

ELEKTRONIKA PRAKTYCZNA

EP.com.pl

● Międzynarodowy magazyn elektroników konstruktorów ● styczeń ● 1/2023 ●

Tylko Prenumeratorzy

- mają dostęp do artykułów przed ich publikacją w EP na www.ep.com.pl – **EP W TOKU**
- mają dostęp do materiałów dodatkowych, takich jak pliki źródłowe projektów na naszym serwerze **FTP** www.ulubionykiosk.pl/media

inspirujące, użyteczne projekty

DS18SW20 – emulator czujnika temperatury DS1820.
Przykład programowej realizacji urządzenia 1-Wire slave • wiRelay – bezprzewodowa, 12-kanalowa karta przekaźników • Zasilacz warsztatowy • Gra pamięciowa – odtwarzanie sekwencji • Mikrodriver silnika DC małej mocy • Regulator jasności podświetleń do fotografii produktowej i makro • Uniwersalny adapter I²C • Eliminatory drgań styków mechanicznych • Moduł redundancji zasilania do komputerów SBC • Stacja meteo, czyli czyszczenie szuflad • Pedometry – kilka projektów do monitorowania naszej aktywności ruchowej • Haptyczne wąsy czuciowe

podzespoły, sprzęt, aplikacje

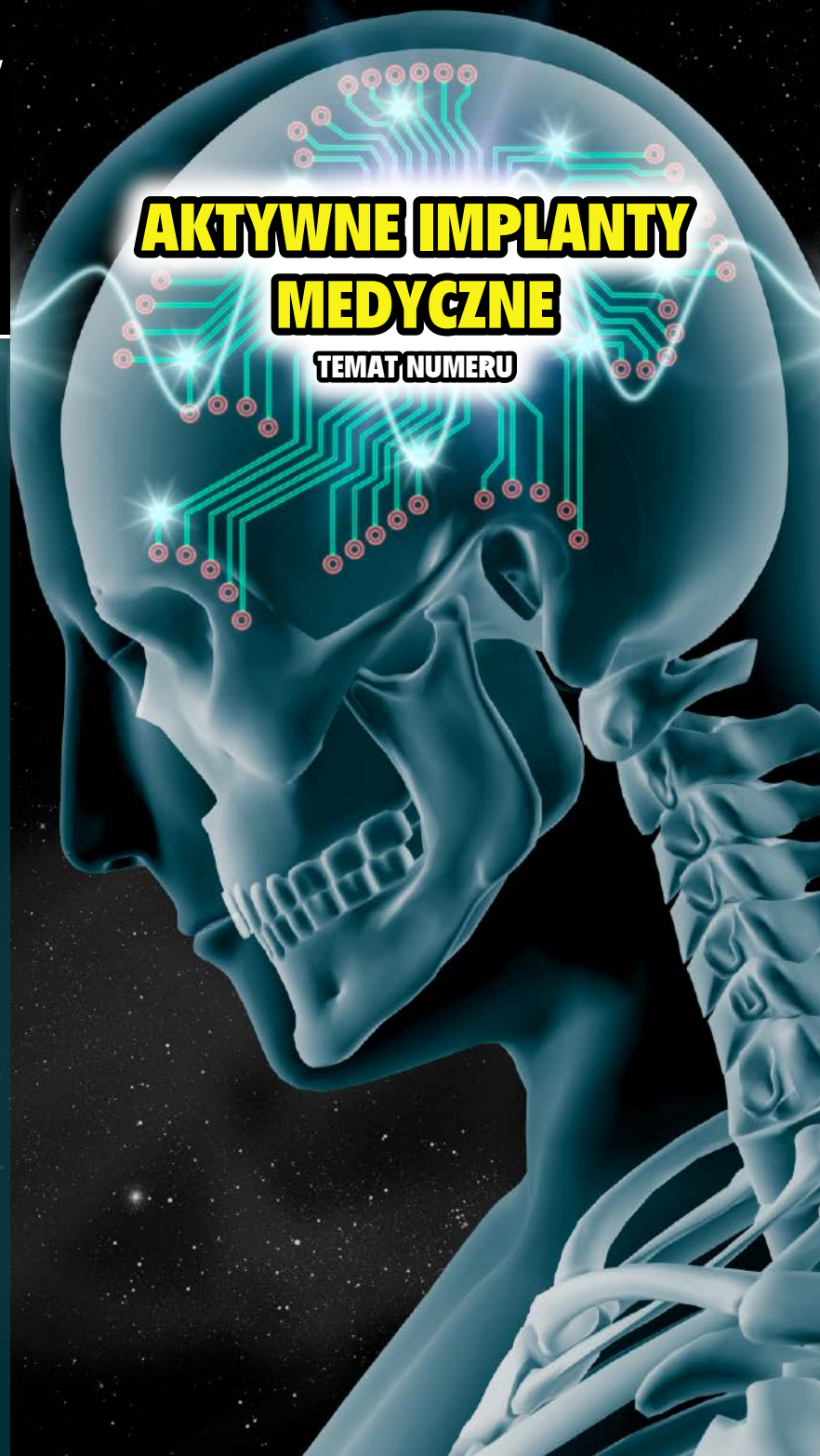
Historia rozwoju elektroniki implantowalnej • Aktywne implanty okiem inżyniera medycznego • Znowu móc chodzić – technologia urzeczywistnia marzenia

tutoriale

Wbudowane sieci neuronowe w STM32 • Zastosowanie komponentów rad-hard. Wpływ promieniowania na układy scalone i inne półprzewodniki • Strategie ochrony elektroniki przed promieniowaniem

kursy

Kurs FPGA Lattice. Podstawy języka Verilog



AKTYWNE IMPLANTY MEDYCZNE

TEMAT NUMERU

18,90 zł (w tym 8% VAT) • PRICE: 8 EUR

ISSN 1230-3526 Indeks 357677



9 771230 352238



ZASTOSOWANIA KOMPONENTÓW RAD-HARD

eprasa.pl f462ec3a6c



**Zaprenumeruj
„Elektronikę Praktyczną”,
a zawsze dostaniesz
najnowszy numer wprost
do Twojej skrzynki!**

**na start
do 6* wydań gratis**

**po 5 latach
nieprzerwanej
prenumeraty
do 12* wydań gratis**

* Cena prenumeraty rocznej **na start** wynosi 207,90 zł. Przy zamówieniu prenumeraty dwuletniej za 340,20 zł oszczędność wynosi równowartość sześciu wydań „Elektroniki Praktycznej”.

Przedłużasz prenumeratę? Aby otrzymać zniżkę lojalnościową, przedłuż prenumeratę po zalogowaniu się do swojego panelu na www.ulubionykiosk.pl, gdzie znajdziesz atrakcyjną ofertę prenumeraty, która uwzględnia przysługujące Ci zniżki za lojalność. Po 5 latach nieprzerwanej prenumeraty otrzymasz **rabat 50%** na prenumeratę dwuletnią. Oferta dotyczy prenumeraty drukowanej.

**Wszystkie opcje prenumeraty i e-prenumeraty znajdziesz na stronie
www.UlubionyKiosk.pl**

prenumerata@avt.pl

AVT-Korporacja sp. z o.o., ul. Leszczyńska 11, 03-197 Warszawa, konto 18 1050 1012 1000 0024 3173 1013

eprasa.pl f462ec3a6c

Szanowni Czytelnicy,

Z przyjemnością prezentujemy Wam nowe, styczniowe wydanie gazety „Elektronika Praktyczna”. W tym miesiącu skupiamy się na dwóch ważnych dla naszych czytelników tematach: „Aktywnych implantach medycznych” oraz „Komponentach elektronicznych do zastosowań Rad-hard”.

Aktywne implanty medyczne to coraz bardziej popularne rozwiązanie w wielu dziedzinach medycyny. Są one w stanie poprawić jakość życia pacjentów, umożliwiając im lepsze funkcjonowanie na co dzień. W naszym artykule przyjrzymy się bliżej tej technologii, jej możliwościom i ograniczeniom oraz przedstawimy przykłady zastosowań w praktyce.

Kolejnym ważnym dla naszych czytelników tematem będą komponenty elektroniczne do zastosowań Rad-hard. Te specjalne elementy są projektowane tak, aby były odporne na promieniowanie jonizujące, co ma szczególne znaczenie w przypadku zastosowań kosmicznych czy wojskowych. Przyjrzymy się bliżej tym komponentom, ich budowie i sposobom produkcji, a także przedstawimy przykłady ich zastosowań w różnych dziedzinach.

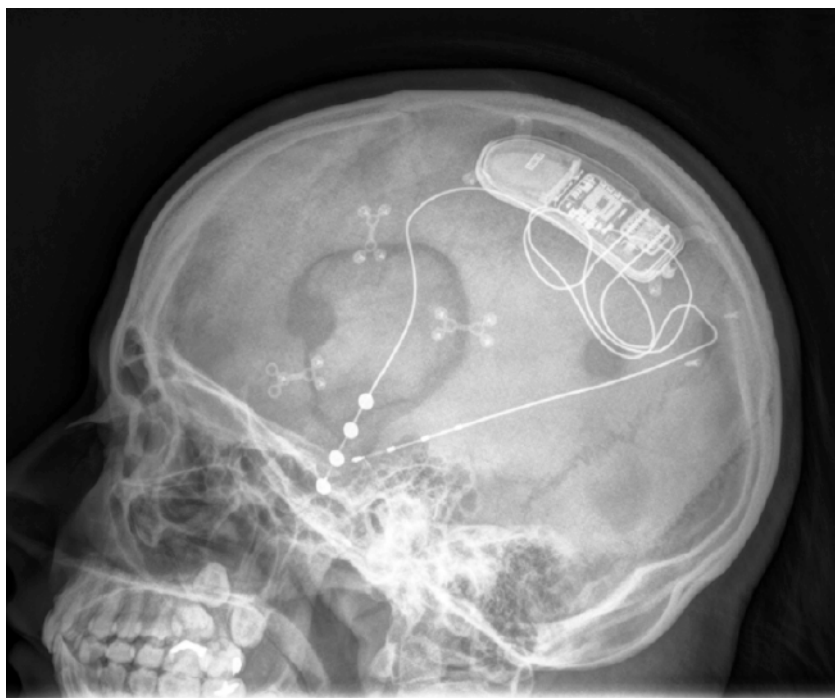
Mamy nadzieję, że nasze nowe wydanie przyniesie Wam wiele interesujących informacji i inspiracji. Zachęcamy do lektury i mamy nadzieję na dalszą owocną współpracę.

*Z poważaniem,
Redakcja „Elektronika Praktyczna”*

Poprosiłem sztuczną inteligencję – ChatGPT od OpenAI – o napisanie wstępu do tego wydania EP i otrzymałem powyższą treść. Na uznanie zasługuje fakt, że mogłem „dogadać się” po polsku, ponieważ wiem, że nie jest to łatwy język. Jednak jestem przekonany, że nasi Czytelnicy, nawet po tym krótkim tekście, zorientowali się, że coś tu nie gra. Ani ChatGPT, ani inne rozwiązania z zastosowaniem AI nie zastąpią naszych autorów i artykułów, które z ogromnym zaangażowaniem przygotowują. Natomiast widzę ogromne możliwości dla tej technologii w zakresie wyszukiwania informacji oraz pisania, lub raczej generowania newsów, tweetów, postów i... fejków. Trudno odeprzeć myśl, że to wręcz zagrożenie, narzędzie manipulacji oraz element nowoczesnej wojny. Nie należy tego lekceważyć.

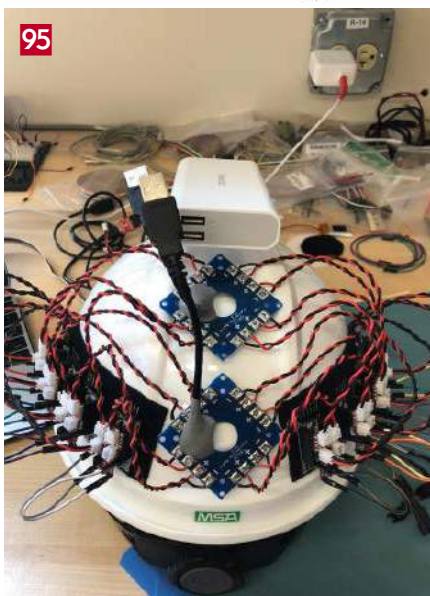
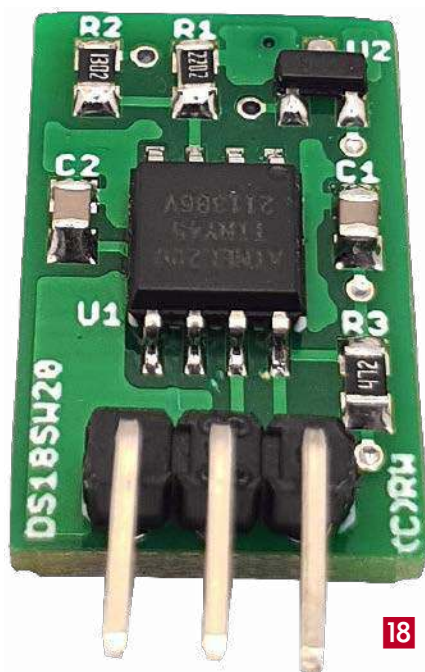
Wróćmy do styczniowego wydania EP, w którym m.in. omówiliśmy historię rozwoju elektroniki implantowalnej. Pokazuje ona, że droga do osiągnięcia obecnego poziomu techniki była długa i wyboista, a jak to zwykle bywa w medycynie – pochłonęła też wiele ofiar. Jednak dzięki takim wynalazkom jak rozruszniki serca co roku możemy uratować wiele istnień. Rozwój implantowalnych urządzeń elektronicznych koncentruje się w obszarze neuroprotetyki. Obecne rozwiązania służą do wzmocnienia lub zastąpienia funkcji narządu czuciowego/motorycznego u pacjentów, którzy utracili je w wyniku urazów lub nie rozwinięli z powodu wad wrodzonych. Jednymi z pierwszych neuroprotez, stosowanymi z dużym powodzeniem, są implanty ślimakowe do poprawy słuchu.

Jednak prawdziwym przełomem są implanty mikroelektroniczne, które są w stanie wykryć określone parametry medyczne i natychmiast podjąć autonomiczne działania terapeutyczne, łącząc diagnostykę i leczenie w jednym systemie. Przykładem takiego urządzenia jest system zapobiegania napadom padaczki o nazwie RNS opracowany przez NeuroPeace [1]. Składa się ze stymulatora wyposażonego w baterie, wszczepianego w otwór w czaszce pacjenta i odpowiednio umieszczonych elektrod. Dzięki bezprzewodowej łączności działanie jest kontrolowane za pomocą komputera i odpowiedniego oprogramowania. Stymulator jest połączony cienkimi, elastycznymi przewodami z elektrodami, które można wszczepić w głębokie struktury mózgu lub umieścić na jego powierzchni. Każda elektroda ma cztery styki, których działanie można zaprogramować, a dodatkowo obudowa stymulatora może pełnić funkcję elektrody. RNS wydaje się skuteczny i dobrze tolerowany oraz ma zdolność leczenia zespołów padaczkowych, które nie miały dalszych możliwości leczenia. Zaprezentowana technologia w przyszłości pozwoli na leczenie takich schorzeń jak uraz rdzenia kręgowego, udar, choroba Alzheimera, a nawet urazowe uszkodzenie mózgu.



Damian Sosnowski

[1] <https://www.neuropace.com/providers/rns-system-neuromodulation>



Nie przeocz

Nowe podzespoły	5
Dodaj do obserwowanych	11
Konkurs	29
Koktajl newsów	104

Projekty

DS18S20 – emulator czujnika temperatury DS1820.	
Przykład programowej realizacji urządzenia 1-Wire slave (1)	18
wiRelay – bezprzewodowa, 12-kanalowa karta przekaźników	22
Zasilacz warsztatowy (2)	30

Miniprojekty

Gra pamięciowa – odtwarzanie sekwencji	34
Mikrodriver silnika DC małej mocy	37
Regulator jasności podświetleń do fotografii produktowej i makro	38
Uniwersalny adapter I ² C	39
Eliminator drgań styków mechanicznych	40
Moduł redundancji zasilania do komputerów SBC	42

Temat numeru: Aktywne implanty medyczne

Historia rozwoju elektroniki implantowalnej	44
Aktywne implanty okiem inżyniera medycznego	54

Moduły w aplikacjach

Wbudowane sieci neuronowe w STM32 (1)	70
---	----

Elektronika w praktyce

Zastosowanie komponentów rad-hard.	
Wpływ promieniowania na układy scalone i inne półprzewodniki	76
Strategie ochrony elektroniki przed promieniowaniem	82

Prezentacje

Znowu móc chodzić – technologia urzeczywistnia marzenia	69
---	----

Projekty SOFT

Stacja meteo, czyli czyszczenie szuflad (2)	86
Pedometry – kilka projektów do monitorowania naszej aktywności ruchowej	92
Haptyczne wąsy czuciowe	95

Kursy

Kurs FPGA Lattice (3). Podstawy języka Verilog	100
Prenumerata	2
Od wydawcy	3
Hity następnego numeru	107

nowe podzespóły

Z kilkuset nowości wybraliśmy te, których nie wolno przeoczyć. Bieżące nowości można śledzić na www.elektronikaB2B.pl



48,5-calowe wyświetlacze LCD-TFT dla branży transportowej

W ofercie firmy Unisystem dostępne są nowe modele wyświetlaczy LCD-TFT typu 4851-Y opracowane przez Litemax. To ekrany o przekątnej 48,5 cala, które doskonale sprawdzają się w branży transportowej, także przy montażu na zewnątrz. Oferują rozdzielczość 1920×360 przy proporcjach obrazu 16:3. Zapewniają doskonałą czytelność treści, także w warunkach intensywnego oświetlenia naturalnego i sztucznego – jasność (3000 cd/m²), kontrast (7700:1) oraz kąty obserwacji (89/89/89/89°). Zakres temperatur pracy to –20...60°C.

Producent przygotował dwa warianty wyświetlacza: SSF4851-Y oraz SSH4851-Y, różniące się między sobą akcesoriami, które standardowo dołączane są przez producenta.

Modele SSF wyposażone są jedynie w płytke sterującą podświetleniem LED. Ograniczeniem tej wersji jest komunikacja wyłącznie za pomocą jednego interfejsu (LVDS).

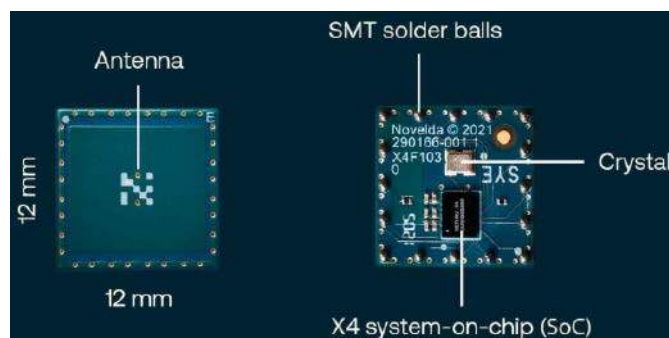
Modele SSH są rozbudowanym wariantem SSF, w których poza płytka sterującą podświetleniem LED, występuje również AD Board umożliwiający transmisję obrazu i dźwięku za pomocą innych interfejsów (DVI, HDMI i DP).

W ofercie dostępne są również rozwiązania SSD, które można nazywać monitorami – w tych modelach cała elektronika umieszczona jest w obudowie. W tabeli zestawiono porównanie poszczególnych parametrów modeli SSF4851-Y oraz SSH4851-Y.

Wyświetlacze SSF4851-Y i SSH4851-Y sprawdzają się zarówno w aplikacjach wewnątrz pojazdów i pomieszczeń, jak i przy montażu

na zewnątrz, choć w zależności od warunków panujących w środowisku konieczne może być dodanie elementów ogrzewających lub ochładzających elektronikę (np. grzałek lub wiatraków). Co więcej, oba rozwiązania mogą być również umieszczane w lokalizacjach nasłonecznionych. Zastosowana w nich technologia hiTNI chroni ciekłe kryształy przed efektem przegrzewania, jednocześnie niwelując zjawisko występowania czarnych plan. Kształt rozwiązań z serii 4851-Y kojarzy się z aplikacjami transportowymi, ze względu na występujące w tej branży ograniczenia przestrzeni, np. we wnętrzach pojazdów. Z tego samego powodu wyświetlacze szerokoformatowe są stosowane w systemach informacji pasażerskiej, w tym jako podwieszane ekrany na stacjach i przystankach, na których prezentowane są bieżące informacje dotyczące podróży. Ponadto, wyświetlacze SSF4851-Y i SSH4851-Y z powodzeniem można wdrażać również w innych branżach, np. jako niestandardowy nośnik treści reklamowych w handlu.

<https://unisystem.pl>



Zaawansowany czujnik zbliżeniowy do wykrywania obecności człowieka

Czujnik zbliżeniowy UWB Proximity Sensor firmy Novelda to najbardziej zaawansowany na rynku komponent tego typu do wykrywania obecności człowieka. Może wykrywać nawet najdrobniejsze ruchy, w tym oddech i bicie serca, automatycznie aktywując np. laptop i wyłączając go po odejściu użytkownika. Nadaje się do zastosowań m.in. w elektronice użytkowej i systemach inteligentnego budynku, zwiększając bezpieczeństwo danych i ograniczając straty mocy m.in. w systemach zarządzania oświetleniem, ogrzewaniem i klimatyzacją. UWB Proximity Sensor zapewnia szeroki kąt widzenia (180° w poziomie i w pionie) i umożliwia programowanie zakresu pomiarowego od 0,5 do 1,5 m. Jest zamykany w obudowie SMD o wymiarach 12×12×2 mm. Może być umieszczany za panelami z tworzywa sztucznego, szkła lub tkaniny, co nie wpływa na jego parametry pracy. Zawiera komplet elementów w torze przetwarzania sygnału oraz własną antenę, zapewniając łatwą i tanią integrację w urządzeniach docelowych. Może współpracować z dowolnym mikrokontrolerem host wyposażonym w 8 kB pamięci Flash i 1 kB pamięci RAM. UWB Proximity Sensor pracuje z napięciem zasilania 1,8...3,3 V przy poborze mocy poniżej 20 mW. Komunikuje się z mikrokontrolerem za pomocą interfejsu I²C lub SPI.

www.novelda.com

Parametr	SSF4851-Y	SSH4851-Y
rozmiar	48,5"	48,5"
rozdzielczość	1920×360 px	1920×360 px
obszar aktywny	1209,6×226,8 mm	1209,6×226,8 mm
proporcje	16:3	16:3
jasność	3000 cd/m ²	3000 cd/m ²
kontrast	7700:1	7700:1
kąty obserwacji (°)	89/89/89/89	89/89/89/89
interfejs	LVDS	LVDS/DVI, HDMI, DP
pobór mocy	71 W	75 W
wymiary zewnętrzne	1232,4×249,6×22,8 mm	1232,4×249,6×22,8 mm
wymiary ramki	11,4×11,4×11,4×11,4 mm	11,4×11,4×11,4×11,4 mm
waga	5,4 kg	5,4 kg
zakres temperatur pracy	–20...60°C	–20...60°C
dodatkowe	hiTNI	hiTNI



Przełącznik kontaktronowy SPST NO do przetwarzania sygnałów o częstotliwości do 3 GHz

Miniaturowy przełącznik kontaktronowy 113RF firmy Pickering Electronics umożliwia przełączanie sygnałów o częstotliwości do 3 GHz w szybkich systemach w.c.z. Jest to przełącznik o konfiguracji 1 Form A (SPST NO), charakteryzujący się małymi stratami wtrąconymi i małymi gabarytami (12,5×6,6×3,7 mm), nadający się idealnie do systemów o gęstym upakowaniu podzespołów. Zawiera wewnętrzny ekran magnetyczny, zapobiegający interferencjom sygnałów pomiędzy sąsiednimi elementami. Jego żywotność przekracza 250 milionów cykli. Model 113RF charakteryzuje się obciążalnością do 10 W/0,5 A, maksymalnym napięciem przełączanym 100 V, rezystancją kontaktu 0,12 Ω i rezystancją izolacji 1012 Ω . Jest odporny na uderzenia i wibracje do odpowiednio 50 g i 20 g. Występuje w wersjach o napięciu znamionowym i rezystancji cewki 3 V/100 m Ω i 5 V/300 m Ω . Może pracować w zakresie temperatury otoczenia od -20 do +85°C.

www.pickeringrelay.com

Czujnik ciśnienia bezwzględny odporny na wodę i składniki organiczne zawarte w gazach

Alps Alpine rozpoczyna masową produkcję nowego czujnika ciśnienia bezwzględnego HSPPAD143C o szczelnej konstrukcji, wyróżniającego się dużą odpornością na wodę i składniki organiczne zawarte w gazach. Jest on polecany do zastosowań w gazomierzach, a dzięki miniaturowej konstrukcji, również we wszelkiego typu urządzeniach przenośnych. HSPPAD143C charakteryzuje się dużą odpornością chemiczną, osiągniętą dzięki przeprojektowanej osłonie z żywicy oraz zapewnia większą odporność na wstrząsy w porównaniu z wcześniejszym modelem HSPPAD143A. Wymiary (3,1×3,1×2,6 mm) i kształt czujnika pozostają niezmiennic. Wodoodporność uzyskuje się dzięki zastosowaniu O-ringu.

Czujnik HSPPAD143C pracuje w zakresie temperatury otoczenia od -40 do +85°C i w zakresie napięcia zasilania 1,7...3,6 V, pobierając typowo 1,8 μ A prądu przy częstotliwości próbkowania 1 Hz. Charakteryzuje się zakresem pomiarowym od 300 do 2000 hPa, 17-bitową rozdzielczością, dokładnością ± 2 hPa i poziomem szumu 0,026 hPa. Komunikacja z mikroprocesorem odbywa się za pośrednictwem interfejsu I²C.

www.alpsalpine.com



Seria bezpieczników dużej mocy do samochodowych systemów ładowania

Bezpieczniki nowej serii POWrFUSE PF-K firmy Bourns zostały zaprojektowane specjalnie do zastosowań w wysokonapięciowych systemach ładowania i zarządzania zasilaniem. Spełniają wymogi normy ISO 8820-8 dotyczącej wkładek bezpiecznikowych stosowanych w pojazdach EV/HEV. Pracują z napięciem zasilania do 500 VDC i zapewniają zdolność wyłączenia do 20 kA @ 500 VDC. Ich zakres dopuszczalnej temperatury pracy rozciąga się od -55 do +125°C.

W ramach serii PF-K dostępne są warianty o prądach znamionowych 15, 30, 40 i 50 A z dwoma opcjami montażu (śrubowy, PCB). Zostały one przetestowane w wewnętrznym laboratorium firmy Bourns, zgodnie z motoryzacyjnymi standardami jakości pod kątem stabilności w trudnych warunkach środowiskowych, cyklicznych obciążeń uderowych i charakterystyk elektrycznych.

www.bourns.com

Rezystory chipowe dużej mocy o tolerancji od 0,1% produkowane na podłożach AlN

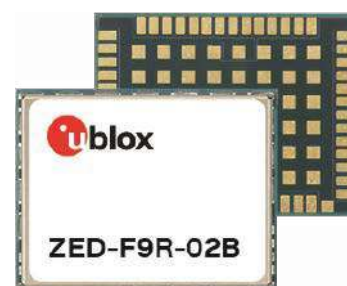
TT Electronics wprowadza na rynek nową serię cienkowarstwowych rezystorów chipowych TFHP, produkowanych na ceramicznych podłożach z azotku glinu o 6-krotnie większej przewodności cieplnej od podłoży aluminiowych. Charakteryzują się one dużą niezawodnością, wąskim przedziałem tolerancji i dużą mocą znamionową, wynoszącą 2 W dla wersji w obudowach 1206 i 6 W dla wersji 2512. Nadają się doskonale do zastosowań m.in. w precyzyjnych zasilaczach, wzmacniaczach mocy i przemysłowych systemach sterowania. Rezystory TFHP zapewniają węższy przedział tolerancji od innych odpowiedników grubowarstwowych, wynoszący 0,1%. Mogą pracować w temperaturze otoczenia od -55 do +155°C. Ich współczynnik TCR wynosi 25 ppm/°C. Dostępne są wersje o rezystancji od 50 Ω do 30 k Ω .

www.ttelectronics.com



Precyzyjny moduł GNSS do zastosowań wymagających centymetrowej dokładności

Firma u-blox wydała nową aktualizację oprogramowania firmware, zwiększającą dokładność pozycjonowania precyzyjnego modułu GNSS ZED-F9R. Obsługa usług korekcji położenia SPARTN 2.0 i QZSS CLAS pozwala na zastosowania układu w aplikacjach wymagających centymetrowej



dokładności, np. w różnego typu pojazdach zautomatyzowanych. Moduł może pracować nawet w trudnych warunkach zewnętrznych, np. w gęstej zabudowie miejskiej. Zawiera inercyjną jednostkę pomiarową do pozycjonowania RTK (Real-Time Kinematic) i wykorzystuje specjalne algorytmy do łączenia danych IMU z pomiarami GNSS, danymi z usług korekcji i modelem dynamiki pojazdu, mające zapewnić centymetrową dokładność nawet w sytuacji, gdy sam GNSS zawiedzie.

ZED-F9R bazuje na 184-kanalowym odbiorniku GNSS u-blox F9, zapewniającym równoczesne śledzenie satelitów z kilku konstelacji GNSS (GPS L1C/A L2C, GLO L1OF L2OF, GAL E1B/C E5b, BDS B1I B2I, QZSS L1C/A L1S L2C). Jego struktura obejmuje oscylator TCXO, zegar RTC, pamięć flash, akcelerometr + żyroskop 3D, diplexer i filtry SAW. ZED-F9R pracuje z napięciem zasilania 2,7...3,6 V, pobierając średnio 85 mA prądu. Jest zamykany w obudowie LGA o wymiarach 22×17×2,4 mm.

Pozostałe cechy to:

- częstotliwość aktualizacji: do 30 Hz,
- dokładność pozycjonowania (RTK): <0,01 m + 1 ppm CEP,
- czułość: -160 dBm w trybie śledzenia,
- interfejsy: 2×UART, USB, SPI, DDC (kompatybilny z I²C),
- zakres temperatury pracy: -40...+85°C,
- współpraca z anteną aktywną,
- kwalifikacja AEC-Q100.

www.ublox.com

Miniaturowy 100-woltowy tranzystor GaN FET o rezystancji RDS(on) równej 3,8 mΩ

EPC2306, najnowszy tranzystor GaN FET z oferty firmy EPC (Efficient Power Conversion), charakteryzuje się małymi stratami przy przewodzeniu i przełączaniu, wynikającymi z rezystancji RDS(on) zredukowanej do zaledwie 3,8 mΩ oraz małych ładunków wewnętrznych



QG, QGD i QOSS. Jest to tranzystor o napięciu BVDSS równym 100 V, dopuszczalnym ciągłym prądzie drenu 48 A i dopuszczalnym prądzie impulsowym 197 A (25°C, 300 μs), zamykany w chipowej obudowie QFN o powierzchni 5×3 mm z wyprowadzeniem radiatora. Jego zakres zastosowań obejmuje wszelkiego typu aplikacje wysokoprądowe o dużej gęstości mocy, w tym systemy zasilające i napędowe oraz wzmacniacze audio klasy D.

