyka przejściowa tranzystora MOSFET o grubości tlenku brankowego 20 nm i długości kanału $3\mu\text{m}$



Rys. 10. Charakterystyki wyjściowe tranzystora MOSFET o grubości tlenku brankowego 20 nm i długości kanału $3\mu\text{m}$

ATLAS daje również możliwość generowania struktur urządzenia i jego zachowań pod wpływem określonych napięć, takich jak: koncentracja dziur i elektronów, rozkład potencjału, rozkład pola elektrycznego i gęstość prądu.

W projektowanym tranzystorze kanał uformowany jest z dziur i jest widoczny na symulacji przedstawionej na rysunku 11.

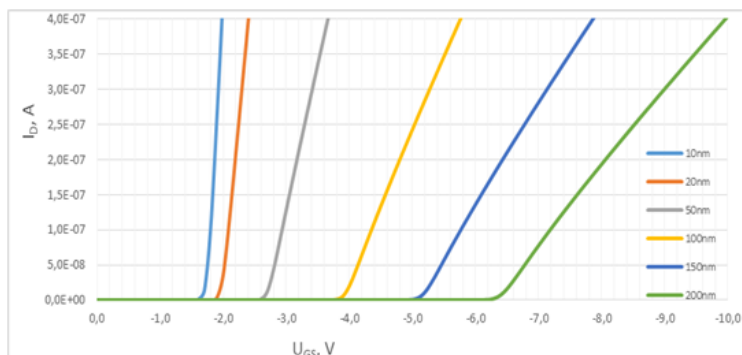


Rys. 11. Rozkład dziur w tranzystorze MOSFET o grubości tlenku ramkowego 20 nm i długości kanału 3 μm przy napięciu -5 V na bramce

Wpływ grubości tlenku bramkowego na napięcie progowe UT

Jednym z podstawowych paramentów technologicznych jest grubość tlenku bramkowego. Wykonano symulacje dla różnych grubości tlenku bramkowego wynoszących 10nm, 20nm, 50nm, 100nm, 150nm i 200nm. Wygenerowano charakterystyki przejściowe dla poszczególnych wartości tlenku bramkowego i przedstawiono je na wspólnym wykresie (rys. 12).

Po porównaniu charakterystyk zaobserwowano, że wraz ze wzrostem grubości tlenku bramkowego wzrasta wartości napięcia progowego.

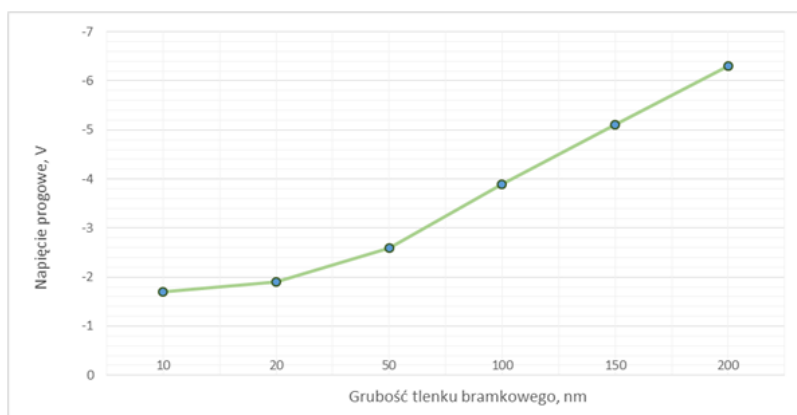


Rys. 12. Zestawienie charakterystyk przejściowych różnych wartości grubości tlenku bramkowego

Otrzymane wyniki przedstawiono w tabeli 2, a zależność pomiędzy napięciem progowym a grubością tlenku bramkowego przedstawiono na rysunku 13.

Tab. 2. Wartości napięcia progowego w zależności od grubości tlenku bramkowego

Grubość tlenku bramkowego, nm	Napięcie progowe, V
10	-1,7
20	-1,9
50	-2,6
100	-3,9
150	-5,1
200	-6,3

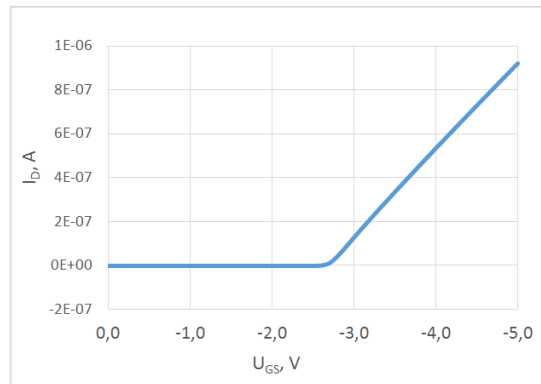


Rys. 13. Wykres zależności wartości napięcia progowego od grubości tlenku bramkowego

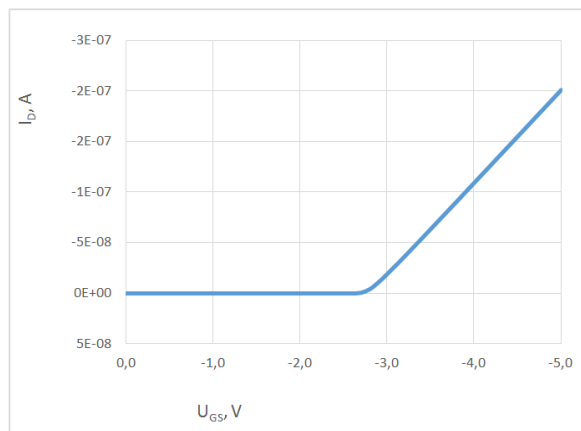
Analiza otrzymanych wyników pozwala na wybranie optymalnej wartości grubości tlenku bramkowego tranzystora MOSFET. Należy jednak wziąć pod uwagę fakt, że dla cienkich warstw tlenku bramkowego istnieje ryzyko jego przebicia, wynikające z jego nierównomiernego rozkładu.

Wpływ modulacji długości kanału.

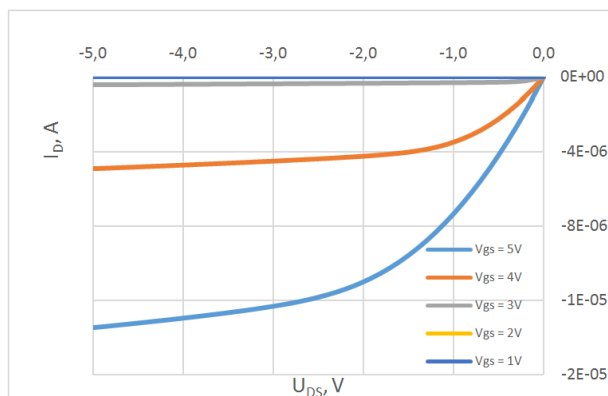
Kolejnym parametrem podlegającym zmianie była długość kanału tranzystora. Symulacje zostały przeprowadzone dla stałej wartości grubości tlenku bramkowego wynoszącej 50nm, długość kanału była zmieniana w zakresie od 3 μm do 10 μm , co 1 μm . Wygenerowano charakterystyki przejściowe i wyjściowe. Na rysunkach 14 i 15 przedstawione zostały charakterystyki przejściowe odpowiednio dla długości kanału 3 μm i 10 μm . Na rysunkach 16 i 17 przedstawione zostały charakterystyki wyjściowe dla długości kanału 3 μm i 10 μm .



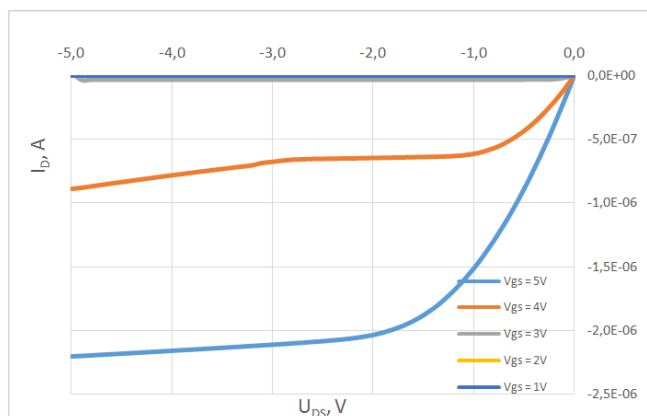
Rys. 14. Charakterystyka przejściowa tranzystora o długości kanału 3 μm



Rys. 15. Charakterystyka przejściowa tranzystora o długości kanału 10 μm



Rys. 16. Charakterystyki wyjściowe tranzystora o długości kanału 3 μm



Rys. 17. Charakterystyki wyjściowe tranzystora o długości kanału $10\mu\text{m}$

Po analizie otrzymanych charakterystyk stwierdzono, że modulacja długości kanału nie ma zauważalnego wpływu na napięcie progowe, lecz ma wpływ na wartość prądu drenu, która zmniejsza się wraz z wydłużeniem się kanału. Przy długości kanału $10\mu\text{m}$, prąd drenu osiągnął wartość czterokrotnie mniejszą niż przy długości kanału $3\mu\text{m}$.

Fakt, że modulacja długości kanału nie ma znaczącego wpływu na napięcie progowe jest bardzo istotny. Oznacza to, że można przyjąć dłuższy kanał, który jest łatwiejszy do wykonania. Na dzień dzisiejszy dostępna dla nas technologia pozwala na produkcję tranzystorów o najmniejszej długości kanału $5\mu\text{m}$ (przy bardzo zwiększonej precyzji możliwe jest zmniejszenie tej wartości do $3\mu\text{m}$), lecz należy uwzględnić potencjalne odchylenia wynikające m.in. z centrowania płytki. Wykonanie tranzystora z dłuższym kanałem eliminuje ten problem.

6. PODSUMOWANIE

W ramach pracy wykonano model tranzystora z węglika krzemu oraz szereg symulacji elektrycznych. Wygenerowane charakterystyki przejściowe i wyjściowe potwierdziły możliwość wykonania tranzystora opierającego się na parametrach technologicznych zastosowanych przy diodzie p-n. Szereg symulacji pozwolił na sprawdzenie wpływu podstawowych parametrów na pracę tranzystora. Określone w obliczeniach numerycznych grubości tlenku bramkowego oraz długości kanału pozwolą na ustalenie parametrów procesów technologicznych do wytworzenia tranzystora MOS z węglika krzemu.

LITERATURA

- [1] Zarębski J., *Tranzystory mocy typu MOS z węgliku krzemu*, „Elektronika” 7/2004
- [2] Casady J.B., Johnson R.W., *Status of Silicon Carbide (SiC) as a wide-bandgap semiconductor for high temperature applications*, Solid-State Electronics, “Elsevier” 1996
- [3] Sadow S. E., *Silicon Carbide Biotechnology*, “Elsevier”, 2012
- [4] Waczyński K., Wróbel E.: *Technologie mikroelektroniczne – metody wytwarzania materiałów i struktur półprzewodnikowych*, Gliwice 2006
- [5] Agarwal S. i inni, *SiC Materials and Devices*, Academic Press, 1988
- [6] Kociubiński A., Duk M., Teklińska D., Kwietniewski N., Sochacki M., Borecki M., *Fabrication and characterization of epitaxial 4H-SiC on junctions*, SPIE, 2014
- [7] Kościwicz K., Tymicki E., Graszka K.: *SiC – materiał dla elektroniki*, „Elektronika”, 9/2006
- [8] Jeleński A., *Półprzewodniki szerokopasmowe dla elektroniki i fotoniki*, KKE 2005
- [9] Kociubiński A., Duk M., Korona M., Muzyka K.: *4H-SiC Photodiode Model for DC SPICE Circuit Simulation*, SPIE 2015
- [10] Kociubiński A., Borecki M., Duk M., Sochacki M., Korwin-Pawłowski M. L., *3D photodetecting structure with adjustable sensitivity ratio in UV–VIS range*, “Microelectron. Eng.”, vol. 154, p. 48–52
- [11] Kociubiński A., Duk M., Masłyk M., Kwietniewski N., Sochacki M., Borecki M., Korwin-Pawłowski M., *Technology and characterization of 4H-SiC p-i-n junctions*, *Photonics Applications in Astronomy, Communications, Industry and High-Energy Physics Experiments* 2013, 89030V
- [12] Borecki M., Kociubiński A., Duk M., Kwietniewski N., Korwin-Pawłowski M. L., Doroz P., Szmidt J., *Large-area transparent in visible range silicon carbide photodiode*, *Photonics Applications in Astronomy, Communications, Industry and High-Energy Physics Experiments* 2013, 89030H
- [13] Turek M., Pysznik K., Prucnal S., Drożdżel A., Żuk J., *Źródło jonów dla potrzeb implantacji jonami Al⁺*, „Elektronika” 9/2009

PROJEKT I TECHNOLOGIA ZŁĄCZA SCHOTTKY'EGO Z WĘGLIKA KRZEMU

1. WSTĘP

Wraz ze współczesnym rozwojem technologii wzrasta też zapotrzebowanie na wysokotemperaturowe i wysokonapięciowe zastosowania urządzeń elektro-nicznych, pracujących dodatkowo przy dużych mocach. Powszechnie stosowana technologia krzemowa okazała się być niewystarczająca ze względu na ograniczenia materiałowe krzemu. W związku z tym poszukiwane są nowe rozwiązania technologiczne dla materiałów półprzewodnikowych o szerokim paśmie zabronionym i wysokim krytycznym natężeniu pola elektrycznego, które pozwolą sprostać wysokim wymaganiom pracy w trudnych warunkach. Do materiałów takich należą m.in. węglik krzemu, fosforek galu, azotek galu, oraz diament. Wśród nich największym zainteresowaniem cieszy się węglik krzemu. Nieustannie trwają badania nad jego właściwościami i możliwymi zastosowaniami, dzięki temu jest najlepiej dopracowany technologicznie wśród wymienionych powyżej materiałów, co pozwala na seryjną produkcję elementów na jego podłożu. Do tej pory największe osiągnięcia dotyczą technologii wytwarzania diod, a przede wszystkim diod Schottky'ego. Ich działanie polega na przepływie nośników większościowych, a co za tym idzie – z uwagi na brak nośników mniejszościowych oraz obecność tzw. gorących nośników – przełączanie urządzeń z barierą Schottky'ego następuje bardzo szybko.

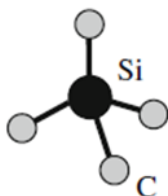
Węglik krzemu jest materiałem problematycznym technologicznie, zwłaszcza w porównaniu z krzemem, ze względu np. na konieczność stosowania dużo wyższych temperatur do jego obróbki, przez co potrzebne są do tego wyspecjalizowane maszyny. Wytworzenie zarówno monokryształu SiC o wymaganej jakości, jak i urządzeń półprzewodnikowych na jego podłożu jest bardzo kosztownym procesem. W celu ograniczenia znacznych wydatków oraz czasu wykorzystuje się symulacje elektryczne pozwalające m.in. przewidzieć rozkład nośników elektrycznych, procesy zachodzące podczas pracy danego urządzenia, jego działanie. Projektowanie i symulacja pozwalają na dokładną analizę oraz ewentualną korektę budowy, czy technologii jeszcze przed wytworzeniem, co znacznie obniża koszty i skraca czas związany z wyprodukowaniem kompletnego urządzenia.

¹ Politechnika Lubelska, Koło Naukowe „SEMICON”

2. WĘGLIK KRZEMU – PARAMETRY I WŁAŚCIWOŚCI

Pomimo, że węgiel krzemu po raz pierwszy został otrzymany w roku 1824, jego potencjał w kategorii materiałów półprzewodnikowych został doceniony dopiero w latach 90. ubiegłego stulecia, wtedy też rozpoczęły się intensywne badania nad jego możliwymi zastosowaniami i właściwościami elektrycznymi [1].

Węgiel krzemu jest związkiem typu IV-IV o budowie krystalicznej tworzonej przez tetraedryczne połączenie czterech atomów węgla i atomu krzemu w centrum silnymi wiązaniami kowalencyjnymi (Rys. 1), co czyni SiC materiałem bardzo stabilnym chemicznie, mechanicznie i termicznie.



Rys. 1. Tetraedryczna budowa komórki elementarnej SiC [2]

Węgiel krzemu wykazuje dwuwymiarową wielopostaciowość zwaną politypizmem. Występuje w ponad 250 strukturach różniących się między sobą ułożeniem kolejnych warstw złożonych z atomów krzemu i węgla [1]. Politypy są pogrupowane w trzy główne kategorie: heksagonalną (H), kubiczną (C) i romboedryczną (R). Każda z nich posiada odmienne właściwości i wartości parametrów elektrycznych, takie jak szerokość pasma zabronionego, czy ruchliwość nośników. Najistotniejsze ze względu na wykorzystanie w praktyce są odmiany 4H-SiC, 6H-SiC oraz 3C-SiC. Struktura 4H-SiC, ze względu na najszerszą przerwę zabronioną (3,2 eV) i znacznie większą ruchliwość elektronów w porównaniu z politypem 6H-SiC, jest najczęściej wykorzystywana do wytwarzania unipolarnych urządzeń elektronicznych.

Możliwości zastosowań węgla krzemu jako podłożowego materiału półprzewodnikowego dla przyrządów, które przetwarzają duże moce i pracują w wysokich temperaturach oraz wysokich częstotliwościach wynikają z jego właściwości, tj.:

- szerokiej przerwy energetycznej
- wysokiej przewodności cieplnej
- wysokiej wartości przebicia pola elektrycznego
- dobrej odporności chemicznej oraz termicznej.

Bardzo istotnym parametrem jest szerokość pasma zabronionego, która w przypadku struktury 4H jest blisko 3 razy większa niż w krzemie. Dzięki temu

węglik krzemu zyskuje nad nim znaczną przewagę pod względem wysokości temperatury i częstotliwości pracy. Urządzenia z niego wykonane są w stanie pracować przy temperaturach, które w teorii dochodzą do 500–600°C [3]. W praktyce ze względu na ograniczenia innych komponentów układu i samego urządzenia (np. odporność obudowy, przewodów) najwyższe dopuszczalne temperatury to ok. 250°C, przy czym zastosowanie krzemu kończy się na około 150°C [4].

Krytyczne pole elektryczne węglika krzemu jest około dziesięciokrotnie większe niż w przypadku krzemu, więc grubość urządzeń z niego wykonanych może być dziesięciokrotnie mniejsza bez ryzyka przebicia. W połączeniu z właściwościami płynącymi z trzykrotnie większego współczynnika przewodności cieplnej (mniej skomplikowane i mniejsze systemy chłodzenia), urządzenia wykonane na podłożu z węglika krzemu mogą być mniejsze i generować mniej strat [5].

Ze względu na szerokość pasma zabronionego, SiC jest odpowiednim materiałem dla wysokich mocy niebieskich diod i detektorów promieniowania UV. W tym celu stosuje się m.in. diody Schottky'ego i diody p-i-n. W połączeniu ze wspomnianymi już zaletami płynącymi z właściwości materiałowych, detektory UV z węglika krzemu nadają się do pracy w warunkach ekstremalnych, jak medycyna lub technika nuklearna, jako urządzenia służące analizom chemicznym i biologicznym, detekcji płomienia oraz wspomnianym już zastosowaniom optoelektronicznym [6–10].

Problemem w stosowaniu węglika krzemu jest opanowanie technologii na poziomie umożliwiającym produkcję seryjną. Obecnie najlepiej dopracowanym przyrządem jest dioda Schottky'ego. Pomimo problemów, węglik krzemu pozostaje jednym z najbardziej obiecujących materiałów w swojej kategorii.

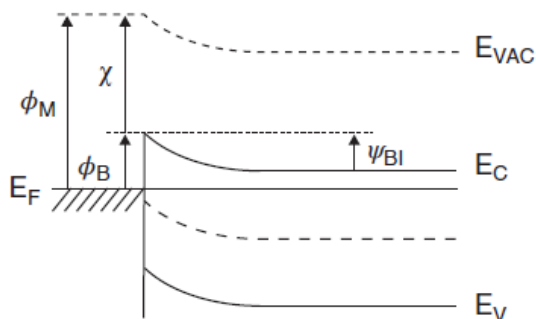
3. DIODA SCHOTTKY'EGO Z WĘGLIKA KRZEMU

Diody z barierą Schottky'ego (SBD, ang. *Schottky Barrier Diode*) dzięki prostocie swojej budowy są atrakcyjne jako urządzenia służące do badań nad nowymi, trudnymi materiałami podłożowymi, jakim jest węglik krzemu.

Kontakt Schottky'ego jest to złącze metal-półprzewodnik (*m-s*). Podłoże składa się z cienkiej, słabo domieszkowanej warstwy epitaksjalnej na mocno domieszkowanym półprzewodniku o tym samym typie przewodnictwa. Na warstwie epitaksjalnej wykonywane są złącza Schottky'ego, a od strony podłoża kontakty omowe.

Bardzo ważną cechą diody jest duża szybkość jej przełączania, która wynika z braku nośników mniejszościowych oraz obecności tzw. nośników gorących. Diody z barierą Schottky'ego wykonane z węglika krzemu najczęściej są wytwarzane na podłożu o politypii 4H typu n w związku z większą ruchliwością elektronów w porównaniu do 6H-SiC [11].

W złączu Schottky'ego funkcja pracy umieszcza poziom Fermiego metalu blisko środka pracy półprzewodnika (rys. 2). Tworzy to barierę dla wstrzykiwania nośników z metalu do pasma półprzewodnika, więc napięcie może przepływać tylko nośnikami większościowymi z półprzewodnika do metalu. Zapewnia to jednokierunkowość prądu i wytwarza prostujący związek prąd – napięcie. Domieszkowanie półprzewodnika w takim złączu nie może być zbyt wysokie, w przeciwnym razie jego pasmo zubożenia będzie zbyt małe i dojdzie do efektu tunelowania elektronów pomiędzy pasmem nośników większościowych półprzewodnika i stanami na tym samym poziomie energetycznym w metalu prowadząc do powstania kontaktu nie prostującego, lecz kontaktu omowego [12].



Rys. 2. Model pasmowy diody Schottky'ego o podłożu typu n [12]. E_{VAC} – poziom próżni; E_C – dolna krawędź pasma przewodnictwa półprzewodnika; E_V – górna krawędź pasma walencyjnego półprzewodnika; E_F – poziom Fermiego metalu; ϕ_M – praca wyjścia metalu; ϕ_B – wysokość bariery Schottky'ego; χ – powinowactwo elektronowe półprzewodnika; ψ_{BI} – potencjał wbudowany złącza

Ze względu na szeroką przerwę zabronioną węgla krzemu przy wytwarzaniu złączy z barierą Schottky'ego stosuje się metale (pojedynczo lub w kombinacjach) o stosunkowo dużej wartości pracy wyjścia elektronów. Ich przykłady przedstawiono w tabeli 1.