EPC2306 jest ponad dwukrotnie tańszym odpowiednikiem wprowadzonego wcześniej na rynek modelu EPC2302 (1,8 mΩ), zamykanym w obudowie o identycznych wymiarach i charakteryzującym się identycznym napięciem BVDSS. Może stanowić jego zamiennik w aplikacjach, w których priorytetem jest cena podzespołów, a RDS(on) ma

REKLAMA

COMPUTER CONTROLS

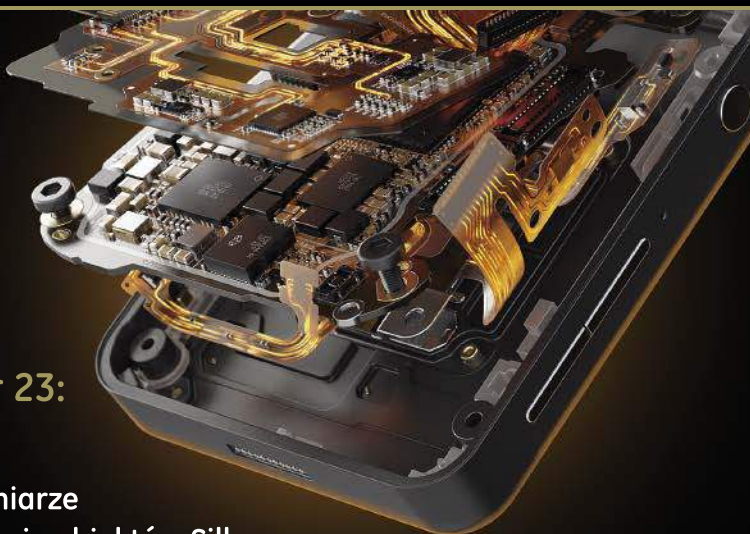
Autoryzowany dystrybutor Altium w Polsce



ALTium
DESIGNER23

W najnowszym Altium Designer 23:

- Obsługa wiązek przewodów
- Porty o automatycznie ustalonym rozmiarze
- Automatyczne wycinanie lub przesuwanie obiektów Silkscreen
- Raport historii projektu
- Komentarze zależne od warstwy



Computer Controls Sp. z o.o.
Bielsko-Biała, ul. Budowlanych 1

tel.: +48 (33) 485 94 90

e-mail: info@ccontrols.pl
www.ccontrols.pl

mniej krytyczne znaczenie. Cena hurtowa EPC2306 wynosi 3,08 USD przy zamówieniach 1000 sztuk.

Firma EPC oferuje też płytkę ewaluacyjną EPC90145 z wyjściowym stopniem mocy zrealizowanym na dwóch tranzystorach EPC2306, skalonym sterownikiem bramek (uP1966E) i niezbędnymi komponentami współpracującymi. Służy ona do badania parametrów przetwornic DC-DC buck i boost o prądzie wyjściowym do 45 A, zrealizowanych na bazie EPC2306. Cena płytki EPC90145 wynosi 200 USD.

www.epc-co.com



Diody prostownicze o krótkim czasie regeneracji do układów korekcji PFC

Do oferty firmy Diodes wchodzi dwie serie diod o krótkim czasie regeneracji, zaprojektowanych do układów korekcji PFC. Są one produkowane w wersjach o napięciu znamionowym od 600 do 1000 V i dopuszczalnym prądzie przewodzenia 1...30 A. Poza krótkim czasem regeneracji do ich zalet należy też małe napięcie przewodzenia, mała rezystancja termiczna, mały prąd upływu oraz odporność na duże impulsy prądowe. Łagodna charakterystyka regeneracyjna zmniejsza poziom generowanych zaburzeń elektromagnetycznych.

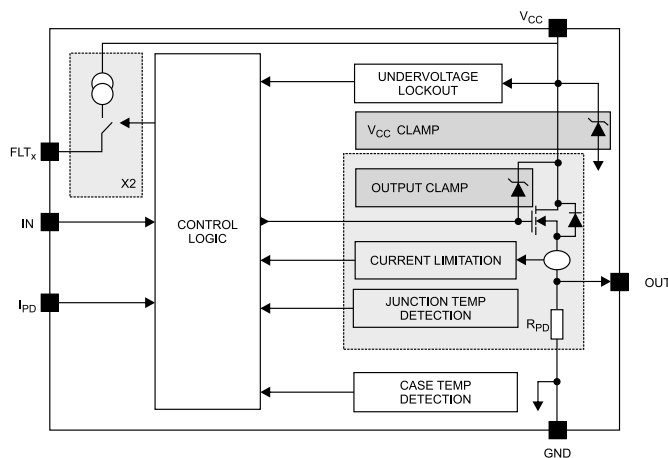
	V_{RRM} (V)	I_F (A)	V_F @ $T_A = 25^\circ C$		I_R @ $T_A = 125^\circ C$		I_R @ $T_A = 25^\circ C$		t_{rr} typ. (ns)	t_{rr} maks. (ns)	I_{FSM} (A)
			typ. (V)	maks. (V)	typ. (μA)	maks. (μA)	typ. (μA)	maks. (μA)			
MURS160A	600	1	1,03	1,25	13	150	5	25	50	35	
MURS160	600	1	1,03	1,25	13	150	5	25	50	35	
MURS360B	600	3	1,05	1,25	33	150	5	35	50	100	
MURS360	600	3	1,05	1,25	33	150	5	35	50	100	
MURS460C	600	4	1,15	1,28	35	250	10	35	50	110	
DTH8E06D	600	8	2,4	2,9	35	400	30	18	25	125	
DTH8L06D	600	8	1,1	1,3	13	200	8	41	70	120	
DTH8L06DNC	600	8	1,2	1,3	13	200	8	41	70	120	
DTH8R06D	600	8	2,4	2,9	35	400	30	18	25	80	
DTH8R06D1	600	8	2,4	2,9	35	400	30	18	25	80	
DTH8S06D	600	8	2,7	3,4	92	200	15	16	21	70	
DTH1206D	600	12	2,4	2,9	30	600	45	21	30	120	
DTH3006D	600	30	2,1	2,4	86	1,000	100	26	45	350	
DTH3006PT	600	30	2,1	2,4	86	1,000	100	26	45	350	
DTH810D	1000	8	1,5	2	16	200	5	65	85	80	

Diody serii MURS (1...4 A) są produkowane w obudowach SMA, SMB i SMC. W przypadku serii DTH (8...30 A) dostępne są cztery warianty obudów: TO-220, ITO-220, TO-247 i TO-252.

www.diodes.com

Szybki przełącznik zasilania do systemów bezpieczeństwa

IPS1025HF to szybki przełącznik zasilania do systemów bezpieczeństwa wymagających krótkiego czasu opóźnienia przy włączaniu/wyłączeniu, pracujący w konfiguracji high-side. Układ reaguje na sygnały włączenia i wyłączenia w czasie krótszym niż 60 μs , umożliwiając systemom ochronnym zapewnienie określonego poziomu nienaruszalności bezpieczeństwa (SIL). Dzięki bardzo szerokiemu zakresowi napięcia wejściowego od 8,65 V do 60 V i tolerancji na przepięcia



wejściowe do 65 V, nadaje się do pracy w ciężkich warunkach przemysłowych. Jego zakres zastosowań obejmuje sterowniki PLC, automaty sprzedające, urządzenia peryferyjne I/O oraz maszyny CNC.

IPS1025HF umożliwia zaprogramowanie dwóch wartości progowych ogranicznika prądu, co pozwala na współpracę z obciążeniami o dużym prądzie początkowym, np. silnikami i obciążeniami pojemnościowymi. Stopień wyjściowy z N-kanalowym tranzystorem MOSFET o typowej wartości $R_{DS(ON)}$ równej 12,5 m Ω zapewnia dużą sprawność energetyczną i małą emisję ciepła. Odporność na impulsy energetyczne do 14 J zwiększa niezawodność przy współpracy z obciążeniami indukcyjnymi.

IPS1025HF zawiera zabezpieczenie podnapięciowe, nadnapięciowe, przeciążeniowe, zwarcie oraz przed rozłączeniem linii GND i VCC. Wbudowane obwody diagnostyczne sygnalizują przeciążenie wyjścia, przegrzanie struktury oraz nadmierną temperaturę obudowy. Układ spełnia wymogi norm IEC 61000-4-2 ESD, IEC 61000-4-4 i IEC 61000-4-5 w zakresie odporności na wyładowania ESD, szybkie stany przejściowe i przepięcia. IPS1025HF jest produkowany w obudowach PowerSSO-24 i QFN48L. Jego ceny hurtowe zaczynają się od 3,51 USD przy zamówieniach 1000 sztuk.

www.st.com

Barometryczny czujnik ciśnienia do obliczania wysokości w pomieszczeniach

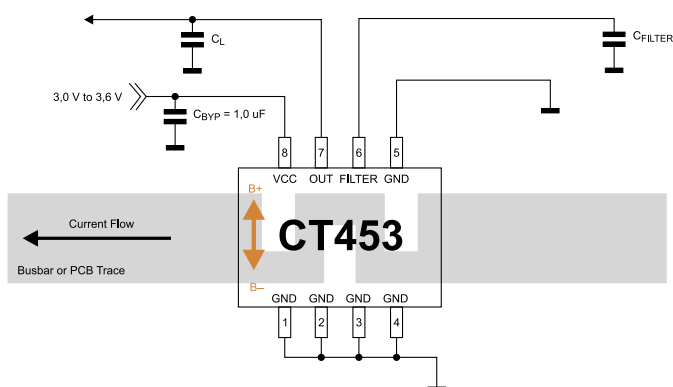
Moduły nawigacyjne w maszynach przemysłowych, takich jak roboty, służą do śledzenia lokalizacji w dużych, wielopiętrowych przestrzeniach wewnętrznych, takich jak fabryki, centra dystrybucyjne, lotniska i garaże. Oprócz szerokości i długości geograficznej, najbardziej zaawansowane modele mierzą również wysokość n.p.m., którą można określić za pomocą czujników ciśnienia barometrycznego. Nowy czujnik ND015A firmy Superior Sensor Technology umożliwia prowadzenie precyzyjnego pomiaru wysokości w pomieszczeniach. Bardzo mały poziom podłogi szumowej pozwala w jego przypadku zapewnić dokładność pomiaru na poziomie ± 150 cm. ND015A zawiera filtry 50/60 Hz, zapewniający dużą dokładność pomiaru w obszarach podatnych na zaburzenia z linii energetycznych lub sieci elektrycznej. Oferuje również funkcję filtrowania cyfrowego, eliminującą wpływ dźwięków, wibracji i szybkich ruchów. ND015A jest obecnie dostępny w ofercie firm Digi-Key i Mouser. Ważniejsze parametry to:

- zakres pomiarowy: 0...15 psia,



- maksymalne ciśnienie bezpieczne: 35 psia,
- okres pomiarów: 2,25 ms,
- pasmo: 1...200 Hz,
- dokładność: 0,1 %FS,
- błąd całkowity (TEB): typ. 0,15 %FS,
- rozdzielczość: 16 bitów,
- PSRR: 0,0005 %FS/mV,
- stabilność długoterminowa: typ. 0,1 %FS/rok.

www.superiorsensors.com



Bezkontaktowe czujniki prądu o dużym stosunku sygnał/szum

Firma Crocus Technology powiększa ofertę bezkontaktowych czujników prądowych o nową rodzinę CT45x, obejmującą czujniki o dużym współczynniku sygnału do szumu, przekraczającym 65 dB. Są to układy oparte na technologii XtremeSense TMR (Tunnel Magneto-Resistive), mogące znaleźć zastosowanie np. w systemach zarządzania akumulatorami, falownikach i konwerterach DC-DC.

Zapewniają odporność na zewnętrzne pola magnetyczne >50 dB, umożliwiając producentom eliminację ekranów.

Nowa oferta obejmuje czujnik o napięciu zasilania 3,3 V (CT453) i dwie wersje 5-woltowe (CT450, CT452). CT452 i CT453 zapewniają tłumienie składowej sumacyjnej.

Wszystkie czujniki rodziny CT45x charakteryzują się pasmem 1 MHz, czasem odpowiedzi <300 ns oraz identycznymi zakresami pomiarowymi od 0...6 mT i ± 6 mT do 0...24 mT i ± 24 mT. Uzyskały kwalifikację AEC-Q100 Grade 1, pozwalającą na zastosowania w elektronice samochodowej. Ich całkowity błąd pomiaru nie przekracza $\pm 1,0\%$ w zakresie temperatury otoczenia od -40 do +125°C.

www.crocus-technology.com

Wzmacniacze szerokopasmowe na zakres częstotliwości pracy 1,4...7,1 GHz

CMX90G301 i CMX90G302 to uniwersalne wzmacniacze szerokopasmowe o zakresie częstotliwości pracy od 1,4 do 7,1 GHz, zamykane w obudowach VQFN-16 o powierzchni 3×3 mm. Układy te wyróżniają się charakterystyką częstotliwościową o dodatnim nachyleniu (odpowiednio +1 dB i +2 dB w paśmie pracy), eliminując konieczność kompensowania strat w systemie, rosnących wraz z częstotliwością. Wymagają dołączenia jedynie kondensatorów na linii zasilania. Ich wejścia i wyjścia są dopasowane do impedancji 50 Ω . Są produkowane w technologii GaAs pHEMT, zapewniającej energooszczędną pracę, małe szумы i duże wzmocnienie. Mogą być zasilane napięciem od 2,7 do 5 V i pobierają typowo 20 mA prądu.



REKLAMA

świat radio

Magazyn wszystkich użytkowników eteru
KRÓTKOFALARSTWO CB RADIOTECHNIKA

przejrzyj i kup na
www.ulubionykiosk.pl

Pozostałe parametry to:

- wzmacnienie: 14,5...15,5 dB,
- NF: 2 dB,
- maks. moc wejściowa: +10 dBm,
- P1dB: +11,5 dBm @ 3,5 GHz,
- OIP3: 21 dBm ($\Delta f=10$ MHz, $f=3,5$ GHz),
- zakres temperatury pracy: $-40...+105^{\circ}\text{C}$.

www.cmlmicro.com

Transile w obudowach cylindrycznych odporne na impulsy energetyczne do 5 kW

Transile nowej serii 5KP-HRA firmy Littelfuse zapewniają odporność na impulsy energetyczne do 5 kW (10/1000 μs) przy cyklu pracy 0,01%. Występują w wersjach jedno- i dwukierunkowych, zamykanych w obudowach z wyprowadzeniami osiowymi o maksymalnych wymiarach $\varnothing 9,1 \times 9,1$ mm. Charakteryzują się krótkim czasem odpowiedzi, nieprzekraczającym 1 ps (dla skoku napięcia od 0 do VBR min.). Ich maksymalna moc rozpraszana przy pracy ciągłej wynosi 8 W.

Transile 5KP-HRA przeszły testy odporności na impulsy udarowe oraz testy wysokotemperaturowe (168 h @ 175°C) i cykli termicznych (20 cykli 15-minutowych $-55...+150^{\circ}\text{C}$). Są produkowane w kilkudziesięciu wersjach o napięciu pracy od 5 do 220 V i napięciu przebicia 7...270 V. Zapewniają odporność na wyładowania ESD do 30 kV (przenoszonych przez kontakt i powietrze), zgodnie z wymogami IEC-61000-4-2. Ich prąd upływu wynosi mniej niż 2 μA dla wersji o napięciu przebicia >12 V.

www.littelfuse.com



Wysokotemperaturowe rezystory z kwalifikacją AEC-Q200 odporne na duże impulsy prądowe

Rezystory chipowe nowej serii CRG2512 z oferty firmy Bourns wyróżniają się odpornością na duże impulsy prądowe. Do ich zalet należy też odporność na siarkę, bardzo szeroki zakres dopuszczalnej temperatury pracy od -55 do $+170^{\circ}\text{C}$ oraz mała indukcyjność resztkowa (≤ 5 nH) i małe szумы, wynikające z zastosowania

metalowego elementu rezystancyjnego. Rezystory te mogą być stosowane jako czujniki natężenia prądu w układach zasilania i napędowych. Są produkowane w chipowych obudowach rozmiaru SMD2512 (metryczny 6432). Występują w wersjach 1- i 2-watowych o rezystancji od 0,11 do 0,68 Ω .

Pozostałe parametry to:

- współczynnik temperaturowy: ± 50 ppm/ $^{\circ}\text{C}$,
- tolerancja: $\pm 1\%$, $\pm 5\%$,
- odporność na udary: 100 g (6 ms),
- odporność na wibracje: 5 g (10...2000 Hz),
- kwalifikacja AEC-Q200.

www.bourns.com



Moduł pomiarowy do testowania pojazdów i akumulatorów w środowisku wysokonapięciowym do 1500 V

Nowy moduł pomiarowy HISO-HV-4 wypełnia lukę na rynku aparatury do testowania pojazdów elektrycznych i pakietów akumulatorowych. Może być użyty do prowadzenia pomiarów w środowiskach wysokonapięciowych (do 1500 V), rozszerzając zakres możliwości testowych dla inżynierów działów badawczo-rozwojowych. HISO-HV-4 rozszerza rodzinę małogabarytowych modułów pomiarowych CANSASfit, bazujących na szynie CAN. Służy do akwizycji danych podczas testów mobilnych, umożliwiając bezpośrednie podłączenie wszystkich typowych sygnałów pomiarowych (napięcie, prąd, temperatura, prędkość obrotowa oraz przemieszczenie i prędkość liniowa).

HISO-HV-4 zapewnia wzmocnioną izolację do 1000 V (CAT II). Cztery kanały z bananowymi gniazdami laboratoryjnymi dostarczają dane pomiarowe przez magistralę CAN przy maksymalnej szybkości próbkowania 1 kSps w paśmie do 400 Hz. Mała, zatrzaskowa obudowa CANSASfit umożliwia bezpośrednie mechaniczne i elektryczne dokowanie do innych modułów fit, a w szczególności do uzupełniających modułów HISO typu HISO8-T i HISO-UT-6, przeznaczonych do współpracy z czujnikami temperatury (RTD, termopara), akcelerometrami MEMS i innymi czujnikami niskonapięciowymi, podłączonymi do linii wysokonapięciowej.

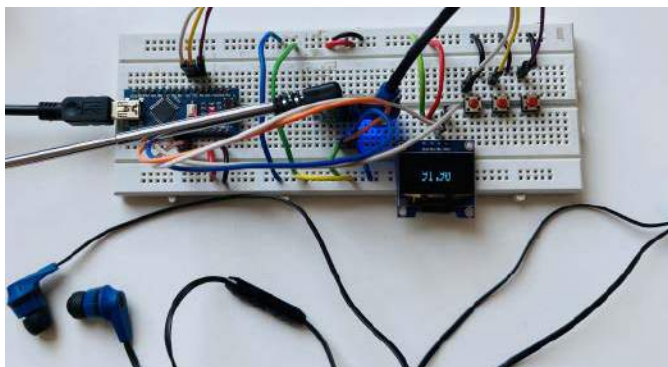
www.imc-tm.com

REKLAMA

www.ep.com.pl/EPwtoku

dodaj do obserwowanych

Przedstawiamy redakcyjny wybór najciekawszych projektów spośród ostatnio anonsowanych w internecie. Są to projekty na różnych etapach realizacji. Warto się zapoznać z projektami zakończonymi i śledzić realizację projektów niegotowych, by czerpać z nich inspirację do własnych prac.



Radio z modułem Arduino

Autor tego projektu konstruuje obecnie odbiornik radiowy z zastosowaniem modułu Arduino Nano. Głównymi komponentami są: płytka Arduino Nano i moduł radiowy FM TEA5767. Moduł radiowy nie ma wbudowanego głośnika, więc podłączony jest zewnętrzny głośnik (słuchawki też będą działały poprawnie). Jednym z największych problemów, z jakimi spotkał się autor podczas budowy, było to, że moduł radiowy FM TEA5767 nie działał poprawnie z płytką Arduino Mega, natomiast działa prawidłowo z płytą Arduino Nano. Obecnie prace koncentrują się na dodawaniu różnych funkcji. Najnowszym dodatkiem są 3 przyciski, które mają przypisane częstotliwości, co pozwala na ich szybką zmianę i ułatwia przełączanie między preferowanymi stacjami. Docelowo dodany ma zostać enkoder obrotowy lub joystick do poruszania się między wszystkimi stacjami.

Ważnym elementem rozbudowy było dodanie ekranu OLED, na którym jest wyświetlana nazwa stacji, którą nasłuchuje radio. Co dziwne, autor napotkał problem podczas oprogramowywania wyświetlacza – funkcja `DisplayFreq()`, która odpowiada za wyświetlanie danych na ekranie powodowała generowanie szumu w module

radiowym. Umieszczenie `DisplayFreq()` wewnątrz pętli `loop()` powodowało ciągły szum w słuchawkach. Nie jest jasne, co powoduje ten szum, ale udało się to rozwiązać wywołując aktualizację częstotliwości tylko w chwili naciśnięcia przycisku.

<https://hackaday.io/project/188135-arduino-radio>



Neutron Pico

Komputer jednopłytkowy Neutron Pico został zaprojektowany do obsługi systemu operacyjnego Neutron, ale ponieważ w istocie jest to tylko Raspberry Pi Pico, można na nim uruchomić wszystko, co normalnie możliwe byłoby do uruchomienia na tym mikrokontrolerze. System ma wyjście wideo z 12-bitowym kolorem, podobnie jak OCS w komputerze Amiga. Dostępny jest DAC wykonany na drabince R2R, ale także dodano bufor wideo RGB zaprojektowany do sterowania obciążeniem 75 Ω , więc obraz powinien być lepszej, jakości niż pochodzący z przeciętnego wyjścia GPIO do VGA. Komputer ma także 16-bitowy stereofoniczny kodek audio i złącze AC'97, a także klasyczne

REKLAMA

HAMMOND

Obudowa miniaturowa 1551W IP68

Dowiedz się więcej: <https://hammfg.com/1551W>

Skontaktuj się z nami, aby otrzymać bezpłatną próbkę ewaluacyjną.
eusaes@hammfg.com +44 1256 812812



różowo-zielono-niebieskie potrójne gniazdo PC98 z tyłu. Wyposażono go w gniazdo kart SD do przechowywania danych, mikrokontroler STM32F0 do sterowania zasilaniem, obwód resetowania i parę portów PS/2. Dostępna jest również magistrala rozszerzeń z SPI oraz I²C do dodawania nowych funkcji. MIDI? SCSI? Autor nie wyklucza dodania żadnego z tych interfejsów w przyszłości.

Pliki projektowe znajdują się na w repozytorium projektu na GitHubie. To także jest ciekawe, gdyż komputer ten jest sprzętem typu open-source, a zespół projektowy używa GitHub Actions do niezawodnej konwersji plików KiCad na schematy w PDF i pliki Gerber, przygotowane do produkcji. Typowe projekty open-source mają w swoim repozytorium już zbudowane pliki wynikowe (pliki PDF, Gerbery, firmware, jako HEX itp.) wraz z kodem źródłowym. Nie daje nam to pewności, czy pliki te zostały przygotowane w poprawny (dla nas) sposób. Ten projekt korzysta z GitHub Actions do generowania tych plików na żądanie. Oznacza to, że można zobaczyć skrypty użyte do ich zbudowania – dosłownie i z definicji dostępne są wszystkie narzędzia i pliki potrzebne do zrobienia tego, co robi GitHub. To samo dotyczy plików PDF z schematami i interaktywnych plików z listami elementów. Nie ma wątpliwości, która wersja jest, która – wszystko jest w dziennikach git, wyraźnie oznaczone każdym unikalnym numerem wersji. To nie tylko Open Hardware. To w pełni odtwarzalny sprzęt.

Dodatkowo, istotnym aspektem projektu, są komentarze na schemacie. Rzadko, kiedy spotyka się z objaśnieniami działania systemu na schemacie, a tutaj autor, w miarę możliwości, starał się, aby schematy opisywały nie tylko, czym coś jest lub robi, ale także, dlaczego. To nie jest tylko opis pomysłu, który stoi za podzespołem, ale może być zasobem edukacyjnym dla każdego, kto chciałby dowiedzieć się więcej o elektronice i o tym, co jest potrzebne, aby utworzyć własną płytę główną kompatybilną z ATX.

Moduł wykorzystuje system operacyjny Neutron OS. Jest to prosty system operacyjny napisany w języku Rus, zainspirowany CP/M i wczesnym MS-DOSem. Jest jednozadaniowy, z płaską przestrzenią pamięci. Zasadniczo ten system operacyjny wystarczy, aby aplikacja została załadowana z karty SD do pamięci RAM, a gdy już tam będzie, może robić, co chce.

<https://hackaday.io/project/180407-neutron-pico>
<https://GitHub.com/Neutron-compute/neutron-pico>



Tetrinsic – czuły na nacisk enkoder suwakowy z dotykowym sprzężeniem zwrotnym i wyświetlaczem TFT

Autor konstrukcji chce zaprojektować urządzenie wejściowe, które może zapewnić więcej niż standardowy przełącznik klawiatury. Przeznaczone jest ono dla Tetent – systemu komputerowego, nad jakim pracuje. Odkąd zaczął pracę nad swoim modulem, doszedł do wniosku, że obecne rozwiązanie (5 oddzielnych przycisków dla palców), jest wadliwe. Zmienił on projekt tak, aby system używał małego silnika

BLDC połączonego z paskiem, który działa również, jako powierzchnia nasadki klawiszy. Zostało to zainspirowane różnymi projektami open-source, takimi jak SmartKnob View itp. Moduł będzie korzystał z czujnika siły i silnika do symulacji różnych suwaków i manipulatorów dotykowych. Silnik jest również potrzebny do maksymalnego skrócenia czasu uczenia się nowego sposobu pisania.

Inne rozwiązania, analizowane przez autora tego urządzenia, to przełączniki wykorzystujące efekt Halla i ich potencjalne podwójne działanie. Istnieją również inne małe przełączniki dotykowe (siła nacisku 0,6 N i 1,6 N) podwójnego działania, ale ich druga siła uruchamiająca jest zbyt duża dla tego zastosowania. Przełączniki Void firmy Riskable mają z kolei konstrukcję, która umożliwia regulowanie siły naciskania bez konieczności podłączania elektrycznego do płytki drukowanej. Kluczem są tutaj magnesy, które są tańsze niż kupowanie przełączników dotykowych o naciskach 0,8 N i 0,5 N, a ponadto pozwalają na podświetlenie za pomocą diod LED, ponieważ potrzebne komponenty są mniejsze.

<https://hackaday.io/project/184180-tetrinsic-gd0041>



Monitor zasilania z interfejsem Ethernetowym i wsparciem dla MQTT

Zaprezentowane urządzenie, służy do monitorowania mocy, prądu, napięcia, częstotliwości, współczynnika mocy i zliczania zużytej energii w sieci energetycznej. Dane na żywo są dostępne na ekranie lub w aplikacji internetowej przyjaznej dla urządzeń mobilnych. Podstawowe funkcje układu to:

- wyświetlanie predefiniowanych metryk na ekranie układu,
- wysyłanie wartości do brokera MQTT,
- automatyczne zerowanie wskaźnika zużycia energii na początku miesiąca,
- pokazywanie wartości energii z poprzedniego miesiąca w interfejsie użytkownika,
- automatycznie wykrywa obecność Home Assistant'a i zapewnia integrację z tym systemem.

Dostępna jest również aplikacja internetowa, która aktualizuje metryki za pośrednictwem WebSockets. Pokazuje ona ostrzeżenie, gdy dana wartość nie mieści się w założonym zakresie, przechowuje zakresy zapisane w pamięci EEPROM i umożliwia ich konfigurowanie w interfejsie użytkownika. System przechowuje również konfigurację MQTT w pamięci EEPROM i umożliwia konfigurowanie jej w interfejsie użytkownika.

Obecnie, jak donosi autor, Power Monitor działa prawie rok bez poważnych problemów. Tak jak się spodziewał, zastosowany wyświetlacz OLED zaczął blaknąć, więc niezbędne jest dodanie wyłączania wyświetlacza na podstawie czujnika światła, ale nie zostanie ono zaimplementowane w tej wersji sprzętowej. Dodano zamiast tego szereg opcji do lepszej integracji z Home Assistant, które szczegółowo opisano i pokazano na stronie z projektem.

<https://hackaday.io/project/183312-power-monitoring>

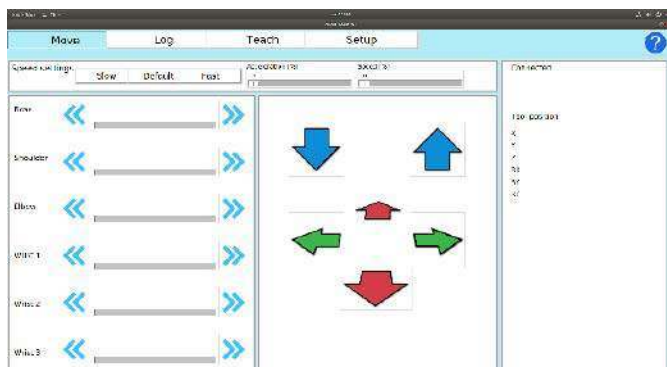
Robotyczne ramie Faze4

Faze4 to 6-osiove ramie robotyczne, wykonane z pomocą druku 3D. Wykorzystuje silniki krokowe i cykloidalne reduktory drukowane w 3D. Elektronika układu jest również dostępna i jak pisze autor projektu, Petar Crnjak, może ona być przez niego również dostarczona w formie złożonego modułu. Budowa tego robota była możliwa, gdy autor systemu opracował działającą przekładnię cykloidalną do wydrukowania w 3D. Następnie zaprojektował wokół niej całe ramie. Zawiera 6 silników krokowych, 3×NEMA 23, 2×NEMA 17 i 1×NEMA 14. Przeguby zawierają głównie cykloidalne przekładnie i pasy, a jeden zamiast tego przekładnię planetarną. Sześć kluczowych przegubów jest bazowane na czujnikach indukcyjnych, pozostałe zawierają krańcówki. Oprogramowanie ramienia działa na Teensy 3.5. Całkowity koszt budowy tego ramienia to około 1000 dolarów.

Cały projekt rozpoczął się, więc od napędów cykloidalnych, które mają następujące zalety:

- są naprawdę łatwe do produkcji w technice druku 3D,
- mogą mieć duży współczynnik redukcji,
- są tanie w budowie,
- mają niskie luzy, nawet przy wykonaniu z druku 3D.

Wszystkie elementy ramienia zostały zaprojektowane przy użyciu oprogramowania Solidworks. W ramieniu występują 3 wersje napędów cykloidalnych o różnych przełożeniach: wersja 11:1, wersja 15:1 i wersja 27:1.



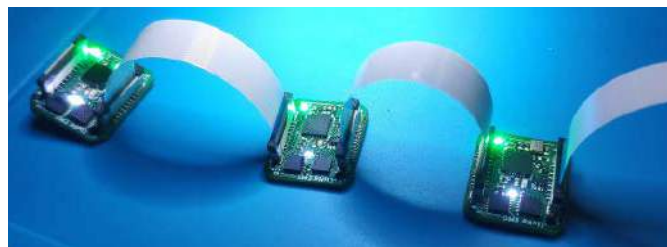
Jeśli chodzi o estetykę, wygląd systemu został zainspirowany modelem LR Mate 200iD firmy FANUC. Celem było również ukrycie wszystkich przewodów w ramieniu, tak jak robi to większość przemysłowych producentów. Jedyne widoczne przewody (lub rury) byłyby tymi dla chwytaka. Waga ramienia wynosi około 14...15 kg, ale można ją zmniejszyć drukując z mniejszym wypełnieniem. Dodatkowo, wszystkie silniki przeniesiono do podstawy ramienia, aby zmniejszyć wagę każdego segmentu. Pomysłem było również skopiowanie ramienia firmy Kuka, ale autor zrezygnował z tego projektu i po prostu wybrał prawdopodobnie najbardziej podstawowy projekt, w którym każdy silnik znajduje się bezpośrednio na przegubie.

Ramie udało się poprawnie uruchomić pod koniec zeszłego roku, ale przed autorem jeszcze sporo pracy nad optymalizacją, dlatego warto śledzić dalszy rozwój tego projektu.

<https://hackaday.io/project/167247-faze4-robotic-arm>

CleverHand – modułowe, przenośne urządzenia do elektromiografii powierzchniowej w systemach o dużej gęstości

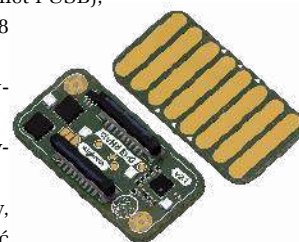
Robotyczne protezy to wysoko zaawansowane urządzenia, które wymagają bardzo drogich komponentów, aby zapewnić akceptowalną sprawność. W związku z tym znalezienie skutecznej protezy



ręki za mniej niż kilkadziesiąt tysięcy euro jest prawie niemożliwe. Aby rozwiązać ten problem, autor opracowuje element takiej inteligentnej ręki – niedrogą alternatywę do elektromiografii powierzchniowej o dużej gęstości (sEMG).

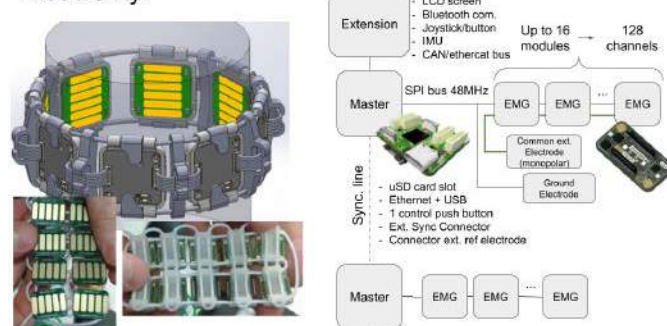
Tak prezentują się wymagania i ich obecny status w projekcie:

- układ mobilny – obecnie trwają prace m.in. nad łącznością Wi-Fi (zaimplementowano tylko Ethernet i USB),
- min. 16 kanałów (dostępne do 128 kanałów – 16×8),
- pasmo powyżej 2 kHz (użytkowano 2,5 kHz),
- rozdzielczość min. 16-bitów (użytkowano do 24 bitów),
- pomiar bipolarny i monopolarny,
- modułowość – można połączyć ze sobą do 16 modułów,
- niski koszt – poniżej 15 euro za moduł,
- w pełni otwarty projekt,
- 8 wbudowanych elektrod,
- elastyczne prowadzenie elektrod pomiarowych,
- wspólna elektroda zewnętrzna,
- cztery diody RGB LED.



Od bardzo prostych analogowych obwodów bazujących na wzmacniaczach operacyjnych, system przeszedł ewolucję do modułowych, wielokanałowych urządzeń cyfrowych. zastosowano interfejs SPI,

Modularity:



REKLAMA

BORNICO Teraz większe MOŻLIWOŚCI
bornico.com.pl

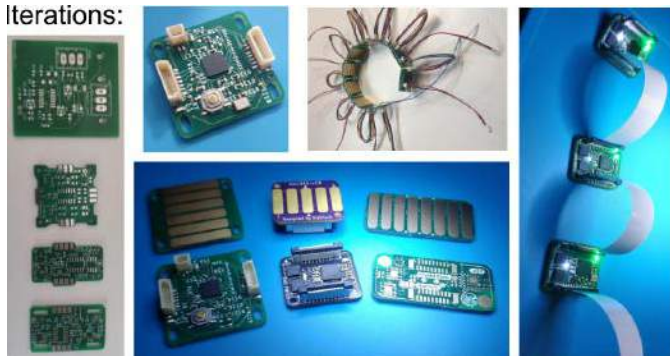
- montaż kontraktowy elektroniki
- projektowanie urządzeń i systemów

Zakład Elektroniczny BORNICO

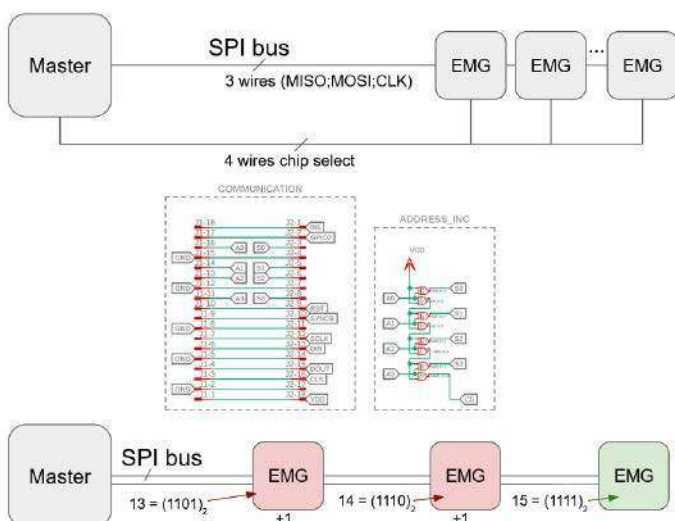
ul. Małczyńska 25
26-600 Radom
tel. +48 48 365 58 22
bornico@bornico.com.pl



Iterations:



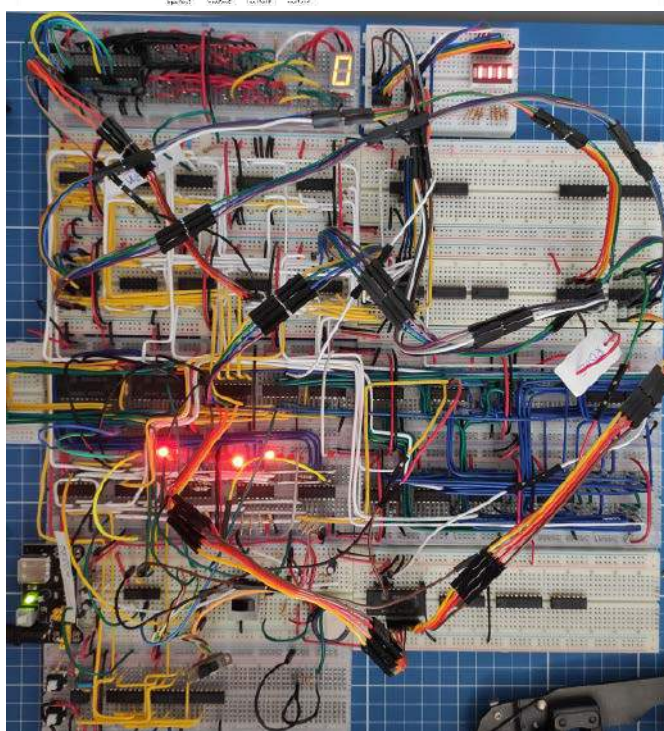
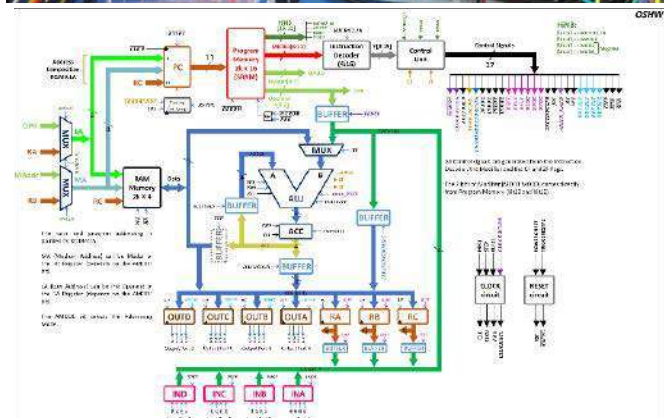
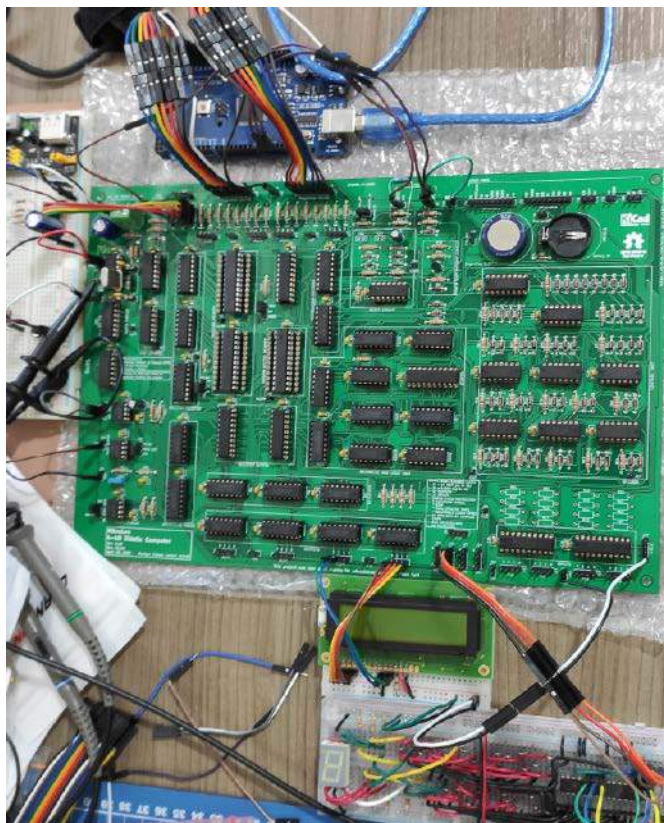
do zapewnienia ergonomicznej i szybkiej transmisji. Magistrale SPI są bardzo szybkie, jednak liczba potrzebnych przewodów wzrasta wraz z liczbą modułów podłączonych do magistrali. W tym projekcie ograniczeniem jest liczba łączy, ponieważ nie można sobie pozwolić na 3+16 linii tylko do komunikacji. Rozwiązaniem jest opracowanie systemu adresowania w celu programowego wybrania modułu. Standardowy system adresowania wymaga zakodowanego na stałe adresu dla każdego modułu. Można to zrobić za pomocą padów do przyłutowania na każdej płytce, aby ustawić ich indywidualny adres.



Niestety system ten nie jest bardzo modułowy, gdyż wymusza zróżnicowanie każdego modułu. Rozwiązaniem, jest adresowanie modułów według ich kolejności na magistrali. System adresowania działa na 4 liniach. Każdy moduł odczytuje 4-bitowy adres obecny na 4 przewodach, a następnie zwiększa o jeden adres dla następnego modułu. Każdy moduł jest aktywowany tylko wtedy, gdy odczytuje 0xF (0b1111, 15). W ten sposób, gdy moduł próbuje zwiększyć wartość do 0xF, aktywuje się bit przeniesienia, który jest podłączony do pinu wyboru chipa (CS) interfejsu SPI. Na schemacie zaprezentowano wizualne objaśnienie działania tego systemu adresowania.

MikroLeo – 4-bitowy mikrokomputer edukacyjny, rozpowszechniany na licencji open source

MikroLeo to 4-bitowy mikrokomputer dydaktyczny – projekt ten został opracowany głównie w celach edukacyjnych. Jak pisze autor, ma to być platforma do nauki dla tych, którzy lubią elektronikę, komputery i programowanie. Podjęto wiele wysiłków, aby ten projekt urzeczywistnić. Skierowany jest do studentów, entuzjastów i każdego innego, kto chce zrozumieć lub poszerzyć swoją wiedzę z zakresu elektroniki i dowiedzieć się, jak działa prosty komputer. Poza tym to także próba uratowania opowieści o początkach rozwoju układów scalonych i komputerów na jednym chipie, zademonstrowania możliwości, jakie miały te maszyny w czasach ich powstawania.



Jest to w pełni otwarty projekt sprzętu i oprogramowania, który można zbudować w domu. Tylko płytkę drukowaną musi być wyprodukowana przez jakąś zewnętrzną firmę na podstawie dokumentacji produkcyjnej, jaką właśnie przygotowuje autor tego projektu. Na razie aktualizacje są dokonywane tylko w repozytorium projektu na GitHubie. Główne cechy systemu to:

- 4-bitowy procesor,
- pamięć programu 2k×16 (do 4k),
- 2k×4 pamięci RAM (do 4k),
- 4 porty wyjściowe (16 wyjść),
- 4 porty wejściowe (16 wejść),
- instrukcje realizowane w czasie pojedynczego cyklu (RISC),
- architektura Harvardzka,
- 3 tryby wykonania:
- krok po kroku,
- z zegarem 3 MHz (dokładna podstawa czasu),
- z zegarem o regulowanej prędkości (1...200 Hz),
- brak MPU, mikrokontrolerów i innych złożonych układów scalonych,
- brak mikro kodu,
- adresowanie pośrednie w celu ułatwienia realizacji podprogramów,
- pamięć programu zaimplementowana w pamięć RAM ułatwia programowanie,
- akceptuje pamięci w rozmiarze 300 i 600 mil (w obudowach DIP),
- superkondensator lub bateria do podtrzymania przechowywania programu w pamięci RAM (dla wersji o niskim poborze mocy),
- może być programowany ręcznie za pomocą przełączników lub przez Arduino/ESP32,
- zbudowany z układów scalonych z rodziny 74HCTxxx zapewniających niskie zużycie energii i kompatybilność z obwodami TTL,
- wszystkie części są do montażu przewlekane, co ułatwia budowę układu w domu,
- wszystkie sygnały sterujące, rejestry i licznik programów są dostępne poprzez złącza dostępne z zewnątrz,
- całość umieszczona jest na pojedynczej płytce drukowanej z dwoma warstwami (wymiary płytki to 295,9×196,9 mm).

Należy pamiętać, że niektóre bufony służą do przeglądania wartości rejestrów w dowolnym momencie, ponieważ projekt ten jest przeznaczony głównie do celów edukacyjnych. Jeśli ten aspekt układu nie byłby istotny, możliwe jest uproszczenie konstrukcji systemu i pozbycie się części elementów.

<https://github.com/edson-acordi/4bit-microcomputer>

<https://hackaday.io/project/188340-mikroleo>

Loko – najmniejszy lokalizator GPS z baterią wystarczającą na 270 dni

Loko to 12-gramowy lokalizator o ekstremalnej żywotności baterii. Jest zasilany baterią radiowy lokalizator GPS o otwartym kodzie źródłowym. To małe, proste, przydatne urządzenie, które wysyła dane nawigacyjne do odbiornika za pośrednictwem interfejsu LoRa typu peer-to-peer. W przeciwieństwie do innych podobnych technologii, LoRa nie ma stałych kosztów subskrypcji

itp. Loko bazuje na komunikacji radiowej i nie wymaga karty SIM ani żadnej miesięcznej opłaty. Działa również w dowolnym miejscu, nawet, jeśli nie ma tam zasięgu 2G/3G/LTE. Jest

on konfigurowany do przesyłania danych nawigacyjnych do istniejących sieci LoRaWAN.

Loko został zaprojektowany z myślą o bardzo niskim zużyciu energii, dzięki czemu może pracować ponad 30 dni na jednym ładowaniu swojej małej, wbudowanej baterii pastylkowej. Nadajnik waży zaledwie 12 gramów i jest nieco większy od palca (30×23 mm). Loko działa z 2 jednostkami: jednostką naziemną (odbiornikiem) i jednostką powietrzną (nadajnikiem). Jednostka Air przesyła dane nawigacyjne do jednostki naziemnej we wstępnie skonfigurowanych odstępach czasu, a jednostka naziemna przesyła otrzymane dane do aplikacji na smartfona przez Bluetooth (obecna wersja tylko na iPhone). Aplikacja automatycznie pokazuje lokalizację dowolnych modułów Loko Air na mapie 2D.

Przykładową aplikacją jest użycie systemu do znalezienia zagubionego drona. Wystarczy podłączyć moduł Loko Air do swojego drona przed wylotem. Jeśli dron zaginie, włączamy jednostkę naziemną Loko, aby go szukać. Gdy jednostka naziemna odbierze sygnał z jednostki Air, jej lokalizacja pojawi się w aplikacji Loko. Zasięg działania systemu jest zależny od wielu czynników i warunków atmosferycznych, ale wynosi on powyżej 5 km i oferuje dokładność wymaganą do odnalezienia zagubionego urządzenia. Istnieje możliwość podłączenia do 30 jednostek powietrznych do jednej jednostki naziemnej.

<https://bit.ly/3GDPudu>

MIDIBUS – system działający z CAN do dystrybucji komunikatów MIDI w syntezytorze modułowym

Jest to prototyp magistrali do dystrybucji komunikatów MIDI w systemach syntezytorów modułowych, takich jak eurorack. Projekt pierwotnie miał być zgłoszeniem do konkursu hackaday Designlab 2021, ale zajęcia na uniwersytecie okazały się zbyt absorbujące, aby możliwe było dokończenie projektu na czas.

Jak podkreśla autor jest to jedynie proponowany transport MIDI z obsługą MIDI 2.0. Te specyfikacje mogą nie być ostateczne. Specyfikacja sprzętu wygląda obecnie następująco:

REKLAMA

ZAJRZYSZ NA TE STRONY

All In One

- Projektowanie i wykonywanie
 - modeli karkasów i obudów na drukarce 3D
 - transformatorów i induktorów
 - prototypów PCB
- Modelowanie 3D modułów i urządzeń
- Projektowanie urządzeń zasilających

Feryster – producent elementów EMC

www.feryster.pl

www.gamma.pl

info@gamma.pl

PODZESPOŁY ELEKTRONICZNE

RACK i Eurocarta 19" Wyposażenie szaf 19"

www.obudowa.pl

Producent obudów dla elektroniki tel. 032-230-2301

www.piekarczyk.pl

części elektroniczne

sprzedaz@piekarczyk.pl tel. 22 599 49 70



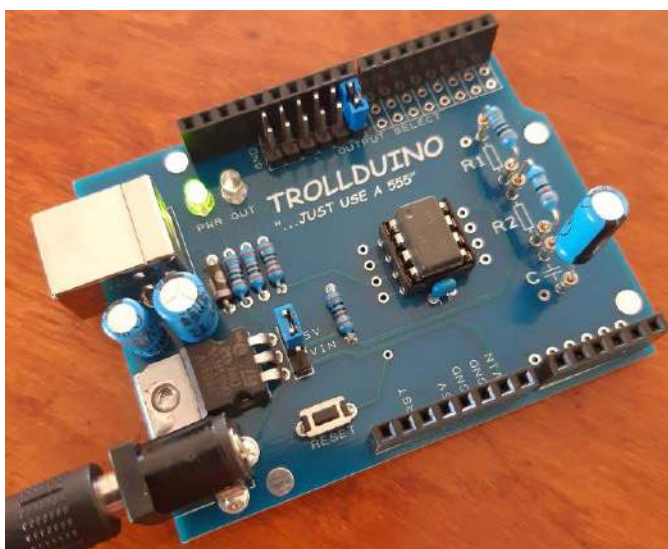


- 4-żyłowy płaski kabel taśmowy do połączenia,
- złącze MicroMatch lub WR-MM,
- najbardziej zewnętrzne piny złącza są podłączone do filtra RC,
- możliwość wyboru sposobu terminacji na każdym module,
- wbudowany kontroler i transceiver CAN FD.

Terminację końcową można włączyć na module za pomocą przełącznika DPDT lub za pomocą zworek. A pozycja końcowa powinna być wyraźnie zaznaczona na płytce drukowanej. Rozdzielone zakończenia jest preferowane w celu zminimalizowania EMR.

Specyfikacja komunikacji: zakłada szybkość magistrali 1 MHz (TSEG1=600 ns, TSEG2=300 ns, Sync=100 ns) i 11-bitowy, standardowy identyfikator (każde urządzenie na magistrali powinno mieć unikalny identyfikator). Komunikaty MIDI są przesyłane przez magistralę przy użyciu formatu UMP. Pole ID ramki CAN zawiera identyfikator urządzenia nadawczego do celów arbitrażu. Te identyfikatory powinny być generowane losowo. Ponadto urządzenie powinno wygenerować nowy identyfikator, jeśli otrzyma ramkę z własnym identyfikatorem.

<https://hackaday.io/project/182092-midibus>



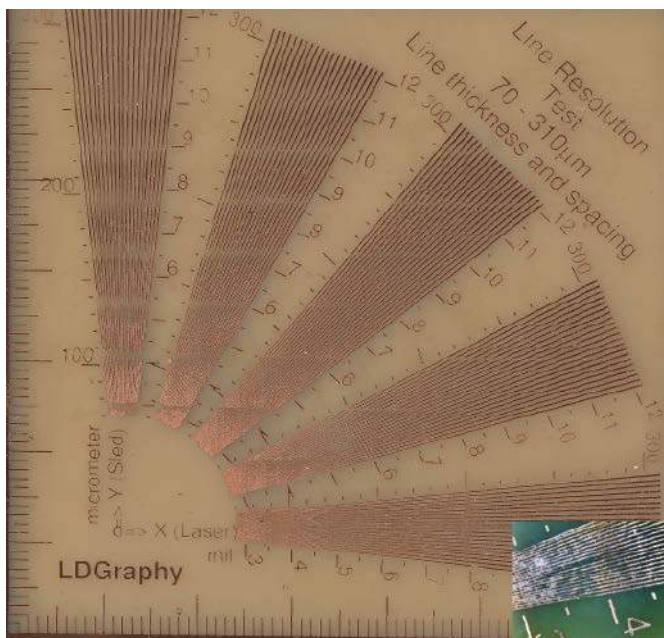
Trollduino

Projekt jest odpowiedzią na komentarze, które często widzi się na forach internetowych dla elektroników – „To głupie! Po co używać Arduino, skoro można zrobić to samo na 555?”. Niekoniecznie w tym brzmieniu, ale można nazwać to reprezentatywnym komentarzem do dowolnego postu z projektem, który bazuje na dowolnym module Arduino. Odpowiedzią na te komentarze jest moduł z układem 555 w tradycyjnym formacie Arduino UNO. Teraz mamy szansę

wreszcie sprostać wyzwaniu z tych komentarzy – zbudować układ z legendarnym timerem zamiast modułu Arduino z mikrokontrolerem. To idealny projekt dla bardziej oldschoolowych majsterkowiczów i miłośników prostych rozwiązań analogowych.

Nie trzeba uczyć się krzykliwego, mylącego, nowego, wymyślnego IDE. Programowanie analogowe jest proste i łatwe dzięki dwóm rezystorom i kondensatorowi za pomocą zintegrowanego prostego interfejsu programowania analogowego (ISAPI). Piny układu są kompatybilne z większością shieldów Arduino (funkcjonalność może być ograniczona).

<https://hackaday.io/project/176844-trollduino-v10>



Skanner laserowy Prism

Ten otwarty projekt to projektor laserowy o wysokiej rozdzielczości, odpowiedni do produkcji płytek drukowanych (PCB) lub drukowania 3D. Głowica laserowa może być używana do pisania na odpowiednio przygotowanym podłożu. Autor obecnie pracuje nad dodaniem czytania z podłoża. Celem tego projektu jest opracowanie głowicy laserowej do druku 3D lub wytwarzania PCB, która wykorzystuje obracający się pryzmat i jest przy tym łatwa w montażu. Obecnie, do testów, używany

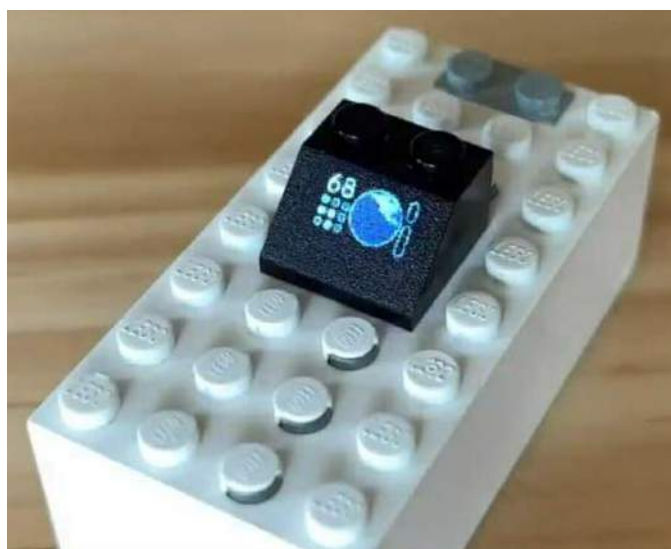
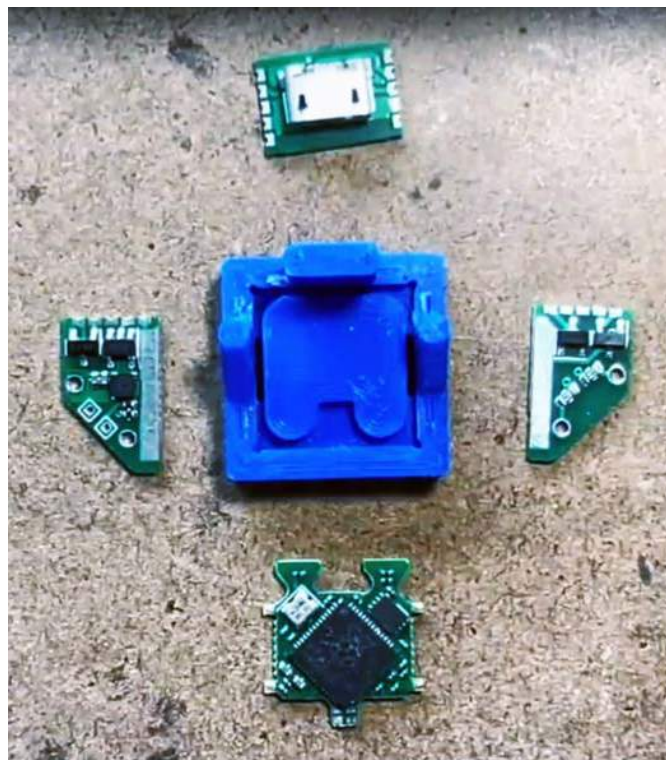


jest papier cyjanotypowy, ponieważ można go wywołać wodą. System wykorzystuje laser 405 nm o mocy 500 mW i pryzmat wirujący z prędkością do 21000 RPM, co przekłada się na przemieszczanie plamki lasera z prędkością do 34 metrów na sekundę. Rozdzielczość układu zależna jest od kilku czynników:

- rozmiaru plamki (FWHM) – średnica 25 mikrometrów,
- błędu skanera krzyżowego – 40 mikrometrów (błąd prostopadły do linii skanowania),
- dokładność stabilizacji kierunku skanowania – 2,2 mikrometra (wyłączenie i następnie włączenie głowicy skanującej),
- jitter – 35 mikronów (błąd równoległy do linii skanowania).

Maksymalna długość linii z skanera wynosi 24 mm. Sercem sterownika urządzenia jest FPGA ICE40UP5K-SG48, oprogramowany w pakiecie narzędzi Icestorm. Firestarter cape, sterownik lasera, wykorzystuje trzy sterowniki silników krokowych TMC2130, sterownik PWM i sterowanie wentylatorem. Obraz można wyeksportować np. z Kicada do SVG i przekonwertować na instrukcję dla głowicy laserowej. Wynik naświetlenia na płytce drukowanej ze zbliżeniem pokazano na zdjęciu. Można osiągnąć rozdzielczość poniżej 100 mikronów. Rozdział między ścieżkami jest całkiem dobry na papierze cyjanotypowym, ale wciąż trochę nierówny na PCB.

<https://hackaday.io/project/21933-prism-laser-scanner>



Kłoczek Lego z wbudowanym ekranem OLED i mikrokontrolerem RP2040

James Brown (Ancient) zbudował mały komputer w klocku LEGO, wykorzystując do tego mikrokontroler Raspberry Pi RP2040 i 0,42-calowy wyświetlacz OLED. W końcu minifigurki LEGO mają dostęp do odpowiedniego dla ich wzrostu prawdziwego komputera. Obecnie James nie opublikował jeszcze wielu informacji, jak samodzielnie odtworzyć kompilację, ale udostępnił kod napisany w Micropythonie,



aby wyświetlać obrazy w skali szarości na monochromatycznym ekranie OLED.

Na temat projektu więcej można dowiedzieć się z filmu, zamieszczonego w materiale źródłowym. Na filmie autor pokazuje, jak zmieścić komputer z Raspberry Pi RP2040 w klocku LEGO. Małutki projekt składa się z pięciu płytek drukowanych, które są lutowane z pomocą wydrukowanych w 3D uchwytów. Następnie lutowany jest wyświetlacz OLED i po upewnieniu się, że wszystko działa poprawnie, cały układ zatapia się w żywicy, odlanej na kształt klocka. System można następnie zasilac zasilaczem dla klocków LEGO lub z portu USB, który jest odsłonięty w elemencie. Przez port ten można również programować wbudowany mikrokontroler. I, jeśli ktoś się zastanawia... tak, na mikrokontrolerze może działać Doom.

<https://bit.ly/3Z8gdPg>

REKLAMA

OBWODY DRUKOWANE
Produkcja, Projektowanie, Montaż

Certyfikat Underwriters Laboratories 94V-0 E480148 TYPE 1 Zakład produkcyjny: 05-660 Warka ul. M. Ropelewskiej 17 tel. 22 781 63 95 22 761 95 80 fax. 22 781 63 95 w 23 www.elmax.waw.pl elmax@elmax.waw.pl	Płytki jednostronne Płytki dwustronne Płytki na podłożu aluminium Płyty czołowe FR4	Serie dowolne Prototypy Maksymalny wymiar płytek 1w 630 mm
	Dokumentacja technologiczna Dokumentacja konstrukcyjna Trawione szablon SMD	Montaż elektroniki Krótkie terminy Wykonania super ekspresowe
	Aktywny kalkulator prototypów na stronie internetowej	Pokrycie Sn lub SnPb inne na życzenie Maski, opisy montażowe w różnych kolorach

**Podstawowe parametry:**

- zakres pomiarowy temperatury: -40...125°C
- rozdzielczość pomiaru temperatury: 0,5°C
- dokładność pomiaru temperatury: ±2°C
- obsługiwane rozkazy: SEARCH_ROM, READ_ROM, MATCH_ROM, SKIP_ROM, CONVERT_T oraz READ_SCRATCHPAD,
- napięcie zasilania: 2,7...5 V, pobierany prąd: 6 mA.