Tab. 1. Wartości prac wyjścia przykładowych metali stosowanych w diodach Schottky'ego [13]

Metal	Praca wyjścia, eV
tytan (Ti)	4,6
nikiel (Ni)	4,9
platyna (Pt)	5,3
złoto (Au)	5,4
iryd (Ir)	5,6

3.1. Model i symulacja diod Schottky'ego z węglika krzemu

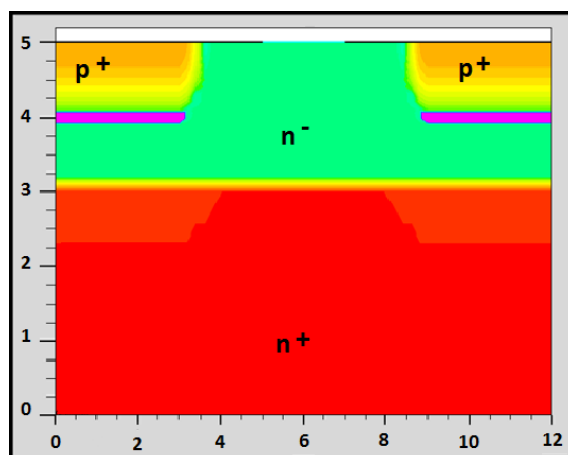
Symulacje elektroniczne mają na celu ułatwienie zaprojektowania urządzenia obniżając koszty i skracając czas jego wytworzenia poprzez wyliczanie wartości np. prądu i napięcia świadczących o jakości pracy i reakcji urządzenia na zmienne warunki napięciowe, również te nietypowe – problematyczne do zmierzenia w rzeczywistości, dając szansę na korektę założeń i planowanych procesów.

W celu przeprowadzenia symulacji diody z barierą Schottky'ego użyto oprogramowania Silvaco TCAD. Wykorzystano narzędzie ATLAS, które jest w stanie przewidzieć działanie półprzewodnikowych urządzeń elektronicznych, rozkład ich domieszek oraz reakcji w nich zachodzących podczas pracy również w nietypowych warunkach napięciowych.

Przygotowane zostały modele dwóch najczęściej stosowanych konstrukcji diod Schottky'ego: struktura z pierścieniem ochronnym oraz struktura planarna. Do otrzymania złącza m-s wykorzystano cztery metale (tytan, nikiel, platyna, iryd), o wartościach pracy wyjścia zgodnych z danymi w tabeli 1, po czym porównano ich wpływ na właściwości diody.

Pierwszym zaprojektowanym i zasymulowanym modelem był model złącza na podłożu typu n z pierścieniem ochronnym w postaci obszaru domieszkowanego typu p (rys. 3). Jego przekrój składa się z następujących elementów:

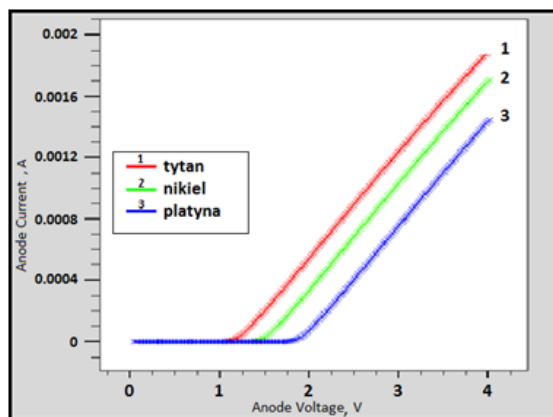
- 1) podłoże typu n^+ 4H-SiC o koncentracji domieszki $1 \times 10^{20} \text{ cm}^{-3}$;
- 2) warstwa epitaksjalna typu n^- 4H-SiC o koncentracji domieszki $1 \times 10^{16} \text{ cm}^{-3}$;
- 3) dwa, symetryczne obszary p^+ o koncentracji domieszki $1 \times 10^{19} \text{ cm}^{-3}$;
- 4) warstwa metalu o wysokiej pracy wyjścia (tytan, nikiel, platyna).



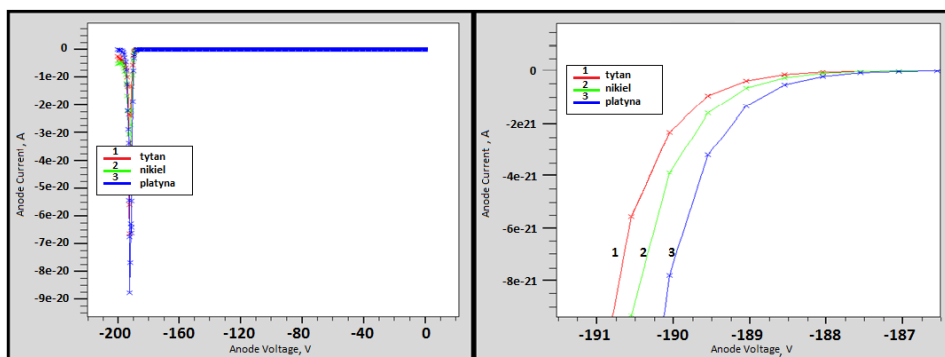
Rys. 3. Rozkład domieszek w zaprojektowanej diodzie Schottky'ego z ochronnym pierścieniem p^+

Wygenerowane zostały charakterystyki w kierunku przewodzenia oraz w kierunku zaporowym (Rys. 4-5). Przeprowadzone symulacje prezentują

zależność parametrów diody od zastosowanego metalu w złączu m-s z barierą Schottky'ego. Wyższa wartość pracy wyjścia metalu powoduje wzrost wartości napięcia przewodzenia U_f oraz spadek wartości napięcia przebicia U_b . Nie wykazano wpływu na wartość prądu wstecznego diody I_r . Wartości parametrów zostały przedstawione w tabeli 2.



Rys. 4. Zestawienie charakterystyk I – U w kierunku przewodzenia dla modelu diody z pierścieniem p^+ dla trzech wariantów metalizacji anody



Rys. 5. Zestawienie charakterystyk I – U w kierunku zaporowym dla modelu diody z pierścieniem p^+ dla trzech wariantów metalizacji anody

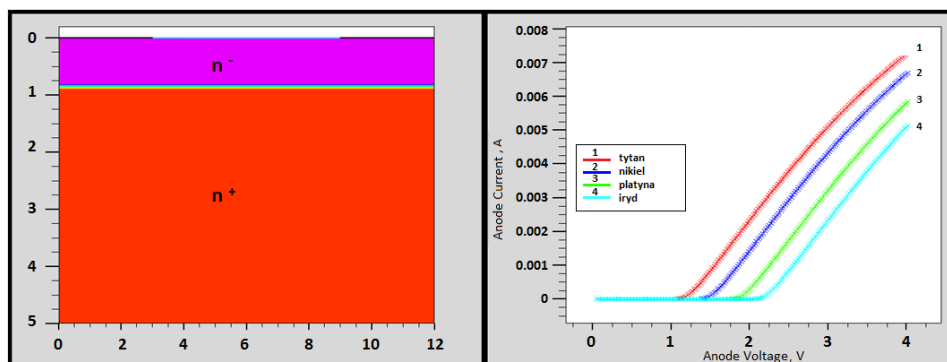
Tab. 2. Wartości napięcia przewodzenia oraz napięcia przebicia dla modelu diody z pierścieniem p^+ w zależności od zastosowanego metalu

Metal	Praca wyjścia eV	U_f V	U_b V
tytan (Ti)	4.6	1.25	-189.9
nikiel (Ni)	4.9	1.5	-189.6
platyna (Pt)	5.3	1.95	-189.2

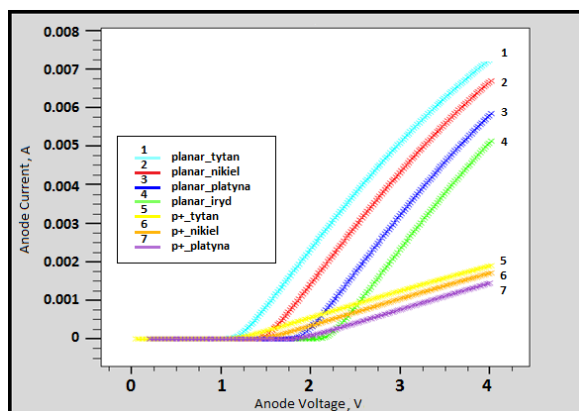
Planarna konstrukcja modelu złącza m-s z barierą Schottky'ego zawierała następujące elementy:

- 1) podłoże typu n^+ 4H-SiC o koncentracji domieszki $1 \times 10^{20} \text{ cm}^{-3}$;
- 2) warstwa epitaksjalna typu n^- 4H-SiC o koncentracji domieszki $1 \times 10^{18} \text{ cm}^{-3}$;
- 3) warstwa epitaksjalna typu n^- 4H-SiC o koncentracji domieszki $1 \times 10^{16} \text{ cm}^{-3}$;
- 4) warstwa metalu o wysokiej pracy wyjścia (tytan, nikiel, platyna, iryd).

Wykazano analogiczną zależność parametrów diody od pracy wyjścia metalu, jak w przypadku konstrukcji z pierścieniem ochronnym p^+ (Rys. 6). Ponadto, na rysunku 7 przedstawiono porównanie charakterystyk $I(U)$ w kierunku przewodzenia dwóch rozważanych konstrukcji struktury. Nie wykazano różnic pomiędzy wartościami napięcia przewodzenia U_f w obu przypadkach, natomiast zależność prądowo-napięciowa dla struktury planarnej okazała się wyraźnie silniejsza.



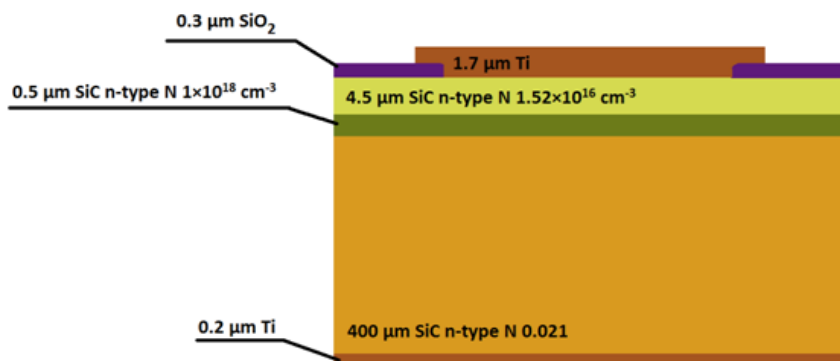
Rys. 6. Rozkład domieszek oraz zestawienie charakterystyk $I-U$ w kierunku przewodzenia modelu diody planarnej dla czterech wariantów metalizacji anody



Rys. 7. Porównanie charakterystyk $I-U$ w kierunku przewodzenia diody planarnej oraz diody z pierścieniem p^+ dla różnych metali anody

3.2. Charakteryzacja diod Schottky'ego z węgla krzemu

Badaniom poddane zostały diody z barierą Schottky'ego o konstrukcji planarnej. Złącza wytworzono wykorzystując 3-calowe płytki węgla krzemu o politypii 4H wyprodukowane przez firmę Cree [14] zawierające podłoże typu n o rezystywności równej $0,021\Omega\text{cm}$, warstwę epitaksjalną typu n o grubości $0,5\mu\text{m}$ i domieszkowaniu na poziomie $1\times 10^{18}\text{cm}^{-3}$ oraz warstwę epitaksjalną typu n o grubości $4,5\mu\text{m}$ i koncentracji domieszki $1,52\times 10^{16}\text{cm}^{-3}$. Pierwiastkiem pełniącym rolę domieszki w podłożu był fosfor. Procesy technologiczne przy wytwarzaniu omawianych złączy zostały przeprowadzone podczas prac nad wytworzeniem złączy p-n z węgla krzemu [15–18]. Schematyczny przekrój wytworzonych diod Schottky'ego przedstawiony został na rysunku 8. W przypadku obu kontaktów (anody i katody) do wykonania metalizacji wykorzystano tytan.

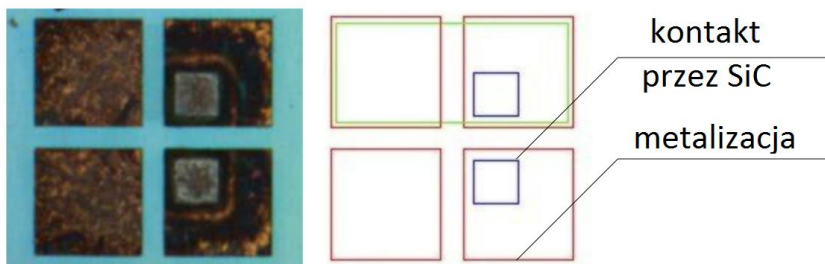


Rys. 8. Schematyczny przekrój wytworzonych diod z barierą Schottky'ego z węgla krzemu

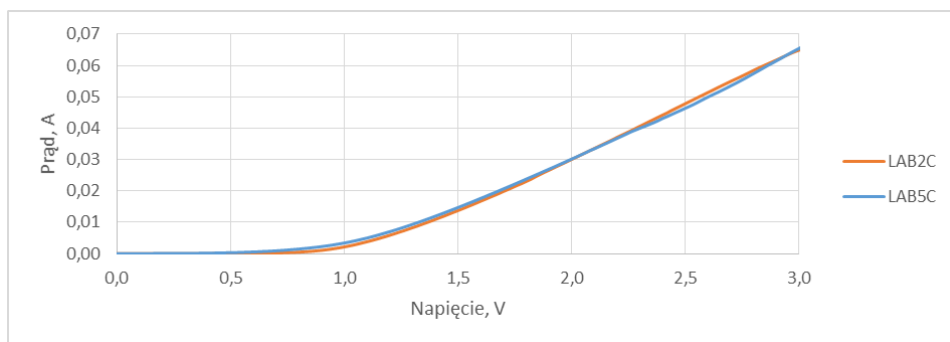
Zaprojektowano i wytworzono kwadratowe struktury: LA_B2C ($80\mu\text{m} \times 80\mu\text{m}$) (rys. 9) oraz LA_B5C ($140\mu\text{m} \times 140\mu\text{m}$). Pomiaru zostały wykonane w temperaturze pokojowej, w ciemności, przy użyciu urządzenia pomiarowego KEITHLEY SMU 236 w zakresie 0–3 V w kierunku przewodzenia oraz 0–10 V w kierunku zaporowym (rys. 10–12).

Na podstawie wyznaczonych charakterystyk prądowo-napięciowych stwierdzono typowe działanie diod oraz stabilność ich pracy w zakresie pomiarowym. Napięcie przewodzenia struktur wynosi w przybliżeniu 1 V, a więc jest zgodne z wartościami oczekiwanymi na podstawie przeprowadzonych symulacji elektrycznych dla złączy z barierą Schottky'ego o konstrukcji planarnej z metalizacją tytanową wykonanego z węgla krzemu. Pomiaru przy polaryzacji zaporowej diody zostały wykonane zaledwie do -10 V, ponieważ tylko niewielkie wartości napięć są odpowiednie dla planowanego zastosowania diod, które będą służyć jako urządzenia wykrywające promieniowanie ultrafioletowe [17]. Wartość prą-

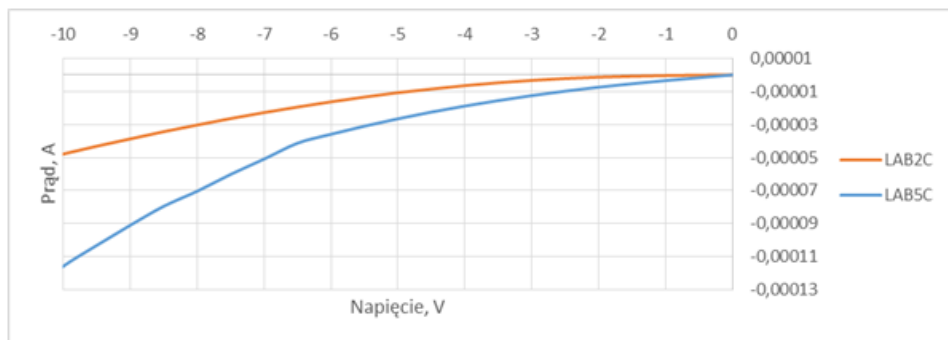
du wstecznego (ciemnego) okazała się być zbyt wysoka i wyniosła około 10^{-4} A dla -10 V.



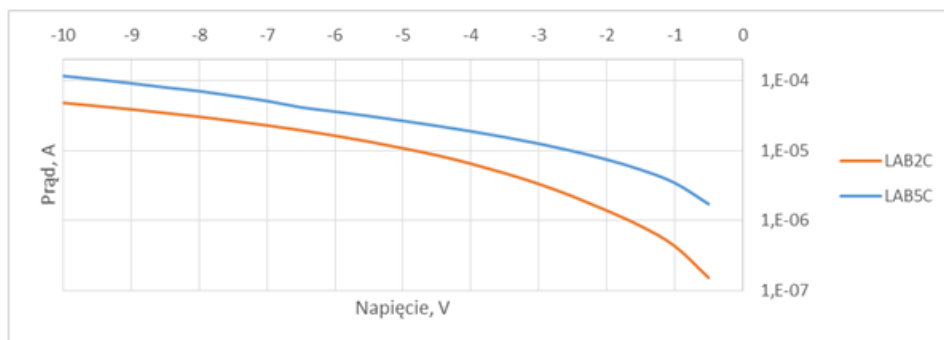
Rys. 9. Widok z góry wytworzonej oraz zaprojektowanej diody Schottky'ego (LA_B2C, $80\ \mu\text{m} \times 80\ \mu\text{m}$)



Rys. 10. Charakterystyki I-U w kierunku przewodzenia wytworzonych struktur



Rys. 11. Charakterystyki I-U w kierunku zaporowym wytworzonych struktur



Rys. 12. Charakterystyki półlogarytmiczne I-U w kierunku zaporowym wytworzonych struktur

4. PODSUMOWANIE

Węglik krzemu ze względu na swoje właściwości posiada duży potencjał jako materiał podłożowy urządzeń elektronicznych. Przeszkodą jest technologia, która pomimo wielu podobieństw do dobrze opanowanej technologii krzemowej, jest od niej bardziej złożona i problematyczna. Znacznym ułatwieniem może być stosowanie symulacji elektrycznych podczas projektowania struktur, których analiza pozwala na ograniczenie kosztów i czasu wytwarzania urządzeń z węglika krzemu.

W związku z planowanym zastosowaniem wytworzonych diod z barierą Schottky'ego jako detektorów promieniowania UV, najważniejszym parametrem struktur jest ich prąd ciemny osiągany przy małych wartościach napięć w polaryzacji zaporowej diody. Opierając się na wynikach przeprowadzonych symulacji elektrycznych, wartość prądu otrzymanych struktur jest zbyt wysoka. Może być to rezultat obecności defektów na powierzchni półprzewodnika, a ich wpływ może zostać zminimalizowany poprzez modyfikację wybranych procesów technologicznych, np. podczas przygotowania powierzchni materiału podłożowego.

LITERATURA

- [1] Kościwicz K., Tymicki E., Graszka K., SiC – materiał dla elektroniki. „Elektronika”, 22–27, 9/2006.
- [2] Fan, J. Chu P. K., *Silicon Carbide Nanostructures. Fabrication, Structure and Properties*, Springer, Switzerland (2014).
- [3] Barlik R., Rąbkowski J., Nowak M., *Przyrządy półprzewodnikowe z węglika krzemu (SiC) i ich zastosowania w energoelektronice*. „Przegląd Elektrotechniczny”, 11/2006, s. 1–8

- [4] Janke W., Węglik krzemu w energoelektronice – nadzieje i ograniczenia. „Przegląd Elektrotechniczny”, 11/2011, s. 41–46
- [5] Yonn S. P., *SiC Materials and Devices*, vol. 52, Academic Press, 1998.
- [6] Borecki M., Doroz P., Szmidt J., Korwin-Pawłowski M.L., Kociubiński A., Duk M., *Sensing Method and Fiber Optic Capillary Sensor for Testing the Quality of Biodiesel Fuel*. The Fourth International Conference on Sensor Device Technologies and Applications, 25–31 Sierpnia 2013, Barcelona, Hiszpania, 19–24, 2013.
- [7] Borecki M., Korwin-Pawłowski M.L., *Optical capillary sensors for intelligent classification of microfluidic samples*. Pod red. Teik-Cheng Lim: *Nanosensors - Theory and applications in industry, healthcare and defense*. CRC Press Boca Raton FL., 215–245, 2011.
- [8] Borecki M., Korwin-Pawłowski M.L., Beblowska M., Szmidt J., Jakubowski J., *Optoelectronic Capillary Sensors in Microfluidic and Point-of-Care Instrumentation*. “Sensors” 10, 2010, p. 3771–3797
- [9] Kociubiński A., Duk M., Korona M., Muzyka K., *4H-SiC Photodiode Model for DC SPICE Circuit Simulation*. Proc. of SPIE 9662, 96620V-1–96620V-7, 2015
- [10] Kociubiński A., Duk M., Teklińska D., Kwietniewski N., Sochacki M., Borecki.: *Fabrication and characterization of epitaxial 4H-SiC pn junctions*. Proc. Of SPIE 9228, 922804-1 – 922804-6, 2014.
- [11] Baliga J., *Advanced power rectifier concepts*. Springer, 2009.
- [12] Kimoto T., Cooper J. A., *Fundamentals of Silicon Carbide Technology. Growth Characterization Devices and Applications*. Singapur, John Wiley & Sons, 2014.
- [13] Bieniek T., Stęszewski J., Sochacki M., Szmidt J., *Symulacje elektryczne diod Schottky'ego oraz tranzystorów RESURF JFET i RESURF MOSFET na podłożach z węglika krzemu (SiC)*. „Elektronika”, 11–15, 7–8/2008.
- [14] www.cree.com [dostęp dnia 27.04.2017].
- [15] Kociubiński A., Duk M., Korona M., Muzyka K., *4H-SiC Photodiode Model fro DC SPICE Circuit Simulation*, Proc. of SPIE 9662, 96620V-1–96620V-7, 2015.
- [16] Kociubiński A., Duk M., Teklińska D., Kwietniewski N., Sochacki M., Borecki M., *“Fabrication and characterization of epitaxial 4H-SiC pn junctions*, Proc. of SPIE 9228, 922804-1–922804-6, 2014.
- [17] Kociubiński A., Duk M., Masłyk M., Kwietniewski N., Sochacki M., Borecki M., Korwin-Pawłowski M.L., *Technology and characterization of 4H-SiC p-i-n junctions*, Proc. of SPIE 8903, 89030V-1–89030V-6, 2013.
- [18] Kociubiński A., Borecki M., Duk M., Sochacki M., Korwin-Pawłowski M.L., *3D photodetecting structure with adjustable sensitivity ratio in UV–VIS range*, *Microelectron. Eng.*, vol. 154, p. 48–52, 2016.

PROJEKT I TECHNOLOGIA KONDENSATORÓW GRZEBIENIOWYCH Z MIEDZI

1. WSTĘP

Ostatnie lata przyniosły szczególnie widoczne zmiany w zakresie pomiarowych aplikacji medycznych. Wynika to między innymi z rozwoju innych dziedzin, w tym mikroelektroniki umożliwiającej budowę wysoce wyspecjalizowanych, zminiaturyzowanych systemów elektronicznych. Rosnąca świadomość z korzyści wynikający z zastosowania medycyny wymusza przeprowadzanie badań na żywych organizmach. Konieczność zrozumienia zmian zachodzących poziomie molekularnym komórki pod wpływem zjawisk fizjologicznych, jak również związków farmakologicznych i toksykologicznych, wpłynęła na rozwój nowoczesnych technik badawczych, takich jak np. testy oparte na pomiarze impedancji.

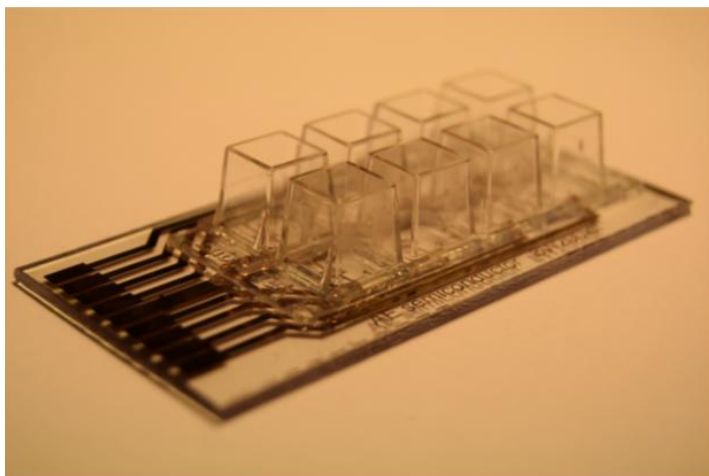
Niniejsza publikacja została poświęca zagadnieniom związanym z technologią struktur do zaawansowanego systemu pomiarowego ECIS[®] (ang. *Electric cell-substrate impedance sensing*) umożliwiającego obserwowanie aktywności badanych komórek za pomocą analizy zmieniających się parametrów elektrycznych monitorowanych w czasie rzeczywistym.

2. SYSTEM POMIAROWY ECIS[®]

System do pomiaru impedancji komórek w warunkach fizjologicznych w czasie rzeczywistym ECIS[®] jest techniką pozwalającą na badanie aktywności komórek oraz zmian w ich strukturze i morfologii. Metoda ta została zastosowana po raz pierwszy w 1984 roku przez dr. Ivara Giaever oraz dr. Charlesa R. Keese konkurując z metodami badań elektryczności komórek z wykorzystaniem mikroskopów [1, 2, 5].

Technika pomiaru impedancji w czasie rzeczywistym umożliwia analizę cyklu życia komórki, od adhezji, poprzez wzrost do uzyskania konfluentnej hodowli, aż do obumarcia komórki. Dzięki możliwości zbadania rekacji komórki na różne stymulanty metoda ECIS[®] jest wykorzystywana w obszarze nowoczesnych badaniach medycznych prowadzących do stworzenia nowych leków pozwalając na zastąpienie testów przeprowadzanych na zwierzętach.

¹ Politechnika Lubelska, Koło Naukowe „SEMICON”



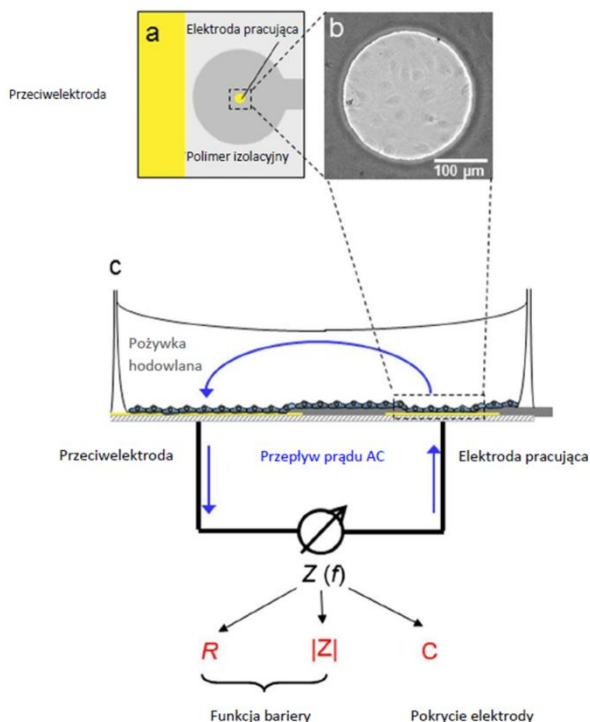
Rys. 1. Matryca pomiarowa dla systemu ECIS®

2.1. Budowa i zasada działania systemu pomiarowego ECIS®

System pomiaru impedancji komórek w czasie rzeczywistym zbudowany jest w oparciu o układ elektrod z cienkiej warstwy metalizacji osadzonych na biokompatybilnym podłożu (poliwęglan, Lexan™ lub PET). Hodowla komórkowa wraz z niezbędną pożywką nanoszona jest bezpośrednio na powierzchnię elektrod, a polimer izolujący ogranicza obszar badania do określonej geometrii [1]. Konstrukcja najbardziej podstawowego układu pomiarowego składa się z pojedynczych okrągłych elektrod o średnicy $250\mu\text{m}$ oraz współpracujących z nimi elektrod przeciwnych. Badana hodowla komórek rozwijająca się na powierzchni układu wpływa na przepływ prądu pomiędzy elektrodami poprzez izolacyjną właściwość błon komórkowych. Każda zmiana w strukturze komórek ma znaczący wpływ wyniki pomiarów [1].

2.2. Pomiar impedancji komórek

Elektrodę badawczą, na której osadzono jedynie pożywkę hodowlaną bez żadnych komórek, można w przybliżeniu opisać jako szeregowo połączony układ kondensatora i rezystora. Rezystancja omowa systemu określana jest przez efekt zwężenia elektrody, jak również przez jonowy ładunek badanego roztworu. Polaryzacja powierzchni elektrody z naniesionym roztworem jonowym jakim jest pożywka hodowlana, która jest zależna od częstotliwości, prowadzi do czasowego gromadzenia się jonów na granicznej międzyfazowej powierzchni elektrodowo-elektrolitycznej. Efektem powyższego zjawiska jest powstanie podwójnej warstwy elektrycznej zbliżonej elektrycznie do kondensatora [1, 6].

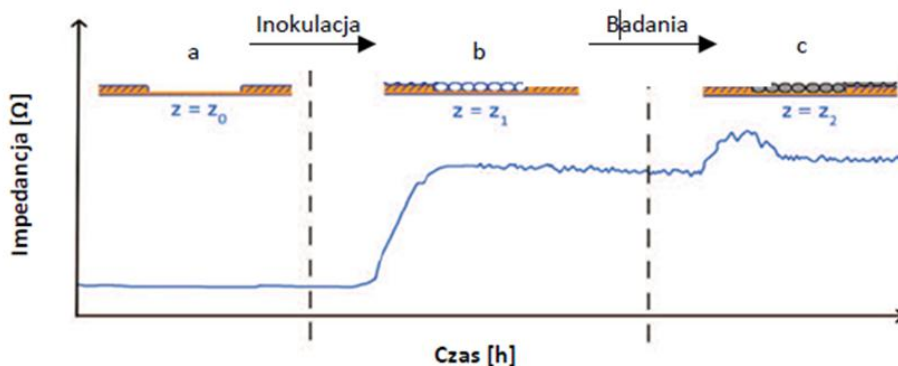


Rys. 2. Ideowy obwód elektryczny systemu pomiarowego ECIS[®]: a) rzut z góry pojedynczej elektrody, b) elektroda pokryta komórkami, c) zasada pomiaru [1]

Naniesienie komórek na elektrody badawcze na matrycy ECIS[®] powoduje w pewnym zakresie częstotliwości (od 100Hz do 100kHz) wzrost impedancji. Rosnące na powierzchni elektrody komórki stają się izolatorami, których rezystancja jest zależna od morfologii, adhezji oraz liczby komórek na elektrodzie. Ideowa charakterystyka impedancji w funkcji czasu pozwala na analizę wpływu zmian morfologicznych na zmianę impedancji. Elektrode badawczą, na której osadzono komórki można w przybliżeniu opisać układem równoważnym do układu bez komórek, z dodatkowo dołączoną impedancją powstałą z warstwy komórkowej [1–4].

Wyszczególnić można 3 charakterystyczne stany:

- a) niepokryta komórkami elektroda – stała wartość impedancji;
- b) inokulacja komórek na powierzchnię elektrody, wzrost komórek do momentu uzyskania konfluentnej hodowli – wzrost impedancji aż do momentu uzyskania odpowiedniego ustabilizowanego poziomu;
- c) badania wpływu różnych stymulant na hodowlę komórek – monitorowanie zmian impedancji w stosunku do poziomu konfluentnej hodowli.



Rys. 3. Charakterystyka impedancji badanej hodowli komórek w funkcji czasu [2]

Komórki hodowlane naniesione na matrycę ECIS[®] wpływają na przepływ prądu w elektrodzie. Zmianie ulega impedancja na granicy elektrody i roztworu badanego, której charakter zależy od częstotliwości prądu użytego w ramach badań. W badaniu większości komórek, wzrost impedancji komórek w stosunku do badań nad samą pożywką hodowlaną odnotowywany jest w szerokim spektrum częstotliwości od kilkuset do kilkuset tysięcy Hz. Niemniej jednak główny zakres częstotliwości używany w badaniach impedancji to od 10kHz do 40kHz [1, 2, 7].

Technologia do pomiaru impedancji komórek w warunkach fizjologicznych w czasie rzeczywistym ECIS[®] oferuje możliwość monitorowania takich parametrów cyklu życia komórki jak:

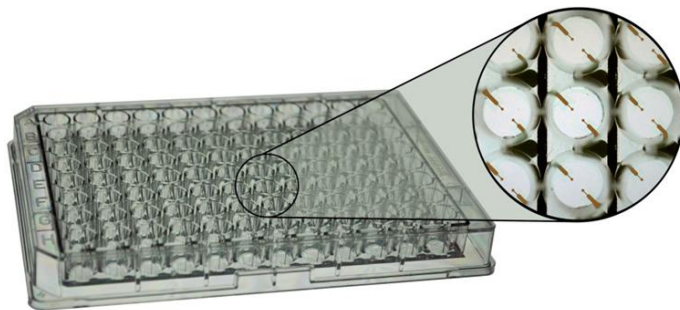
- adhezja i rozprzestrzenianie się
- proliferacja i ruchliwość
- zróżnicowanie biologii wzrostu
- funkcje bariery
- cytotoxyczność.

3. MATRYCE POMIAROWE I ELEKTRODY

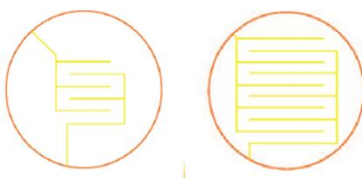
W badaniach pomiaru impedancji komórek w czasie rzeczywistym z zastosowaniem systemu ECIS[®], hodowla komórkowa wraz z niezbędną pożywką nanoszona jest na powierzchnię elektrod. Są one wykonane z cienkiej warstwy metalizacji osadzonej na podłożu z poliwęglanu (LexanTM lub PET) lub włókna szklanego. W praktyce badania przeprowadza się na matrycach pomiarowych, które są standaryzowanymi układami zawierającymi wiele elektrod. Do podłoża doklejone jest specjalne uformowane studnie wykonane z polistyrenu, które odseparowuje pojedyncze elektrody tworząc przestrzeń dla badanej

hodowli. Można wyróżnić dwa typy matryc, na których znajduje się 8 lub 96 elektrod danego typu. W systemie ECIS[®] zastosowanie znalazło wiele typów elektrod różniących się między sobą kształtem oraz powierzchnią metalizacji. Typ elektrody jej kształt oraz powierzchnia, wpływa na czułość pomiarów parametrów elektrycznych badanych z zastosowaniem technologii ECIS[®]. Małe elektrody charakteryzują się wyższą czułością wobec zmian zachodzących pomiędzy elektrodą i warstwą komórek.

a)

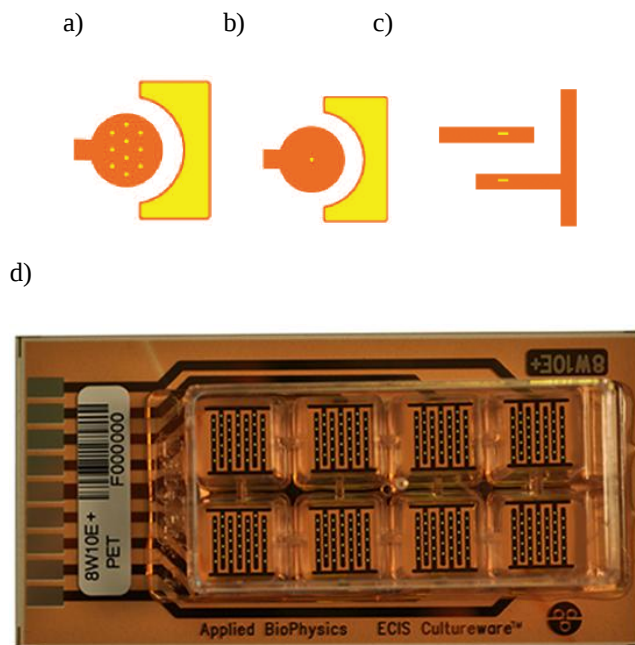


b)



Rys. 4. a) Matryca 96W1E, b) schemat elektrody [2]

Zaletą ich stosowania jest umożliwienie badania pojedynczych komórek oraz możliwość mikroskopowego obserwowania komórek odpowiedzialnych za zmiany impedancji. Zastosowanie małych elektrod pozwala wytworzenia dużego pola elektrycznego przez prąd przemienny niskiego napięcia, niezbędnego w przeprowadzeniu elektroporacji lub zniszczenia komórki podczas eksperymentu. Jednakże, wraz ze zmniejszeniem wielkości elektrody wielkość populacji komórek, z której dokonywany jest pomiar, zmniejsza się, zapewniając mniej statystyczne pokrycie badanych komórek. Ponadto, przy bardzo małych elektrodach występują problemy związane z impedancją wynikającą z pojemności w obrębie ścieżek łączących elektrody z wyprowadzeniami do połączenia z elektroniką ECIS[®]. Wpływ na wyniki pomiarów elektrycznych badanej hodowli ma również liczba pojedynczych elektrod w danej studni pomiarowej [1, 2].



Rys. 5. Schemat elektrody a) 8W10E, b) 8W10E, c) 8W2LE, kolor żółty – złoto, kolor pomarańczowy – warstwa izolacyjna, d) matryca 8W10E+ [2]

4. METALIZACJA

W technologii półprzewodnikowej wykorzystuje się wiele materiałów takich, jak miedź, tytan, złoto lub aluminium. Do wykonania połączeń tworzących układy elektroniczne takie jak elektrody wykorzystywane w technologii ECIS[®] konieczne jest zastosowanie surowców cechujących się bardzo dobrym przewodnictwem prądu. Właściwość ta jest niezbędna do uzyskania miarodajnych wyników badań pozbawionych błędów wynikających ze strat związanych z przekazaniem sygnałów elektrycznych [8, 9].

Właściwości materiałów które są wykorzystywane w technologii półprzewodnikowej powinny spełniać następujące wymagania:

- niska rezystancja,
- duża szybkością układu,
- łatwy do formowania,
- łatwe trawienie do generowania wzorów (litografia),
- stabilny w otoczeniu utleniającym,
- stabilność mechaniczna – dobra adhezja i małe naprężenia,
- gładka powierzchnia po naniesieniu,

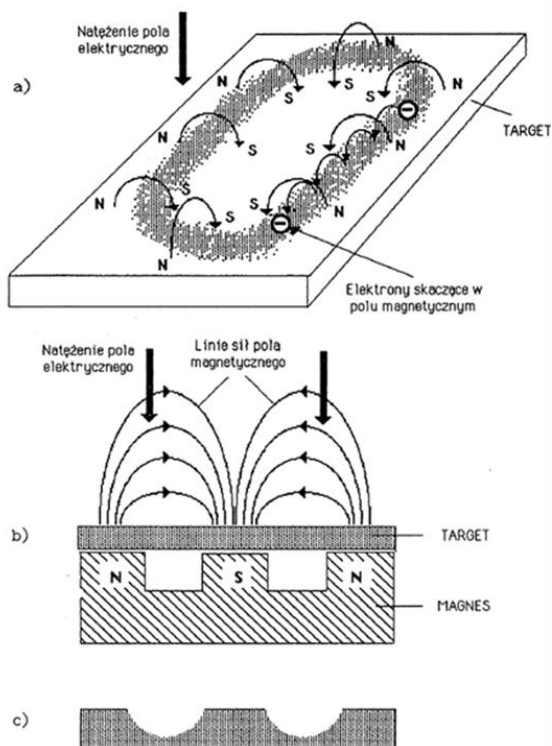
- stabilność podczas procesów technologicznych, w tym wygrzewania w wysokiej temperaturze, utlenianie suchego i mokrego, getterowania (usuwania gazów), pasywacji,
- nie reagowanie chemiczne z metalem końcowym – zwykle aluminium,
- nie powinien zanieczyszczać przyrządu, podłoża lub urządzeń i wyposażenia technologicznego,
- dobre właściwości przyrządu (np. charakterystyki omowe) i żywotność (odporność na starzenie),
- w przypadku kontaktu omowych - niska rezystancja połączenia, minimalne wnikanie w warstwę przyłączeniową, minimalna elektromigracja,
- w przypadku stosowania z organizmami żywymi (lub pochodne – płyny organiczne, krew, etc.) – biokompatybilność.

4.1. Miedź

W aplikacjach medycznych, takich jak system do pomiaru impedancji komórek w warunkach fizjologicznych w czasie rzeczywistym ECIS[®], dobór odpowiednich materiałów, z których wykonane są elementy pomiarowe jest bardzo ważny. Mnogość zastosowań tego rodzaju urządzeń jak również fakt, iż prace badawcze niejednokrotnie odbywają się na żywych komórkach ma wpływ na selekcję pożądanych parametrów. W tabeli 4.1 zestawiono ze sobą parametry takich materiałów jak miedź, złoto i aluminium [8–11]. Niniejsza praca została poświęcona technologii kondensatorów grzebieniowych wykonanych z miedzi przeznaczonych do badań przy użyciu systemu ECIS[®]. Analiza parametrów pozwoliła na wybranie miedzi jako materiału umożliwiającego pomiary parametrów elektrycznych komórek. Jest to materiał o bardzo niskiej rezystancji, tym samym ma on bardzo dobrą przewodność prądu. Pozwala to na otrzymanie dokładnych wyników w badaniach nawet bardzo małych zmian elektrycznych zachodzących w badanych komórkach

4.2. Sputtering

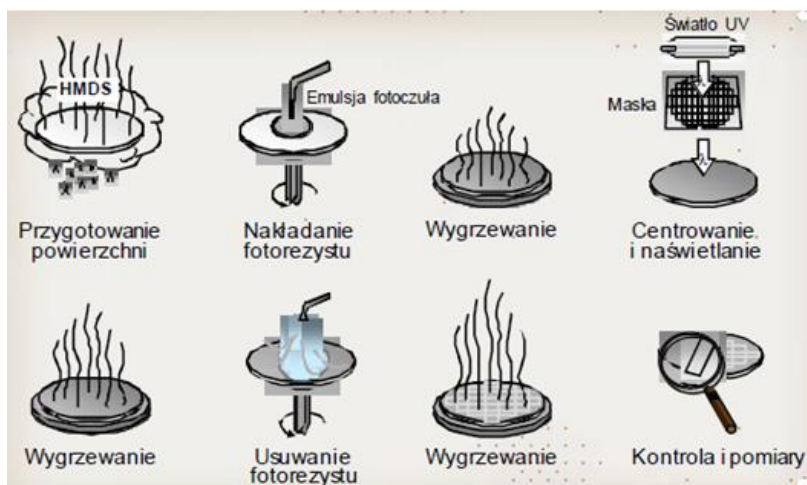
Rozpylanie jonowe, nazywane często rozpylaniem magnetronowym, jest to technika fizycznego osadzania cienkich powłok z fazy gazowej, która powstała na początku XX wieku. Umożliwia ona osadzanie metali oraz ich stopów. Bazuje ona na zjawisku odparowania pojedynczych cząstek materiału osadzanego ze źródła pod wpływem energii zjonizowanego przez silne pole elektryczne gazu. Cechą charakterystyczną rozpylania jonowego jest konieczność przeprowadzenia procesu osadzania w próżni [9, 10].



Rys. 6. Proces rozpylania magnetronowego: a) oddziaływanie pola magnetycznego na cząsteczki gazu, b) przedstawienie linii pola magnetycznego przechodzące przez target, c) źródło po wielu cyklach osadzania [12]

5. PROCESY TECHNOLOGICZNE MAJĄCE NA CELU ODWZOROWYWANIE KSZTAŁTÓW

W technologii MEMS do uzyskania pożądanych struktur, takich jak kondensatory, rezystor lub cewki, konieczne jest zastosowanie procesów, za pomocą których możliwe jest odwzorowanie kształtów w osadzonym materiale. W technologii rozróżniane są metody pośrednie i bezpośrednie, w których do uzyskania oczekiwanej architektury przyrządu, konieczne jest zastosowanie specjalnej emulsji [12]. Przykładem metody bezpośredniej jest trawienie skanowaną wiązką laserową. Kształt struktury jest osiąganym poprzez odwzorowanie maski fotolitograficznej lub jest bezpośrednio odtwarzany przez urządzenia sterujące z modelu matematycznego. Cechuje się ona bardzo niską wydajnością i jest niezwykle rzadko spotykana w produkcji urządzeń typu MEMS [33, 13].



Rys. 7. Algorytm wykonania procesu fotolitografii

Etapy procesu litografii przedstawiono na rysunku 7 odbywa się według opisanego poniżej algorytmu:

- naniesienie warstwy fotorezystu
- naświetlanie
- wywołanie fotorezystu
- trawienie materiału
- usunięcie zbędnego fotorezystu.