*** Uwaga!** Elektroniczne zestawy do samodzielnego montażu. Wymagana umiejętność lutowania! Podstawową wersją zestawu jest wersja [B] nazywana potocznie KIT-em (z ang. zestaw). Zestaw w wersji [B] zawiera elementy elektroniczne (w tym [UK] – jeśli występuje w projekcie), które należy samodzielnie wlutować w dołączoną płytkę drukowaną (PCB). Wykaz elementów znajduje się w dokumentacji, która jest podlinkowana w opisie kitu. Mając na uwadze różne potrzeby naszych klientów, oferujemy dodatkowe wersje:

Dodatkowe materiały do pobrania ze strony www.ulubionykiosk.pl/media

- AVTS952 eT – wielokanałowy, bezprzewodowy system pomiaru temperatury (EP 9/2022)
- AVTS949 Energooszczędny termometr LED (EP 8/2022)
- AVTS892 Energooszczędny termometr z kalibracją (EP 10/2021)
- AVTS635 Bezprzewodowy, energooszczędny system pomiaru temperatury (EP 8-9/2018)
- AVT1999 2-kanałowy termometr MIN-MAX z alarmem (EP 8/2018)
- AVTS623 4-kanałowy termometr z interfejsem Wi-Fi (EP 4/2018)
- AVTS566 THPStation – rozbudowany termometr z Wi-Fi (EP 1/2017)
- AVTS535 Termometr 2-kanałowy z interfejsem Bluetooth (EP 4/2016)
- AVTS518 Termometr bezprzewodowy (EP 11/2015)
- AVT1863 Termometr z interfejsem Bluetooth (EP 8/2015)
- AVT1790 Termometr XXL (EP 2/2014)
- AVTS489 8-kanałowy termometr z alarmem i wyświetlaczem LCD (EP 11/2013)
- AVTS420 Wielopunktowy termometr z rejestracją (EP 10/2013)
- AVT1734 Termometr do wędzarni (EP 4/2013)
- AVTS373 Tlogger – rejestrator temperatury (EP 12/2012)
- AVT1705 Moduł do pomiaru temperatury z interfejsem RS485 (EP 9/2012)

- wersja [C] – zmontowany, uruchomiony i przetestowany zestaw [B] (elementy wlutowane w płytkę PCB),
 - wersja [A] – płytka drukowana bez elementów i dokumentacji.
- Kity, w których występuje układ scalony wymagają zaprogramowania, mają następujące dodatkowe wersje:
- wersja [A+] – płytka drukowana [A] + zaprogramowany układ [UK] i dokumentacja,
 - wersja [UK] – zaprogramowany układ.

Nie każdy zestaw AVT występuje we wszystkich wersjach! Każda wersja ma załączony ten sam plik PDF! Podczas składania zamówienia upewnij się, którą wersję zamawiasz! <http://sklep.avt.pl>

W przypadku braku dostępności na stronie sklepu osoby zainteresowane zakupem płytek drukowanych (PCB) prosimy o kontakt via e-mail: kity@avt.pl.

W ofercie AVT*

AVT5965

DS18S20 – emulator czujnika temperatury DS1820

Przykład programowej realizacji urządzenia 1-Wire slave (1)

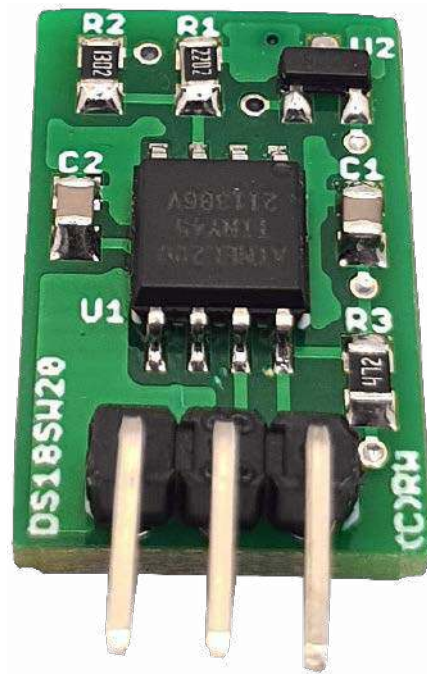
Pomysł na ten projekt narodził się w dość nietypowych okolicznościach. Otóż, w jednym z urządzeń dostarczonych przez klienta byłem zmuszony wymienić uszkodzony termometr scalony pod postacią bardzo popularnego układu typu DS1820 firmy Maxim Integrated (dawniej Dallas Semiconductor) pracującego pod kontrolą interfejsu 1-Wire. Ku mojemu wielkiemu zaskoczeniu dość szybko okazało się, że zdobycie tego rodzaju, jakby nie patrzeć prostego, elementu graniczy z cudem, a ceny u dystrybutorów sięgają astronomicznych kwot rzędu 50 do nawet 120 zł za sztukę (!).

To niewyobrażalne, aby element tego typu osiągał tak kuriozalne ceny, gdy jeszcze jakiś czas temu można było go kupić za równowartość 1\$. Rozumiem bieżące problemy z dostępnością jakichkolwiek półprzewodników (a zwłaszcza mikrokontrolerów) lecz nic nie uzasadnia tak absurdalnej sytuacji. Na pocieszenie można dodać, że znacznie przystępniejsze, powiedziałbym normalne, ceny ma inny termometr tego producenta a mianowicie DS18B20. Niestety oba wspomniane peryferia nie są bezpośrednimi odpowiednikami, gdyż różnią się dokładnością, a co najważniejsze mają inną organizację wbudowanej pamięci (scratchpad) przechowującej wartość mierzonej temperatury a co za tym idzie nie można ich stosować zamiennie. I właśnie wtedy do głowy wpadł mi dość oryginalny, jak mi się wydaje, pomysł skonstruowania własnego termometru DS1820, którego projekt nazwałem dość sugestywnie, a mianowicie DS18SW20 (SW od słowa *software*).

Jednak zaznaczam, że nie będę realizował całej dostępnej funkcjonalności termometru DS1820 (choć zapewne z 90%) a skupię się wyłącznie na tych możliwościach, które zazwyczaj potrzebne są w systemach mikroprocesorowych. Zaprezentuję jednak szczegółowy opis implementacji w języku C, w związku z czym zainteresowani Czytelnicy bez problemu będą mogli dodać brakujące funkcje.

Magistrala 1-Wire

Zanim jednak przejdę do opisu samego urządzenia, jak i szczegółów implementacyjnych, nie sposób choćby skrótowo nie przypomnieć informacji na temat bardzo interesującej magistrali 1-Wire, tym razem jednak z punktu widzenia układu podrzędnego typu Slave. Tak jak wspomniano wcześniej, komunikacja na magistrali 1-Wire odbywa się przy udziale wyłącznie jednego przewodu (stąd nazwa interfejsu) oznaczonego jako DQ, który



może jednocześnie pełnić rolę przewodu zasilającego w konfiguracji tzw. zasilania pasywnego (de facto nieobsługiwane przez nasze urządzenie DS18SW20). W przypadku magistrali 1-Wire, tak jak w przypadku większości interfejsów szeregowych, transmisja przebiega w konfiguracji Master↔Slave.

Układ nadrzędny (Master) steruje wyszukiwaniem i adresowaniem układów podrzędnych (Slave), steruje przepływem danych oraz

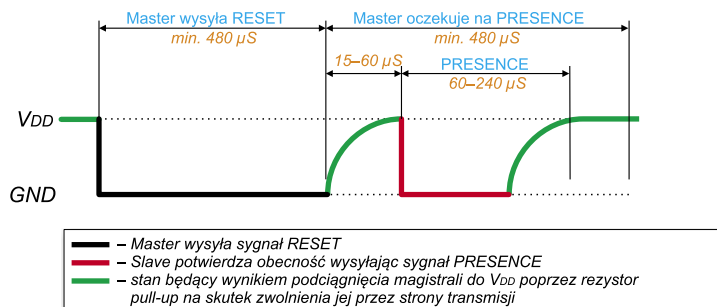
generuje sygnał zegarowy (inicjuje wysyłanie i odbieranie danych). Dane przesyłane są synchronicznie z prędkością do 16,3 kbps w trybie standard oraz do 115 kbps w trybie *overdrive*. Należy szczególnie podkreślić, że przesłanie każdego bitu informacji niezależnie od kierunku transmisji inicjowane jest wyłącznie przez układ Master za pomocą wygenerowania opadającego zbocza sygnału (ściągnięcie magistrali do logicznego „0” przez czas z zakresu 1...5 μ s). Po wystąpieniu takiego zbocza sygnału układ Slave podejmuje różne działania, których scenariusz zależy od oczekiwanego kierunku transmisji. Tego typu organizacja protokołu transmisji zapewnia prawidłową synchronizację przesyłanych danych bez potrzeby stosowania dodatkowych linii sterujących.

Minimalny czas trwania pojedynczego bitu jest ściśle określony i wynosi $60 \mu\text{s} + 1 \mu\text{s}$ na tak zwany czas odtworzenia zasilania (*recovery time*). Wyznacza on maksymalną prędkość transmisji w trybie standard ($1/61 \mu\text{s} = 16,3 \text{ kbps}$). Co ważne, każde z urządzeń podłączonych do magistrali musi mieć wyjście typu otwarty dren lub otwarty kolektor, a linia danych połączona jest do zasilania przez rezystor podciągający o typowej wartości 4,7 k Ω , co w stanie bezczynności powoduje utrzymywanie się stanu wysokiego na tej linii zapewniającego zasilanie urządzeń podrzędnych (jeśli pracują w trybie zasilania pasywnego).

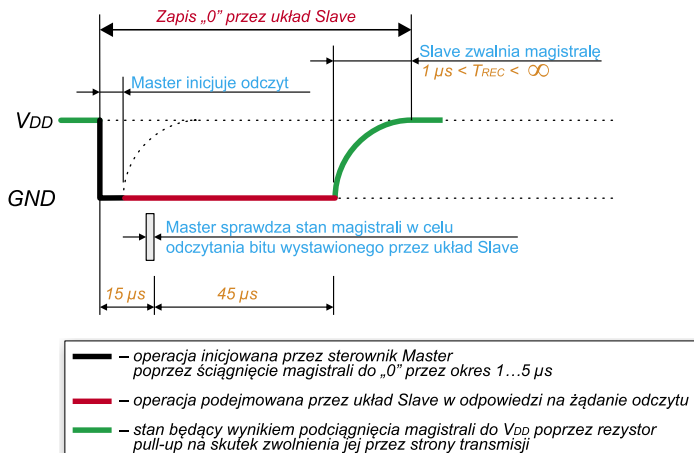
Sama magistrala nie ma ustalonego formatu danych a sposób przesyłania informacji zależy od konfiguracji i właściwości układów podrzędnych. Przesyłane słowa są zawsze jednobajtowe a jako pierwszy transmitowany jest bit mniej znaczący. Dodatkową i jedną z najważniejszych cech urządzeń z interfejsem 1-Wire, o czym wspomniano na wstępie, odróżniającą je jednocześnie np. od urządzeń standardu I²C, jest unikatowy, ośmiobajtowy adres zapisany w pamięci ROM peryferium. Adres ten jest niepowtarzalny i właściwy tylko i wyłącznie pojedynczemu układowi scalonemu (dla elementów produkowanych przez firmę Maxim/Dallas zapisywany jest na etapie produkcji). Najmniej znaczący bajt tego adresu zawiera kod rodziny układów (Family code), kolejne 6 bajtów zawiera unikatowy kod konkretnego egzemplarza (właściwy adresu układu) a najbardziej znaczący bajt zawiera sumę kontrolną CRC8 (*Cyclic Redundancy Check*). Suma ta wyliczana jest na podstawie poprzednich siedmiu bajtów i jest ustalana na etapie produkcji (służy do kontroli poprawności transmisji).

Protokół transmisji danych interfejsu 1-Wire definiuje kilka, podstawowych sygnałów sterujących i stanów pracy magistrali:

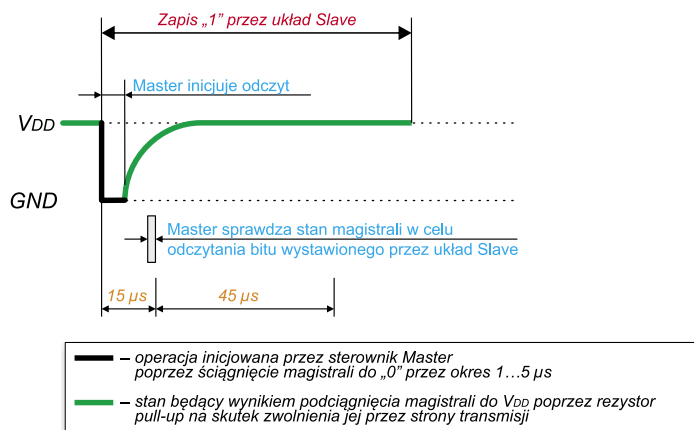
- sygnał *Reset*, wysyłany przez układ Master, będący żądaniem zgłoszenia się układów Slave,



Rysunek 1. Procedura inicjalizacji magistrali 1-Wire



Rysunek 2. Sekwencja sygnałów sterujących obrazująca operację zapisu logicznego „0” wykonywaną przez układ Slave a inicjowaną przez układ Master



Rysunek 3. Sekwencja sygnałów sterujących obrazująca operację zapisu logicznej „1” wykonywaną przez układ Slave a inicjowaną przez układ Master

- sygnał *Presence*, wysyłany przez układy Slave, będący potwierdzeniem obecności tych układów na magistrali danych,
- zapis logicznej „1” i „0”,
- odczyt logicznej „1” i „0”.

Najlepszym sposobem na zrozumienie zależności czasowych dla poszczególnych stanów pracy magistrali jest zobrazowanie tych sygnałów widziane z perspektywy układu Slave. Na **rysunku 1** pokazano sekwencję inicjalizacji magistrali 1-Wire, która to, jak wspomniano wcześniej, umożliwia układowi Master wykrycie podłączonych do magistrali układów Slave. Sekwencja ta rozpoczyna się poprzez wysłanie przez układ Master sygnału *Reset* (ściągnięcie magistrali do masy przez czas 480 μ s)

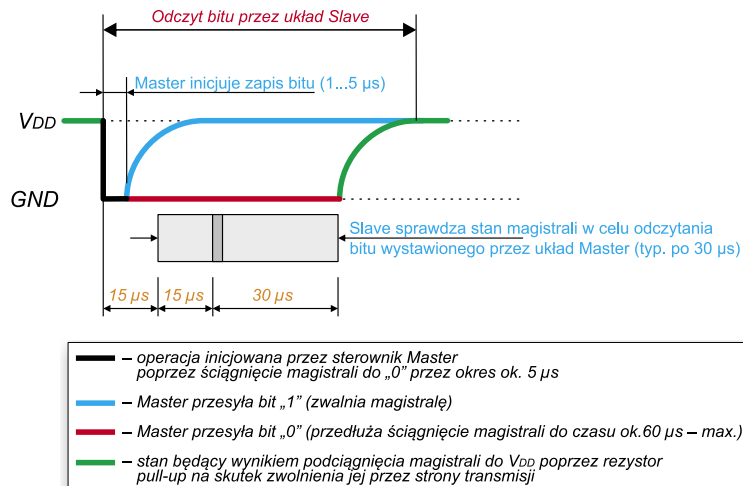
i po odczekaniu czasu 15...60 μ s, odpowiedzią układów Slave poprzez wysłanie sygnału *Presence* (ściągnięcie magistrali do masy przez czas 60...240 μ s).

Na **rysunkach 2 i 3** pokazano z kolei sekwencje sygnałów sterujących obrazujące operację zapisu danych (logicznego „0” i „1”) przez układ Slave na magistralę 1-Wire a będącą wynikiem żądania odczytu ze strony układu Master (ściągnięcie magistrali do masy przez czas 1...5 μ s). Dalej, na **rysunku 4** pokazano sekwencję sygnałów sterujących obrazującą operację odczytu danych wykonywaną przez układ Slave a będącą wynikiem zapisu tychże danych dokonanego przez układ Master (jak wcześniej, inicjowaną poprzez ściągnięcie magistrali do masy przez czas 1...5 μ s).

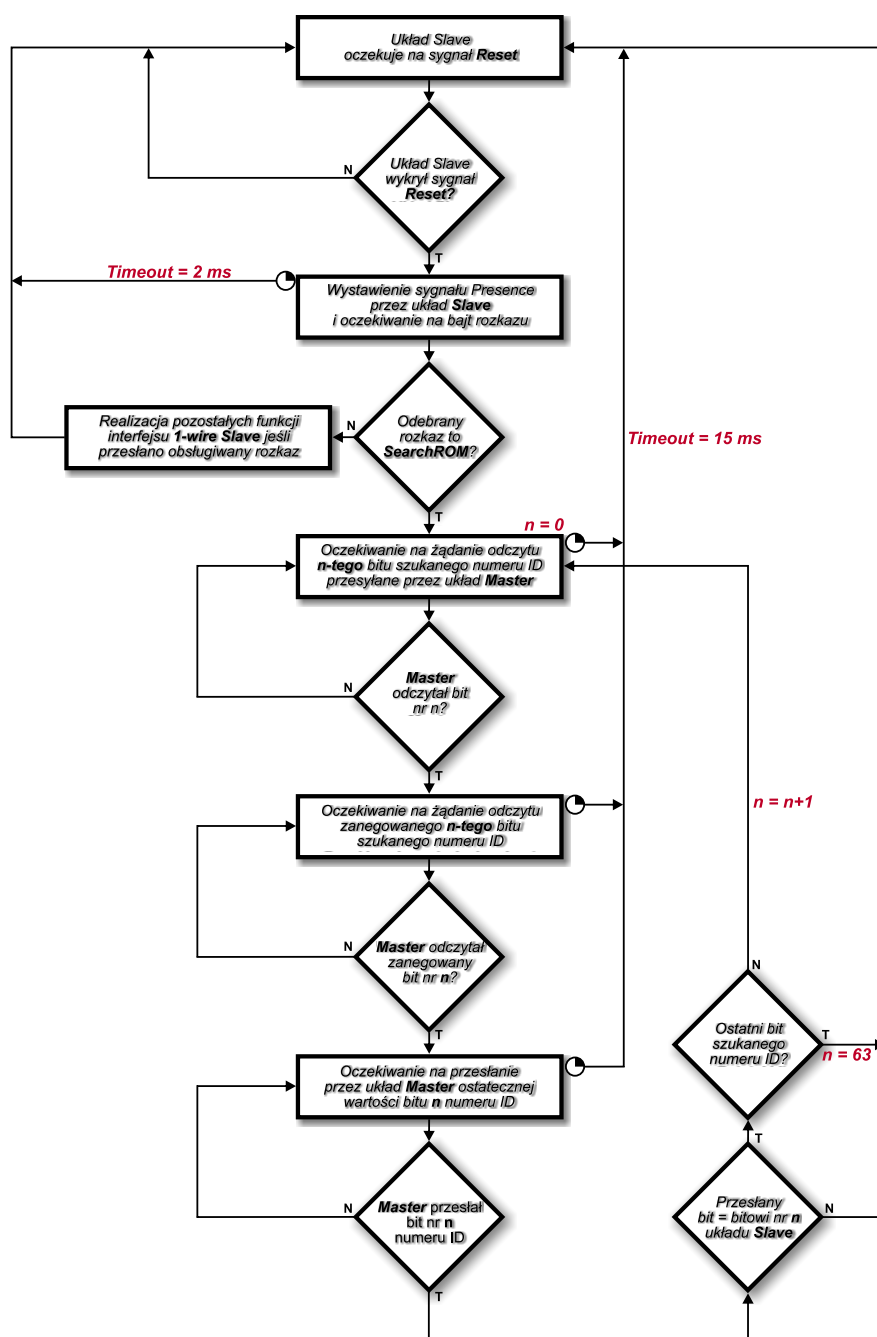
Na koniec, na **rysunku 5**, zaprezentowany jest graf funkcjonalny badając najciekawszego mechanizmu charakterystycznego wyłączenie dla magistrali tego typu a przeznaczonego do wyszukania adresów wszystkich układów do niej przyłączonych nazywany SEARCH ROM. Należy bowiem pamiętać, że układ Master nie musi znać adresów przyłączonych do magistrali układów, w związku z czym musi istnieć pewny mechanizm na ich ustalenie. Operacja taka inicjowana jest poprzez wysłanie przez układ Master rozkazu 0xF0 (zwanego SEARCH ROM). Następnie każdy układ Slave wystawia na magistralę wartość pierwszego, najmniej znaczącego bitu swojego numeru ID (oczywiście zapis inicjowany jest zawsze przez układ Master). Mając na uwadze specyfikę interfejsu 1-Wire, polegającą na tym, że wszystkie układy podłączone są do tej samej magistrali danych należy pamiętać, że na odebraną komendę przeszukiwania odpowiedzą dokładnie w tym samym czasie wszystkie układy Slave w związku, z czym rezultat wystawienia bitów docierający do układu Master jest iloczynem logicznym stanów wyjść wszystkich tych układów (tzw. *wired and*). Kolejnym krokiem jest wystawianie przez układy Slave zanieganowanego, pierwszego bitu swojego numeru ID (oczywiście jak zwykle, na żądanie układu Master). W tym momencie układ Master decyduje o przyjętej wartości pierwszego bitu adresu i wystawia go na magistralę, a układy Slave odczytują go. Jeśli przesłana przez układ Master wartość (w tym trzecim kroku) pierwszego bitu numeru ID jest zgodna z rzeczywistą wartością pierwszego, najmniej znaczącego bitu numeru ID wybranego układu Slave, układ ten kontynuuje proces aż do przebrnięcia przez wszystkie 64 bity numeru. Jeśli taka zgodność nie występuje, układ Slave pozostaje nieaktywny aż do następnego sygnału Reset i żądania wyszukiwania numerów ID (rozkazu 0xF0). W ten prosty sposób układ Master, po przejściu przez każde 64 bity numeru ID, dysponuje każdorazowo kolejnym, znalezionym numerem ID układu pracującego na magistrali 1-Wire. Oczywiście opis ten pokazano z punktu widzenia układu Slave, gdyż algorytm dla układu Master jest nieco bardziej skomplikowany (bazuje na algorytmie tzw. drzewa binarnego).

Budowa i działanie

W tym miejscu dysponujemy już niezbędną wiedzą z zakresu protokołu 1-Wire przejdźmy zatem do schematu ideowego naszego urządzenia, który pokazano na **rysunku 6**. Jak widać, zaprojektowano bardzo prosty, wręcz trywialny, system mikroprocesorowy, którego sercem jest niewielki mikrokontroler ATtiny25 taktowany wewnętrznym oscylatorem 8 MHz odpowiedzialny za realizację całej, założonej funkcjonalności.



Rysunek 4. Sekwencja sygnałów sterujących obrazująca operację odczytu stanu magistrali 1-Wire wykonywaną przez układ Slave a inicjowaną przez układ Master



Rysunek 5. Sekwencja sygnałów sterujących obrazująca mechanizm znajdowania numerów ID układów Slave (z punktu widzenia tychże układów)

Wykaz elementów, kupuj na stronie sklep.avt.pl (Warszawa, ul. Leszczyńska 11, tel. +48222578451, e-mail: handlowy@avt.pl)

Rezystory: (SMD0805)

R1: 22 kΩ 1%

R2: 13 kΩ 1%

R3: 4.7 kΩ

Kondensatory: (SMD0805)

C1, C2: 100 nF ceramiczny X7R

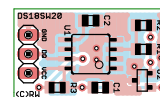
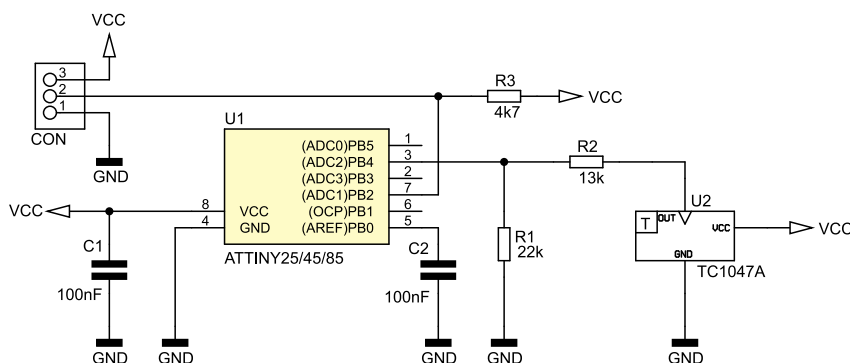
U2: TC1047A (SOT-23)

Półprzewodniki:

U1: ATtiny25/45/85 (SOIC-8)

Pozostałe:

CON: złącze GOLDPIN kątowe 3×1 pin



Rysunek 7. Schemat montażowy urządzenia DS18S20

Rysunek 6. Schemat ideowy urządzenia DS18S20

Jako element pomiarowy (termometr) zastosowano scalony przetwornik temperatura-napięcie pod postacią układu scalonego TC1047A produkcji Microchip cechujący się doskonałą liniowością (z nachyleniem 10 mV/°C) oraz akceptowalną dokładnością na poziomie $\pm 2^\circ\text{C}$. Nie jest to co prawda dokładność emulowanego termometru typu DS1820 (lub DS18S20), ale moim zdaniem jest jak najbardziej akceptowalna, zaś sam układ jest łatwo dostępny i niedrogi. Co oczywiste, do odczytu wskazań termometru użyto wbudowanego w strukturę mikrokontrolera 10-bitowego przetwornika ADC (kanał ADC2) oraz wewnętrznego napięcia odniesienia o wartości 1,1 V. Odczytywane 10-bitowe dane z przetwornika ADC przeskalowano w taki sposób, aby uzyskać rozdzielczość pomiaru równą $0,5^\circ\text{C}$ oraz zakres mierzonych temperatur rzędu $-40\dots+125^\circ\text{C}$ (jak

dla emulowanego termometru). Jako, że poziom napięć wyjściowych układu TC1047A (100...1750 mV) dla pełnego zakresu mierzonych temperatur nie mieści się w zakresie dopuszczalnych napięć wejściowych przetwornika ADC (w przypadku korzystania z wewnętrznej referencji 1,1 V) zastosowano prosty dzielnik napięcia pod postacią rezystorów R1/R2 (koniecznie dokładnych – 1%) niwelujący tę niedogodność.

Montaż i uruchomienie

W tym momencie przejdźmy do szczegółów dotyczących montażu układu. Schemat płytki PCB urządzenia DS18S20 pokazano na **rysunku 7**. Jak widać zaprojektowano bardzo kompaktowy (12×20 mm) dwustronny obwód drukowany z wyłącznym montażem elementów SMD. Montaż urządzenia rozpoczynamy

od przylutowania półprzewodników (o nieproblematycznych obudowach), następnie lutujemy elementy bierne a na samym końcu złącze GOLDPIN. Poprawnie zmontowane i zaprogramowane urządzenie nie wymaga żadnego uruchamiania i powinno działać już po podłączeniu napięcia zasilającego.

Jedynym procesem, jaki możemy wykonać jest zaprogramowanie opcjonalnego adresu urządzenia Slave w pamięci EEPROM mikrokontrolera. W przypadku problemów z dostępnością przetwornika TC1047A u dystrybutorów elektroniki bez wahania możemy go zamówić na stronie producenta układu (z resztą w niższej cenie, niż u dystrybutorów) lub otrzymać za darmo (i to kilka sztuk) w ramach programu bezpłatnych próbek (sample) z czego skorzystałem podczas montażu urządzenia.

Tyle w kwestiach sprzętowych, których nie ma nazbyt wiele, gdyż cała magia dzieje się w oprogramowaniu. W drugiej części artykułu, która ukaże się w kolejnym wydaniu EP, dokładnie omówimy najważniejsze fragmenty kodu programu sterującego urządzeniem.

Robert Wołgajew, EP

REKLAMA



O projektach, miniprojektach, projektach soft i na wiele innych tematów dyskutuj na forum.ep.com.pl

**Podstawowe parametry:**

- komunikacja z potwierdzaniem poprawności transmisji,
- moduł nadawczy zasilany z baterii CR2016 (teoretyczny czas pracy na jednej baterii: 10 lat),
- maksymalny prąd obciążenia (tryb uśpienia/nadawanie): 0,2 μ A / 22,6 mA,
- zasilanie odbiornika: 230 V AC, ok. 2 W,
- 12 wyjść przełącznikowych,
- maksymalne obciążenie wyjść: 5 A/250 VAC
- częstotliwość pracy transceivera: 868 MHz
- zasięg w terenie otwartym: ok. 100 m

*** Uwaga!** Elektroniczne zestawy do samodzielnego montażu. Wymagana umiejętność lutowania! Podstawową wersją zestawu jest wersja [B] nazywana potocznie KIT-em (z ang. zestaw). Zestaw w wersji [B] zawiera elementy elektroniczne (w tym [UK] – jeśli występuje w projekcie), które należy samodzielnie wstawić w dołączoną płytkę drukowaną (PCB). Wykaz elementów znajduje się w dokumentacji, która jest podlinkowana w opisie kitu. Mając na uwadze różne potrzeby naszych klientów, oferujemy dodatkowe wersje:

Dodatkowe materiały do pobrania ze strony www.ulubionykiosk.pl/media

- | | |
|---------|---|
| AVT5960 | Przełącznik elektromagnetyczny sterowany optoelektronicznie (EP 11/2022) |
| AVT5876 | Energooszczędny przełącznik bistabilny (EP 8/2021) |
| AVT5794 | Moduł przełącznikowy z gasikami (EP 8/2020) |
| AVT5710 | 8-kanalowy moduł przełącznikowy z USB (EP 8/2019) |
| AVT5682 | Przełącznik elektromagnetyczny 230 V sterowany optoelektronicznie (EP 6/2019) |
| AVT5632 | Moduł przełączników z interfejsem USB (EP 3/2019) |
| AVT5588 | Sterownik-timer z 8 przełącznikami (EP 6/2017) |
| AVT1916 | Konfigurowalny przełącznik 4-kanalowy (EP 8/2016) |
| AVT1890 | Moduł przełączników z USB (EP 6/2016) |
| AVT1815 | 4-kanalowy przełącznik sterowany dowolnym pilotem IR (EP 8/2014) |
| AVT5368 | Programowalny moduł przełączników (EP 11/2012) |
| AVT5353 | Moduł przełączników z interfejsem USB (EP 7/2012) |
| AVT1691 | Uniwersalny moduł przełącznikowy (EP 8/2012) |
| AVT1659 | 8-kanalowy miniaturowy moduł przełączników (EP 1/2012) |
| AVT1656 | Uniwersalny moduł wykonawczy (EP 12/2011) |
| AVT1560 | 8-kanalowa karta przełączników (EP 2/2010) |

- **wersja [C]** – zmontowany, uruchomiony i przetestowany zestaw [B] (elementy wstawiane w płytkę PCB),
 - **wersja [A]** – płytka drukowana bez elementów i dokumentacji.
- Kit, w którym występuje układ scalony wymagający zaprogramowania, mają następujące dodatkowe wersje:
- **wersja [A+]** – płytka drukowana [A] + zaprogramowany układ
 - **[UK]** i dokumentacja,
 - **wersja [UK]** – zaprogramowany układ.

Nie każdy zestaw AVT występuje we wszystkich wersjach! Każda wersja ma załączony ten sam plik PDF! Podczas składania zamówienia upewnij się, którą wersję zamawiasz!
<http://sklep.avt.pl>

W przypadku braku dostępności na stronie sklepu osoby zainteresowane zakupem płytek drukowanych (PCB) prosimy o kontakt via e-mail: kity@avt.pl.

W ofercie AVT*
AVT5966



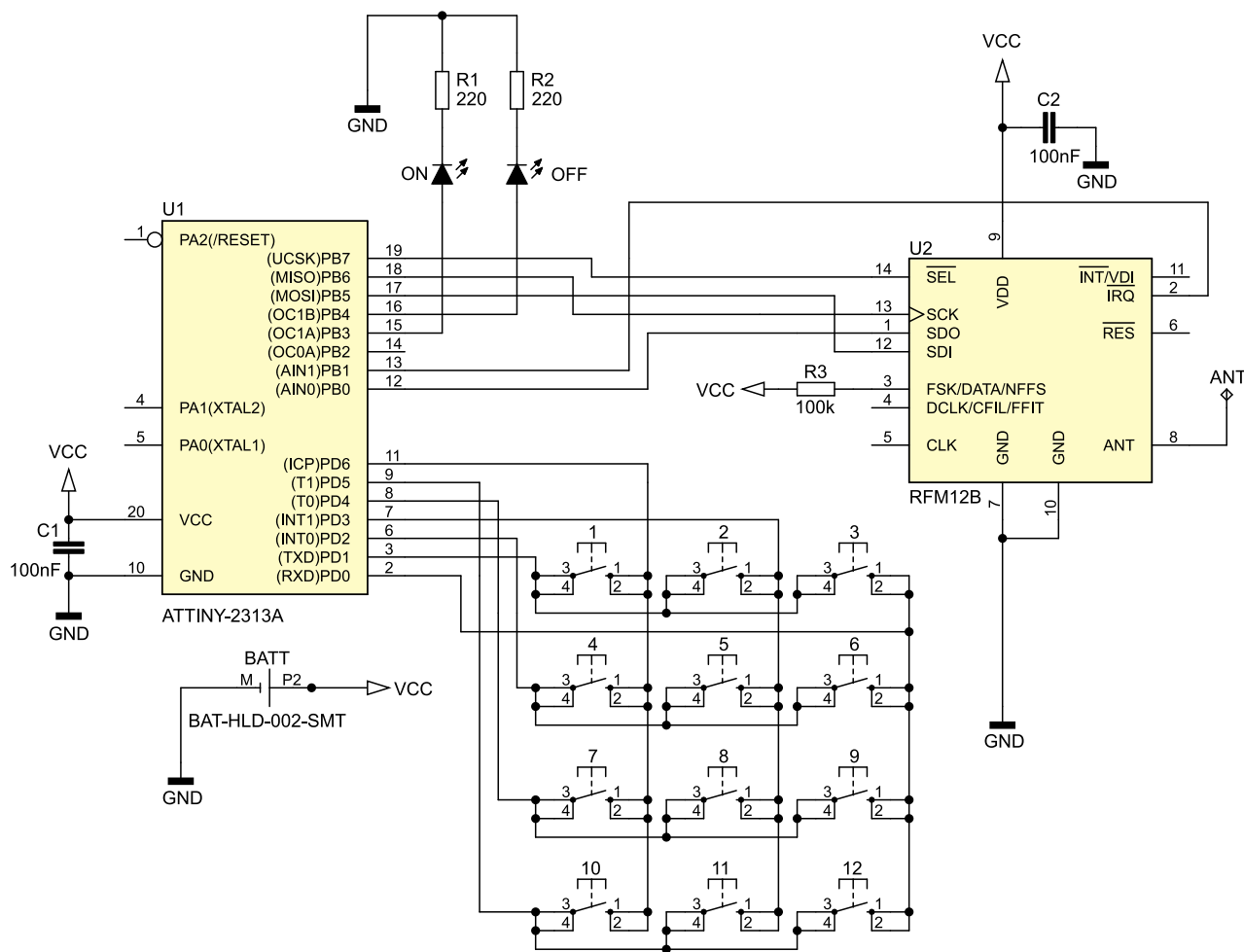
wiRelay – bezprzewodowa, 12-kanalowa karta przełączników

Tym razem, zupełnie inaczej, niż miało to miejsce zazwyczaj, pomysł na niniejszy projekt narodził się w głowie...mojego kolegi z redakcji. Zaproponowano mi zaprojektowanie 12-kanalowej, bezprzewodowej karty przełączników, gdyż, jak się okazuje, tego rodzaju urządzenia dość często znajdują zastosowanie w gospodarstwach domowych, zwłaszcza wtedy, gdy zamieszkujemy domek jednorodzinny. Co prawda w swoim portfolio mam już urządzenie o podobnej funkcjonalności, lecz tym razem postanowiłem zbudować system o większej uniwersalności i lepszych cechach użytkowych.

Kompletny system składa się z nadajnika i odbiornika rozkazów sterujących. Głównym problemem implementacyjnym było znalezienie sprawdzonego i energooszczędnego modułu do bezprzewodowej transmisji danych. Co oczywiste, zależało mi na tym,

aby nadajnik tego systemu pracował na zasilaniu bateryjnym i charakteryzował się minimalnym poborem mocy zapewniającym długą pracę urządzenia. Ponieważ w moim poprzednim projekcie energooszczędnego systemu pomiaru temperatury z EP 8/2018 [1]

i EP 9/2018 [2] zdobyłem duże doświadczenie w zakresie obsługi bardzo ciekawych modułów RFM12B pracujących w paśmie 433, 868 lub 915 MHz (w zależności od wersji), należących do bogatej oferty modułów radiowych produkowanych przez firmę HopeRF, zdecydowałem się na ich zastosowanie również w tej konstrukcji. Moduły, o których mowa stanowią kompletne rozwiązanie toru radiowego nadawczo-odbiorczego. Dostarczają wygodny interfejs komunikacyjny SPI pozwalający na przeprowadzenie pełnej konfiguracji elementu w ramach dostępnej szerokiej palety ustawień jak i sterowanie komunikacją radiową. Nie będę powtarzał informacji dotyczących specyfikacji i obsługi tych peryferiów, gdyż takowe



Rysunek 1. Schemat ideowy nadajnika systemu wiRelay

zamieściłem w ramach wspomnianych wcześniej artykułów. Zainteresowanych tymi szczegółami Czytelników odsyłam do wskazanych wcześniej materiałów. Zanim jednak przejdę do schematów urządzenia kilka słów napiszę na temat specyfikacji całego systemu.

Tak, jak wspomniano powyżej, system wiRelay składa się z dwóch modułów komunikacyjnych: nadawczego (pilota), który obsługuje klawiaturę sterującą i odbiorczego (karty przekazników), który steruje wbudowanymi przekaznikami dużej mocy. Moduł nadawczy pracuje na zasilaniu baterijnym w postaci pastylki CR2016 i większość swojego czasu spędza w trybie uśpienia (dla ograniczenia poboru mocy) czekając na zmianę stanu obsługiwanej klawiatury matrycowej. Wspomnianej zmianie stanu towarzyszy wybudzenie

układu (mikrokontrolera i modułu RF), wysłanie komunikatu do karty przekazników, odbiór odpowiedzi i ponowne uśpienie urządzenia. W ten prosty sposób ograniczamy do minimum pobór energii ze źródła zasilania pozwalając na wieloletnią pracę urządzenia.

Uważny Czytelnik dostrzeże pewne ograniczenia i sformułuje związane z nimi zapytania. Otóż bateria CR2016 (3 V) zastosowana w urządzeniu przeznaczona jest do zasilania urządzeń cechujących się bardzo niskim poborem prądu – rzędu ułamków mA, do pojedynczych mA. Nasz układ po wybudzeniu aktywuje nadajnik modułu RFM12B, który w czasie transmisji pobiera prąd rzędu 22,6 mA, co stanowi bardzo duże obciążenie dla niewielkiej baterii zasilającej. Na szczęście transmisja trwa około

20 ms, w związku z czym pobrana energia jest bardzo mała, niemniej jednak takie obciążenie skromnej baterii ma pewne reperkusje. Po pierwsze – z czasem spada jej znamionowa pojemność, napięcie znamionowe oraz wzrasta rezystancja wewnętrzna. Spadek pojemności nie jest jakiś drastycznie wielki, ale można go szacować na 25%, przy spadku napięcia baterii do 2,2 V. Zagadnienie jest naprawdę bardzo ciekawe, w związku z czym zachęcam ambitnych Czytelników do zgłębienia tematu, który został omówiony na stronie [3], gdzie inżynierowie firmy Energizer i Nordic Semiconductor bardzo drobiazgowo wyjaśnili ten interesujący temat w dokumencie o nazwie „High pulse drain impact on CR2032 coin cell battery capacity”. Na szczęście główne podzespoły nadajnika

Wykaz elementów, kupuj na stronie sklep.avt.pl (Warszawa, ul. Leszczyńska 11, tel. +48222578451, e-mail: handlowy@avt.pl)

Karta przekazników

Rezystory: (miniaturowe 1/8 W, raster 0,2")

R1: 22 kΩ

R2...R13: 1,8 kΩ

R14: 100 kΩ

Kondensatory:

C1, C3, C5, C6: 100 nF ceramiczny (raster 0,1")

C2, C4: 100 µF/16 V elektrolityczny (5 mm, raster 2,5 mm)

Półprzewodniki:

U1: LD1117V33 (TO-220)

U2: ATtiny2313A (DIL20)

U3, U4: ULN2803 (DIL18)

U5: RFM12B-866MHz (SMD)

B1: mostek prostowniczy 1 A (okrągły, raster 0,2")

CHN1...CHN12: czerwona dioda LED 3 mm

Pozostałe:

L1: dławik osiowy 10 µH (raster 0,2")

TR1: transformator do druku TE22.5/D/9V

K1...K12: przekaznik AZ921-1A-12DE ZETTLER (JZC-49F-12V)

AC: złącze śrubowe AK500/2

OUT1...OUT3: złącze śrubowe AK500/8

Pilot

Rezystory: (SMD0805)

R1, R2: 220 Ω

R3: 100 kΩ

Kondensatory: (SMD0805)

C1, C2: 100 nF ceramiczny X7R

Półprzewodniki:

U1: ATtiny2313A (SO20)

U2: RFM12B-866MHz (SMD)

ON: czerwona dioda LED (SMD1206)

OFF: zielona dioda LED (SMD1206)

Pozostałe:

BATT: koszyk baterii CR2016 SMT typu BAT-HLD-002-SMT

1...12: microswitch SMD typu KSC411J 70 SH LFS C&K (wysokość 5,2 mm)

OBUDOWA: plastikowa do pilota typu 13120.44 TEK0

pracują już przy niewielkich napięciach zasilających (moduł RF: 2,2 V, mikrokontroler: 1,8 V) i nawet przy 25% spadku pojemności, tego typu aplikacja powinna zapewnić wieloletnią pracę urządzenia.

Aby ocenić jak długo nadajnik pracował będzie na pojedynczej baterii CR2016 należy zastanowić się z jakich etapów składa się jego praca i jakie są wtedy prądy pobierane ze źródła napięcia zasilającego. Przystępując do obliczeń przyjąłem następujący podział cyklu pracy urządzenia:

- etap trybu power-down (uśpienie), który trwa z dużym przybliżeniem 24 h/dobę i podczas którego pobierany jest prąd rzędu 0,2 μ A,
- etap transmisji modułu RFM12B (po wybudzeniu), który trwa średnio 20 ms i podczas którego pobierany jest prąd rzędu 22,6 mA.

Założono ponadto, iż wybudzenie nadajnika, a więc zmiana stanu podłączonej klawiatury następuje 15 razy na dobę. Przy tych założeniach otrzymano teoretyczny, ponad 62-letni czas pracy na pojedynczej baterii CR2016 (o zredukowanej o 25% pojemności) co wydaje się wartością niespotykaną i grubo przekraczającą czas życia samej baterii, standardowo przyjmowany, jako 10 lat.

Uważny Czytelnik dostrzeże pewną subtelność w opisie sposobu działania nadajnika (pilota) systemu wiRelay. Otóż napisałem, że pilot, oprócz nadania rozkazu, odbiera również odpowiedź wysłaną przez kartę przekazników. Tak, to prawda. Oba moduły w istocie przeprowadzają dwustronną komunikację wyłącznie po to, aby po stronie pilota sterującego pozyskać informację na temat wykonania rozkazu sterującego (zmiany stanu sterowanego przekaznika). Do tego rodzaju potwierdzenia służą diody LED zamontowane na obwodzie drukowanym nadajnika (pilota). Oczywiście tego rodzaju rozwiązanie nie daje 100% pewności co do faktu wykonania rozkazu sterującego, gdyż nie trudno wyobrazić sobie sytuację, gdzie rozkaz wysłany przez nadajnik został w istocie wykonany, zaś stosowne potwierdzenie wysłane przez odbiornik nie dotarło do nadajnika przez co zasygnalizował on (zupełnie błędnie) niewykonanie tegoż rozkazu sterującego, lecz moim zdaniem lepsze jest tego rodzaju rozwiązanie, niż brak jakiegokolwiek potwierdzenia po stronie nadajnika. Z resztą mechanizm ten jest niezbędny wyłącznie wtedy, gdy naocznie nie jesteśmy w stanie przekonać się o tym, iż sterowany przez nas przekaznik (a więc podłączony do niego odbiornik) nie zmienił stanu pracy. Sama karta przekazników pozostaje, co oczywiste, w nasłuchu i po otrzymaniu rozkazu sterującego zmienia stan związanego z nim przekaznika (w przypadku braku błędów) po czym odsyła potwierdzenie do odbiornika (pilota). Warto również podkreślić, iż transmisja pakietów danych

między stronami komunikacji oparta jest stosowną sumą kontrolną CRC8 minimalizując ryzyko potencjalnych błędów. Tyle w kwestii funkcjonowania systemu.

Budowa i działanie

Pora na zaprezentowanie schematu ideowego nadajnika, który został pokazany na rysunku 1. Jak widać, zaprojektowano bardzo prosty system mikroprocesorowy, którego sercem jest niewielki mikrokontroler ATtiny2313A (niskonapięciowy) taktowany wewnętrznym oscylatorem RC o częstotliwości 1 MHz sterujący pracą modułu transceivera, dzięki realizacji programowej obsługi interfejsu SPI oraz obsłudze przerwania zewnętrznego Pin Change Interrupt 0 (wyprowadzenie PB1) odpowiedzialnego za mechanizm wysyłania i odbierania danych. Ponadto mikrokontroler obsługuje klawiaturę matrycową utworzoną z przycisków 1...12, której obsługa realizowana jest z użyciem przerwania zewnętrznego Pin Change Interrupt 2 (wyprowadzenia PD0...PD6) inicjującego wybudzenie mikrokontrolera (i modułu RF) oraz obsługuje dwie diody LED (ON i OFF) służące sygnalizacji wykonania wysłanych rozkazów sterujących.

Zgodnie z tym, co napisano wcześniej, nasz nadajnik powinien charakteryzować się minimalnym zapotrzebowaniem na energię elektryczną, jako że jest zasilany niewielką baterią CR2016. W związku z powyższym zastosowano następujące mechanizmy programowo-sprzętowe:

- wyłączono wszystkie nieużywane peryferia mikrokontrolera (komparator analogowy, TIMER1, TIMER0, USI, USART),

Listingi 1. Fragment pliku nagłówkowego związany z obsługą klawiatury matrycowej

```
//Definicje portu klawiatury
#define KEY_PORT PORTD
#define KEY_PIN PIND
#define KEY_DDR DDRD
#define COL1 (1<<PD6)
#define COL2 (1<<PD3)
#define COL3 (1<<PD0)
#define ROW1 (1<<PD1)
#define ROW2 (1<<PD2)
#define ROW3 (1<<PD4)
#define ROW4 (1<<PD5)

//Definicje portu led
#define LED_PORT PORTB
#define LED_DDR DDRB
#define LED_RED PB3
#define LED_GREEN PB4
#define LED_RED_ON
LED_PORT |= (1<<LED_RED)
#define LED_GREEN_ON
LED_PORT |= (1<<LED_GREEN)
#define LED_BOTH_ON
LED_PORT |= ((1<<LED_GREEN)|(1<<LED_RED))
#define LED_BOTH_OFF
LED_PORT &= ~(1<<LED_GREEN)|(1<<LED_RED))

//Definicje rozkazów sterujących
#define RELAY_UNDEFINED 0x00
#define RELAY_ON 0b11110101
#define RELAY_OFF 0b10101111

//Timeout - ms
#define TIMEOUT 100
```

Listing 2. Fragment funkcji main w zakresie konfiguracji obsługi klawiatury matrycowej

```
//Podciągnięcie wierszy klawiatury pod VCC
KEY_PORT |= ROW1|ROW2|ROW3|ROW4;
//Kolumny klawiatury jako wyjścia
//ze stanem 0 - przygotowanie do przerwania
KEY_DDR |= COL1|COL2|COL3;
//Konfiguracja i uruchomienie
//przerwania zewnętrznego obsługi klawiatury
//Zmiana stanu na wierszach klawiatury
//powoduje przerwanie
PCMSK2 = ROW1|ROW2|ROW3|ROW4;
GIMSK |= (1<<PCIE2);
```

Listing 3. Funkcja obsługi przerwania odpowiedzialna za obsługę klawiatury matrycowej

```
//Dane przeznaczone do wysyłki: Rozkaz i numer przekaznika
volatile char dataToSend[2];
//Tablica kolumn
const uint8_t Cols[3] = {COL1, COL2, COL3};
//Tablica wierszy
const uint8_t Rows[4] = {ROW1, ROW2, ROW3, ROW4};

ISR(PCINT2_vect){
    //Wyłączenie przerwania klawiatury
    GIMSK &= ~(1<<PCIE2);
    //Wszystkie kolumny ustawiamy
    //na początek jako wejścia
    KEY_DDR &= ~(COL1|COL2|COL3);

    //Ustalamy numer wciśniętego klawisza
    for(uint8_t Col=0; Col<3; Col++){
        //Wybraną kolumnę ustawiamy,
        //jako wyjściową ze stanem 0
        KEY_DDR |= Cols[Col];
        //Szukamy czy nie zwarto jej
        //do wybranej wiersza
        for(uint8_t Row=0; Row<4; Row++){
            if(!(KEY_PIN & Rows[Row]))
                dataToSend[1] = Col + (Row*3);
        }
        //Wybraną kolumnę ustawiamy,
        //jako wejściową
        KEY_DDR &= ~Cols[Col];
    }

    //Wszystkie kolumny ustawiamy ponownie,
    //jako wyjścia ze stanem 0 - przygotowanie do przerwania
    KEY_DDR |= COL1|COL2|COL3;

    //Opóźnienie by ocenić czy krótkie,
    //czy długie naciśnięcie klawisza
    _delay_ms(300);
    //Sprawdzamy czy klawisz nadal jest naciśnięty
    if(!(KEY_PIN & ROW1) || !(KEY_PIN & ROW2) ||
    !(KEY_PIN & ROW3) || !(KEY_PIN & ROW4))
        dataToSend[0] = RELAY_ON;
    else dataToSend[0] = RELAY_OFF;
}
```


Ustawienia Fuse-bitów obu modułów:

```
CKSEL3...0: 0100
SUT1...0: 10
CKDIV8: 0
CKOUT: 1
```

- wprowadzono mikrokontroler w tryb niskiego poboru mocy Power-down, z którego wybudzany jest wyłącznie poprzez zmianę stanu klawiatury matrycowej. Wybudzony mikrokontroler przeprowadza transmisję danych, po czym przechodzi ponownie w tryb Power-down,
- nieużywany transceiver RFM12B wprowadzany jest każdorazowo w tryb niskiego poboru mocy.

Dzięki tym rozwiązaniom otrzymano wspomniane wcześniej wartości zapotrzebowania na energię dla całego urządzenia oraz teoretyczne czasy jego działania. Na koniec kilka słów komentarza należy się implementacji obsługi klawiatury matrycowej realizowanej w funkcji przerwania zewnętrznego *Pin Change Interrupt 2* (wyprowadzenia PD0...PD6). Zaczniemy zatem od pliku nagłówkowego porządkującego ustawienia sprzętowe w tym zakresie, który to pokazano na **listingu 1**. Dalej, na **listingu 2** pokazano fragment funkcji *main* w zakresie podstawowej konfiguracji obsługi klawiatury matrycowej. I na koniec funkcja obsługi przerwania odpowiedzialna za obsługę klawiatury matrycowej, której ciało pokazano na **listingu 3** (wraz z deklaracją niezbędnych zmiennych globalnych upraszczających ciał funkcji). Funkcja ta ustawia stosowne wartości zmiennej *dataToSend[]*, które transmitowane są do karty przekaźników.

Za samą transmisję danych do karty przekaźników, jak i odbiór i sygnalizację potwierdzenia odpowiedzialny jest fragment funkcji *main* pokazany na **listingu 4**. Z kolei za sygnalizację potwierdzenia wysłanego przez kartę przekaźników odpowiada prosta funkcja pokazana na **listingu 5**.

Tyle w kwestii nadajnika, przejdźmy zatem do schematu ideowego karty przekaźników, który pokazano na **rysunku 2**. Tak jak poprzednio, zaprojektowano bardzo prosty system mikroprocesorowy, którego sercem jest niewielki mikrokontroler ATtiny2313A (niskonapięciowy) taktowany wewnętrznym oscylatorem RC o częstotliwości 1 MHz sterujący pracą modułu transceivera dzięki realizacji programowej obsługi interfejsu SPI oraz obsłudze przerwania zewnętrznego *Pin Change Interrupt 2* (wyprowadzenie PD6) odpowiedzialnego za mechanizm odbierania i wysyłania danych. Ponadto mikrokontroler ten steruje pracą 12 przekaźników dużej mocy a dokonuje tego dzięki zastosowaniu scalonych sterowników pod postacią układów ULN2803 (U3, U4). W odróżnieniu jednak od układu nadajnika, w ramach aplikacji odbiornika zaprojektowano niewielki, kompletny układ

zasilający przeznaczony do podłączenia do sieci 230 VAC. Było to niezbędne z uwagi na fakt, iż moduł odbiornika steruje pracą przekaźników, z których każdy w czasie załączenia pobiera prąd rzędu 10 mA. Pomimo tego sumaryczny pobór mocy urządzenia w przypadku załączenia wszystkich przekaźników nie przekracza 2 W.

Przejdźmy zatem to szczegółów implementacyjnych, jeśli chodzi o program obsługi aplikacji. Prawdę mówiąc nie ma tego zbyt wiele, gdyż wspomniany program jest niezmiernie prosty. Na początek fragment funkcji *main* odpowiedzialny za odbiór rozkazu sterującego i odesłanie odpowiedzi pokazany na **listingu 6**. I na sam koniec funkcja odpowiedzialna za sterowanie przekaźnikami, której ciało pokazano na **listingu 7**.

Listing 4. Fragment funkcji *main* odpowiedzialny za transmisję danych do karty przekaźników, jak i odbiór i sygnalizację potwierdzenia

```
//Wychodzimy z trybu powerDown modułu RFM12B
RFM12bPowerUp();

//Uruchomienie przerwania odpowiedzialnego
// za nadawanie i odbieranie
GIFR |= (1<<PCIF0);
GIMSK |= (1<<PCIE0);

//Wysyłamy ramkę danych do odbiornika - 11ms
RFM12bStartTx((char*) dataToSend, 2, 0xFF);
//Czekamy na zakończenie transmisji
while(RFM12B.Status != PACKET_SENT);
delay_us(100);

//Przygotowujemy się na odbiór potwierdzenia
RFM12bStartRx();

Timeout = 0;
while(1){
    //Odebrano poprawną ramkę danych:
    //Size|Rozkaz|CRC8
    if(RFM12B.Status == NEW_PACKET){
        //Sprawdzamy, czy odebrane potwierdzenie
        //jest takie samo, jak wysłany rozkaz
        if(RFM12B.Buffer[1] == dataToSend[0])
            blinkLED(dataToSend[0]);
        else blinkLED(RELAY_UNDEFINED);
        break;
    }
    //Odebrano błędną ramkę danych: sygnalizujemy błąd
    else if(RFM12B.Status == FAULTY_PACKET){
        blinkLED(RELAY_UNDEFINED);
        break;
    }
    //Sprawdzamy czy upłynął maksymalny czas
    //na odbiór potwierdzenia
    if(++Timeout > TIMEOUT){
        blinkLED(RELAY_UNDEFINED);
        break;
    }
    delay_ms(1);
}

//Wyłączenie przerwania odpowiedzialnego
//za nadawanie i odbieranie
GIMSK &= ~(1<<PCIE0);

//Wprowadzamy moduł RFM12B
//w tryb powerDown
RFM12bPowerDown();
```

Listing 5. Funkcja odpowiedzialna za sygnalizację potwierdzenia wysłanego przez kartę przekaźników

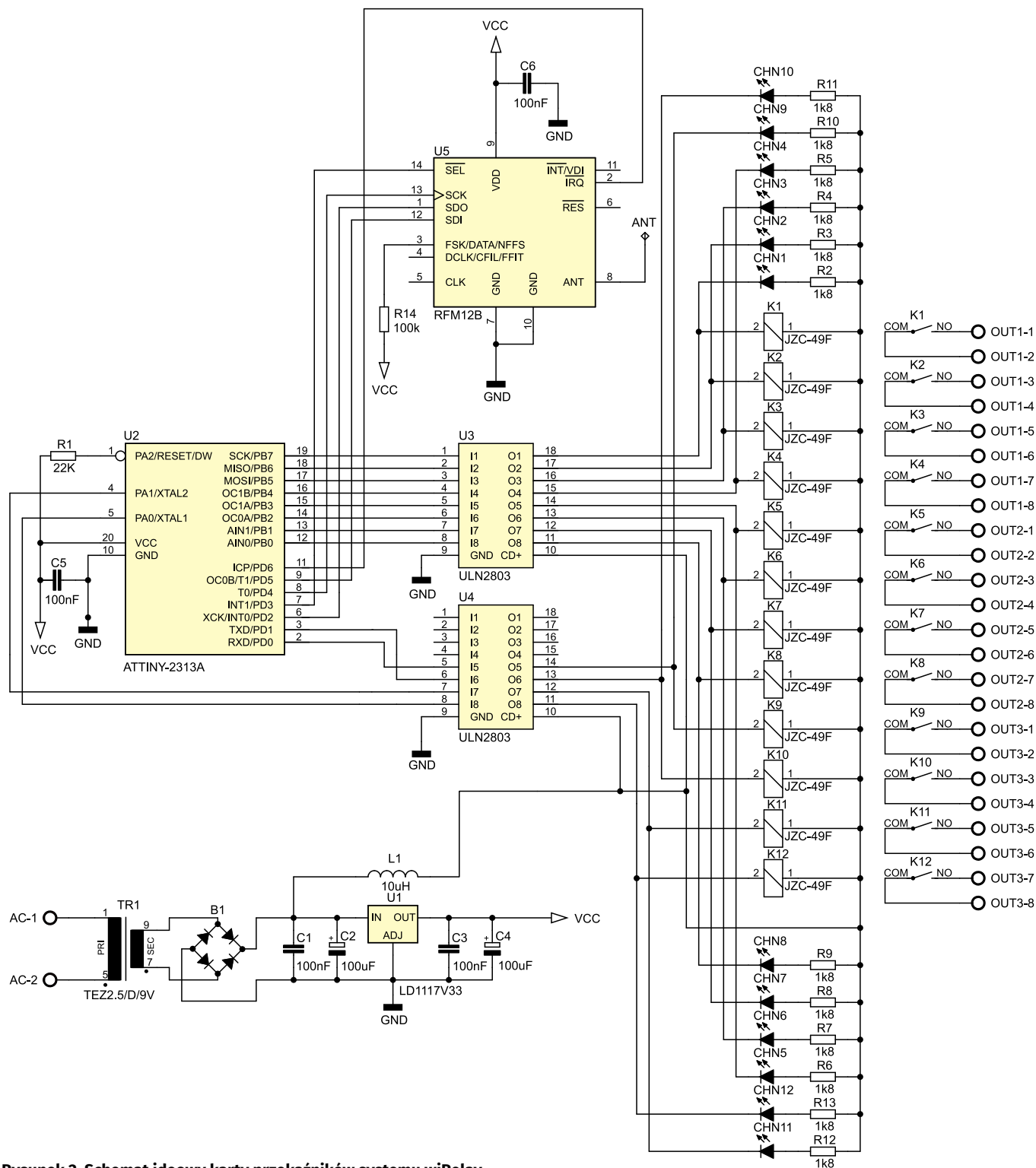
```
void blinkLED(uint8_t Led){
    switch(Led){
        case RELAY_ON: LED_RED_ON; break;
        case RELAY_OFF: LED_GREEN_ON; break;
        case RELAY_UNDEFINED: LED_BOTH_ON; break;
    }
    delay_ms(200);
    LED_BOTH_OFF;
}
```

Listing 6. Fragment funkcji *main* odpowiedzialny za odbiór rozkazu sterującego i odesłanie odpowiedzi

```
//Uruchomienie i konfiguracja RFM12B,
//w tym interfejsu SPI oraz przerwania nadawania/odbioru
RFM12bInit(0xFF);
//Przygotowujemy moduł na odbiór danych
RFM12bStartRx();

sei();
while(1){
    //Odebrano poprawną ramkę danych:
    //Size|Rozkaz|Przekaźnik|CRC8
    if(RFM12B.Status == NEW_PACKET){
        //Sterujemy przekaźnikiem
        ctrlRelay(RFM12B.Buffer[1], RFM12B.Buffer[2]);
        //Odsyłamy potwierdzenie
        //w postaci odebranego rozkazu
        Ack = RFM12B.Buffer[1];

        delay_ms(1);
        //Wysyłamy potwierdzenie
        //do odbiornika - 10ms
        RFM12bStartTx(&Ack, 1, 0xFF);
        //Czekamy na zakończenie transmisji
        while(RFM12B.Status != PACKET_SENT);
        delay_ms(1);
        //Przygotowujemy moduł
        //na kolejny odbiór danych
        RFM12bStartRx();
    }
    //Odebrano błędną ramkę danych:
    //restartujemy procedurę odbiorczą,
    //bo inaczej nastąpiłaby blokada odbioru
    else if(RFM12B.Status == FAULTY_PACKET)
        RFM12bStartRx();
}
```



Rysunek 2. Schemat ideowy karty przekaźników systemu wiRelay

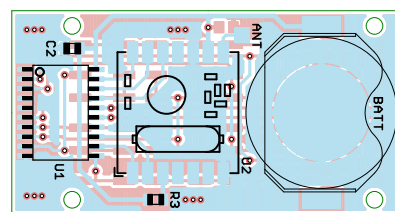
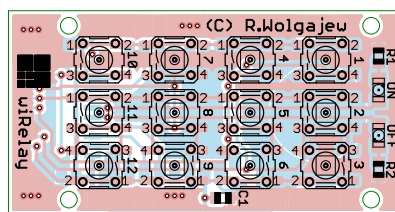
Tyle w kwestiach implementacyjnych dotyczących obu układów.