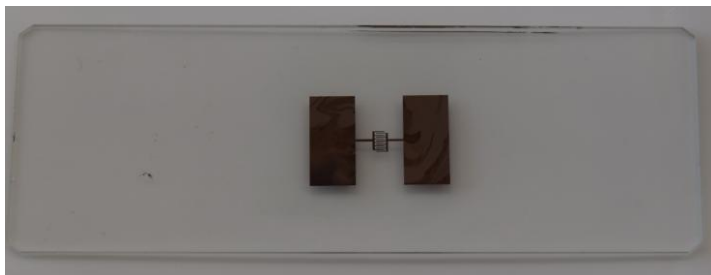
6. WYKONANIE KONDENSATORÓW GRZEBIENIOWYCH Z MIEDZI

Celem, który został założony do realizacji w ramach niniejszej pracy było opracowanie technologii wykonania kondensatorów grzebieniowych z metalizacją miedzi. Plan projektu zakładał wykonanie prób z wykorzystaniem różnych podłoży. Ze względu na zróżnicowane właściwości, mnogość zastosowań w aplikacjach elektronicznych i medycznych, jak również niejednorodne reakcje na różne stymulanty, zdecydowano się na zastosowanie podłoży wykonanych z następujących materiałów:

- szkło mikroskopowe,
- poliester,
- politereftalan etylu – PET,
- poliwęglan – PC,
- PCB.

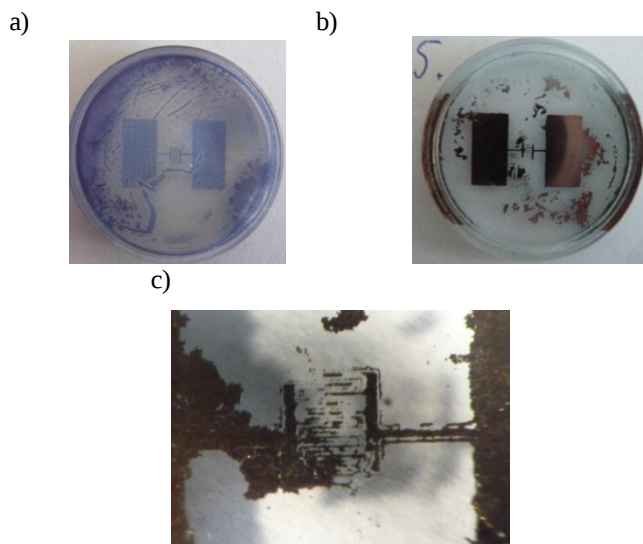
Prace badawcze zostały rozpoczęte od wykonania kondensatorów grzebieniowych na podłożu ze szkła laboratoryjnego. Materiał ten został wybrany ze względu na jego obojętność chemiczną. Po wykonaniu serii kondensatorów za-

uważano że na powierzchni szklanej miedź cechuje się niskim poziomem adhezji. Utrudnia to możliwość wykonania struktury o odpowiedniej jakości pozbawionej zniekształceń.

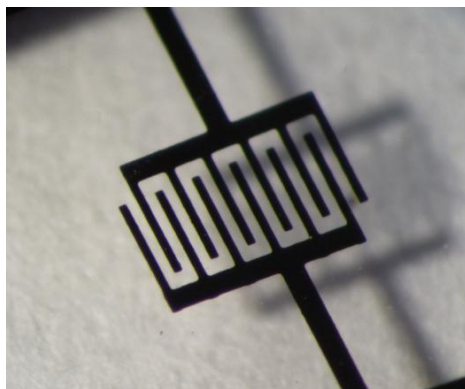


Rys. 8. Kondensator grzebieniowy na podłożu szklanym

Kolejnym podłożem, na którym zostały wykonane kondensatory był poliestr. Płytki, na których zostały wytworzone struktury były w formie niewielkich okrągłych kapsli. Było to podłoże, które zapewniało odpowiedni kształt do osadzenia wewnątrz hodowli. Jednakże, przy zastosowaniu tej metody napotkano na trudności związane z wykonaniem struktury o odpowiedniej jakości. Problemy ze znalezieniem właściwej metody, parametrów trawienia miedzi oraz brak powtarzalności zachodzenia tego procesu pozwoliło na zakwalifikowanie poliestru jako materiału nieodpowiedniego do wykonania podłoża.

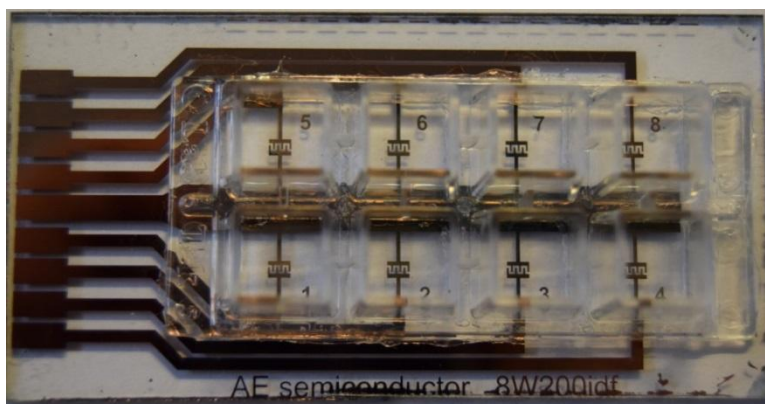


Rys. 9. a) Defekt podłoża poliestru, b), c) struktura kondensatora całkowicie zniszczona w procesie wytrawiania



Rys. 10. Prawidłowo wykonany kondensator grzebieniowy

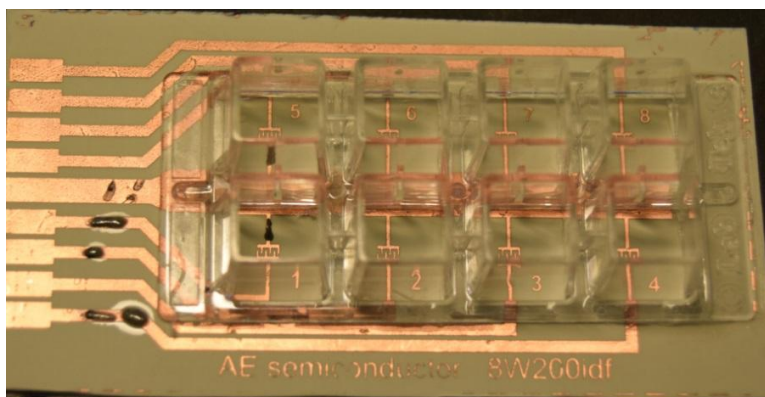
Kolejne prace badawcze obejmowały wykonanie kondensatorów grzebieniowych z metalizacji miedzi na podłożu wykonanym z poliwęglanu – PC. Materiał ten został wybrany, ponieważ jest on wykorzystywany do tworzenia matryc wykonywanych komercyjnie. Wymiary płyki zostały dostosowane do możliwości przeprowadzenia badań za pomocą techniki ECIS[®]. Po przeprowadzeniu serii prób wykonania kondensatorów na poliwęglanie za pomocą różnych technik odwzorowania kształtu stwierdzono, że jest to materiał pozwalający na wykonanie struktur o odpowiedniej jakości (rys. 11).



Rys. 11. Gotowa matryca badawcza z kondensatorami grzebieniowymi z szerokością jednego palca 100 μm

Równolegle do wykonania struktur na podłożu PC trwały próby wykonania kondensatorów grzebieniowych osadzonych na podłożu wykonanym z PCB (rys. 12). Materiał ten został wybrany z powodu jego powszechnego zastosowania w elektronice. Struktury otrzymane na tym podłożu miały być

wykorzystane do pomiarów jej właściwości elektrycznych. Kondensatory otrzymane metodami analogicznymi do tych wykorzystanych na PC okazały się bardzo dobrej jakości, co pozwoliło na przeprowadzenie na nich badań substancji biologicznych (rys. 13).



Rys. 12. Kondensatory grzebieniowe na podłożu PCB



Rys. 13. Stacja pomiarowa systemu ECIS

7. PODSUMOWANIE

W ramach niniejszych prac osiągnięty został założony cel jakim było wykonanie kondensatorów grzebieniowych wykonanych z miedzi służących do przeprowadzania pomiarów w systemie ECIS[®]. Przeprowadzono analizę zagadnienia, jakim jest pomiar impedancji komórek w warunkach fizjologicznych w czasie rzeczywistym. W toku wykonywanego projektu zostały wykonane takie procesy, jak osadzanie warstwy metalizacji z wykorzystaniem metody rozpylania magnetronowego oraz odwzorowanie kształtu na otrzymanej powierzchni metodą subtraktywną i addytywną. Zaprojektowano oraz wykonano maski

technologiczne konieczne do przeprowadzenia procesu fotolitografii. Przeprowadzono szereg testów mających na celu wypracowanie metod i parametrów trawienia metalizacji miedzi na podłożach wykonanych z różnych materiałów. Ostatecznie gotowe matryce zostały wykorzystane do przeprowadzenia testów monitorowania funkcji życiowych komórek w czasie rzeczywistym za pomocą systemu ECIS®.

LITERATURA

- [1] Stolwijk J. A., Matrougui K., [i in.], *Impedance analysis of GPCR – mediated changes in endothelial barrier function: overview and fundamental considerations for stable and reproducible measurements*, “Pflugers Arch – Eur J Physiol” 2015, nr 467, p. 2193
- [2] Applied BioPhysics, *Product Guide*, Corporate Headquarters: 185 Jordan Road Troy, NY 12180 1-866-301-ECIS (3247)
- [3] Prendecka M., Mlak R., Małecka-Massalska T., [i in.], *Effect of exopolysaccharide from Ganoderma applanatum on the electrical properties of mouse fibroblast cells line L929 culture using an electric cell-substrate impedance sensing (ECIS) – Preliminary study*, “Annals of Agricultural and Environmental Medicine” 2016, vol. 23, nr 2, s. 293
- [4] Matthew R. Pennington, Gerlinde R. Van de Walle, *Electric Cell-Substrate Impedance Sensing To Monitor Viral Growth and Study Cellular Responses to Infection with Alphaherpesviruses in Real Time*, mSphere 2017, vol. 2
- [5] Oleś A., *Metody doświadczalne fizyki ciała stałego*. WNT, Warszawa, 1998
- [6] McAdams ET, Lackemeier A, McLaughlin JA, Macken D, Jossinet J, *The linear and nonlinear electrical properties of the electrode-electrolyte interface*, “Biosensors & bioelectronics” 1995, nr 10, s. 74
- [7] Wegener J, Keese CR, Giaever I., *Electric cell-substrate impedance sensing (ECIS) as a noninvasive means to monitor the kinetics of cell spreading to artificial surfaces*, “Exp Cell Res.” 2000, nr 259(1), s. 158
- [8] Dr. Lim Soo King, *Microelectronic Fabrication Metallization*, UTAR 2012.
- [9] Szczepański Z., Okoniewski S., *Technologia i materiałoznawstwo dla elektroników*. WSiP, Warszawa, 2007
- [10] R. Rosenberg, D. C. Edelstein, C.-K. Hu and K. P. Rodbell, *Copper metallization for high performance silicon technology*, IBM T. J. Watson Research Center, Yorktown Heights, New York 10598
- [11] Leszek Golonka, *Zastosowanie ceramiki LTCC w mikroelektronice*, Wrocław 2001
- [12] Beck R., *Technologia krzemowa*, PWN, Warszawa, 1991
- [13] Gary S. May, Simon M. Sze, *Fundamentals of Semiconductor Fabrication*, Wiley, 2004

BADANIE WYTRZYMAŁOŚCI POŁĄCZEŃ DRUTOWYCH DLA WYBRANYCH METALIZACJI

1. WSTĘP

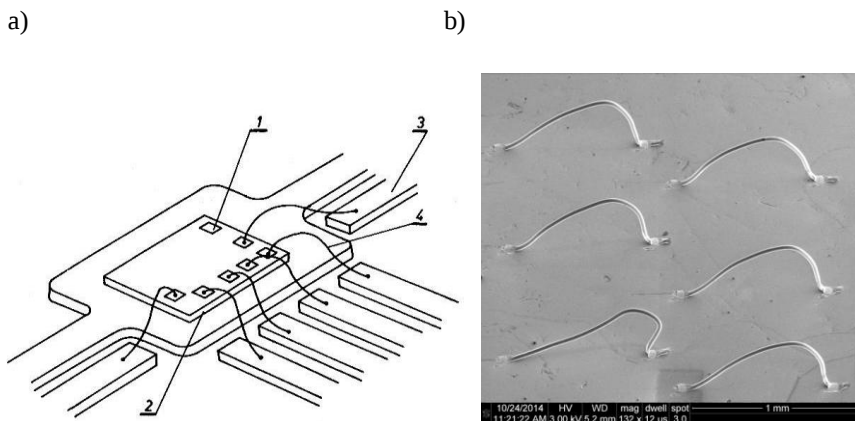
Jakość połączeń drutowych wykonanych w technologii półprzewodnikowej ma znaczący wpływ na wydajność i niezawodność układów scalonych. Duże zagęszczenie oraz mnogość połączeń ma odzwierciedlenie w koszcie produkcji całego urządzenia. Dlatego niezbędne jest badanie istniejących obecnie połączeń na styku drutu i podłoża. Ciągły rozwój tej dziedziny elektroniki jest związany z wykorzystaniem coraz bardziej skomplikowanych rozwiązań konstrukcyjnych oraz coraz lepszych materiałów. Istnieje wiele metod określających jakość połączeń drutowych, jednak najwięcej informacji dostarczają metody mechaniczne. Najpopularniejszą z nich, ze względu na łatwość stosowania oraz dużą ilość informacji jakie można otrzymać, jest metoda zrywania drutu [1–3].

W pracy przedstawiono dwie metody montażu drutowego – ultrakompresję oraz ultratermokompresję. Omówiona została ich zasada działania, wady i zalety oraz możliwość zastosowań. Zaplanowano i wykonano szereg połączeń powyższymi metodami, o różnych wartościach parametrów łączenia, stosując drut złoty i aluminiowy o średnicy 25 μm oraz podłoża z metalizacją złotą, aluminiową, cynową oraz miedzianą. Następnie przeprowadzono badanie wytrzymałości połączeń metodą zrywania drutu. Efektem pracy są zestawy wartości parametrów zgrzewania dla przebadanych złączy.

2. MONTAŻ DRUTOWY – METODY I ZASTOSOWANIE

Pierwsze urządzenie do montażu drutowego zostało zaprojektowane w 1957 roku przez firmę Bell Laboratories. Montaż drutowy (ang. *wire bonding*) jest techniką wykonywania połączeń elektrycznych pomiędzy dwoma polami kontaktowymi struktury półprzewodnikowej (rys. 1). Połączenia wykonywane są za pomocą cienkiego drutu oraz kombinacji ciepła, ciśnienia i/lub energii ultradźwiękowej. Spajanie dwóch metalicznych powierzchni pozostających w fazie stałej zachodzi poprzez doprowadzenie ich do bezpośredniego kontaktu, a następnie przez zachodzący proces dyfuzji atomów lub wzajemnego łączenia się elektronów [1, 4, 5].

¹ Politechnika Lubelska, Koło Naukowe „SEMICON”



Rys. 1. Połączenia drutowe w elemencie półprzewodnikowym: a) schemat [6] 1 – pole kontaktowe układu scalonego, 2 – struktura półprzewodnikowa, 3 – zewnętrzne pole kontaktowe, 4 – połączenie drutowe; b) zdjęcie z mikroskopu elektronowego

Powszechnie stosowanymi metodami montażu drutowego są ultrakompresja i ultratermokompresja, których charakterystykę przedstawiono w tabeli 1. Metody ultrakompresji i ultratermokompresji różnią się również kształtem wykonywanych złączy, użytymi narzędziami, szybkości wykonywania połączeń oraz stosowanych drutów (Tab. 2).

Tab. 1. Metody wykonywania połączeń drutowych [7, 8]

Metoda	Nacisk	Temperatura	Energia ultradźwiękowa	Drut	Kontakt
termokompresja	duży	300–500 °C	nie	Au	Al, Au
ultrakompresja	mały	25 °C	tak	Au, Al	Al, Au
ultratermokompresja	mały	100–150 °C	tak	Au, Cu	Al, Au

Od czasu wynalezienia tej techniki, została ona znacząco rozwinięta i nadal jest szeroko stosowana w przemyśle elektronicznym. Na jej sukces składa się szereg zalet, tj. wysoki stopień automatyzacji, duża elastyczność, precyzyjna kontrola parametrów łączenia, powtarzalność, niskie koszty oraz łatwość oceny i analizy wykonanych połączeń.

Tab. 2. Charakterystyka metod montażu drutowego [7, 8]

Złącze	Metoda	Narzędzie	Drut	Kontakt	Szybkość
kulkowe	TK, UTK	kapilara	Au, Cu	Al, Au	10 połączeń/s (UTK)
klinowe	UK, UTK	sonotroda	Au, Al	Al, Au	4 połączenia/s

TK – termokompresja, UK – ultrakompresja, UTK – ultratermokompresji

Montaż drutowy znajduje zastosowanie we wszystkich obudowach wykorzystywanych w mikroelektronice, szczególnie w obudowach typu CSP, BGA, PQ-FP, COB. Istnieją również pewne ograniczenia, np. ograniczona gęstość połączeń oraz duże opóźnienia w przesyłaniu sygnału ze względu na długość połączeń [1, 3, 9].

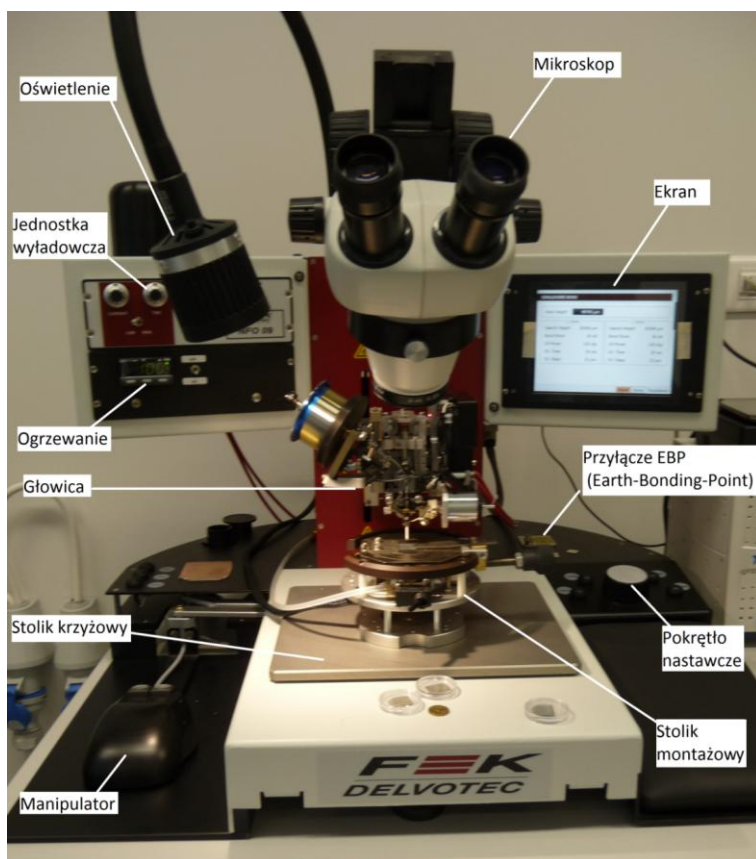
2.1. Bonder 53xx – urządzenie do montażu drutowego

W pracy wykorzystano Bonder 53xx BDA (ang. *Ball Deep Access*) firmy F&K Delvotec (Rys. 2). Ma on możliwość pracy w dwóch wariantach, ultrakompresji i ultratermokompresji. Jest wyposażony w specjalne końcówki bondujące, sonotrodę – do połączeń ultrakompresyjnych oraz kapilarę – do połączeń ultratermokompresyjnych. Zmiana metody wykonywania połączeń ogranicza się do wymiany końcówki oraz drutu, co trwa około 20 minut [10].

Wszystkie podzespoły bondera 53xx BDA zostały zmontowane w jedną całość na płycie będącej jednocześnie jego podstawą. Głównymi elementami urządzenia są:

- głowica z narzędziem wykonującym połączenie (kapilarą lub sonotrodą)
- przesuwany stolik krzyżowy, którym można manipulować w 2 osiach
- podgrzewany stolik montażowy
- sterownik i regulator temperatury
- generator ultradźwiękowy
- regulowany mikroskop firmy Leica
- komputer sterujący.

Bonder działa w oparciu o system operacyjny Windows XP Embedded, który jest wykorzystywany w różnych aplikacjach przemysłowych z dedykowaną aplikacją firmy F&K Delvotec sterującą przebiegiem procesu łączenia. Oferowany interfejs aplikacji pozwala użytkownikowi na łatwy dostęp do parametrów urządzenia i programu. Bezpieczną i wydajną obsługę umożliwia interaktywne obsługa operacji programowania

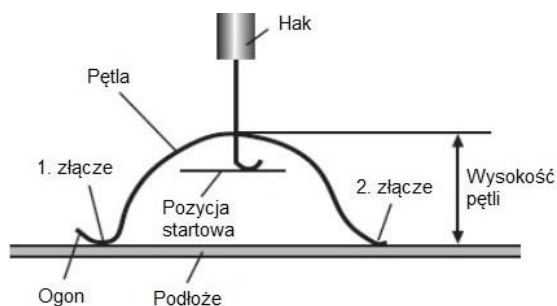


Rys. 2. Bonder 53xx BDA firmy F&K Delvotec

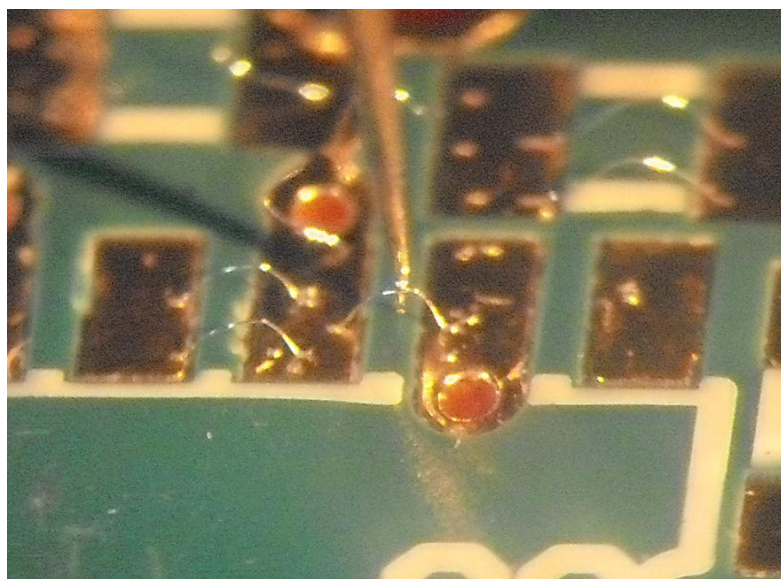
2.2. Metoda zrywania drutu

Największy wpływ na niezawodność pracy układów elektronicznych ma proces pakowania układów scalonych. Jest to proces najbardziej kosztowny, dlatego ważne jest jego zoptymalizowanie. Podstawowym czynnikiem określającym jakość połączeń jest ich wytrzymałość mechaniczna. Istnieje wiele sposobów oceny ich jakości, a najpopularniejszą z nich jest metoda zrywania drutu [1, 4].

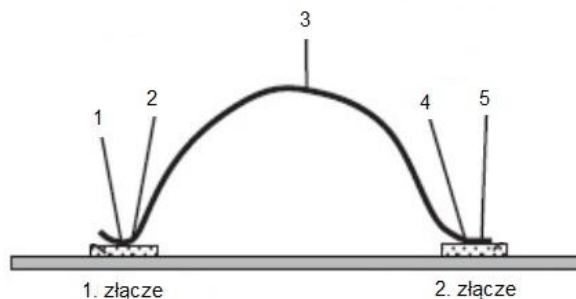
Zrywanie połączenia (ang. *bond pull test*) jest niszczącą metodą określającą wytrzymałość mechaniczną połączenia (rys. 3–4). Zrywanie realizowane jest poprzez wprowadzenie haczyka pod pętlę drutu, a następnie unoszenie go do góry, do momentu zerwania połączenia. Dokonany zostaje pomiar siły ciągu haka, a jej wartość wyświetlona na ekranie urządzenia [1]. Zerwanie połączenia może nastąpić w jednym z pięciu obszarów przedstawionych na rysunku 5.



Rys. 3. Istota zrywania połączeń drutowych [11]



Rys. 4. Zdjęcie przed zerwaniem połączenia haczykiem umieszczonym pod drutem z Au



Rys. 5. Lokalizacja możliwych obszarów zerwania połączenia. 1 – oderwanie pierwszego złącza (kulkowego lub klinowego) od pola kontaktowego; 2 – przerwanie drutu w miejscu karbu przy pierwszym złączu; 3 – przerwanie drutu w miejscu przyłożenia haczyka; 4 – przerwanie drutu w miejscu karbu przy drugim złączu; 5 – oderwanie drugiego złącza od pola kontaktowego [11]

Najczęstszymi powodami zerwań są mikropęknięcia i przewężenia drutu bezpośrednio za uformowanym złączem. Jest to skutkiem wyginania drutu podczas formowania pętli oraz naprężeń w wyniku oddziaływania karbu, który powstaje w momencie wykonania złącza. Odrywanie się złączy od podłoża może być spowodowane niedostateczną czystością metalizacji lub drutu. Najkorzystniejsza sytuacja podczas zrywania występuje wtedy, gdy drut zostaje przzerwany w miejscu przyłożenia haczyka. Wytrzymałość połączenia w takim przypadku jest równa wytrzymałości drutu na zrywanie, a jego niezawodność jest na bardzo wysokim poziomie [1, 4, 12].

3. PROJEKT POŁĄCZEŃ DRUTOWYCH

W ramach niniejszej pracy zaplanowano i wykonano połączenia drutowe metodą ultrakompresji i ultratermokompresji. Do ich realizacji wykorzystano Bonder 53xx BDA. Przygotowano cienkie warstwy metalizacji na podłożu krzemowym oraz PCB. Wykonano połączenia różniące się wartościami parametrów zgrzewania, metodą zgrzewania oraz doбором użytych materiałów. Następnie przeprowadzono kontrolę optyczną powyższych połączeń.

3.1. Planowanie eksperymentu

Wykorzystując urządzenie Bonder 53xx BDA zostały zrealizowane połączenia na cienkich metalizacjach ze złota, aluminium, miedzi oraz cyny.

Podłożem o złotej metalizacji był laminat szklano-epoksydowy FR4 z 35 μ m warstwą miedzi pokrytą złotem chemicznym. Jako podłoże o cynowej metalizacji wykorzystano laminat pokryty warstwą cyny techniką grubowarstwową. Kolejnymi podłożami były płytki krzemowe ze 100nm warstwą miedzi oraz 1 μ m warstwą aluminium. Warstwy zostały wykonane metodą rozpylania magnetyronowego.

Do wykonania połączeń ultrakompresyjnych wykorzystano sonotrodę z węgliku wolframu o długości 1 cala i płaskim profilu stopy oraz zastosowano drut aluminiowy o średnicy 25 μ m z domieszką 1% krzemu. Wartość siły zerwania drutu mieści się w przedziale od 14cN do 16cN [10, 13].

Do ultratermokompresji wykorzystano ceramiczną kapilarę o długości 16 mm oraz 30° kącie stożka. Użyto drut złoty o średnicy 25 μ m, którego siła zerwania wynosi ponad 14cN.

Zaplanowano wykonanie 14 serii połączeń metodą ultrakompresji (12 na podłożu Al, 2 na Au) oraz 110 serii połączeń ultratermokompresyjnych (36 na Cu, Al, Sn, 2 na Au), przy czym każda seria zawiera co najmniej 5 połączeń.

Wykonując połączenia metodą ultrakompresji zastosowano wartości parametrów zawierających się w następujących zakresach:

- siła docisku: bez zmian (25cN),
- moc ultradźwięków: 85–100dig (co 5dig),

- czas działania ultradźwięków: 35–45ms (co 5 ms),
- temperatura zgrzewania: bez zmian (25°C).

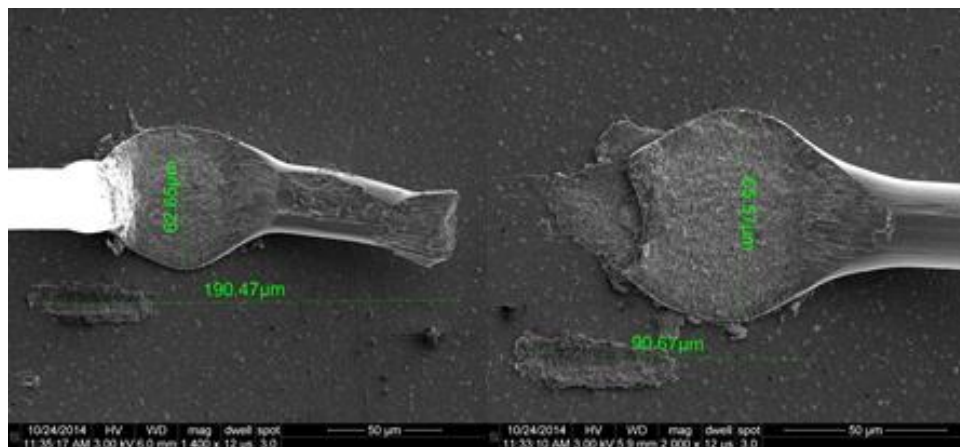
Natomiast wykonując połączenia metodą ultratermokompresji zastosowano wartości parametrów zawierających się w następujących zakresach:

- siła docisku: 30–50cN (co 10cN),
- moc ultradźwięków: 130–170dig (co 20dig),
- czas działania ultradźwięków: 25–40ms (co 5ms),
- temperatura zgrzewania: bez zmian (100°C).

Badanie połączeń metodą zrywania odbyło się wykorzystując dedykowaną zrywarkę Micro Pull Tester MPT-3 firmy Mechatronika.

3.2. Wykonanie połączeń drutowych metodą ultra kompresji

Wykonanie 14 serii połączeń odbyło się z wykorzystaniem wcześniej zaplanowanych parametrów zgrzewania, drutu aluminiowego oraz podłoża o metalizacji złotej i aluminiowej. Zwiększanie mocy i czasu trwania ultradźwięków powodowało coraz większe odkształcenie drutu w miejscu karbu. Może to powodować znaczące osłabienie połączenia w tym obszarze. Wymiary i kształt złączy połączenia ultrakompresyjnego przedstawiono na rysunku 6.

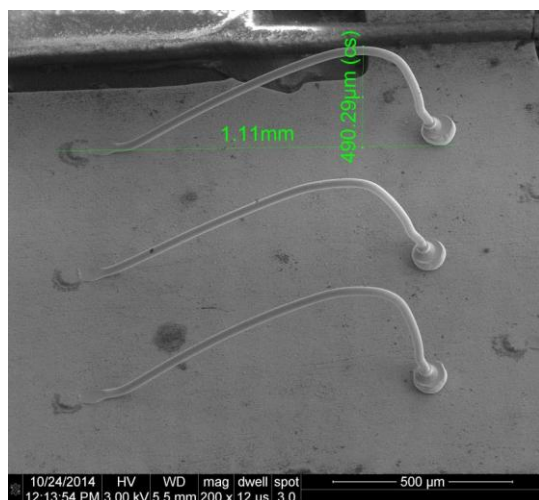


Rys. 6. Wymiary i kształt złączy klinowych

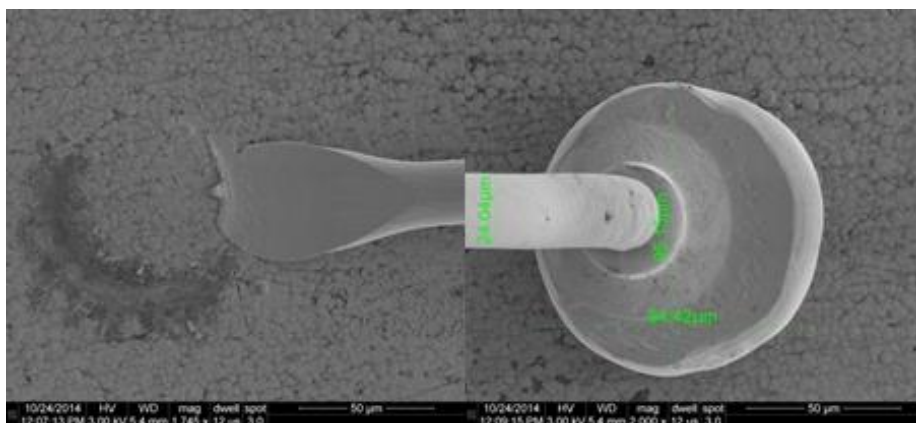
3.3. Wykonanie połączeń metodą ultra termokompresji

Wykonano 114 serii połączeń ultra termokompresyjnych wykorzystując drut złoty i podłoża o metalizacji złotej, aluminiowej, cynowej oraz miedzianej.

W menu bondera wybrano podstawowy kształt pętli o długości 1mm oraz o wysokości 500µm. Długość całego połączenia wynosiła około 1,1mm (Rys. 7). Średnica złącza kulkowego wynosiła około 95 µm. Wymiary i kształt złączy został przedstawiony na rysunku 8.

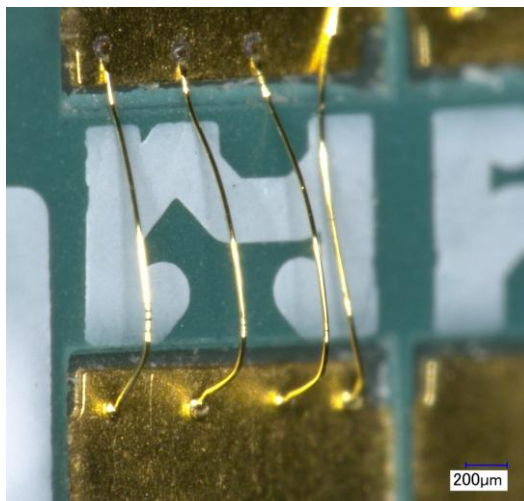


Rys. 7. Wymiary połączenia ultratermokompresyjnego



Rys. 8. Złącze klinowe i kulkowe

Połączenia wykonane na metalizacji złotej i aluminiowej zostały wykonane zgodnie z zaplanowanym eksperymentem (Rys. 9). W przypadku połączeń wykonywanych na podłożach Cu i Sn, ze względu na trudność w uzyskaniu poprawnych połączeń, okazało się konieczne skorygowanie parametrów zgrzewania, głównie zwiększenia ich wartości. Dodatkowo, dla metalizacji cynowej było to zwiększenie temperatury podłoża ze 100°C do temperatury 170°C. Podczas wykonywania połączeń dochodziło do częstego zapychania się kapilary drutem. W celu jej udrożnienia wykonywano trawienie w wodzie królewskiej.



Rys. 8. Połączenia ultratermokompresyjne na złotej metalizacji

4. TEST ZRYWANIA POŁĄCZEŃ

Jakość połączenia jest określona siłą wiązania drutu z metalizacją pola kontaktowego. Metoda zrywania drutu jest podstawową metodą badania wytrzymałości połączeń drutowych. Ponadto, ważne są również inne czynniki, na które warto zwrócić uwagę, tj. wyginanie się pętli drutu, mikropęknięcia czy deformacja złączy [1, 4]. Siła zerwania zależy głównie od średnicy drutu. Zakładając, że kąt α zawierający się między podłożem a drutem tworzącym obydwie złącza będzie taki sam, dla 25 μm drutu, minimalna wartość siły zerwania pojedynczego złącza wynosi ponad 4cN, a dla wartości średniej siły, więcej niż 50% siły zerwania samego drutu [14].

Wyniki dla zrywania połączeń wykonanych metodą ultrakompresji przedstawiono w tabeli 3.

Tabela 3. Wyniki zerwań połączeń wykonanych metodą ultrakompresji

Płytki krzemowa (drut – Al, podłoże – Al)					Płytki krzemowa (drut – Al, podłoże – Al)				
Bond Force	US. Power	US. Time	F	nr zerwania	Bond Force	US. Power	US. Time	F	nr zerwania
cN	dig	ms	mN	–	cN	dig	ms	mN	–
25	85	35	74	2	25	95	35	61	2
25	85	40	53	2	25	95	40	59	2
25	85	45	65	2	25	95	45	52	2
25	90	35	56	2	25	100	35	53	2

25	90	40	67	2	25	100	40	52	2
25	90	45	57	2	25	100	45	48	2
PCB (druć – Al, podłoże – Au)									
Bond Force	US. Power	US. Time	F	nr zerwania					
cN	dig	ms	mN	–					
25	100	40	42	2					
20	100	45	34	2					

Bond Force – siła zgrzewania; US. Power – moc ultradźwięków; US. Time – czas działania ultradźwięków; *F* – siła zerwania

We wszystkich połączeniach nastąpiło zerwanie oznaczone numerem 2, czyli w miejscu karbu. Na podstawie uzyskanych wyników można stwierdzić, że zwiększenie czasu trwania ultradźwięków oraz ich mocy powodowało osłabienie połączeń. Tylko połączenia o najniższych wartościach parametrów, o sile zerwania 74mN, spełniły warunek wytrzymałości przedstawiony powyżej. Takie połączenie można uznać za połączenie wysokiej jakości.

Połączenia wykonane na złotej metalizacji wymagały zwiększenia czasu trwania ultradźwięków oraz ich mocy w celu uzyskania pełnego połączenia. Mimo to, siła zerwania tych połączeń jest bardzo niewielka, co dyskwalifikuje je z zastosowań. Połączenie materiałów Al–Au jest narażone na powstawanie niekorzystnych dla złączy związków międzymetalicznych. Ze względu na szybszą dyfuzję złota niż aluminium, powstaje krucha faza, która prowadzi do degradacji złącza [1, 4]. W tabeli 4 oraz 5 przedstawiono uśrednione wyniki zerwań uzyskane z testu zrywania połączeń ultratermokompresyjnych.

Tab. 4. Wyniki zerwań połączeń ultratermokompresyjnych

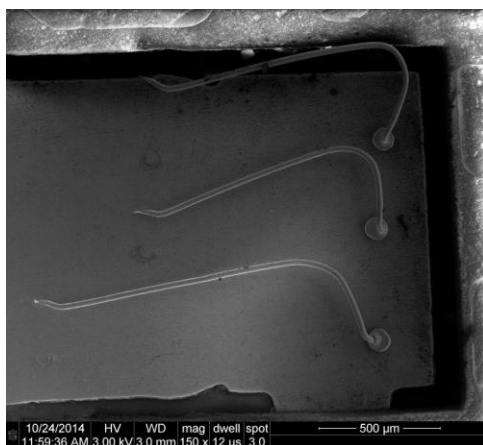
Laminat (druć – Au, podłoże – Sn)					
Bond Force	US. Power	US. Time	F	Temperatura podłoża	nr zerwania
cN	dig	ms	mN	°C	–
40	150	35	45	150	5
50	150/170	40/45	39	170	5
60	150	40	25	170	5

Tab. 5. Wyniki zerwań połączeń wykonanych metodą ultratermokompresji

Płytki krzemowa (drut – Au, podłoże – Al)					PCB (drut – Au, podłoże – Au)				
Bond Force	US. Power	US. Time	F	nr zerwania	Bond Force	US. Power	US. Time	F	nr zerwania
cN	dig	ms	mN	–	cN	dig	ms	mN	–
30	130	30	113	2	30	150	25	52	5
30	150	35	95	2	30	150	30	44	5
30	170	40	112	2	Laminat (drut – Au, podłoże – Cu)				
Bond Force	US. Power	US. Time	F	nr zerwania	Bond Force	US. Power	US. Time	F	nr zerwania
50	130	30	97	5	cN	dig	ms	mN	–
50	150	35	87	5	30	130	25	31	5
50	170	40	94	2	30	130	35	39	5
70	130	25	50	5	30	130	40	28	5
70	150	30	35	5	30	170	35	35	5
70	170	35	55	5	30	170	40	32	5
70	170	40	74	5	30	170	40	32	5

Połączenia wykonane drutem złotym na aluminiowej metalizacji charakteryzowały się zerwaniem w miejscu oznaczonym numerem 2 oraz 5 (rys. 5). Zerwanie połączenia w miejscu numer 2 jest spowodowane działaniem karbu i inicjowaniem mikropęknięć w tym obszarze. Zastosowanie większej siły docisku powodowało uszkodzenie drugiego złącza połączenia, a tym samym malała siła potrzebna na jego zerwanie (zerwanie nr 5). Przy zastosowaniu mniejszej siły docisku uzyskano połączenia o bardzo dużej wartości siły zerwania, bliskiej wytrzymałości samego drutu. Takie połączenia odznaczają się bardzo wysoką niezawodnością pod względem mechanicznym [4].

Zrywanie połączeń na metalizacji cynowej przeprowadzono dla trzech z jedenastu wykonanych połączeń. Spowodowane było to niską jakością złączy, które samoistnie odrywały się od podłoża. W kilku przypadkach również kształt pętli nie pozwalał na wprowadzenie haczyka w celu zerwania połączenia. Zerwania nastąpiły w miejscu oznaczonym numerem 5 (rys. 5). Przykład takiego zerwania przedstawiono na rysunku 10. Połączenia charakteryzowały się bardzo niską wytrzymałością, co dyskwalifikuje je z zastosowań.



Rys. 10. Zerwane połączenie w miejscu oznaczonym numerem 5

Połączenia zrealizowane na płytce PCB o metalizacji Au oraz Cu charakteryzowały się występowaniem zerwania w miejscu oznaczonym numerem (5) (rys. 5). Całkowita utrata kontaktu z podłożem następowała przy użyciu niewielkiej siły ciągu haczyka. Podobnie jak w przypadku połączeń na metalizacji cynowej, takie połączenia nie są akceptowane. Mogło być to spowodowane niedostatecznym oczyszczeniem podłoża, użyciem zbyt niskich wartości parametrów wiązania, twardej metalizacji, z którą trudniej zachodzi proces dyfuzji oraz zbyt niskiej temperatury podgrzania płytki [4].

Uzyskany na drodze eksperymentu, opracowany zestaw wartości parametrów zgrzewania do wykonywania połączeń metodą ultra kompresji i ultra termokompresji na różnych metalizacjach, przedstawiono w tabeli 6 i tabeli 7.

Tab. 6. Zestaw wartości parametrów zgrzewania dla połączeń wykonanych metodą ultrakompresji

Płytki krzemowa (drut – Al, podłoże – Al)				Płytki krzemowa (drut – Al, podłoże – Al)			
Bond Force	US. Power	US. Time	Temperatura zgrzewania	Bond Force	US. Power	US. Time	Temperatura zgrzewania
cN	dig	ms	°C	cN	dig	ms	°C
25	85	35	25	25	95	35	25
PCB (drut–Al, podłoże–Au)							
Bond Force	US. Power	US. Time	Temperatura zgrzewania				
cN	dig	ms	°C				
25	100	40	25				

Tab. 7. Zestaw wartości parametrów zgrzewania dla połączeń wykonanych metodą ultratermokompresji

Płytką krzemowa (drut – Au, podłoże –Al)				PCB (drut – Au, podłoże – Au)			
Bond Force	US. Power	US. Time	Temperatura zgrzewania	Bond Force	US. Power	US. Time	Temperatura zgrzewania
cN	dig	ms	°C	cN	dig	ms	°C
30	130	30	100	30	150	25	100
Laminat (drut – Au, podłoże – Cu)				Laminat (drut – Au, podłoże – Sn)			
Bond Force	US. Power	US. Time	Temperatura zgrzewania	Bond Force	US. Power	US. Time	Temperatura zgrzewania
cN	dig	ms	°C	cN	dig	ms	°C
30	130	35	100	40	150	35	150

5. WNIOSKI

W ramach niniejszej pracy zaplanowano i wykonano szereg połączeń drutowych metodami ultrakompresji i ultratermokompresji na podłożach o metalizacji Au, Al, Cu, Sn, wykorzystując drut złoty i aluminiowy. Zbadano wytrzymałość mechaniczną połączeń drutowych metodą zrywania drutu. Na podstawie otrzymanych wyników można stwierdzić, że połączeniami o najmniejszej wytrzymałości okazały się połączenia wykonane drutem złotym na miedzianej metalizacji, których siła zerwania wynosiła 39cN. Połączenia o takiej wytrzymałości nie spełniają wymagań dotyczących możliwości ich zastosowania. Najlepsze połączenia uzyskano stosując drut złoty na aluminiowej metalizacji, o sile zerwania równej 113cN. Na podstawie kontroli optycznej defektów struktury połączeń stwierdzono, że istotny wpływ na wytrzymałość połączeń miało oddziaływanie karbu, powodujące mikropęknięcia w obszarze złącza. Finalnie opracowano zestaw wartości parametrów zgrzewania dla testowanych podłoży do uzyskania wysokiej jakości połączeń elektrycznych.

LITERATURA

- [1] Harman G., *Wire Bonding in Microelectronics*. McGraw-Hill, Nowy Jork 2010, 13–153.
- [2] Allen J. J., *Micro Electro Mechanical System Design*. CRC Press, Boca Raton 2005, s. 339–362.
- [3] Szczepański Z., Okoniewski S. *Technologia i materiałoznawstwo dla elektroników*. WSiP, Warszawa 2007, s. 216–225.
- [4] Wang C., Sun R., *The quality test of wire bonding*. *Modern Applied Science*, Tom 3, 2009, nr 12, s. 50–56.