Montaż i uruchomienie

Przejdźmy zatem do schematu montażowego nadajnika, który pokazano na **rysunku 3**. Jak widać zaprojektowano bardzo niewielki obwód drukowany z wyłącznym montażem elementów SMD po obu stronach laminatu. Projektując nadajnik systemu wiRelay, jak i jego obwód drukowany chciałem, aby docelowe urządzenie wyposażone było w gęstą i niewielką obudowę,

przez co etapem wyjściowym w procesie projektowania było znalezienie atrakcyjnej wizualnie i łatwo dostępnej obudowy. Zdecydowałem się na zastosowanie smukłej,

plastikowej obudowy do pilota typu 13120.44 firmy TEK0 w wersji bez żadnych przycisków sterujących, gdyż zakładałem zastosowanie dedykowanych switchy i nawiercenie



Rysunek 3. Schemat płytki PCB nadajnika systemu wiRelay

Uwaga! Na płycie modułu karty przekaźników zamontowano kompletny zasilacz łącznie z transformatorem zasilanym napięciem sieciowym 230 V AC oraz zamontowano elementy będące na potencjale tego napięcia. Istnieje niebezpieczeństwo porażenia prądem elektrycznym, co może stanowić zagrożenie dla życia i zdrowia użytkowników. W związku z tym, montaż układu oraz jego uruchomienie należy wykonywać pod nadzorem osoby z odpowiednimi uprawnieniami lub mającej niezbędną wiedzę i doświadczenie. Natomiast gotowe urządzenie musi być właściwie zabezpieczone – np. umieszczone w nieprzewodzącej obudowie.

stosownych otworów (na osie switchy i diody LED) w płycie czołowej (górnej) obudowy. W związku z powyższym cały projekt laminatu podporządkowany został wymiarom zastosowanej obudowy. Co więcej, z uwagi na fakt, iż zastosowany typ obudowy umożliwia zamontowanie w nim płytki z elementami o maksymalnej, sumarycznej grubości ok. 7,2 mm musiałem zdecydować się na zastosowanie bardzo niskich switchy SMD oraz wyjątkowo niskiego koszyeczka baterii zasilającej. Poniekąd z tego tytułu wynikła konieczność zastosowania cienkiej baterii CR2016, zamiast zdecydowanie bardziej popularnej CR2032 (o ponad 2 krotnie większej pojemności). Osoby, które nie planują zastosowania obudowy, o której mowa powyżej, mogą nie zważać na poniższe ograniczenia.

Przejdźmy zatem do szczegółów montażowych dotyczących nadajnika. Montaż obwodu drukowanego nadajnika rozpoczynamy od warstwy BOTTOM, na której przylutujemy mikrokontroler, następnie lutujemy moduł RFM12B, dalej elementy bierne a na końcu koszyczek baterii zasilającej. Następnie przechodzimy na warstwę TOP, gdzie w pierwszej kolejności przylutujemy elementy bierne i diody LED a na końcu switchy SMD. Jeśli chodzi o typ zastosowanych diod LED i kwestie obudowy to zakładam dwie możliwości rozwiązania tego problemu. Pierwsza to zastosowanie diod LED SMD (tak, jak przewidziano na płycie nadajnika) i umieszczenie w wywierconych otworach w obudowie samych główek zwykłych

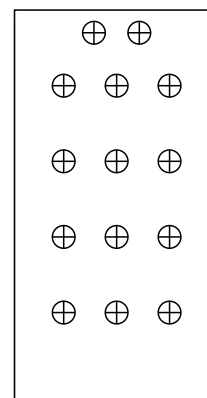
diod LED o średnicy 3 mm podświetlanych niejako przez wlotowane pod spodem diody SMD, zaś drugie rozwiązanie to wlotowanie w miejsce diod SMD typowych diod LED o średnicy 3 mm z odpowiednio przyciętymi i wygiętymi końcówkami. Na koniec,

do tak przygotowanej płytki przylutujemy antenę nadawczą (punkt lutowniczy po stronie BOTTOM) w postaci kawałka przewodu o długości około 17 cm (może być odpowiednio zwinięty). Wygląd zmontowanego nadajnika zarówno od strony warstwy TOP jak i strony warstwy BOTTOM pokazano na **fotografii 1**.

Na koniec kilka słów na temat przygotowania wybranej obudowy. Jak można się domyślić patrząc na obwód drukowany nadajnika, w górnej części jego obudowy (tzn. tej bez otworu na śrubę montażową) należy wywiercić 14 otworów o średnicy 3 mm: 12 otworów dla osi zastosowanych switchy oraz dwa otwory na diody LED. Aby usprawnić proces trasowania otworów przygotowano specjalny szablon, który pokazano na **rysunku 4** (w skali 1:1). Zanim jednak skorzystamy z szablonu, o którym mowa powyżej z wnętrza górnej części obudowy

Listing 7. Funkcja odpowiedzialna za sterowanie przekaźnikami

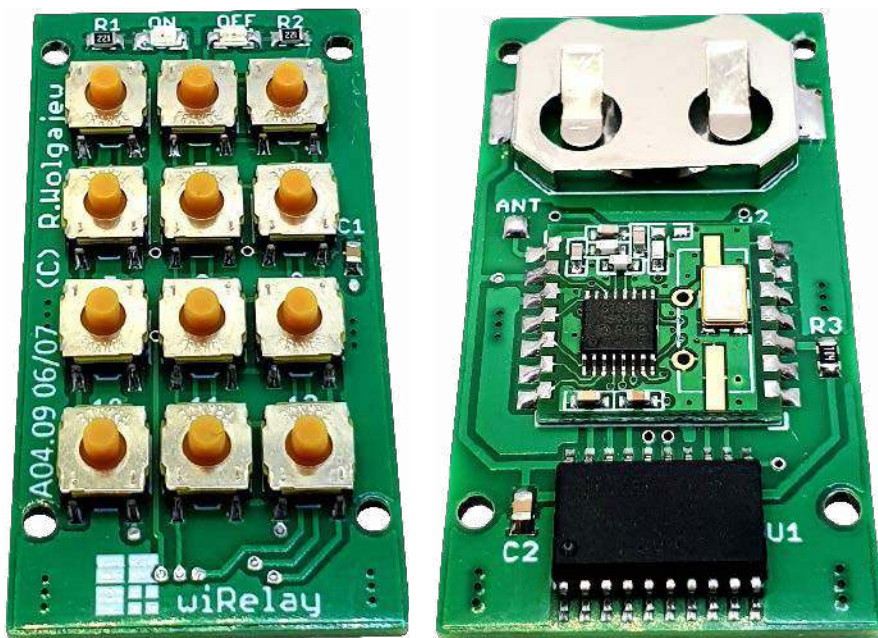
```
void ctrlRelay(uint8_t Command, uint8_t Relay){
  //0...7
  if(Relay < 8){
    if(Command == RELAY_ON)
      PORTB |= (1<<(7-Relay));
    else PORTB &= ~(1<<(7-Relay));
  }
  //8, 9
  }else if(Relay < 10){
    if(Command == RELAY_ON)
      PORTD |= (1<<(Relay-8));
    else PORTD &= ~(1<<(Relay-8));
  }
  //10, 11
  }else{
    if(Command == RELAY_ON)
      PORTA |= (1<<(11-Relay));
    else PORTA &= ~(1<<(11-Relay));
  }
}
```



Rysunek 4. Szablon ułatwiający wykonanie otworów w obudowie nadajnika

(tzn. tej bez otworu montażowego pod śrubę) należy usunąć (przy pomocy ostrego nożyka) wszystkie nadlewki pokazane na **fotografii 2** po to, aby ułatwić umieszczenie szablonu wewnątrz obudowy w celu dokładnego trasowania otworów, oraz aby zwolnić dodatkowe miejsce na obwód drukowany wraz z elementami. Po wykonaniu wspomnianych czynności szablon umieszczamy od środka górnej części obudowy, następnie trasujemy przy jego pomocy miejsca otworów, a na końcu wiercimy je wiertłem do metalu o średnicy 3 mm. Z kolei we wnętrzu dolnej części obudowy (tzn. tej z otworem montażowym pod śrubę) musimy również dokonać pewnych modyfikacji, które mają na celu zwiększenie miejsca na obwód drukowany wraz z elementami. W tym przypadku, podobnie jak poprzednio, usuwamy nadlewki pokazane na **fotografii 3**, które to stanowią podporę zintegrowanych plastikowych tulejek dystansowych uniemożliwiając głębsze osadzenie obwodu drukowanego.

Po takich modyfikacjach obudowa nasza jest gotowa do umieszczenia w niej zmontowanego obwodu drukowanego nadajnika systemu wiRelay. Co prawda po takich



Fotografia 1. Wygląd zmontowanego nadajnika systemu wiRelay od strony TOP i BOTTOM

zabiegach osie zastosowanych mikroprzełączników będą wystawać ponad płaszczyzną obudowy w sposób minimalny, jednak jest to jak najbardziej pożądane, gdyż ograniczy to możliwość przypadkowego wciśnięcia przycisku sterującego. Można jednak nieco ułatwić ich obsługę poprzez nawiercenie stosownych otworów montażowych nieznacznie większym wiertłem, na przykład o średnicy 4 mm.

Przechodzimy do schematu montażowego karty przełączników, który pokazano na **rysunku 5**. Tym razem zaprojektowano dużo większy obwód drukowany, który integruje w sobie kompletny, transformatorowy zasilacz sieciowy i zbudowany jest w zdecydowanej większości z elementów przewlekanych. Montaż obwodu drukowanego karty przełączników rozpoczynamy od przyłutowania modułu RFM12B, następnie lutujemy pozostałe elementy półprzewodnikowe, dalej elementy bierne, zaś na samym końcu elementy mechaniczne oraz transformator do druku. Diody LED możemy umieścić na stosownej długości tulejach dystansowych by ich główki wystawały ponad płaszczyzną obudowy przełączników. Do tak przygotowanej płytki przyłutowujemy antenę odbiorczą w postaci kawałka przewodu o długości około 17 cm (może być odpowiednio zwinięty). Wygląd zmontowanej karty przełączników od strony warstwy TOP pokazano na **fotografii 4**.

Obsługa urządzenia

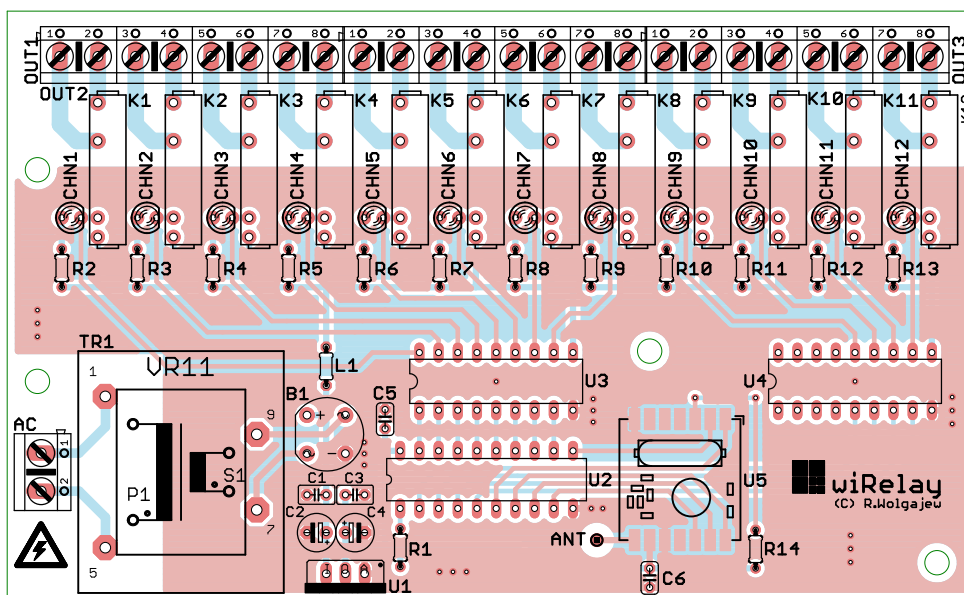
Sposób obsługi systemu wiRelay tak, jak można by tego oczekiwać, jest niezmiernie prosty. Włączeniu pilota systemu wiRelay towarzyszy krótkie (1 s) załączenie obu diod LED, które ma na celu sprawdzenie tych elementów sygnalizacyjnych. Dalej, krótkie naciśnięcie wybranego switcha umieszczonego na nadajniku powoduje wyłączenie powiązanego z nim przełącznika (na karcie przełączników), co powinno zostać potwierdzone poprzez chwilowe zapalenie się diody OFF (zielonej). Długie naciśnięcie wybranego switcha umieszczonego na nadajniku powoduje z kolei włączenie powiązanego z nim przełącznika (na karcie przełączników), co powinno zostać potwierdzone poprzez chwilowe zapalenie się diody ON (czerwonej). Nieodebranie potwierdzenia ze strony karty przełączników (w maksymalnym czasie 100 ms) powoduje chwilowe zapalenie się obu diod LED. W takim wypadku nie mamy pewności, co do stanu sterowanego przełącznika, chyba że możemy wzrokowo ocenić stan karty przełączników korzystając z zamontowanych tam diod LED, których świecenie jest równoznaczne z załączeniem



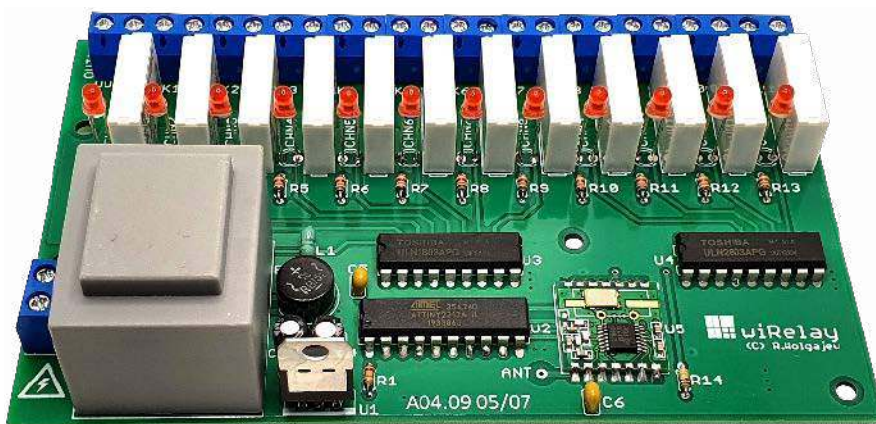
Fotografia 2. Wygląd wnętrza górnej części obudowy z zaznaczeniem nadlepek przeznaczonych do usunięcia



Fotografia 3. Wygląd wnętrza dolnej części obudowy z zaznaczeniem nadlepek przeznaczonych do usunięcia



Rysunek 5. Schemat płytki PCB karty przełączników systemu wiRelay



Fotografia 4. Wygląd zmontowanej karty przełączników systemu wiRelay od strony warstwy TOP

odpowiedniego przełącznika lub zwyczajnie mamy możliwość wzrokowej oceny faktu załączenia do pracy sterowanego urządzenia.

Robert Wołgajew, EP

Odnośniki:

- [1] <https://bit.ly/3C8HLBI>
- [2] <https://bit.ly/3jAKj4H>
- [3] <https://bit.ly/3Ol42zY>

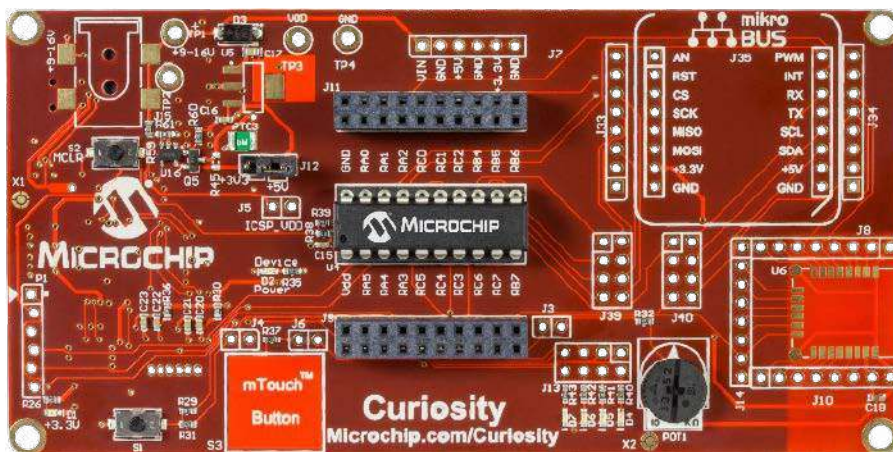
Wygraj uniwersalną płytkę ewaluacyjną Microchip Curiosity Development Board

Płytkę rozwojową Curiosity od firmy Microchip (DM164137) to ekonomiczna, w pełni zintegrowana platforma uruchomieniowa przeznaczona dla początkujących konstruktorów i programistów. Została zaprojektowana do współpracy ze zintegrowanymi środowiskami programistycznymi MPLAB X i MPLAB Code Configurator firmy Microchip. Zawiera zestaw niezbędnych komponentów oraz wbudowany programator/debugger, dzięki czemu nie wymaga dodatkowego sprzętu do rozpoczęcia prototypowania aplikacji z mikrokontrolerami.

W centralnym obszarze płytki rozwojowej znajduje się złącze, które umożliwia zamontowanie 8-, 14- i 20-wyprowadzeniowych nowoczesnych mikrokontrolerów z rodziny PIC. Charakteryzują się one wydajnością, wysoką energooszczędnością i bogatym zestawem urządzeń peryferyjnych niezależnych od rdzenia – CIP (Core Independent Peripherals). Umożliwiają one integrację różnych funkcji aplikacji w jednym mikrokontrolerze, upraszczając projekt i utrzymując niskie zużycie energii, a jednocześnie gwarantując łatwą, programową rekonfigurację.

Płyta zawiera podstawowe elementy interfejsu użytkownika – przełączniki, czujniki pojemnościowe mTouch, potencjometr i diody LED. Wszystkie wyprowadzenia mikrokontrolera dostępne są na złączach o rastrze 2,54 mm, kompatybilnych z popularnymi złączami szpilkowymi gold-pin i przewodami do płytek stykowych. Dodatkowo, ogromny zestaw akcesoriów jest dostępny za pośrednictwem umieszczonego na płytce miejsca dla modułu w standardzie MikroElektronika mikroBUS. W podobny sposób można rozszerzyć płytkę o komunikację Bluetooth Low Energy za pomocą modułu Microchip RN4020. Odpowiednie miejsce na płytce umożliwia wlotowanie tego modułu.

Płytkę jest wykonywana w dwóch wersjach:

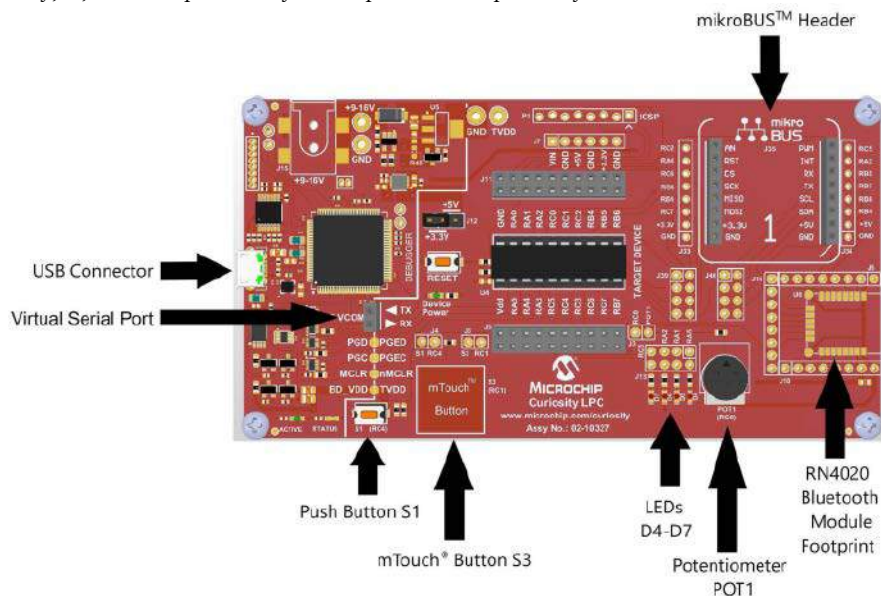


- Rev 2 – w której blok programatora/debuggera został umieszczony na warstwie BOTTOM płytki, tak jak pokazano na fotografii tytułowej,
 - Rev 4 – blok programatora/debuggera został umieszczony na warstwie TOP płytki, tak jak pokazano na **rysunku 1**.
- Obie wersje oferują taką samą funkcjonalność.
- Aby mieć szansę na wygranie płytki ewaluacyjnej Microchip Curiosity Development

Board (DM164137), lub aby otrzymać kupon rabatowy 15% i bezpłatną wysyłkę, należy wypełnić formularz zgłoszeniowy na stronie: <https://bit.ly/3WLoCGs>

Szczegółowe informacje na temat płytki rozwojowej Microchip Curiosity Development Board można znaleźć na: <https://bit.ly/3YRMQrx>

Dokumentacja uruchomieniowa wraz opisem przykładowego projektu jest dostępna na: <https://bit.ly/3PUWzt9>



Rysunek 1. Rozmieszczenie podstawowych komponentów na płytce (wersja płytki – Rev 4)

**Podstawowe parametry:**

- płynnie regulowany zasilacz napięcia stałego w zakresie 0...28 V,
- płynnie regulowana wydajność prądowa 0...3 A,
- podgląd ustawionej granicy zadziałania ograniczenia prądowego na wyświetlaczu LCD,
- wyświetlanie aktualnej wartości napięcia wyjściowego i pobieranego z wyjścia prądu na wyświetlaczu LCD,
- wskazywanie trybu pracy (stałe napięcie/stały prąd) na wyświetlaczu LCD oraz diodzie LED,
- automatyczne załączanie wentylatora schładzającego radiator,
- odłączanie napięcia wyjściowego po wykryciu przegrzania tranzystorów mocy,
- zasilanie napięciem 230 V.

* **Uwaga!** Elektroniczne zestawy do samodzielnego montażu. Wymagana umiejętność lutowania! Podstawową wersją zestawu jest wersja [B] nazywana potocznie KIT-em (z ang. zestaw). Zestaw w wersji [B] zawiera elementy elektroniczne (w tym [UK] – jeśli występuje w projekcie), które należy samodzielnie wlutować w dołączoną płytkę drukowaną (PCB). Wykaz elementów znajduje się w dokumentacji, która jest podlinkowana w opisie kitu. Mając na uwadze różne potrzeby naszych klientów, oferujemy dodatkowe wersje:

Dodatkowe materiały do pobrania ze strony www.ulubionykiosk.pl/media

- Modułowy zasilacz warsztatowy (EP 5/2022)
- Regulowany zasilacz warsztatowy – RPS-02 (EP 4/2022)
- AVT5915 Zasilacz 5 V/1 A z szerokim zakresem napięć wejściowych (EP 1/2022)
- AVT5908 Beztransformatorowy impulsowy zasilacz sieciowy (EP 12/2021)
- AVT5872 Regulowany zamiennik stabilizatora 78xx (EP 7/2021)
- AVT1990 Regulowany zasilacz do płytek stykowych (EP 8/2018)
- Precyzyjny regulowany zasilacz stabilizowany (EP 2/2018)
- AVT5585 Zasilacz laboratoryjny 0...30 V/5 A ze sterowaniem mikroprocesorowym (EP 12/2017-1/2018)
- Multizasilacz (EP 10/2017)
- AVT1976 Precyzyjny, regulowany zasilacz uniwersalny 1,5...32 V/3 A (EP 8/2017)
- AVT3172 Praktyczny zasilacz warsztatowy (EP 5/2017)
- AVT1946 Zasilacz napięcia symetrycznego z LM27762 (EP 2/2017)
- AVT1895 Uniwersalny moduł zasilający (EP 10/2016)

- **wersja [C]** – zmontowany, uruchomiony i przetestowany zestaw [B] (elementy wlutowane w płytkę PCB),
 - **wersja [A]** – płytkę drukowaną bez elementów i dokumentacji.
- Kity, w których występuje układ scalony wymagający zaprogramowania, mają następujące dodatkowe wersje:
- **wersja [A-1]** – płytkę drukowaną [A] + zaprogramowany układ
 - [UK] i dokumentacja,
 - **wersja [UK]** – zaprogramowany układ.

Nie każdy zestaw AVT występuje we wszystkich wersjach! Każda wersja ma załączony ten sam plik PDF! Podczas składania zamówienia upewnij się, którą wersję zamawiasz! <http://sklep.avt.pl>

W przypadku braku dostępności na stronie sklepu osoby zainteresowane zakupem płytek drukowanych (PCB) prosimy o kontakt via e-mail: kity@avt.pl.

W ofercie AVT*

AVT5963

Zasilacz warsztatowy (2)

Nikogo nie trzeba przekonywać, że zasilacz z regulacją napięcia, z ograniczeniem prądu wyjściowego oraz kompletem zabezpieczeń jest fundamentalnym elementem wyposażenia pracowni elektronika. W poprzednim wydaniu „Elektroniki Praktycznej” (EP 12/22) zaprezentowaliśmy pierwszą część opisu projektu zasilacza warsztatowego, w którym zawarty był schemat ideowy konstrukcji, wraz z wyjaśnieniem działania poszczególnych bloków. Druga część artykułu zawiera dokładny opis montażu i uruchomienia oraz wskazówki dotyczące obsługi tego ciekawego urządzenia.

Montaż

Układ został zmontowany na dwóch dwustronnych płytkach drukowanych. Przednia ma wymiary 69×49 mm, natomiast dolna 100×110 mm. Ich schematy zostały pokazane odpowiednio na **rysunku 3** i **rysunku 4**. W odległości 3 mm od krawędzi obu płytek znajdują się po cztery otwory montażowe, każdy o średnicy 3,2 mm, zaś na środku dolnej płytki jest jeszcze jeden, piąty otwór.

Montaż obu płytek proponuję rozpocząć od elementów o najmniejszej wysokości obudowy, czyli rezystorów i diod. Na płycie przedniej należy wstrzymać się z lutowaniem potencjometrów P1 i P2 oraz diod LED, ponieważ są to elementy związane z obudową. Pod mikrokontroler proponuję zastosować podstawkę. Na samym końcu można przykręcić wyświetlacz LCD1, używając do tego tulei dystansowych M3 o długości



Pierwsza część znajduje się pod adresem: <https://ulubionykiosk.pl/media>

12 mm, a następnie wlutować złącza do wyświetlacza. Zmontowaną według tego planu płytkę można zobaczyć na **fotografii 2** i **fotografii 3**. Z kolei płytkę dolną można uzbrajać w elementy według własnej wygody, ale bez tranzystorów T8 i T9 oraz termistora NTC1. Te należy przykręcić do radiatora, a dopiero potem przylutować do laminatu, o czym dalej. Na tym etapie płytka dolna powinna wyglądać tak, jak na **fotografii 4**.

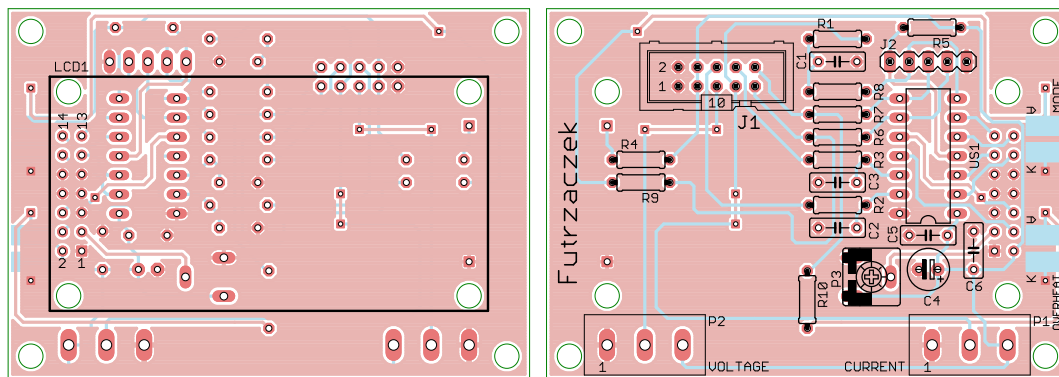
Następnym etapem powinno być zamontowanie płytek oraz pozostałych podzespołów w obudowie. Układ prototypowy zamknięto w standardową obudowę Z17W,

tworząc zwartą konstrukcję. Na tylnej ścianie, w połowie jej wysokości, znalazł się wentylator oraz złącze zasilania zintegrowane z bezpiecznikiem. Wentylator został umieszczony 60 mm od lewej krawędzi obudowy, zaś złącze zasilające 60 mm od prawej, licząc od geometrycznego środka tych elementów do najdalej wysuniętej krawędzi panelu tylnego. Objaśnia to **rysunek 5**. Szczegóły obrazuje **fotografia 5**, gdzie można zobaczyć te dwa podzespoły osadzone w obudowie. Na wentylatorze przykręcono osłonę.

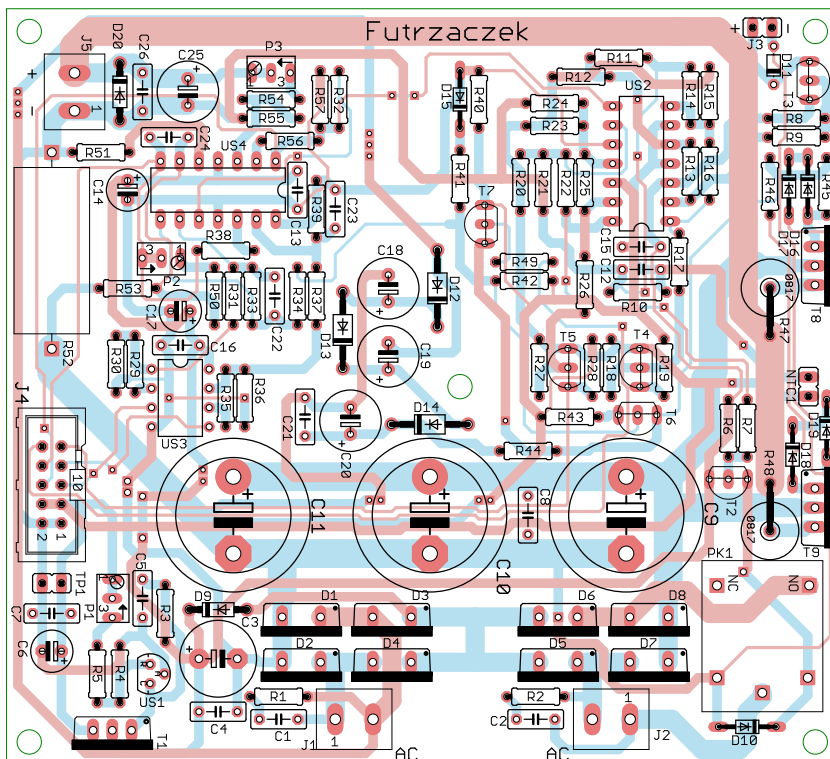
Na płycie przedniej znajduje się nieco więcej podzespołów. W odległości 25 mm od prawej

krawędzi, w połowie wysokości przedniego panelu obudowy, jest otwór o średnicy 12 mm na wyłącznik sieciowy. Płytkę przednią jest umieszczona tak, że dolna krawędź wyświetlacza wypada w połowie wysokości tegoż panelu. Wokół tego otworu są zlokalizowane cztery, o średnicy 3 mm każdy, którymi płytka przednia jest przykręcona do panelu czołowego obudowy za pośrednictwem tulei dystansowych M3 o długości 15 mm. Pod wyświetlaczem, 20 mm poniżej środka wysokości płyty czołowej, są dwa otwory na potencjometry, każdy o średnicy 8 mm. Z kolei zaciski wyjściowe zostały umieszczone 30 mm i 70 mm od lewej krawędzi płyty czołowej, 20 mm nad środkiem wysokości tej płyty. Średnica tych otworów powinna być dostosowana do użytych zacisków. Szczegóły można znaleźć na rysunku 5, z kolei zmontowany panel przedni można zobaczyć na fotografii tytułowej.

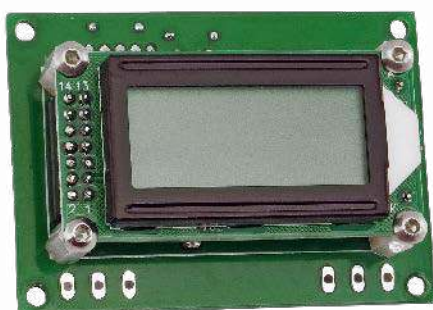
Ostatnim etapem jest montaż podzespołów na spodzie obudowy. W radiatorze trzeba wywiercić trzy otwory, każdy o średnicy 3 mm i każdy w odległości 15 mm od krawędzi radiatora. Jeden powinien znaleźć się w połowie jego szerokości, zaś dwa pozostałe po 16 mm od niego – jeden w lewo, drugi w prawo. Środkowy służy do przykręcenia termistora (śrubą M3 z podkładką) a boczne do tranzystorów. Pod wszystkie te elementy należy zastosować pastę termoprzewodzącą. Podkładki izolacyjne nie są konieczne.



Rysunek 3. Schemat przedniej płytki PCB



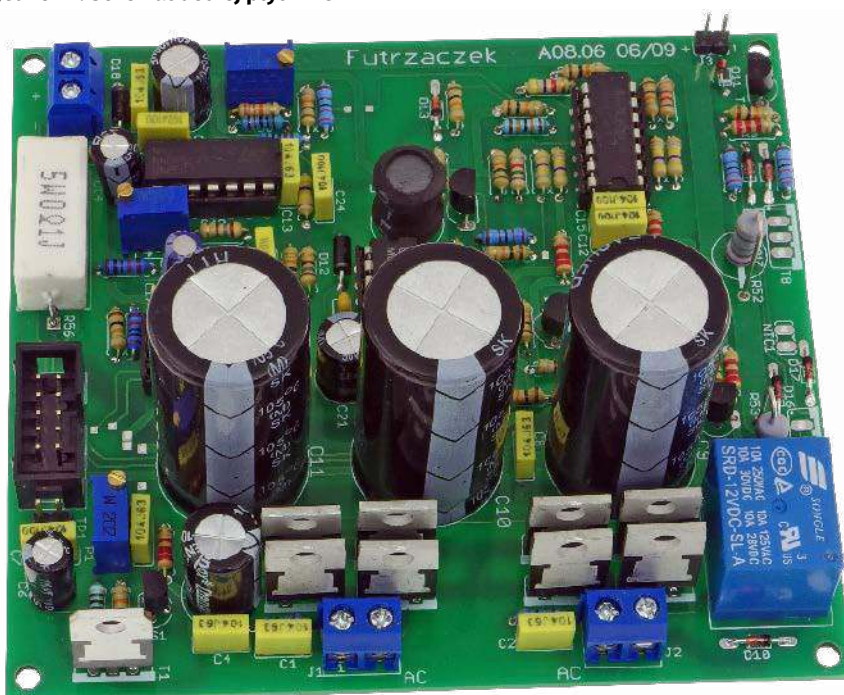
Rysunek 4. Schemat dolnej płytki PCB



Fotografia 2. Zmontowana płytka przednia – widok z przodu



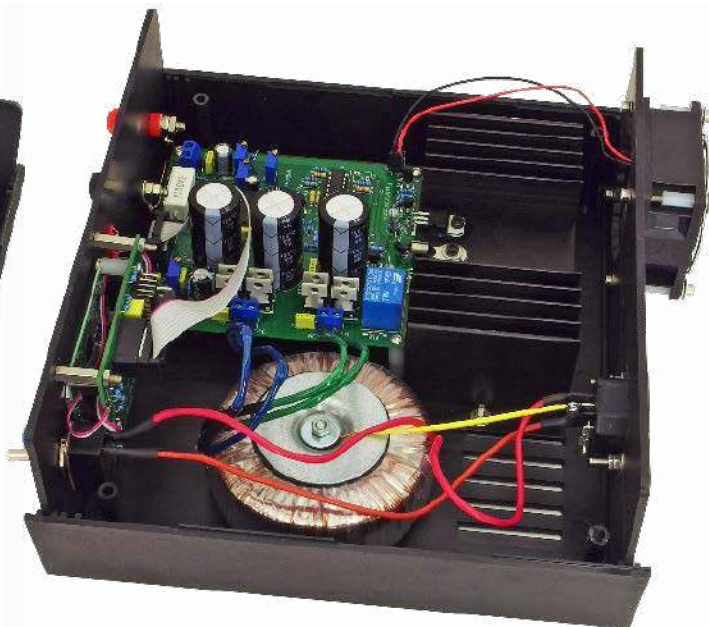
Fotografia 3. Zmontowana płytka przednia – widok od tyłu



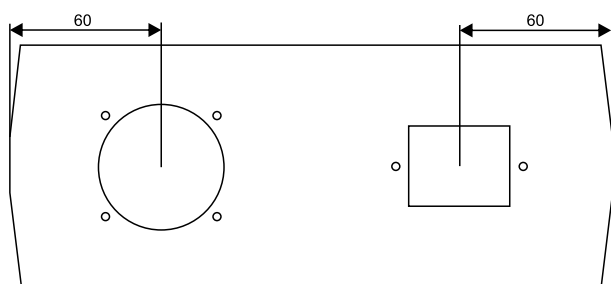
Fotografia 4. Zmontowana płytka dolna



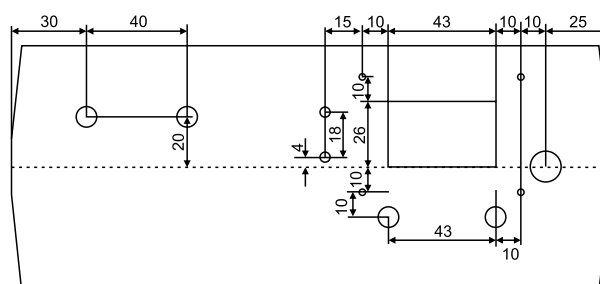
Fotografia 5. Wygląd wnętrza zasilacza pokazujący podzespoły osadzone na płycie tylnej i płycie przedniej oraz sposób zamontowania elementów do radiatora



Fotografia 6. Wygląd wnętrza zasilacza pokazujący podzespoły osadzone na płycie tylnej i płycie przedniej oraz sposób montażu podzespołów na spodzie obudowy



Rysunek 5. Lokalizacja podzespołów na tylnej płycie obudowy



Rysunek 6. Lokalizacja podzespołów na przedniej płycie obudowy

W łapki radiatora, które służą jego przykręceniu do spodu obudowy, należy zrobić po jednym otworze o średnicy 3 mm dla śrub M3. Najlepiej, aby wypadły one w środku każdej z łapek. Tak przygotowany radiator wraz z osadzonymi na nim elementami trzeba przykręcić do spodniej części obudowy (z której polecam usunąć plastikowe kołeczki dystansowe), dosunięty maksymalnie do lewej ścianki i odsunięty od tylnej krawędzi o 25 mm.

W płytkę dolną należy wkręcić tuleje dystansowe o długości 25 mm. Tę należy dosunąć możliwie blisko radiatora i ustawić tak, aby otwory montażowe pod tranzystory pokrywały się z wystającymi z radiatora nóżkami tranzystorów. Po przykręceniu płytki dolnej do spodu obudowy (czterema śrubami M3) można zagiąć wyprowadzenia tranzystorów i termistora, po czym włożyć je w przeznaczone dla nich otwory na płycie dolnej. Na końcu trzeba je przylutować.

Ostatnią rzeczą, która wymaga trwałego umieszczenia w obudowie, jest transformator. Powinien być zlokalizowany po prawej stronie obudowy, blisko prawej ścianki. Należy ustawić go tak, aby nie kolidował z radiatorem tudzież z jakimkolwiek

innym podzespołem. Do jego przykręcenia potrzebna będzie śruba M5 lub o większej średnicy, zbyt cienka może nie wytrzymać tak dużej masy tego elementu.

Szczegóły tego etapu montażu można zobaczyć na **fotografii 6**. Metalowy talerzyk mocujący transformator należy połączyć z zaciskiem przewodu ochronnego złącza IEC. Pod łby śrub i nakrętek, które dotyczą plastiku obudowy, proponuję stosować podkładki by rozłożyć równomiernie wywierany nacisk i nie uszkadzać tworzywa przy wkręcaniu.

Ostatnim etapem jest podłączenie. Uzwojenie pierwotne transformatora należy podłączyć do gniazda IEC za pośrednictwem wyłącznika. Uzwojenia wtórne (najlepiej w tej samej kolejności, choć nie jest to wymóg krytyczny) do zacisków AC na płycie dolnej. Diody LED sygnalizujące przegrzanie (dolna) i przejście w tryb źródła prądowego (górna) należy przylutować do pól lutowniczych na płycie przedniej. Wentylator trzeba podłączyć do podwójnego złącza szpilkowego J3. Zostały zaciski wyjściowe – te trzeba połączyć grubym przewodem z zaciskami napięcia wyjściowego płytki dolnej. Taśmy IDC10 na razie nie podłączamy.

Uruchomienie

Po włożeniu bezpiecznika w przeznaczone dla niego miejsce w złączu IEC oraz włączeniu zasilania, należy podłączyć woltomierz do wyprowadzeń TP1. Potencjometrem P1 trzeba doprowadzić do tego, aby na tych wyprowadzeniach pojawiło się napięcie o wartości równej 5 V. Im dokładniej ten etap wykonamy, tym dokładniejsze będą późniejsze wskazania wyświetlacza, warto więc użyć miernika o możliwie wysokiej dokładności. Jeżeli to się udało, trzeba wyłączyć zasilanie, podłączyć taśmę IDC10 pomiędzy dolną i przednią płytką po czym ponownie włączyć zasilanie.

Na etapie uruchamiania jest konieczne zaprogramowanie pamięci Flash mikrokontrolera dostarczoną wsadą oraz



Rysunek 7. Szczegóły ustawienia bitów zabezpieczających



Fotografia 7. Wygląd płyty przedniej w czasie normalnej pracy



Fotografia 8. Wygląd płyty przedniej po zadziałaniu układu ograniczenia prądowego

zmiana jego bitów zabezpieczających na takie nowe wartości: Low Fuse = 0xE2, High Fuse = 0xDC. Szczegóły zostały pokazane na **rysunku 7**, który zawiera widok okna konfiguracji tych bitów z programu BitBurner. W ten sposób zostanie wyłączony prescaler sygnału zegarowego oraz włączy się Brown-Out Detector, który wprowadzi mikrokontroler w stan zerowania, jeżeli jego napięcie zasilające spadnie poniżej 4,3 V. To znacznie zmniejsza ryzyko zawieszenia się układu podczas uruchamiania.

Po ustawieniu kontrastu potencjometrem P3 na płycie przedniej, naszym oczom powinien ukazać się wyświetlacz jak na **fotografii 7**. Oto, co zawiera:

- w lewym górnym rogu jest aktualne napięcie wyjściowe,
- w prawym górnym rogu mamy aktualny tryb pracy zasilacza (CV – stałe napięcie, CC – stały prąd),
- lewy dolny róg jest zarezerwowany na prąd aktualnie pobierany z zacisków wyjściowych,
- w prawym dolnym rogu mamy przybliżony próg zadziałania ograniczenia prądowego.

Jeżeli nasz zasilacz prawidłowo reaguje na obracanie osi potencjometrów na płycie czołowej, pozostało skalibrować dwa ostatnie punkty.

Do zacisków wyjściowych należy podłączyć woltomierz i używając potencjometru P3 na płycie dolnej, doprowadzić do pojawienia się napięcia 28 V przy pokrętle regulacji napięcia ustawionym na maksimum. Potem ten sam przyrząd należy przełączyć w tryb amperomierza (uwaga, wystąpi zwarcie!), ustawić pokrętkę ogranicznika prądowego na maksimum i tak ustawić potencjometr P2, by popłynął prąd o natężeniu 3 A. Tak, jak poprzednio, im dokładniej ta regulacja zostanie wykonana, tym lepiej będzie działał wbudowany przyrząd pomiarowy. Efekt zadziałania ograniczenia prądowego został pokazany na **fotografii 8**. Załącza się górna dioda LED, na wyświetlaczu pojawia się napis CC a napięcie wyjściowe jest obniżane na tyle, by utrzymać zadane natężenie prądu. W podglądzie ustawionego natężenia prądu nie jest dokonane zaokrąglenie matematyczne, a obcięcie młodszej cyfry – stąd prąd 0,31 A oraz

0,38 A będą pokazywane jako 0,3 A – ponieważ to ma być jedynie wskaźnik.

Z pomiarem prądu jest związana jeszcze jedna kwestia. Przez rezystor pomiarowy stale przepływa prąd o natężeniu kilkunastu miliamperów, co jest związane z rozdzieleniem mas na płycie dolnej. Ten prąd powodował, że wbudowany amperomierz stale pokazywał prąd 0,01 A lub 0,02 A – co jest mało estetyczne. Dlatego wskazania tego amperomierza zaczynają się dopiero od 0,03 A, dla mniejszych prądów pokazuje 0,00 A.

Obwody sterujące zasilacza same decydują o tym, kiedy należy włączyć wentylator. W normalnych warunkach nie wiąże się to z jakimkolwiek komunikatem. Jeżeli jednak doszłoby do przegrzania, układ wyłącza napięcie wyjściowe, załącza dolną diodę LED na panelu przednim, a na wyświetlaczu pokazuje się, w dolnej linii, napis **OVERHEAT** zamiast natężenia prądu. Po ostudzeniu, normalna praca jest przywracana automatycznie.

Michał Kurzela, EP

REKLAMA



Elektronika Praktyczna
@ElektronikaPraktyczna

Strona główna

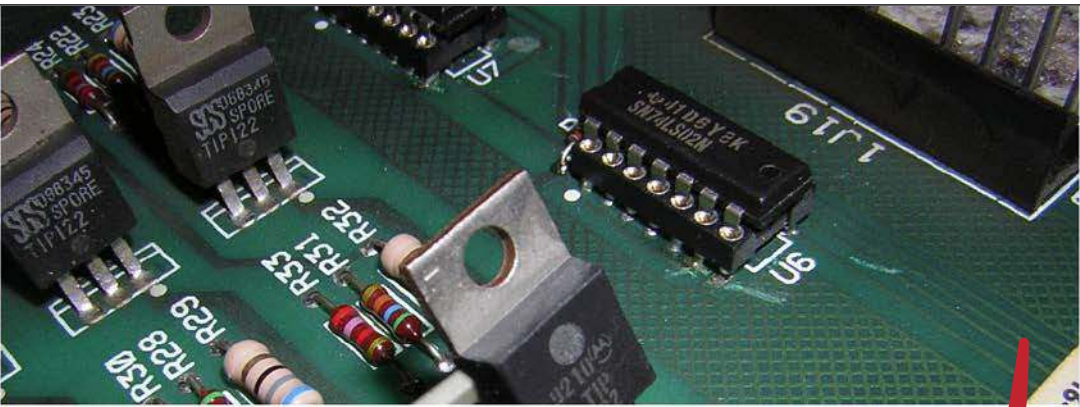
Posty

Filmy

Zdjęcia

Informacje

Społeczność



Lubię to!

Udostępnij

Zaproponuj zmiany

Wyślij e-mail

Wyślij wiadomość

Jesteśmy z Wami w kontakcie, również
w mediach społecznościowych

<https://www.facebook.com/ElektronikaPraktyczna>

**Podstawowe parametry:**

- załączanie generowanej losowo sekwencji diod LED (jednej z czterech),
- wydłużanie sekwencji o jeden krok po każdej w pełni prawidłowej odpowiedzi,
- możliwość regulacji tempa gry, czyli czasu załączania się kolejnych diod LED,
- wprowadzanie sekwencji poprzez wciskanie dużych przycisków,
- automatyczne przejście do uśpienia po około 10 s od ostatniej akcji użytkownika,
- pobór prądu około 120 nA w stanie uśpienia, 850 μ A podczas oczekiwania i do 5 mA podczas wyświetlania,
- zasilanie napięciem stałym 1,8...5 V,
- wbudowane gniazdo baterii CR2032.

***Uwaga!** Elektroniczne zestawy do samodzielnego montażu. Wymagana umiejętność lutowania! Podstawową wersją zestawu jest wersja [B] nazywana potocznie KIT-em (z ang. zestaw). Zestaw w wersji [B] zawiera elementy elektroniczne (w tym [UK] – jeśli występuje w projekcie), które należy samodzielnie wylutować w dołączoną płytkę drukowaną (PCB). Wykaz elementów znajduje się w dokumentacji, która jest podlinkowana w opisie kitu. Mając na uwadze różne potrzeby naszych klientów, oferujemy dodatkowe wersje:

- **wersja [C]** – zmontowany, uruchomiony i przetestowany zestaw [B] (elementy wylutowane w płytkę PCB),
 - **wersja [A]** – płytka drukowana bez elementów i dokumentacji.
- Kity, w których występuje układ scalony wymagający zaprogramowania, mają następujące dodatkowe wersje:
- **wersja [A-1]** – płytka drukowana [A] + zaprogramowany układ [UK] i dokumentacja,
 - **wersja [UK]** – zaprogramowany układ.

Dodatkowe materiały do pobrania ze strony www.ulubionykiosk.pl/media

- | | |
|---------|--|
| AVT5910 | Gra „kto pierwszy” (EP 1/2022) |
| AVT5639 | Gra elektroniczna „Snake” (EP 9/2018) |
| AVT5592 | Gra elektroniczna „Sudoku” (EP 7/2017) |
| AVT5554 | Gra elektroniczna „Snake” (EP 11/2016) |
| AVT1651 | Gra „Kto pierwszy ten lepszy” (EP 11/2011) |
| AVT723 | Uniwersalna gra zręcznościowa (EP 6/2004) |
| AVT5028 | Elektroniczna gra w kości (EP 8/2001) |
| AVT5014 | Gra zręcznościowa (EP 5/2001) |

Nie każdy zestaw AVT występuje we wszystkich wersjach! Każda wersja ma załączony ten sam plik PDF! Podczas składania zamówienia upewnij się, którą wersję zamawiasz! <http://sklep.avt.pl>

W przypadku braku dostępności na stronie sklepu osoby zainteresowane zakupem płytek drukowanych (PCB) prosimy o kontakt via e-mail: kity@avt.pl.

W ofercie AVT*

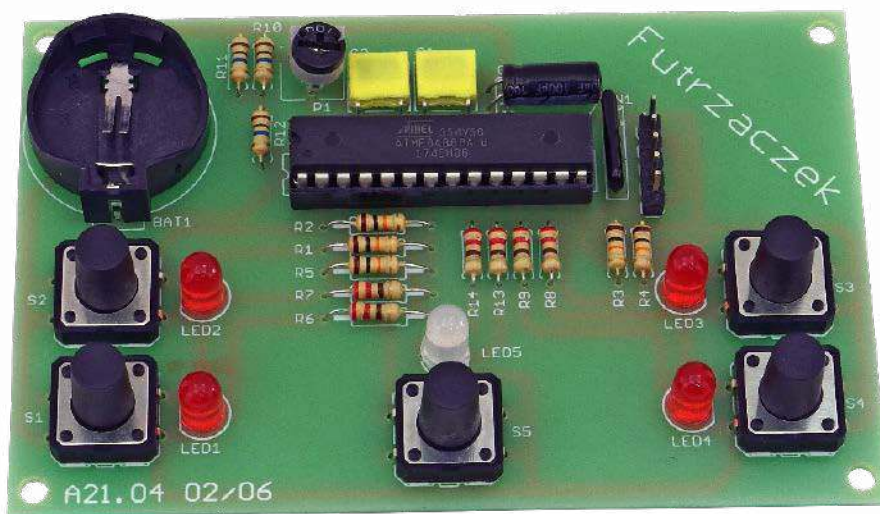
AVT5967

Gra pamięciowa – odtwarzanie sekwencji

Gry i zabawki towarzyszą nam od pierwszych tygodni życia. Ich główną rolę jest rozwijanie naszych zdolności. Jedną z nich jest nasza pamięć, którą powinniśmy ćwiczyć bezustannie na różne sposoby. Ten nieskomplikowany układ może nam w tym pomóc. Zwarta konstrukcja pozwala na trzymanie tej gry przy sobie – w plecaku, schowku samochodowym albo w torbie.

Nasza pamięć może zaskakiwać swoją pojemnością przez wiele lat, o ile regularnie jej używamy. W tym celu powstało wiele gier i zabaw, które wymagają jej wysilania. Kojarzenie symboli na wyłożonych wcześniej kartach to tylko jedna z możliwości. Część z nich wymaga obecności drugiej osoby, co bywa problematyczne. Na szczęście, pomóc nam może elektronika, która będzie z matematyczną precyzją podchodzić do udzielanych odpowiedzi.

Ta niewielka gra ma formę podręcznej konsoli, obsługiwanej obiema dłońmi, na kształt znanych przenośnych konsol do gier. Jej obsługa jest bardzo prosta: po włączeniu zaczyna pokazywać sekwencję załączających się diod. Zaczyna się zawsze od jednej diody, po czym czeka na wcisnięcie odpowiadającego jej przycisku. Jeżeli odpowiedź jest prawidłowa, układ wydłuża sekwencję



– poprzednio załączona dioda jest wskazywana ponownie, po czym dokładana jest kolejna. Jakikolwiek błąd powoduje powrót do początku.

W ćwiczeniu pamięci pomaga fakt, że sekwencje są generowane całkowicie losowo. Prawdopodobieństwo dodania określonej diody (jednej spośród czterech) wynosi 25%. Nie sposób zatem nauczyć się kolejnych kroków sekwencji, ponieważ po przegranej jest ona budowana całkowicie od nowa.

Budowa i działanie

Schemat ideowy omawianego urządzenia znajduje się na **rysunku 1**. Głównym komponentem sterującym pracą układu jest mikrokontroler ATmega88PA-PU

z 8-bitowym rdzeniem AVR. Ma wystarczającą liczbę konfigurowalnych wyprowadzeń i może pracować z napięciem zasilania już od 1,8 V. Przez większość czasu znajduje się w stanie uśpienia, z którego wybudza go przerwanie od zmiany stanu na wyprowadzeniu PD2. Nie realizuje zadań krytycznych czasowo, wobec czego częstotliwość zegara jest stabilizowana przez wbudowany układ oscylatora RC. Rezystory zawarte w drabince RN1 podciągają wejścia służące programowaniu ISP, w tym wejście zerujące, do dodatniego potencjału zasilania, co zmniejsza ryzyko nieprawidłowego zadziałania układu spowodowane ładunkami elektrostatycznymi lub zakłóceniami elektromagnetycznymi.

Wykaz elementów, kupuj na stronie sklep.avt.pl (Warszawa, ul. Leszczynowa 11, tel. +48222578451, e-mail: handlowy@avt.pl)

Kondensatory:

C1, C3: 100 nF raster 5 mm MKT
C2: 100 μ F 25 V raster 2,5 mm

Rezystory: (THT o mocy 0,25 W)

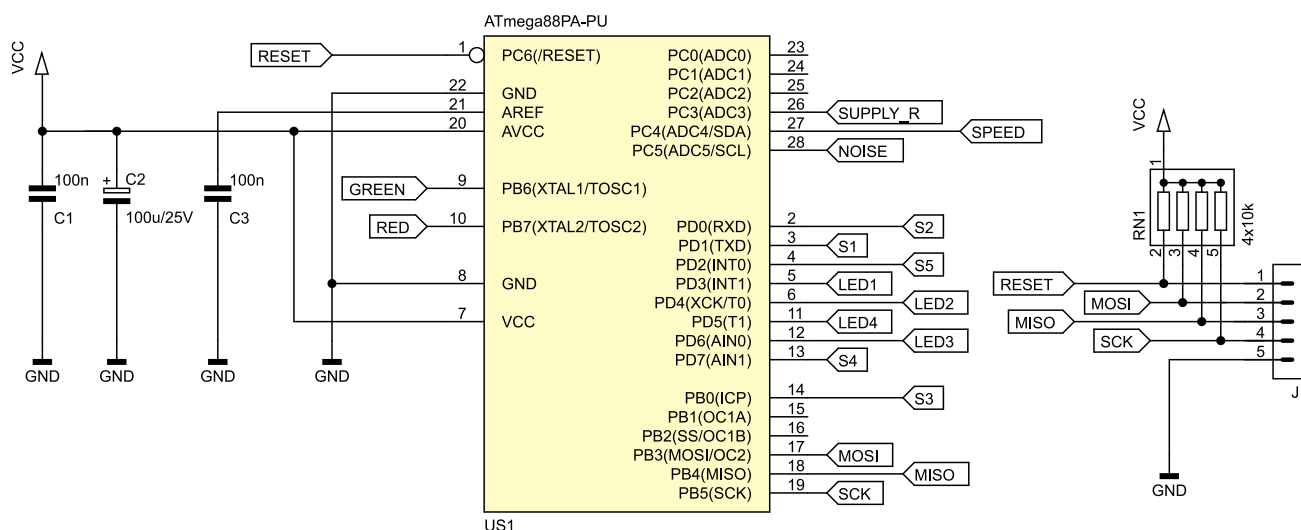
R1...R5: 10 k Ω
R6...R9, R13, R14: 220 Ω
R10...R12: 10 M Ω

RN1: 4 $\times$ 10 k Ω SIL5P1: 100 k Ω montażowy leżący**Półprzewodniki:**

LED1...LED4: czerwone 5 mm np. LED F5 R
LED5: zielona/czerwona 5 mm ze wspólną katodą np. LED F5 R/G-1
U1: ATmega88PA-PU DIP28

Pozostałe:

BAT1: koszyk baterii CR2032 THT leżący (KOSZYK BAT 6)
J1: goldpin 5 pin męski 2,54 mm kątowy THT
S1...S5: mikroswitch 12 $\times$ 12 9 mm THT (MIKROSW TS)
Jedna podstawka DIP28 wąska



W układzie znajduje się pięć przycisków monostabilnych. Cztery z nich – S1...S4 odpowiadają diodom LED, które wskazują sekwencję do zapamiętania, a S5 służy do uruchamiania gry. Wewnętrzne rezystory podciągające mikrokontrolera zostały połączone równolegle z zewnętrznymi R1...R5 co zwiększa pewność prawidłowego zadziałania układu w przypadku np. dotknięcia palcem ścieżki na powierzchni płytki drukowanej.

Wbudowany w mikrokontroler przetwor-
nik analogowo-cyfrowy jest wykorzystywany
do generowania losowych liczb oznaczających

kolejną diodę w sekwencji oraz do ustalania tempera rozgrywki. Załączenie zasilania części analogowej odbywa się po wystawieniu

Hurtownia elementów elektronicznych "AKSOTRONIK" zaprasza do swojego sklepu internetowego.
Zaloguj się i kupuj ON-LINE na naszej stronie:

WWW.AKSOTRONIK.COM.PL

**Magnesy neodymowe
oraz ferrytowe**
Ceny od 0.10zł

**Przełączniki klawiszowe
wodoszczelne, pyłoszczelne**
Ceny od 2.40zł

**Prowadniki
do przewodów**
Ceny od 11.00zł

**Kostki elektryczne
zachwowe**
Ceny od 0.22zł

**Druły oporowe
od 0.16 do 0.81mm**
Ceny od 5.70zł

**Złącza hermetyczne
Superseal**
Ceny od 1.10zł kpl

**Szczotki węglowe
do elektronarzędzi**
Ceny od 2.60zł/kpl

Pudełka/organizery
Ceny od 0.95zł

**Przełączniki do elektronarzędzi
zwykłe i elektromagnetyczne**
Ceny od 7.00zł

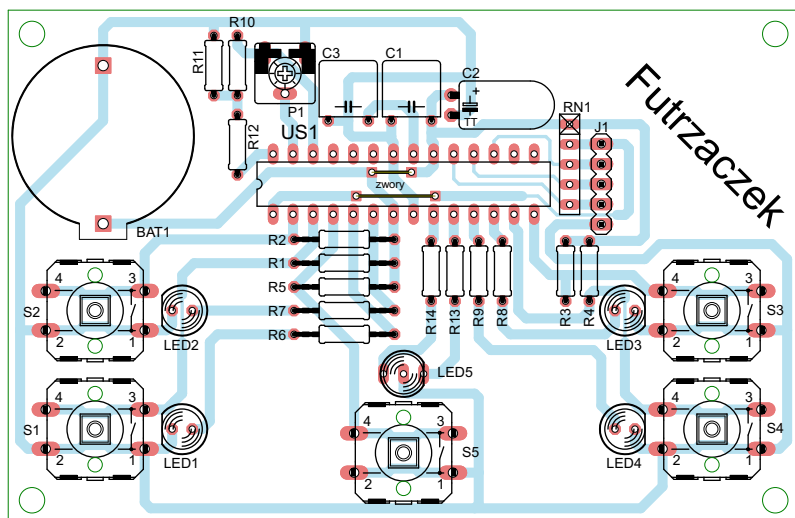
**Zestawy śrubek M2, M3
z nakrętkami i podkładkami**
Ceny od 2.50zł

Aksonotronics
ELEMENTY ELEKTRONICZNE

Uwaga!!! Powyższe ceny dotyczą zakupów minimalnych ilości hurtowych, poprzez nasz sklep internetowy.

W swoje ofercie posiadamy m.in.: poliprowadniki (diody, układy scalone, tranzystory, triaki, elementy optoelektroniczne), elementy dystansowe, złącza, przełączniki, elementy akustyczne, rezystory, kondensatory, kwarce, podstawki, moduły Arduino

Zapraszamy do kontaktu: **INFO@aksotronik.com.pl**, tel.: (22) 783-20-51



Rysunek 2. Schemat montażowy i wzór ścieżek płytki

stanu wysokiego na wyprowadzenie PC3 mikrokontrolera. Wtedy dokonywany jest pomiar napięcia pomiędzy ślizgaczem P1 i masą układu, co służy zadaniu szybkości załączania się kolejnych diod w sekwencji.