- [5] May G. S., Sze S. M. i inni, *Fundamentals Of Semiconductor Fabrication*. WILEY, Nowy Jork 2004, s. 233–235
- [6] Adamczewska J., *Procesy technologiczne w elektronice półprzewodnikowej*. WNT, Warszawa 1980
- [7] Lai Z., Liu J., *The Nordic Electronics Packaging Guideline*,
<http://extra.ivf.se/ngl/A-WireBonding/ChapterA.htm> (11 czerwca 2017)
- [8] Small Precision Tools SPT. <http://www.smallprecisiontools.com/products-and-solutions/chip-bonding-tools> (11 czerwca 2017)
- [9] Chen W.-K., *VLSI Technology*. CRC Press, Boca Raton 2003, s. 285-291
- [10] Bonder 53xx Ball Deep Access, User Guide v3.0. 2012
- [11] *Montaż elektroniczny i systemy testujące*.
http://home.agh.edu.pl/~maziarz/Lecture_1-wire-bonding-en.pdf (11 czerwca 2017)
- [12] Lizak T., Kociubiński A., *Reliability tests of ultrasonic and thermosonic wire bonds*. Proc. of SPIE, vol. 9662, Wilga 2015
- [13] *Heraeus Materials Technology*. <http://heraeus-contactmaterials.com> (11 czerwca 2017)
- [14] Lee K., Özkök M., Schmitz S., *Gold Wire Bonding Performance and Reliability of ENEPIG Surface Finishes*. Atotech, Berlin 2011.

PROJEKT I TECHNOLOGIA KONDENSATORÓW GRZEBIENIOWYCH DO MONITOROWANIA HODOWLI KOMÓREK

1. TECHNIKA MONITOROWANIA IMPEDANCJI KOMÓREK PODCZAS HODOWLI

ECIS[®] (ang. *Electric Cell-substrate Impedance*) jest metodą pomiaru impedancji komórek w czasie rzeczywistym. Pozwala analizować ich aktywności odnoszące się do struktury i morfologii, zdolności do ich podziału, przemieszczania się oraz wielu innych zachowań kierowanych przez cytoszkielet. Technika monitorowania impedancji komórek została wykorzystana w wielu badaniach, będąc znakomitą alternatywą dla prac prowadzonych przy zastosowaniu mikroskopów, w których wyniki otrzymuje się na podstawie obserwacji. Metoda ta stosowana jest do diagnozowania inwazyjnego charakteru komórek nowotworowych, funkcji barierowej komórek śródbłonna, toksyczności testów *in vitro* będących środkiem zastępczym dla testów na zwierzętach oraz transdukcji receptorów sprzężonych z białkami G (GPCR – ang. *G Protein-Coupled Receptor*) w celu odkrywania nowych leków. System ECIS[®] pozwala określić m.in. przebieg życia komórki od jej wyhodowania po uśmiercenie, cykl rozwoju i rozpowszechniania się do momentu uzyskania hodowli w pełni konfluentnej, jak również określić kontrakcje na różnorodne reagenty [1, 8].

1.1. Pomiar impedancji komórek

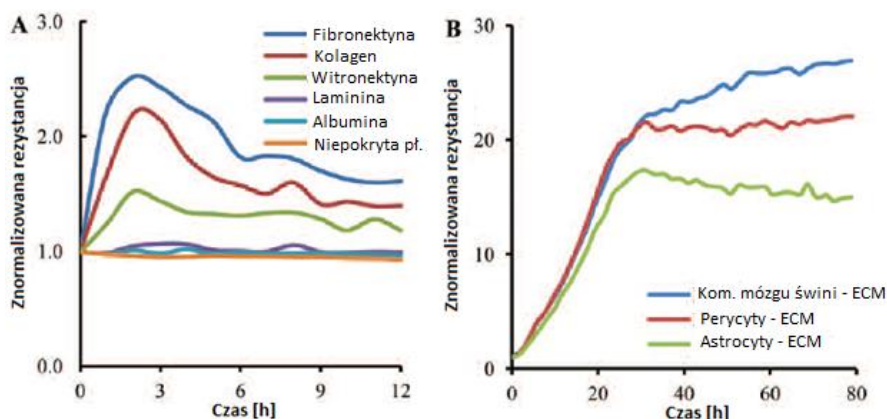
Badanie polega na hodowli komórek, które łączą się z elektrodami ECIS[®] na dnie studzienek. Po przymocowaniu się do powierzchni, komórki zaczynają się rozprzestrzeniać, co powoduje zwiększenie ich obszaru na granicy z elektrodami. To z kolei skutkuje zmianą mierzonej impedancji. Ciągły pomiar w czasie rzeczywistym zapewnia zarejestrowanie dynamikę zmian parametrów elektrycznych (impedancja, rezystancja, pojemność), jak również określenie ich wartości końcowej [1, 6, 8].

Tradycyjne opracowywanie wyników badań komórek dostarcza informacji jedynie o liczbie komórek przyłączonych do każdej powłoki macierzy pozakomórkowej ECM (ang. *extracellular matrix*). Przetwarzając dane za pomocą systemu ECIS[®], można otrzymać informacje dotyczące siły z jaką komórki są przyłączone do podłoża. Przejrzysty charakter elektrod umożliwia, poprzez weryfi-

¹ Politechnika Lubelska, Koło Naukowe „SEMICON”

kację optyczną, określenie wartości rezystancji na komórkę bądź na obszar komórek. Długoterminowa analiza pozwala badać nie tylko wpływ białek ECM na przyłączanie do powierzchni elektrod, ale także przedstawiać skutki funkcjonalne działania różnych czynników na ECM [1, 6, 8].

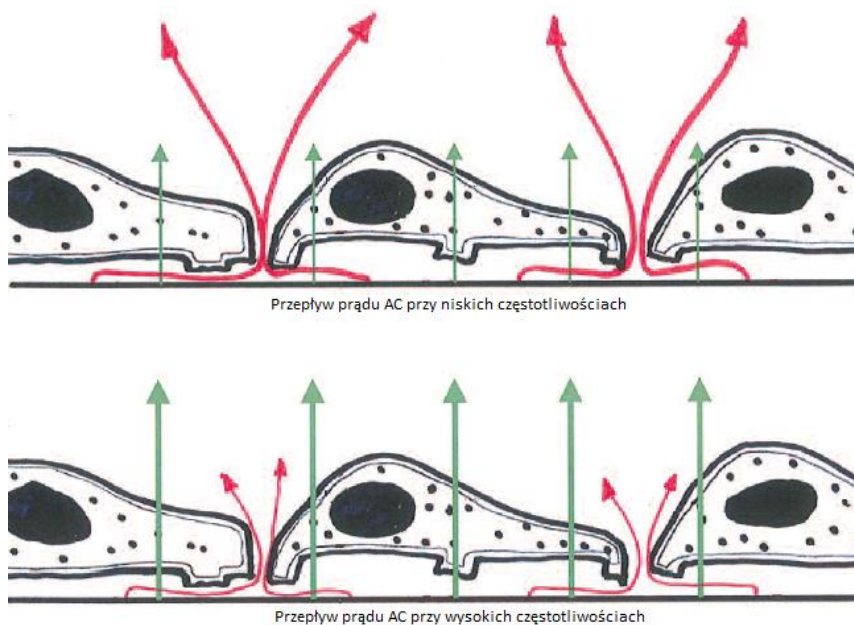
Na rysunku 1 przedstawiono tempo rozprzestrzeniania się komórek mięśni gładkich naczyń nerkowych (A) i śródbłónka naczyń włosowatych mózgu (B). W pierwszym przypadku przed rozpoczęciem hodowli, elektrody zostały poddane działaniu wskazanych białek. W drugim przypadku przed rozpoczęciem rejestracji zmian, perycyty, astrocyty czy również komórki śródbłónka mózgowego narastały swobodnie w elektrodach ECIS[®]. Następnie komórki poddano trypsynizacji i w nowych studzienkach zmierzono tempo ich rozprzestrzeniania i przyłączania [1, 6, 8].



Rys. 1. Wyniki pomiaru impedancji komórek mięśni gładkich naczyń nerkowych (A) i śródbłónka naczyń włosowatych mózgu (B) [1]

1.2. Wpływ częstotliwości sygnału na parametry elektryczne komórek

Częstotliwość z jaką płynie prąd przemienny elektrodami pokrytymi komórkami ma kluczowe znaczenie w systemie ECIS[®]. Przy względnie niskich częstotliwościach (< 2kHz) prąd przepływa pod i pomiędzy sąsiednimi komórkami. Natomiast przy stosunkowo wysokich częstotliwościach (> 40kHz) następuje przepływ bezpośrednio przez błony komórkowe. Na rysunku 2 przedstawiono schemat przepływu prądu dla niskich i wysokich częstotliwości [1, 8, 12].

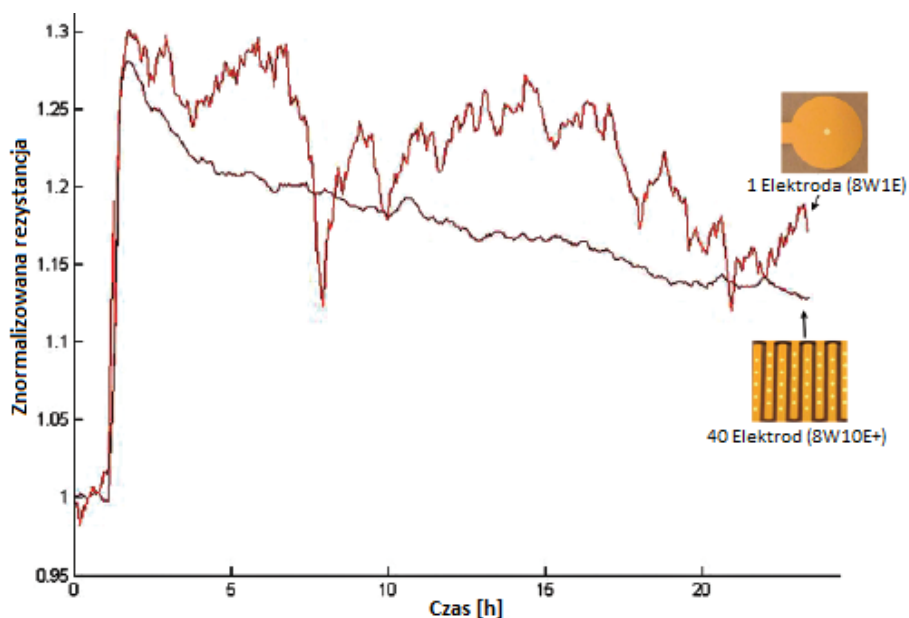


Rys. 2. Schemat przepływu prądu dla niskich i wysokich częstotliwości [1]

1.3. Wpływ konstrukcji elektrod na zachowanie komórek

W elektrodach z niewielką powierzchnią całkowitą ich układ w studzienkach powoduje, że praktycznie cały prąd przechodzi przez pojedynczą warstwę komórki. Pozwala to na analizę uzyskanych danych za pomocą dedykowanego oprogramowania w systemie ECIS[®] służącego do modelowania, w celu zarejestrowania funkcji barierowej, pojemności błony komórkowej, a także odstępów między błoną bazową a elektrodą. Niewielka powierzchnia elektrody umożliwia uzyskanie stosunkowo dużego pola elektrycznego przy względnie niskim prądzie AC. Pozwala to na częściową elektroporację lub całkowite uśmiercenie komórek podczas doświadczeń migracyjnych. Dodatkowo, małe elektrody gwarantują możliwość monitorowania niekontrolowanych zmian morfologicznych pojedynczych lub małych populacji komórek (< 100) w skali nano. Podczas, gdy elektrody większe lub wielokrotne dostarczają szacunkową odpowiedź morfologiczną zbioru komórek (> 1000) [1, 8].

Na rysunku 3 przedstawiono reakcję połączonych warstw komórkowych z domieszką świeżej pożywki. Przedstawiona jest znormalizowana zmiana impedancji w funkcji czasu. Początkowa wartość dla elektrody oznaczonej symbolem 1E wynosi 11,5 k Ω , zaś dla elektrody oznaczonej symbolem 10E+ 1,3k Ω [1, 8].



Rys. 3. Wykres przedstawiający wyniki pomiaru impedancji komórek hodowanych na elektrodach typu 1E oraz 10E+ [1]

2. ZASTOSOWANIE TYTANU DO WYTWORZENIA ELEKTROD

Tytan w medycynie znalazł swoje zastosowanie już w latach 40. XX. w., jako materiał implantologiczny. Wczesne eksperymenty przeprowadzone na zwierzętach przez Bothea i Leventhala wykazały, że tytan posiada znakomitą zgodność tkankową i nawet w dużych dawkach nie wpływa negatywnie na organizmy. Jako znakomity materiał do zastosowań medycznych, dzięki odporności na korozję i dobrym właściwościom mechanicznym został szeroko rozpowszechniony w latach siedemdziesiątych. Czysty tytan charakteryzuje się małą przewodnością cieplną, względnie małą gęstością, dosyć niskim modułem sprężystości, umiarkowaną wytrzymałością, dobrą odpornością na korozję w różnych środowiskach oraz wysoką reaktywnością. W tabeli przedstawiono wybrane właściwości pierwiastka w zestawieniu z innymi metalami [10, 11].

Tytan jest materiałem szczególnie wykorzystywanym w ortopedii ze względu na wysoką wytrzymałość i niski moduł sprężystości. Jednakże tytan ma niską odporność na zużycie poprzez ścieranie ze względu na niewielką twardość. Relatywnie słabo rozwinięte właściwości tribologiczne sprzyjają obróbce powierzchni w celu ulepszenia ich właściwości. Wykorzystywane są różne metody obróbki, w tym użycie powłoki PVD (TiN, TiC), implantacji jonów (N^+), obróbki cieplnej (azotowanie, dyfuzja i utwardzanie) oraz obróbki laserowe z TiC [3, 5, 7].

Tab. 1. Właściwości tytanu w porównaniu z innymi metalami

	Ti	Al	Fe	Ni
Gęstość [g/cm^3]	4.5	2.7	7.9	8.9
Temp. topnienia [$^{\circ}\text{C}$]	1670	660	1538	1455
Przewodność cieplna [W/mK]	15-22	221-247	68-80	72-92
Moduł sprężystości [GPa]	115	72	215	200
Reaktywność z tlenem	Bardzo duża	Duża	Niska	Niska
Odporność na korozję	Bardzo duża	Duża	Niska	Średnia
Cena	Bardzo wysoka	Średnia	Niska	Wysoka

Tytan i jego stopy są uważane za materiały znakomite w zastosowaniach biokompatybilnych. Są stosunkowo obojętne i mają dobrą odporność na korozję uwarunkowaną występowaniem tlenku na ich powierzchni. Tytan łatwo adsorbuje białka z płynów biologicznych. Stwierdzono, iż wchłania albuminę, lamininę V, glikozaminoglikany, kolagenazę, fibronektynę, białka dopełniające i fibrynogen. Powierzchnia tytanu wspiera także wzrost komórek i ich różnicowanie. Istnieje wiele prac poświęconych badaniu interakcji komórek z powierzchniami tytanu. Po wszczępieniu materiału do ludzkiego ciała, początkowo na implantach, pojawiają się zauważane neutrofile i makrofagi, przyczyniające się do powstania załączka nowej kości. Pierwszą reakcją po umieszczeniu wewnątrz organizmu materiału jest adsorpcja białka. Następnie neutrofile i makrofagi w interakcji z cytokinami zwalniają fibroblasty i wpływają na proces kapsułkowania ciała obcego. Tytan, znajdując się w niewielkiej odległości od kości, w odpowiednich warunkach wpływa na zmineralizowanie się komórek. Nie jest jednak bezpośrednio połączony z kością, będąc oddzielony od niej warstwą organiczną. Wiązanie jakie powstaje jest związane z osteointegracją i mechanicznemu związkowi pomiędzy chropowatością powierzchni tytanu i porach występujących w tkance kostnej. W celu zapewnienia wiązania biologicznego tytanu z kością, prowadzone są badania nad modyfikacją powierzchni w celu poprawy przewodnictwa kostnego lub lepszej bioaktywności tytanu [3, 5].

2.2. Zastosowania biomedyczne tytanu

Wcześniejsze zastosowania tytanu w urządzeniach medycznych, chirurgicznych i stomatologicznych opierały się na postępach medycyny podczas II wojny światowej, co wynikało z bardzo surowych wymagań stawianych przez przemysł lotniczy i wojskowy. Zwiększone wykorzystanie tytanu i jego stopów jako biomateriałów zależało od ich niższego modułu, lepszej biokompatybilności i odporności na korozję w porównaniu z bardziej konwencjonalnymi stopami ze stali nierdzewnej i kobaltu. Właściwości te stanowiły siłę napędową do szybkiego wprowadzania nowoczesnych stopów Ti. Zastosowania tytanu i jego stopów można klasyfikować według ich funkcji biomedycznych [3, 5].

2.2.1. Zastosowania kardiologiczne i sercowo naczyniowe

Tytan i stopy tytanu są powszechnie wykorzystywane w implantach układu sercowo-naczyniowego z powodu ich unikalnych właściwości. Wczesnymi przykładami zastosowań były protezy zastawkowe serca, elementy ochronne w rozrusznikach serca, sztuczne serca i urządzenia krążeniowe. Ostatnio zwrócono uwagę na użycie stopu niklu i tytanu (NITINOL) kształtowego w urządzeniach wewnątrznaczyniowych, takich jak stenty i cewki okluzyjne. Atutem stosowania tytanu w układzie sercowo-naczyniowego jest fakt, że jest on wytrzymały, obojętny odczynowo i niemagnetyczny. Dodatkowo, podczas rezonansu magnetycznego, tytan wchodząc w reakcje wspomaga obrazowanie zmian fizjologicznych oraz patologicznych. Zaś wadą jest niewystarczająca radioaktywność w drobniejszych strukturach. W rozrusznikach sercach i urządzeniach wspomagających krążenie, tytan stosowany jest zarówno w częściach mechanicznych pomp, jak i w elementach będących w kontakcie z krwią. Aparaty wykonane w całości z tytanu nie są wykorzystywane, głównie z powodu problemów związanych z krzepnięciem krwi przy powierzchni urządzenia [3, 5, 7].

2.2.2. Zastosowanie tytanu w procesie osteosyntezy

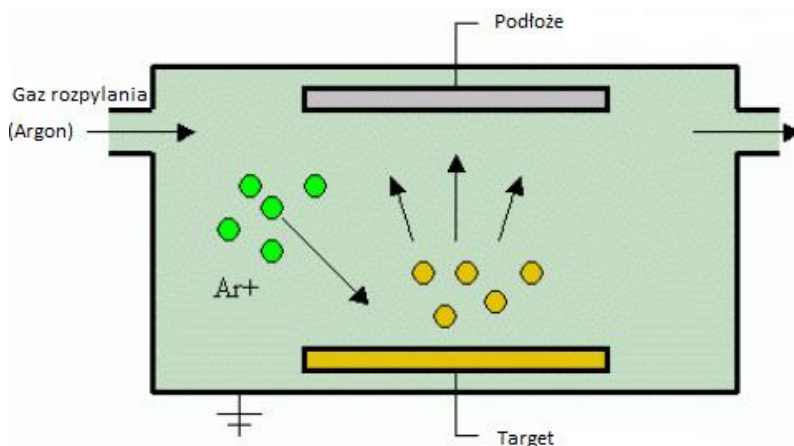
Oprócz sztucznych kości, stawów i implantów stomatologicznych, stopy tytanu wykorzystuje się w osteosyntezie, która jest jedną z chirurgicznych metod leczenia złamanych kończyn. Dzięki niej można przywrócić funkcje utracone podczas wypadku. Tytan i jego stopy są atrakcyjnymi materiałami w implantach osteosyntezy, ponieważ spełniają wszystkie wymagania stawiane w tej metodzie leczenia. Typowe implanty do osteosyntezy obejmują śruby kostne, płyty kostne, implanty szczękowo-twarzowe itp. Śruby kostne są stosowane jako pojedyncze wkręty do bezpośredniego mocowania kostnego lub jako śruby łączące, ścisiskające wywierające nacisk na szczelinę powstałą w czasie pęknięcia. Ponadto, korzysta się z nich podczas mocowania płyt lub innych elementów implantacyjnych do kości. Płyty kostne stosowane są prawie we wszystkich głównych obszarach szkieletowych jako mostki, a nawet jako wewnętrzne stabilizatory. Tytan i jego stopy o szorstkich powierzchniach (powierzchnie ściernie, plazmowe, trawione itp.) lub powierzchniami bioaktywnymi poprawiają osteointegrację, ponieważ są ściśle związane z kością, redukując w ten sposób względne ruchy powodujące wydłużenie procesu gojenia się kości [3, 5].

3. OSADZANIE TYTANU ZA POMOCĄ ROZPYLANIA MAGNETRONOWEGO

Proces rozpylania magnetronowego zwanego również jonowym (ang. *sputtering methods*), polega na uwolnieniu atomów z materiału osadzanego za pomocą cząstek gazu obojętnego (najczęściej argonu), na które działa energia kinetyczna. Powstające w obszarze plazmy cząsteczki bombardują materiał źród-

dłowy tzw. target, który zazwyczaj ma postać płaskiego walca. Wybite z niego atomy osadzają się na całej powierzchni komory próżniowej, w której odbywa się proces, a także na znajdujących się wewnątrz płytkach podłożowych (Rys. 4). Takie działanie powoduje nierównomierne zużywanie się targetu, zanieczyszczenie komory i zużywanie większej niż potrzebna ilości materiału. Jednak powstająca na podłożach warstwa metalizacji posiada równomierną grubość i strukturę na całej powierzchni podłoża [2, 16].

Najważniejszymi parametrami, które charakteryzują warstwy wytworzone tą metodą jest chemiczny skład, grubość, rezystywność i jednorodność warstwy. O tym, który parametr jest najważniejszy, decyduje docelowe zastosowanie uzyskanej warstwy. Przed wyborem podłoża należy sprawdzić jego kompatybilność z nanoszonym materiałem oraz to, czy wchodzi w reakcje chemiczne z badanymi komórkami i pożywką używaną do ich hodowli [2, 8, 13, 16].

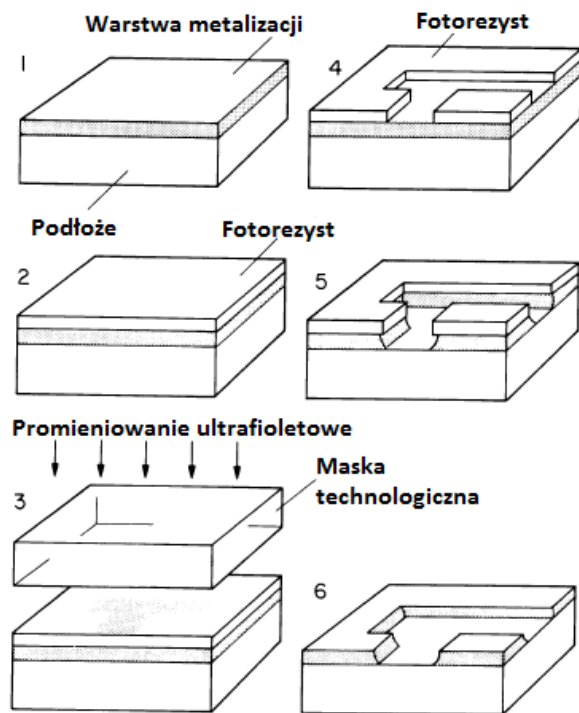


Rys. 4. Schemat komory próżniowej do rozpylania magnetronego [15]

4. PROCES FOTOLITOGRAFII

Wytwarzanie układów scalonych wymaga zastosowania metody tworzenia precyzyjnych i dokładnych wzorów na podłożach krzemowych, które wyznaczają m.in. obszary domieszkowania lub obszary wewnętrznych połączeń. Na rysunku 5 przedstawiono kolejne kroki procesu fotolitografii, który pozwala odwzorować zaprojektowany kształt. Wymaga to wykonania maski technologicznej z odpowiednim dla danej warstwy wzorem układu scalonego, a następnie przeniesieniu go na warstwę emulsji światłoczułej przy użyciu światła ultrafioletowego o długości fal od 350 do 430nm. Obecnie stosuje się kilka metod litografii, a należą do nich fotolitografia UV, litografia elektronowa, litografia rentgenowska i litografia jonowa. O ostatecznym wyborze metody decyduje nie tylko zaplecze technologiczne, lecz również efektywność kosztowa związana

z produkcją układów scalonych. Obecnie proces litografii stanowi mniej niż 10% kosztów gotowego wyrobu. Jeżeli jednak zostanie wybrana technologia wymagająca dużego nakładu w sprzęt o niskiej przepustowości, proces mógłby stać się głównym czynnikiem kształtującym cenę gotowego układu [4, 9].

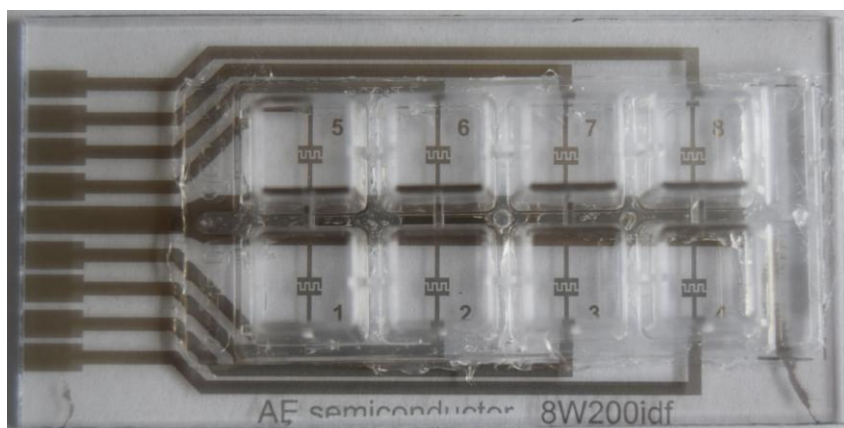


Rys. 5. Schemat procesu fotolitografii [9]

Najczęstszą techniką fotolitograficzną jest druk kontaktowy. Wymaga ona trzymania maski tuż nad powierzchnią płytki i wizualnego dopasowania jej do poprzedniego wzoru znajdującego się na płytce. Po osiągnięciu wyrównania maska zostaje dociśnięta do twardego styku z warstwą fotorezystu, który następnie poddany jest działaniu światła ultrafioletowego przechodzącego przez maskę. Metoda może być zmodyfikowana przez wprowadzenie wolnej przestrzeni pomiędzy maską a płytką. Technika ta, znana jako druk miękkiego kontaktu lub druk zbliżeniowy, niweluje uszkodzenia maski i płytki spowodowane kontaktem, na rzecz gorszej rozdzielczości naświetlanego wzoru. Nową ścieżką rozwoju fotolitografii jest naświetlanie projekcyjne, w którym obraz rzutowany jest na płytkę za pośrednictwem systemu optycznego. Technika ta poprawia rozdzielczość obrazu, a dzięki niewystępowaniu styku nie uszkadza masek ani podłoża [4, 9].

5. KONSTRUKCJA I TECHNOLOGIA ELEKTROD Z TYTANU

Pierwszym etapem prac było zaprojektowanie kształtu elektrody do pomiaru impedancji. Zdecydowano, że do pierwszych testów zostaną wykorzystane elementy grzebieniowe umieszczone po 8 sztuk na pojedynczym podłożu współpracującym z systemem ECIS[®]. Następnie, zaprojektowane zostały maski technologiczne, pozwalające na wykonanie procesu fotolitografii. Odległości pomiędzy poszczególnymi palcami kondensatora oraz ich szerokość wynosi 200 μm . Kolejnym krokiem był wybór podłoża, na którym zostanie osadzona metalizacja. Po wielu eksperymentach badających możliwości stosowania cieczy do trawienia na różnych podłożach, wybrano poliwęglan, który był obojętny dla koniecznych płynów technologicznych na dowolnym etapie wytwarzania oraz monitorowania hodowli komórek. Na podłoże napyłono metodą rozpylania magnetronowego warstwę tytanu o grubości 50nm, napyłarką NANO 36[™] firmy Kurt J. Lesker[®] (rys.6), po czym wykonano proces fotolitografii metodą pozytywową. Na warstwę metalizacji rozprowadzano emulsję, która po naświetleniu ulega wypłukaniu w odpowiednim wywoływaczu. Ostatni etap polegał na wytrawieniu odsłoniętej warstwy metalizacji (rys. 7). Etap ten wymagał dobrania zestawu parametrów trawienia metalizacji. A największym problemem, rozwiązywanym metodą doświadczalną, było zachowanie równowagi między trawieniem metalizacji a fotorezystu, gdyż zbyt silne stężenie substancji trawiących tytan powodowały utratę kształtu elektrod lub wręcz ich całkowity zanik.



Rys. 7. Podłoże z poliwęglanu z tytanowymi elektrodami przeznaczonymi do monitorowania parametrów elektrycznych podczas hodowli komórek



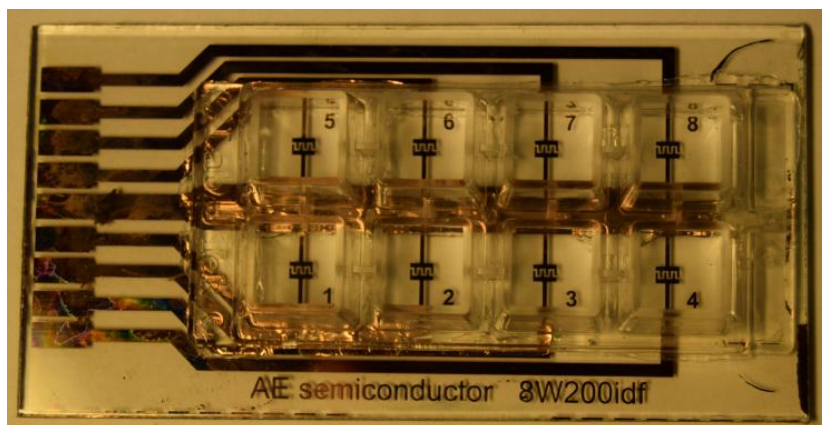
Rys. 3.4. Napyłarka NANO 36™[14]

Po otrzymaniu zadowalających struktur, zmierzono rezystancje doprowadzeń, która wynosiła ok. $20\text{k}\Omega$. Była to zbyt wysoka wartość, która uniemożliwia przeprowadzenie badań pomiaru impedancji komórek. Tak duża rezystancja doprowadzeń wynika najprawdopodobniej z efektu utlenienia się tytanu podczas trawienia. Wartość oczekiwana powinna być jak najniższa, jednocześnie nie może przekraczać kilkuset omów, ze względu na rezystancję pożywki stosowaną do hodowli komórek wynoszącą, przy używanych objętościach studzienek, ok. $2\text{k}\Omega$.

6. KONSTRUKCJA I TECHNOLOGIA ELEKTROD Z MIEDZI

Równolegle prowadzono prace pozwalające na wytworzenie struktur zawierającymi kondensatory grzebieniowe wykonane w warstwie miedzi. Wykonano przyrządy zarówno na podłożu z poliwęglanu, z wykorzystaniem osadzania w procesie rozpylania magnetronowego warstw miedzi o grubościach $0,1\text{--}1,0\mu\text{m}$ (rys. 8), jak i na klasycznym podłożu PCB (ang. *Printed Circuit Board*) z warstwą miedzi o grubości $35\mu\text{m}$.

Rezystancje poszczególnych doprowadzeń nie przekraczały kilku omów, w związku z tym przeprowadzono pierwsze testy z zastosowaniem systemu ECIS®. Podłoża zostały zamontowane w specjalnym uchwycie (rys. 9) umożliwiającym połączenie elektryczne do elektrod sprzężone z dedykowanym oprogramowaniem.

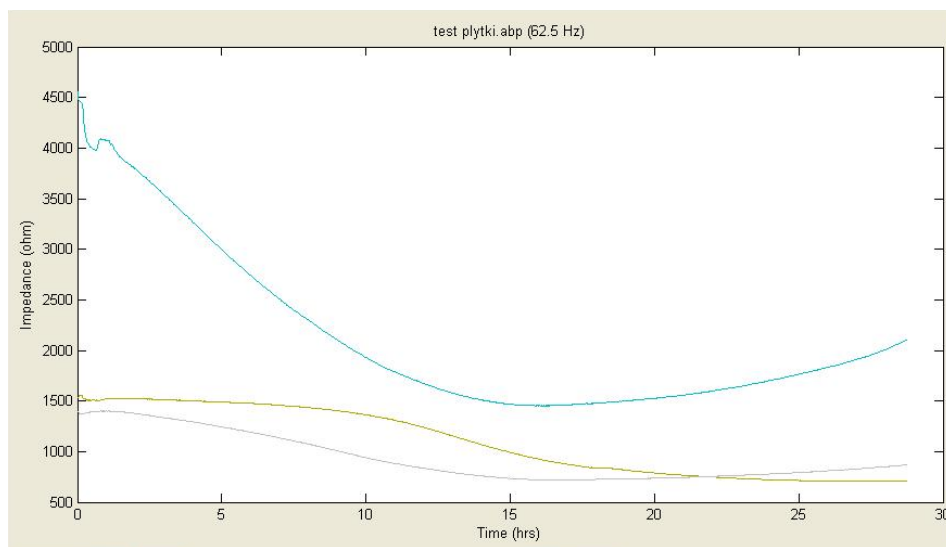


Rys. 8. Podłoże z poliwęglanu z miedzianymi elektrodami przeznaczonymi do monitorowania parametrów elektrycznych podczas hodowli komórek

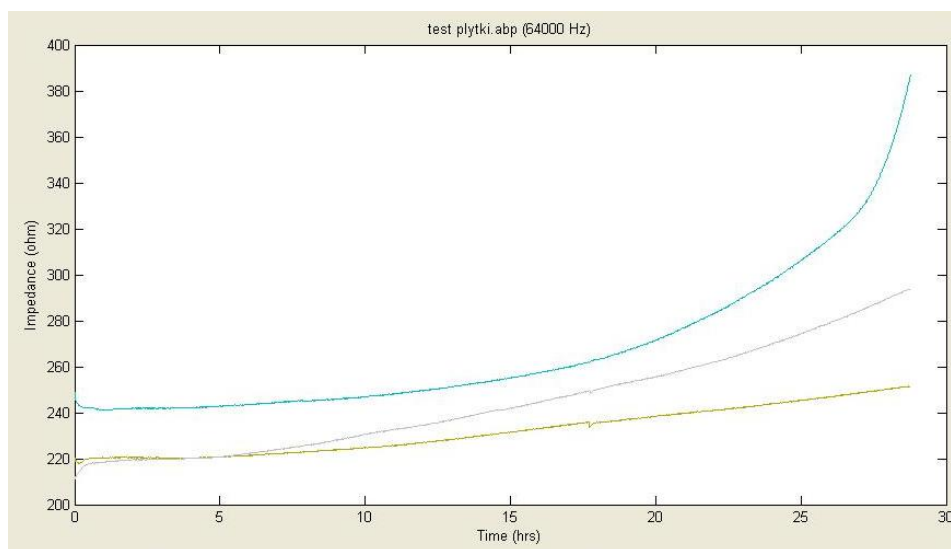


Rys. 9. Zestaw elektrod umieszczonych w uchwycie systemu ECIS®

Poprawne wyniki udało się uzyskać tylko na podłożu PCB, ze względu na reakcję miedzi z pożywką, która powodowała uszkodzenia cienkich warstw metalizacji. Trzy studzienki zostały zalane pożywką zawierającą dodatkowo antybiotyk oraz surowicę i po 24 godzinach w jednej z nich umieszczono fibroblasty. Hodowlę monitorowano przez 28 godzin, w trakcie której wykonano pomiary impedancji komórek zanurzonych w pożywce. Na rysunkach 10 i 11 przedstawiono uzyskane charakterystyki impedancji w funkcji czasu przy skrajnych częstotliwościach 62,5Hz oraz 64kHz. Linia niebieską oznaczono studzienkę zawierającą fibroblasty, natomiast linie beżowa i szara wyznaczają impedancję odniesienia – studzienki zawierającą tylko pożywkę (bez hodowli komórek).



Rys. 10. Wyniki pomiaru impedancji dla częstotliwości 62,5 Hz podczas hodowli komórek (linia niebieska) oraz pomiaru impedancji odniesienia dla pożywki bez hodowli (linie beżowa i szara)



Rys. 11. Wyniki pomiaru impedancji dla częstotliwości 64 kHz podczas hodowli komórek (linia niebieska) oraz pomiaru impedancji odniesienia dla pożywki bez hodowli (linie beżowa i szara)

7. PODSUMOWANIE

W ramach pracy zaprojektowano i wykonano struktury zawierające kondensatory grzebieniowe z tytanu i miedzi na podłożu z poliwęglanu oraz PBC. Konstrukcja przyrządów umożliwia zamontowanie ich w systemie ECIS[®] do monitorowania parametrów elektrycznych podczas hodowli komórek. Uzyskano poprawne struktury na podłożu PCB z warstwą miedzi o grubości 35 μm , które posłużyły do przetestowania ich pracy podczas hodowli komórek. Otrzymane wyniki wskazują, że przyrządy zostały wykonane poprawnie pod względem technologicznym i elektrycznym.

W dalszych etapach prac zaplanowano wykonać struktury na poliwęglanie z grubszą warstwą tytanu (> 500nm) oraz innymi metalami np. nikiel, aluminium oraz chromonikielina.

LITERATURA

- [1] *Applied BioPhysics*, Product Guide, Corporate Headquarters: 185 Jordan Road Troy, NY 12180 1-866-301-ECIS (3247)
- [2] Beck R., *Technologia krzemowa*, PWN, Warszawa, 1991
- [3] Liua Xuanyong, b, Paul K. Chub, Chuanxian Dinga, *Surface modification of titanium, titanium alloys, and related materials for biomedical applications*, "Materials Science and Engineering" R 47 (2004) p. 49–121

- [4] Napieralska M., Jabłoński G., *Research and Training Action for System on Chip Design IST-2000-30193*, Podstawy Mikroelektroniki Łódź 2002
- [5] Niinomi Mitsuo, *Properties of biomedical titanium alloys*, *Materials Science and Engineering*, "Mechanical" A243 (1998) p. 231–236
- [6] Prendecka M., Mlak R., Małecka-Massalska T., *Effect of exopolysaccharide from Ganoderma applanatum on the electrical properties of mouse fibroblast cells line L929 culture using an electric cell-substrate impedance sensing (ECIS) – Preliminary study*, „Annals of Agricultural and Environmental Medicine” 2016, vol. 23, nr 2, s. 293–297
- [7] Rack H.J., Qazi J.I., *Titanium alloys for biomedical applications*, *Materials Science and Engineering C* 26 (2006) 1269–1277
- [8] Stolwijk J. A., Matrougui K., *Impedance analysis of GPCR – mediated changes in endothelial barrier function: overview and fundamental considerations for stable and reproducible measurements*, „Pflugers Arch – Eur J Physiol” 2015, nr 467, s. 2193–2218
- [9] Thompson L.F., *An Introduction to Lithography*, Publication Date: May 3, 1983, <http://pubs.acs.org> [dostępny 01.06. 2017]
- [10] Veiga C., Davim J.P. and Loureiro A.J.R., *Properties and applications of titanium alloys*, “Rev. Adv. Mater. Sci.” 32 (2012) 133–148
- [11] Wang Kathy, *The use of titanium for medical applications in the USA*, *Materials Science and Engineering*” A213 (1996) I34–I37
- [12] Wegener J, Keese CR, Giaever I., *Electric cell-substrate impedance sensing (ECIS) as a noninvasive means to monitor the kinetics of cell spreading to artificial surfaces*, „Exp Cell Res.” 2000, nr 259(1), s. 158–166
- [13] Montaż elektroniczny i systemy testujące: http://home.agh.edu.pl/~maziarz/Lecture_1-wire-bonding-en.pdf (zasoby z dnia 01.06. 2017)
- [14] http://www.lesker.com/newweb/vacuum_systems/deposition_systems_pvd_nano36_sputter.cfm#, (zasoby z dnia 01.06. 2017)
- [15] http://www.mif.pg.gda.pl/homepages/maria/pdf/MF_warstwy.pdf, (zasoby z dnia 01.06. 2017)
- [16] <http://www.pcb007.com/pages/zone.cgi?a=78806> (zasoby z dnia 01.06. 2017).

WYBRANE PROBLEMY W DIAGNOSTYCE OBWODU STEROWANIA SILNIKIEM TDI

1. WSTĘP

Wzrost wymagań dotyczących norm emisji spalin wymusił na producentach aut zastosowanie elektronicznych systemów nadzorujących pracę silnika w celu zmniejszenia zanieczyszczeń emitowanych do atmosfery. W ramach standardu OBD pojawiła się możliwość monitorowania parametrów pracy jednostki napędowej w czasie rzeczywistym, a także odczyt zapisanych w pamięci sterownika błędów powstałych wskutek uszkodzeń [1].

Niejednokrotnie diagnostyka ograniczona jedynie do odczytu kodów DTC zawartych w jednostce sterującej nie gwarantuje usunięcia usterki. Dzieje się tak z powodu ograniczeń w monitorowaniu pracy mechanicznych elementów wykonawczych. Ponadto ze względu na liczne powiązania pomiędzy czujnikami a sterownikami następuje sytuacja, gdzie niesprawność jednego czujnika rzutuje na większą ilość obwodów sterowania. Aby poprawnie zdiagnozować realną przyczynę usterki, należy wykonać dodatkowe odczyty poszczególnych wartości mierzonych. Pozwoli to na trafną diagnozę uszkodzenia.

2. METODYKA BADAŃ

Na potrzeby badań zasymulowano usterki układu EGR oraz układu doładowania poprzez uszkodzenie obwodu sterowania poszczególnych zespołów sterujących. Następnie przebadano wpływ usterek na poszczególne parametry pracy silnika.

2.1. Obiekt badań

Obiektem badań było Audi A4 B5 z 1999 roku (Rys. 1), wyposażone w silnik wysokoprężny o pojemności 1.9 TDI. Jednostka generuje moc 90KM przy 4000obr/min oraz moment 202Nm przy 1900obr/min.

¹ Politechnika Lubelska, Elektrotechnika



Rys. 1 Obiekt badań

2.2. Narzędzia diagnostyczne

Do badań wykorzystano komputer przenośny z systemem Windows, oprogramowanie diagnostyczne VCDS Lite [2] w wersji 1.2 oraz interface diagnostyczny z wtykiem OBD-II [3] (rys. 2).



Rys. 2 Interface diagnostyczny

2.3. Usterka układu recyrkulacji spalin

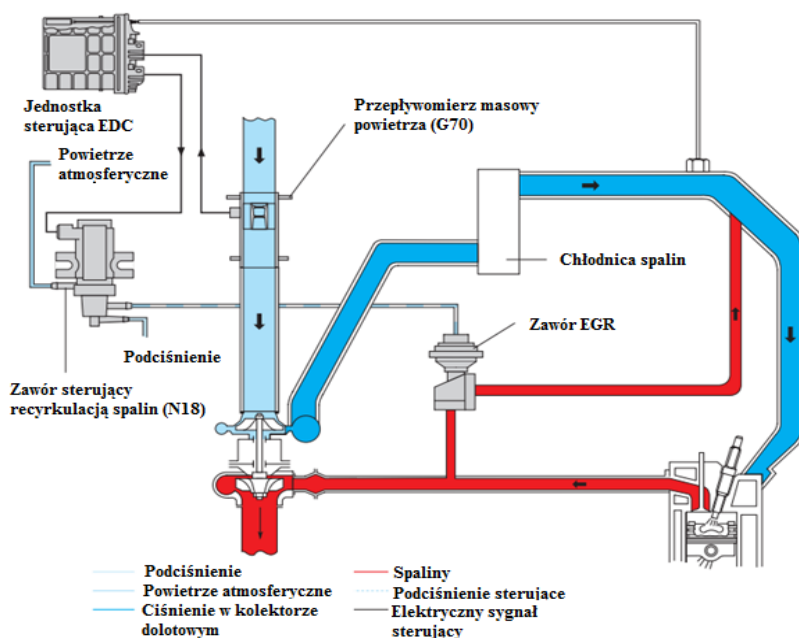
Układ EGR (Rys. 3) odpowiada za obniżenie poziomu emisji spalin silnika w dolnych partiach obrotów. Usterki, jakie mogą wystąpić w obwodzie sterowania można podzielić na dwie grupy:

Diagnostowalne bezpośrednio:

- przerwanie wiązki elektrycznej pomiędzy sterownikiem a zaworem recyrkulacji spalin
- uszkodzenie zaworu sterującego recyrkulacją spalin lub przepływomierza
- brak sygnału z przepływomierza masowego powietrza.

Diagnozowalne pośrednio:

- przerwanie węży podciśnieniowych
- uszkodzenie pompy vacuum
- zanieczyszczenie zaworu nagarem powodujące unieruchomienie go w pozycji zamkniętej bądź otwartej.



Rys. 3 Schemat układu recyrkulacji spalin [4]

2.4. Usterki układu doładowania

Układ doładowania (rys. 4) dostarcza odpowiednią ilość powietrza do komory spalania dzięki czemu silnik jest w stanie wygenerować większą moc. Usterki układu doładowania można podzielić na:

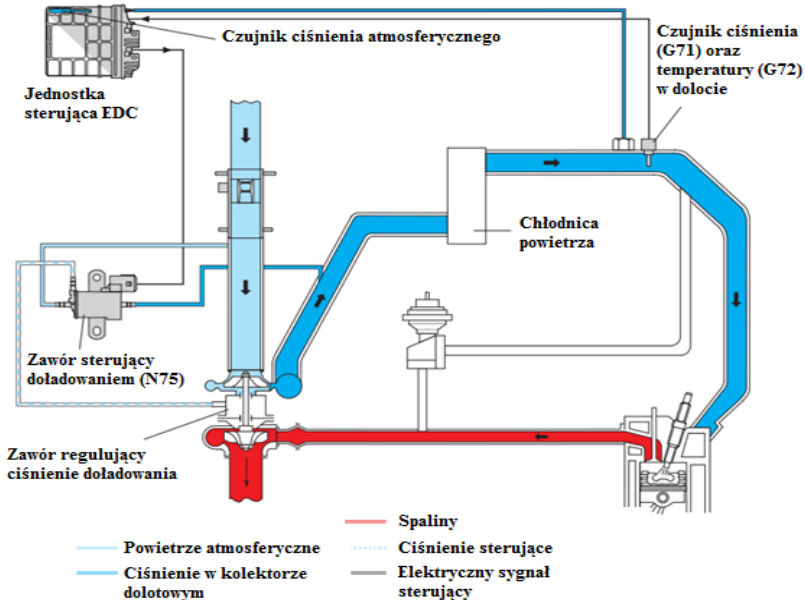
Diagnozowalne bezpośrednio:

- przerwanie wiązki elektrycznej pomiędzy sterownikiem a zaworem sterującym ciśnieniem doładowania
- uszkodzenie zaworu sterującego doładowaniem
- uszkodzenie czujnika ciśnienia lub czujnika wysokości.

Diagnozowalne pośrednio:

- przerwanie węży ciśnieniowych
- unieruchomienie zaworu Wastegate bądź kierownic przepływu (turbo-sprężarka VNT)
- uszkodzenie siłownika pneumatycznego

- mechaniczne uszkodzenie koła kompresji, wirnika turbosprężarki
- rozszczelnienie układu dolotowego
- zapieczenie zaworu EGR w pozycji otwartej.



Rys. 4 Schemat układu doładowania [4]

2.5. Przebieg badań

Przed przystąpieniem do badań należy sprawdzić, czy zostały spełnione następujące warunki:

- temperatura płynu chłodzącego minimum 80°C
- wyłączone wszelkie odbiorniki elektryczne jak np. radio, nawiew, ogrzewanie tylnej szyby, lusterek
- wyłączony układ klimatyzacji
- jeżeli w pamięci sterownika znajdują się błędy, należy je najpierw zapisać za pomocą programu, następnie usunąć. Jeśli mimo kasowania błędy pozostają w pamięci, oznacza to usterkę danego układu sterującego, co bezpośrednio przekłada się na wyniki pomiarów.

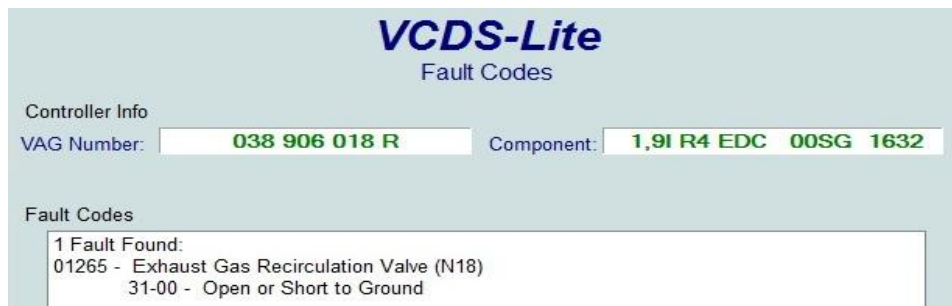
Badania podzielono na trzy części:

- analiza błędów DTC zapisanych w pamięci sterownika
- próba statyczna w tym analiza nastaw podstawowych – wykonywana podczas pracy silnika na biegu jałowym
- próba dynamiczna – przeprowadzona w warunkach drogowych przy pełnym obciążeniu silnika.

3. WYNIKI BADAŃ

3.1. Usterki układu recyrkulacji spalin

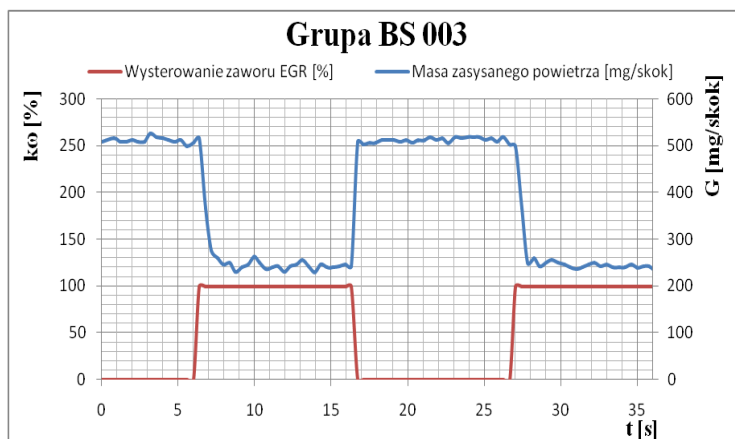
W przypadku uszkodzenia układu EGR występuje jedynie błąd związany z przzerwaniem wiązki elektrycznej pomiędzy sterownikiem a modulatorem podciśnienia. Wszelkie inne mechaniczne błędy nie są wykrywane przez sterownik (rys. 5).



Rys. 5 Błąd układu EGR wykryty przez sterownik silnika

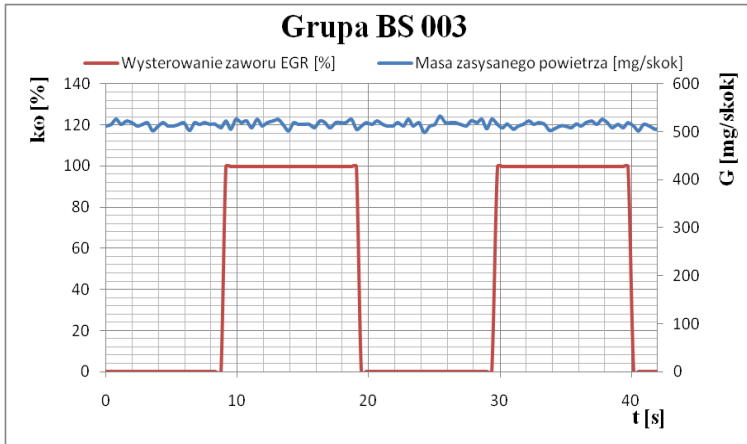
Analiza wartości mierzonych w bloku 003 podczas próby statycznej pozwala zaobserwować, czy przepływomierz masowy powietrza oraz zawór recyrkulacji spalin pracuje poprawnie.

W zależności od procentu wypełnienia impulsu sterującego otwarciem zaworu ERG zmienia się masa zasysanego powietrza, gdyż spaliny trafiają do kolektora ssącego. W ten sposób spada ilość przepływającego świeżego powietrza co przedstawia rysunek 6.



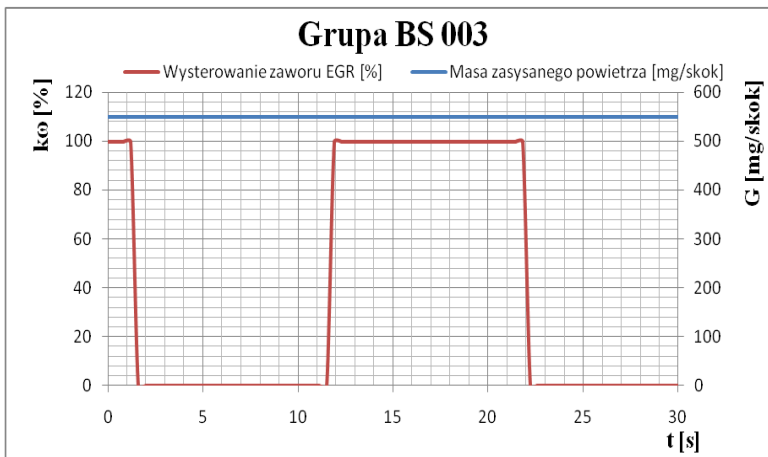
Rys. 6 Poprawnie działający układ EGR

W przypadku mechanicznego uszkodzenia obwodu sterowania, masa zasysanego powietrza nie zmienia się pomimo zmian wysterowania zaworu sterującego recyrkulacją spalin (Rys. 7).



Rys. 7 Uszkodzony układ EGR

Jeśli uszkodzeniu ulegnie przepływomierz, układ EGR nie będzie funkcjonował mimo zmian wysterowania zaworem sterującym recyrkulacją (rys. 8). Ponadto nie obserwuje się wahań wartości zasysanego powietrza typowych dla sprawnego przepływomierza.



Rys. 8 Uszkodzony przepływomierz masowy powietrza

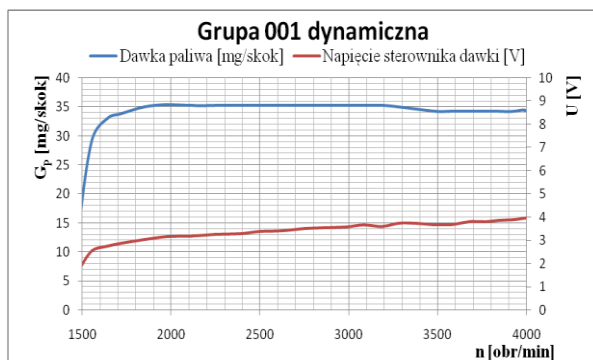
3.2. Usterki układu doładowania

Gdy uszkodzeniu ulegnie układ doładowania, sterownik jest w stanie rozróżnić więcej awarii, gdyż sam układ bazuje się na większej ilości czujników niż układ EGR. W tabeli 1 przedstawiono przykładowe błędy wraz opisami.

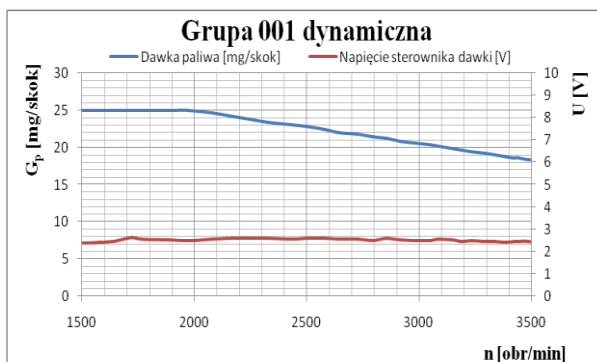
Tab. 1. Przykładowe błędy wykrywane przez sterownik

Nr błędu	Opis
17965	Charge pressure control positive deviation
01262	Solenoid Valve for Boost Pressure Control N75 – Open or Short to Ground
17564	Manifold Pressure Sensor (G71) – Open/Short to Ground
00528	Barometric Pressure Sensor (F96) – Open or Short to Ground

Do diagnostyki układu doładowania zastosowano próbę dynamiczną, gdyż na biegu jałowym silnik nie jest w stanie wygenerować tyle spalin, by turbosprężarka osiągnęła ciśnienie nominalnej pracy. Grupa 001 przedstawia zależność dawki paliwa w zależności od obrotów silnika (rys. 9).



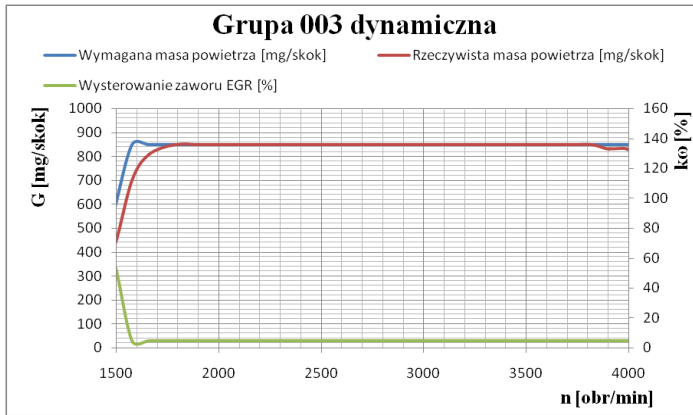
Rys. 9 Dawka paliwa przy braku uszkodzeń



Rys. 9 Dawka paliwa przy uszkodzeniu układu regulacji ciśnienia doładowania

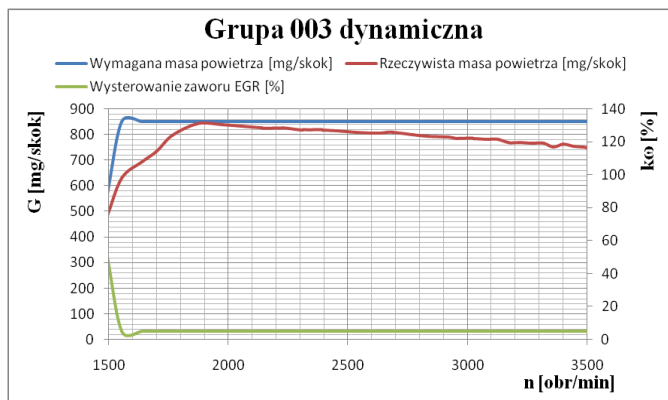
W przypadku uszkodzenia układu sterującego doładowaniem następuje zmniejszenie dawki paliwa w wyższych partiach obrotów, co bezpośrednio wpływa na moc generowaną przez silnik (rys. 10).

Wartość masy zasysanego powietrza w przypadku sprawnej jednostki wynosi około 850mg/ skok od około 1800obr/min, kiedy to turbosprężarka osiąga zadaną wartość doładowania (rys. 11).



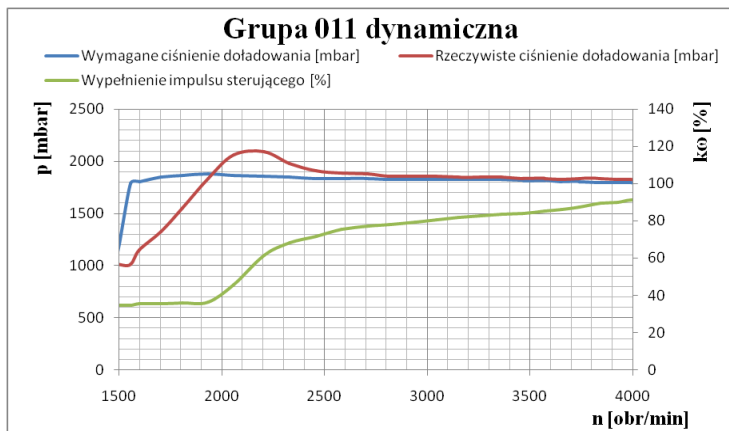
Rys. 10 Masa zasysanego powietrza przy braku uszkodzeń

Uszkodzenie układu sterującego doładowaniem wpływa na obniżenie wartości masy zasysanego powietrza poniżej wartości wymaganej przez mapę. Różnica widoczna jest w pełnym zakresie obrotów i nie wynika z niesprawności przepływomierza masowego powietrza (rys. 12).



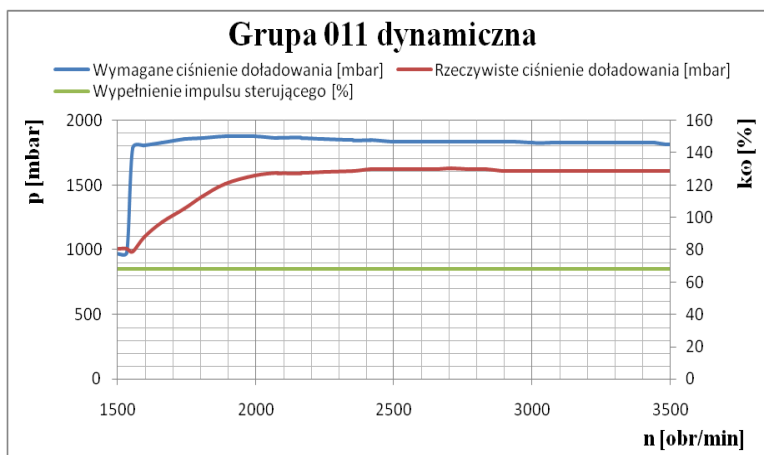
Rys. 11 Masa zasysanego powietrza w przypadku uszkodzenia układu sterowania ciśnienia doładowania

Ciśnienie doładowania sprawnego układu wynosi około 1850mbar (rys. 13) i jest osiągnięte przy około 2000obr/min.



Rys. 12 Ciśnienie doładowania przy braku uszkodzeń

Rzeczywista wartość ciśnienia doładowania w przypadku uszkodzenia znacząco odbiega od wartości zadanej w pełnym zakresie obrotów. Świadczy to jednoznacznie z problemami układu doładowania (rys. 14).



Rys. 13 Ciśnienie doładowania w przypadku uszkodzenia układu sterowania ciśnienia doładowania

5. MODERNIZACJA UKŁADÓW WYKONAWCZYCH

W przypadku siłowników sterowanych pneumatycznie jednostka sterująca nie mogła określić bezpośrednio, czy usterce uległ aktuator czy przewód podciśnieniowy. Zastosowanie nastawników elektrycznych pozwoliło nie tylko uprościć układ sterowania poprzez eliminację pneumatyki, ale także pozwoliło na uzyskanie sygnału zwrotnego o położeniu nastawnika. Dzięki temu sterownik jest w stanie dokładniej określić gdzie występuje błąd, co sprzyja pracy diagnosty przy poszukiwaniu uszkodzeń w silniku.

5.1. Modernizacja zaworu EGR

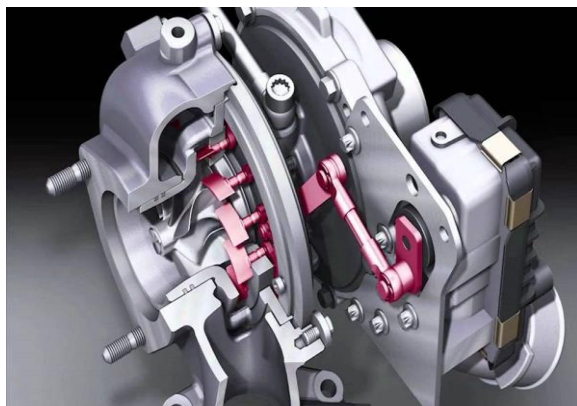
Pneumatyczny zawór recyrkulacji spalin był podatny na zapiekanie się oraz nie wysyłał informacji w postaci elektrycznej odnośnie aktualnej pozycji położenia. Problem ten rozwiązano za pomocą potencjometru, który informował w jakiej pozycji znajduje się zawór, jednak sterowanie nadal bazowało na podciśnieniu. Ostatecznie zbudowano zawór w pełni sterowalny elektronicznie posiadający funkcję samooczyszczania, dzięki czemu zminimalizowano ryzyko zapiekania zaworu. Rysunek 15 przedstawia poszczególne zawory EGR.



Rys. 14 Od lewej, pneumatyczny zawór EGR, pneumatyczny zawór EGR z rozpoznaniem położenia, elektryczny zawór EGR dwudrogowy [5]

5.2. Modernizacja układu sterowania ciśnieniem doładowania turbosprężarki

W przypadku sterowania doładowaniem zastosowano pozycjonowanie położenia kierownic przepływu za pomocą nastawnika elektronicznego (rys. 16). Ponadto stosuje się czujnik prędkości wału kompresora, by zabezpieczyć turbosprężarkę przed nadmiernymi obrotami a w konsekwencji przed uszkodzeniem mechanicznym spowodowanym zbyt dużą prędkością obrotową.



Rys. 15 Turbosprężarka z nastawnikiem elektronicznym [6]

5. PODSUMOWANIE

Diagnostyka komputerowa w znacznym stopniu ułatwia lokalizację usterek elektronicznych jak i mechanicznych pod warunkiem, że zostaną przyjęte odpowiednie kryteria oceny.

W przypadku układu recyrkulacji spalin nie wystarczy odczyt kodów DTC zawartych w sterowniku silnika. Aby określić przyczynę awarii, należy przeprowadzić pomiar statyczny parametrów pracy silnika, w szczególności grupy 003. Pomiar dynamiczny jedynie potwierdza sprawność bądź niesprawność układu recyrkulacji spalin.

W układzie doładowania występuje więcej czujników elektrycznych, a co za tym idzie sterownik jest w stanie rozpoznać więcej błędów. Podczas próby statycznej nie zaobserwowano żadnych odchyleń od wartości zadanych. Aby zdiagnozować poprawność działania układu, niezbędny jest pomiar dynamiczny parametrów pracy silnika.

LITERATURA

- [1] Riehl H.J., *Elektrotechnika i elektronika w pojazdach samochodowych*, Wydaw. Komunikacji i Łączności, Warszawa 2003
- [2] <http://www.ross-tech.com/vcds-lite/download/index.php> (zasoby z 15.08.2015)
- [3] <http://www.viaken.pl/sklep.php?lang=pl&kat=5&prodid=103> (zasoby z 18.11.2015)
- [4] Volkswagen group.: *1.9ltr TDI Industrial Engine Technical status 4/1999*
- [5] Pierburg., *Przepływomierze masowe powietrza*
- [6] <http://dailydriver.pl/porady/technika/jak-dziala-turbosprezarka-ze-zmienna-geometria-kierownicy-vgt/> (zasoby z 16.11.2015)

SPONSORZY I INSTYTUCJE WSPIERAJĄCE
VII SYMPOZJUM NAUKOWE
ELEKTRYKÓW I INFORMATYKÓW



Dziękujemy !

PATRONI VII SYMPOZJUM NAUKOWEGO ELEKTRYKÓW I INFORMATYKÓW

Patronat Lubelskiego Oddziału Stowarzyszenia Elektryków Polskich



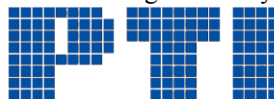
Patronat Jego Magnificencji Rektora Politechniki Lubelskiej



Patronat Dziekana Wydziału Elektrotechniki i Informatyki



Patronat Lubelskiego Koła Polskiego Towarzystwa Informatycznego



Patronat Lubelskiego Oddziału Polskiego Towarzystwa Zarządzania Produkcją