Obwód składający się z rezystorów R10... R12 generuje silnie zaszumione napięcie stałe. Przetwornik wychwytuje chwilową jego wartość, konwertuje, po czym wyłuskuje ostatni, najmłodszy bit. Drugi taki sam pomiar, wykonany po kilku milisekundach, generuje drugi bit. Ich wartość jest zupełnie losowa, bowiem głównych źródłem szumów w tym układzie są szумы termiczne rezystorów, mające rozkład jednostajny (tak zwany szum biały). Dwa zestawione ze sobą bity tworzą liczbę z przedziału 0...3, co odpowiada wprost numerowi diody do załączenia: LED1...LED4.

Zasilaniem dla układu jest pojedyncza bateria BAT1 typu CR2032, której koszyk znajduje się na płytce. Jej pojemność jest na tyle znacząca, że wystarczy na wiele godzin gry oraz wiele miesięcy trwania w uśpieniu. Jednocześnie, cała gra staje się w ten sposób lekka i kompaktowa, ponieważ nie wymaga instalowania dodatkowych źródeł energii lub zasilacza sieciowego.

Montaż i uruchomienie

Układ został zmontowany na jednostronnej płytce drukowanej o wymiarach 100×65 mm, której schemat został pokazany na **rysunku 2**. W odległości 3 mm od krawędzi płytki znalazły się cztery otwory montażowe, każdy o średnicy 3,2 mm.

Montaż proponuję rozpocząć od elementów o najmniejszej wysokości obudowy,

czyli rezystorów, kondensatorów położonych na powierzchni laminatu, diod. Nie wolno zapomnieć o dwóch zwórkach z cienkiego drutu, których miejsce jest pod układem US1. Pod mikrokontroler proponuję zastosować podstawkę, aby ułatwić jego programowanie oraz wymianę w razie uszkodzenia. Zmontowany układ można zobaczyć na fotografii tytułowej.

Mikrokontroler powinien mieć fabryczne ustawienia bitów zabezpieczających. Wbudowany oscylator RC, który taktuje rdzeń sygnałem o częstotliwości 1 MHz oraz wyłączony Brown-Out Detector to prawidłowa konfiguracja w tym układzie. Pamięć Flash trzeba zaprogramować dostarczoną wsadą.

Poprawnie zmontowany i zaprogramowany układ jest gotowy do działania po włożeniu baterii typu CR2032 do koszyka. Napięcie zasilające, o ile będzie pochodziło z zewnętrznego zasilacza, nie powinno przekraczać wartości 5 V i na pewno powinno być dobrze filtrowane, a najlepiej stabilizowane. Układ był projektowany do współpracy z napięciem o wartości około 3 V i w tych warunkach zmierzono pobór prądu: około 120 nA w stanie uśpienia, około 850 µA podczas oczekiwania na wciśnięcie przycisku podczas wprowadzania zapamiętanej sekwencji i około 5 mA podczas świecenia diodą LED. Minimalnym napięciem, przy którym układ działa, jest 1,8 V – wynika to z dokumentacji użytego mikrokontrolera.

Eksplotacja

Jeżeli nasza zabawka ma włożoną baterię, to zapewne znajduje się w stanie uśpienia. Można ją wybudzić poprzez wciśnięcie

przycisku S5. Natychmiast po tym załączy jedną z czterech diod LED umieszczonych przy przyciskach. Prawidłową odpowiedzią będzie wciśnięcie tego przycisku w czasie nie dłuższym niż 10 s – jeżeli w tym czasie nie będzie reakcji, układ powróci do stanu uśpienia. Przejście w ten stan jest sygnalizowane przez diodę LED5, która błyska na czerwono, a potem na zielono.

Założmy, że załączyła się dioda LED3 i chwilę później człowiek wcisnął S3. Świecenie LED5 na zielono potwierdzi, że to był dobry wybór. Zaraz po tym, jak LED5 zgaśnie, załączy się LED3, zgaśnie i załączy się inna dioda, dajmy na to LED1. Po wygaszeniu LED1 trzeba tę sekwencję powtórzyć: najpierw S3, potem S1. Zazielenienie się LED5 wynagrodzi nam włożony wysiłek. Następna sekwencja będzie dłuższa, 3-elementowa: LED3, LED1 i jeszcze jakaś. Podczas wciskania przycisków będziemy widzieć chwilowe załączanie się diod przy nich. To informacja zwrotna, że styki w mikroswitchu zadziałały prawidłowo, a sygnał został prawidłowo zinterpretowany.

Jeżeli jakkolwiek z naszych odpowiedzi będzie nieprawidłowa, czyli wciśnięty przycisk nie będzie odpowiadał diodzie, która się wcześniej załączyła w tej sekwencji, LED5 stanie się czerwona. W ten sposób układ poinformuje nas, że popełniliśmy błąd. Sekwencja zacznie się budować od nowa. Tym razem pierwszą diodą może być, na przykład, LED2.

Potencjometr P1 ustala czas świecenia diod podczas pokazywania sekwencji. Po skręceniu na minimum czas ten wynosi 70 ms, co wymusza silne skupienie uwagi. Z kolei, maksymalny czas świecenia to nieco ponad 1 s. Rozdzielczość regulacji wynosi 1 ms, więc niewielkim wkręciakiem można sobie wyregulować poziom trudności gry. Przypominam, że informacja z P1 jest zbierana wyłącznie po wybudzeniu mikrokontrolera lub po włożeniu baterii w gniazdo BAT1.

Maksymalna długość zbudowanej sekwencji to 256 elementów. Po przekroczeniu tej długości, licznik indeksujący tablicę zapamiętanych elementów przepełnia się i układ wraca na początek zabawy, czyli rozpoczyna od jednej diody. Życzę Czytelnikom takiego sukcesu!

Michał Kurzela, EP

REKLAMA

EPcom.pl

**Podstawowe parametry:**

- dwukierunkowy układ sterowania silnikiem małej mocy,
- obciążalność do 1,76 A,
- napięcie pracy 4,5...18 V,
- zintegrowane zabezpieczenie przeciążeniowe, podnapięciowe i termiczne,
- dostępny tryb obniżonego poboru mocy,
- bez dodatkowych konwerterów może współpracować z logiką 1,8 V, 3,3 V, 5 V.

* **Uwaga!** Elektroniczne zestawy do samodzielnego montażu. Wymagana umiejętność lutowania! Podstawową wersją zestawu jest wersja [B] nazywana potocznie KIT-em (z ang. zestaw). Zestaw w wersji [B] zawiera elementy elektroniczne (w tym [UK] – jeśli występuje w projekcie), które należy samodzielnie wstawić w dołączoną płytkę drukowaną (PCB). Wykaz elementów znajduje się w dokumentacji, która jest podlinkowana w opisie kitu. Mając na uwadze różne potrzeby naszych klientów, oferujemy dodatkowe wersje:

- wersja [C] – zmontowany, uruchomiony i przetestowany zestaw [B] (elementy wstawiane w płytkę PCB),
 - wersja [A] – płytkę drukowaną bez elementów i dokumentacji.
- Kity, w których występuje układ scalony wymagający zaprogramowania, mają następujące dodatkowe wersje:
- wersja [A+] – płytkę drukowaną [A] + zaprogramowany układ [UK] i dokumentacja,
 - wersja [UK] – zaprogramowany układ.

Dodatkowe materiały do pobrania ze strony www.ulubionykiosk.pl/media

- AVT5923 Płynny regulator prędkości i kierunku obrotów silnika 12 V (EP 3/2022)
- AVT5879 Monitor pracy wentylatora (EP 8/2021)
- Projekt 246 Sterownik wentylatora wyciągu łazienkowego (kuchennego) (EP 10/2019)
- AVT5698 Sterownik wentylatorów 12 V dużej mocy (EP 8/2019)
- AVT5653 Sterownik mikrowentylatora (EP 11/2018)
- Uniwersalny driver silnika małej mocy (EP 3/2018)
- AVT1981 Sterownik wentylatora z płynną zmianą obrotów (EP 1/2018)

Nie każdy zestaw AVT występuje we wszystkich wersjach! Każda wersja ma załączony ten sam plik PDF! Podczas składania zamówienia upewnij się, którą wersję zamawiasz! <http://sklep.avt.pl>

W przypadku braku dostępności na stronie sklepu osoby zainteresowane zakupem płytek drukowanych (PCB) prosimy o kontakt via e-mail: kity@avt.pl.

Mikrodriver silnika DC małej mocy

Są urządzenia, których nie da się zrobić prościej. Jednym z nich jest niewielki moduł drivera silnika DC małej mocy na bazie najnowszego komponentu TI – układu DRV8220. W miniaturowej obudowie SOT563-6, z ograniczoną do absolutnego minimum liczbą wyprowadzeń, znajduje się kompletny dwukierunkowy układ sterowania jednym silnikiem.

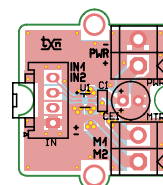
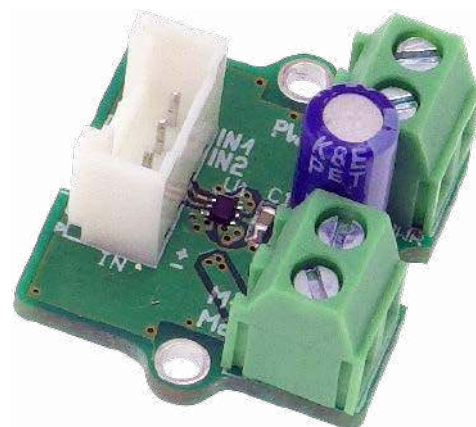
Strukturę wewnętrzną układu DRV8220 pokazano na **rysunku 1**. Zwiera on dwa półmostkowe stopnie mocy MOSFET o obciążalności do 1,76 A i napięciu pracy 4,5...18 V, z wbudowanym układem zabezpieczenia przeciążeniowego, podnapięciowego i termicznego oraz trybem obniżonego poboru mocy. Sposób sterowania dla wejść IN1, IN2 zestawiono w **tabeli 1**.

Budowa i działanie

Schemat modułu pokazano na **rysunku 2**. Aplikacja jest bardzo prosta i składa się z układu DRV8220 i elementów odsprężających. Moduł zasilany jest napięciem stałym ze złącza PWR, które jest filtrowane pojemnościami C1, CE1. Silnik podłączony jest do złącza MTR, sterowanie do złącza IN zgodnego z serią Grove.

Montaż i uruchomienie

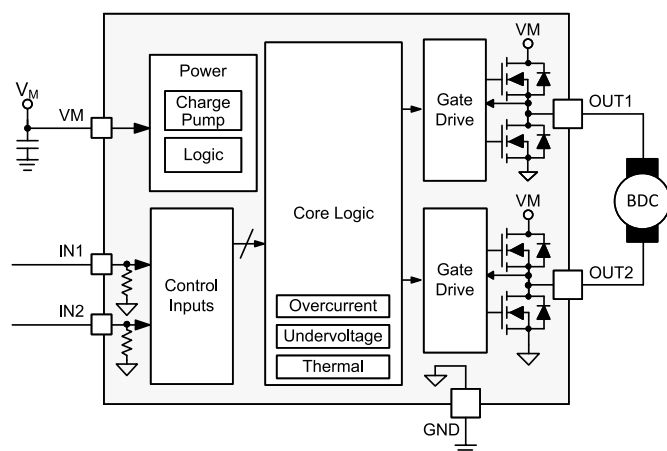
Minimoduł zmontowany jest na dwustronnej płytce drukowanej o wymiarach 20×20 mm. Jej schemat znajduje się na **rysunku 3**. Sposób montażu jest klasyczny i nie wymaga opisu. Moduł nie wymaga uruchamiania. Po podłączeniu silnika do zacisków MTR i zasilania do zacisków PWR, możliwe jest sterowanie jego pracą poprzez zmianę stanów wejść IN1, IN2 zgodnie z wcześniejszym opisem. Napięcie wejściowe na wyprowadzeniach IN1, IN2 nie powinno przekraczać 5,5 V. Moduł bez dodatkowych konwerterów może współpracować z logiką 1,8 V, 3,3 V, 5 V więc praktycznie z wszystkimi mikrokontrolerami i komputerami SBC. Należy pamiętać o prawidłowej filtracji napięcia zasilania



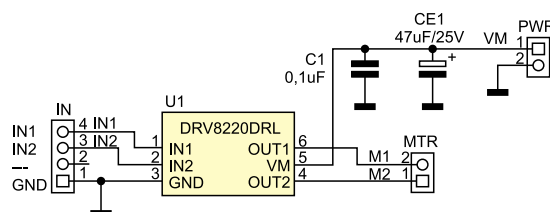
Rysunek 3. Schemat płytki PCB

i dodatkowym jego odsprężaniu, kondensator CE1 filtruje zasilanie samego drivera.

Adam Tatuś, EP



Rysunek 1. Schemat wewnętrzny DRV8220



Rysunek 2. Schemat ideowy modułu

Tabela 1. Sterowanie układ

IN1	IN2	OUT1	OUT2	Opis
0	0	Hi-Z	Hi-Z	Tryb niskiego poboru mocy, stopień mocy wyłaczony
0	1	L	H	Zasilanie silnika w kierunku 1
1	0	H	L	Zasilanie silnika w kierunku 2
1	1	L	L	Hamowanie – wyjścia w stanie niskim

Wykaz elementów, kupuj na stronie sklep.avt.pl (Warszawa, ul. Leszczyńska 11, tel. +48222578451, e-mail: handlowy@avt.pl)

Kondensatory:

C1: 0,1 µF ceramiczny (SMD0603)
CE1: 47 µF/25 V elektrolityczny 5 mm

Półprzewodniki:

U1: DRV8220DRL (SOT-563)

Pozostałe

MTR, PWR: złącze kątowe śrubowe 3,5 mm (DG381-3.5-2)

**Podstawowe parametry:**

- umożliwia niezależną regulację jasności dwóch źródeł światła LED,
- napięcie wyjściowe wynosi 5 V, a maksymalny pobór mocy to 2×2,5 W,
- może być zasilany z powerbanku lub ładowarki USB.

*** Uwaga!** Elektroniczne zestawy do samodzielnego montażu. Wymagana umiejętność lutowania! Podstawową wersją zestawu jest wersja [B] nazywana potocznie KIT-em (z ang. zestaw). Zestaw w wersji [B] zawiera elementy elektroniczne (w tym [UK] – jeśli występuje w projekcie), które należy samodzielnie wstawić w dołączoną płytkę drukowaną (PCB). Wykaz elementów znajduje się w dokumentacji, która jest podlinkowana w opisie kitu. Mając na uwadze różne potrzeby naszych klientów, oferujemy dodatkowe wersje:

- wersja [C] – zmontowany, uruchomiony i przetestowany zestaw [B] (elementy wstawione w płytkę PCB),
 - wersja [A] – płyta drukowana bez elementów i dokumentacji.
- Kity, w których występuje układ scalony wymagający zaprogramowania, mają następujące dodatkowe wersje:
- wersja [A+] – płyta drukowana [A] + zaprogramowany układ [UK] i dokumentacja,
 - wersja [UK] – zaprogramowany układ.

Dodatkowe materiały do pobrania ze strony www.ulubionykiosk.pl/media

- SCL – zaawansowany sterownik oświetlenia schodowego część 1 i 2 (EP 11 i 12/2022)
- Zasilacz power LED średniej mocy (EP 8/2022)
- AVT5921 Włącznik LED z płynnym rozjaśnianiem i ściemnianiem (EP 3/2022)
- AVT5916 Regulator jasności LED sterowany pilotem TV (EP 2/2022)
- Ambient LED controller (EP 1/2022)
- AVT5880 Sterownik 12×LED z interfejsem I²C (EP 8/2021)
- Sterownik LED RGB z układem AL1783 sterowany przez I²C (EP 6/2021)
- AVT5857 Linijowy sterownik LED 3 W (EP 4/2021)
- AVT5839 Minimoduły z driverem I²C do taśm LED RGBW (EP 1/2021)

Nie każdy zestaw AVT występuje we wszystkich wersjach! Każda wersja ma załączony ten sam plik PDF! Podczas składania zamówienia upewnij się, którą wersję zamawiasz! <http://sklep.avt.pl>

W przypadku braku dostępności na stronie sklepu osoby zainteresowane zakupem płytek drukowanych (PCB) prosimy o kontakt via e-mail: kity@avt.pl.

W ofercie AVT*

AVT5968

Regulator jasności podświetleń do fotografii produktowej i makro

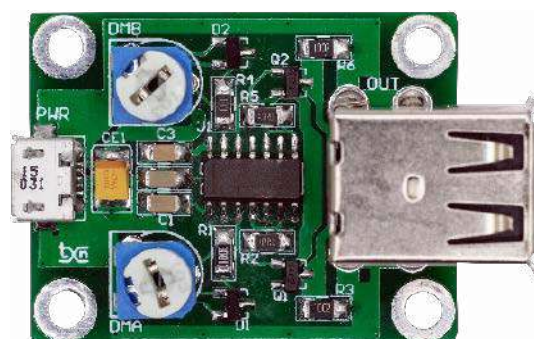
Układ powstał po to, aby ułatwić fotografowanie detali w małym namiocie bezcieniowym. Namiot ma podświetlenie górne, w postaci listew LED oraz dodatkowo został doposażony w podświetlaną podstawkę LED. Oba źródła zasilane są z ładowarki USB, niestety pozbawione są regulacji jasności.

Mały regulator PWM, zasilany z banku energii lub ładowarki USB, umożliwia niezależną regulację jasności dwóch źródeł światła LED. Ułatwia to prawidłowe oświetlenie detalu i uprasza wykonanie poprawnie naświetlonej fotografii.

Budowa i działanie

Schemat ideowy układu jest pokazany na rysunku 1. Moduł składa się z dwóch

identycznych kanałów. Regulator zasilany jest z gniazda USB micro PWR. Podłączony zasilacz musi mieć wydajność min. 1 A. Generator o wypełnieniu regulowanym potencjometrem oznaczonym DMA, zbudowany jest na bramce Schmitta U1A układu HC14. Równolegle połączone bramki U1B, U1C buforują generator i zapewniają wysterowanie tranzystora kluczującego Q1. Przebieg zasilający PWM doprowadzony jest do gniazd OUT. Układ

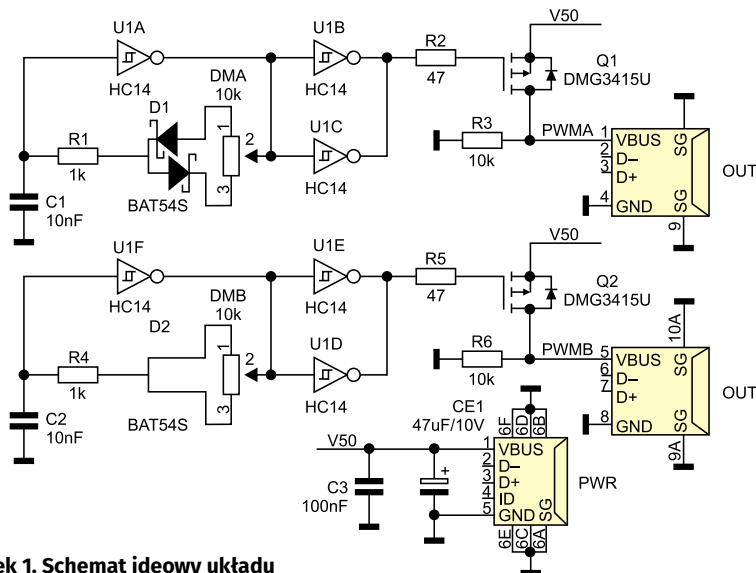


zaprojektowany jest do podświetleń o napięciu 5 V i poboru mocy maksymalnie 2×2,5 W, co wystarcza w opisaney aplikacji.

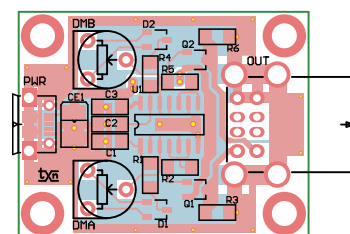
Montaż i uruchomienie

Układ zmontowany jest na niewielkiej dwustronnej płytce drukowanej, której schemat pokazano na rysunku 2. Montaż nie wymaga szczegółowego opisu. Po podłączeniu zasilania do złącza PWR i LED do gniazd OUT, zmieniając położenie potencjometrów powinna być możliwa niezależna regulacja jasności LED. Warto zwrócić uwagę na jakość ładowarki i kabli połączeniowych, aby spadki napięć nie powodowały zakłóceń w regulacji jasności kanałów.

Adam Tatuś, EP



Rysunek 1. Schemat ideowy układu



Rysunek 2. Schemat płytki PCB

Wykaz elementów, kupuj na stronie sklep.avt.pl (Warszawa, ul. Leszczyńska 11, tel. +48222578451, e-mail: handlowy@avt.pl)

Rezystory: (SMD1206, 1%)

R1, R4: 1 kΩ
R2, R5: 47 Ω
R3, R6: 10 kΩ

Kondensatory:

C1, C2: 10 nF (SMD1206)

CE1: 47 μF/10 V tantalowy (SMD3528)

C3: 100 nF (SMD1206)

Półprzewodniki:

D1, D2: BAT54S dioda podwójna (SOT-23)
Q1, Q2: DMG3415U tranzystor MOSFET (SOT-23)
U1: HC14 (SO14)

Pozostałe:

DMA, DMB: potencjometr wzornicowy 10 kΩ 5 mm
OUT: gniazdo USB A podwójne
PWR: gniazdo USB micro (MX105017-0001)

**Podstawowe parametry:**

- wyprowadza sygnały magistrali I²C na złącza SIP/EH/JST/Grove/Qwiic/Gravity oraz śrubowe DG

* **Uwaga!** Elektroniczne zestawy do samodzielnego montażu. Wymagana umiejętność lutowania! Podstawową wersją zestawu jest wersja **[B]** nazywana potocznie KIT-em (z ang. zestaw). Zestaw w wersji **[B]** zawiera elementy elektroniczne (w tym **[UK]** – jeśli występuje w projekcie), które należy samodzielnie włożyć w dołączoną płytkę drukowaną (PCB). Wykaz elementów znajduje się w dokumentacji, która jest podlinkowana w opisie kitu. Mając na uwadze różne potrzeby naszych klientów, oferujemy dodatkowe wersje:

- wersja [C]** – zmontowany, uruchomiony i przetestowany zestaw **[B]** (elementy wlotowane w płytkę PCB),
 - wersja [A]** – płytka drukowana bez elementów i dokumentacji.
- Kity, w których występuje układ scalony wymagający zaprogramowania, mają następujące dodatkowe wersje:
- wersja [A+]** – płytka drukowana **[A]** + zaprogramowany układ **[UK]** i dokumentacja,
 - wersja [UK]** – zaprogramowany układ.

Dodatkowe materiały do pobrania ze strony www.ulubionykiosk.pl/media

- Sterownik mikrokontrolera krokowego do Pi Pico (EP 12/2022)
- Radiomodem ISM do Raspberry Pi Zero (EP 11/2022)
- Moduł LoRa do RPi Pico (EP 9/2022)
- Moduł z wyświetlaczami numitron (EP 8/2022)
- Interfejs aparatury kontrolnej i sygnalizacyjnej standardu M22 do Raspberry Pi (EP 7/2022)
- Sterownik mikrokontrolerów prądu stałego do Rpi Pico (EP 7/2022)

Nie każdy zestaw AVT występuje we wszystkich wersjach! Każda wersja ma załączony ten sam plik PDF! Podczas składania zamówienia upewnij się, którą wersję zamawiasz!
<http://sklep.avt.pl>

W przypadku braku dostępności na stronie sklepu osoby zainteresowane zakupem płytek drukowanych (PCB) prosimy o kontakt via e-mail: kity@avt.pl.

Uniwersalny adapter I²C

Podczas prototypowania układów z interfejsem I²C przy zastosowaniu gotowych modułów bardzo często spotykamy się z koniecznością wykorzystania płytek różnych producentów różniących się nie tylko typem złącza, ale i różnym sposobem wyprowadzania sygnałów I²C. Generuje to całkiem sporą liczbę nietypowych przewodów wymaganych dla realizacji połączenia. Układ adaptera zawiera najbardziej typowe złącza magistrali I²C, co umożliwi zastosowanie standardowych dla danego systemu przewodów, usprawniając codzienne prace warsztatowe.

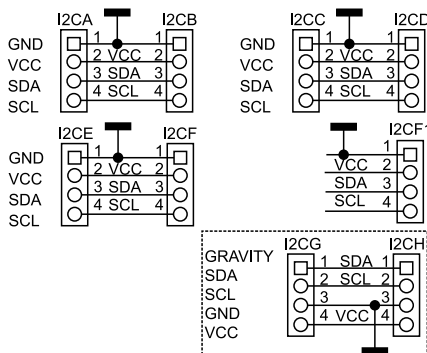


Budowa i działanie

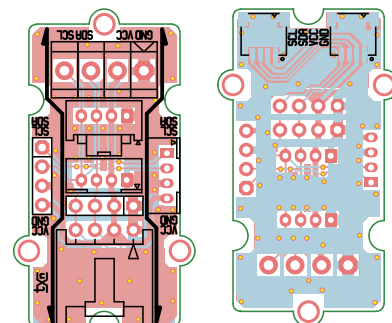
Schemat układu pokazano na **rysunku 1**. Wyprowadza sygnały magistrali I²C pomiędzy złączami SIP/EH/JST/Grove/Qwiic oraz śrubowym DG (I²CA...I²CF1). Dodatkowe złącza SIP/JST (I²CG/I²CH) z przyporządkowaniem sygnałów według standardu Gravity uzupełniają funkcjonalność modułu.

Montaż i uruchomienie

Adapter zmontowany jest na dwustronnej płytce drukowanej, której schemat został pokazany na **rysunku 2**. Montaż jest typowy i nie wymaga opisu.



Adam Tatuś, EP

Rysunek 1. Schemat ideowy adaptera I²C

Rysunek 2. Schemat płytki PCB

Wykaz elementów, kupuj na stronie sklep.avt.pl (Warszawa, ul. Leszczyńska 11, tel. +48222578451, e-mail: handlowy@avt.pl)

I²CA: złącze EH kątowne 2,54 mm
I²CB, I²CH: złącze SIP 2,54 mm

I²CC, I²CG: złącze JST 2 mm
I²CD: złącze śrubowe 3,81 mm (DG381-3.5-4)

I²CE: złącze Grove 2 mm (110990030)
I²CF, I²CF1: złącze JST 1 mm

REKLAMA

Świat projektantów i programistów dla elektroniki w nowej odstonie.

Odwiedź wiecznie młody

ELPORTAL.pl

**Podstawowe parametry:**

- wyjście aktywowane jest po ustabilizowaniu się stanu wejścia na czas dłuższy niż 40 ms
- wejścia zabezpieczone przed skutkami wyładowań ESD,
- komunikacja i odczyt stanu wejść poprzez interfejs I²C,
- moduł dopasowany do płytki Raspberry Pi Zero,
- zasilanie 3,3 V.

* **Uwaga!** Elektroniczne zestawy do samodzielnego montażu. Wymagana umiejętność lutowania! Podstawową wersją zestawu jest wersja [B] nazywana potocznie KIT-em (z ang. zestaw). Zestaw w wersji [B] zawiera elementy elektroniczne (w tym [UK] – jeśli występuje w projekcie), które należy samodzielnie wylutować w dotychczasową płytkę drukowaną (PCB). Wykaz elementów znajduje się w dokumentacji, która jest podlinkowana w opisie kitu. Mając na uwadze różne potrzeby naszych klientów, oferujemy dodatkowe wersje:

- wersja [C] – zmontowany, uruchomiony i przetestowany zestaw [B] (elementy wylutowane w płytkę PCB),
 - wersja [A] – płytka drukowana bez elementów i dokumentacji.
- Kity, w których występuje układ scalony wymagający zaprogramowania, mają następujące dodatkowe wersje:
- wersja [A+] – płytka drukowana [A] + zaprogramowany układ [UK] i dokumentacja,
 - wersja [UK] – zaprogramowany układ.

Dodatkowe materiały do pobrania ze strony www.ulubionykiosk.pl/media

- Sterownik mikrosilnika krokowego dla Pi Pico (EP 12/2022)
- Radiomodem ISM do Raspberry Pi Zero (EP 11/2022)
- Moduł LoRa dla RPi Pico (EP 9/2022)
- Moduł z wyświetlaczami numitron (EP 8/2022)
- Interfejs aparatury kontrolnej i sygnalizacyjnej standardu M22 do Raspberry Pi (EP 7/2022)
- Sterownik mikrosilników prądu stałego do Rpi Pico (EP 7/2022)

Nie każdy zestaw AVT występuje we wszystkich wersjach! Każda wersja ma załączony ten sam plik PDF! Podczas składania zamówienia upewnij się, którą wersję zamawiasz!
<http://sklep.avt.pl>

W przypadku braku dostępności na stronie sklepu osoby zainteresowane zakupem płytek drukowanych (PCB) prosimy o kontakt via e-mail: kity@avt.pl.

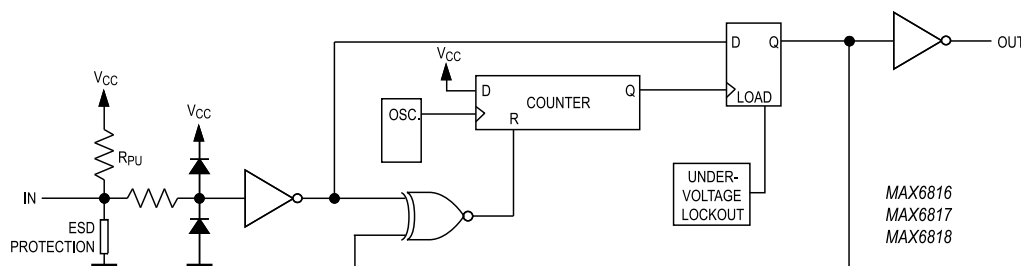
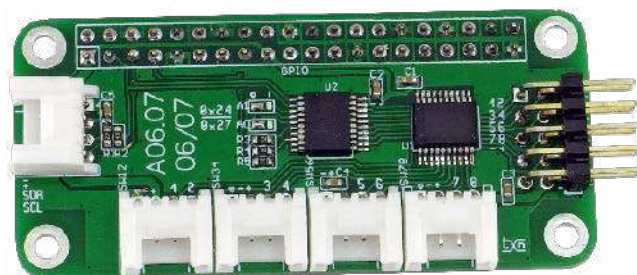
Eliminator drgań styków mechanicznych

Zaprezentowany moduł umożliwia realizację interfejsu do ośmiu styków mechanicznych eliminując konieczność programowego filtrowania ich drgań oraz zabezpieczając wejścia współpracującego układu przed skutkami wyładowań ESD.

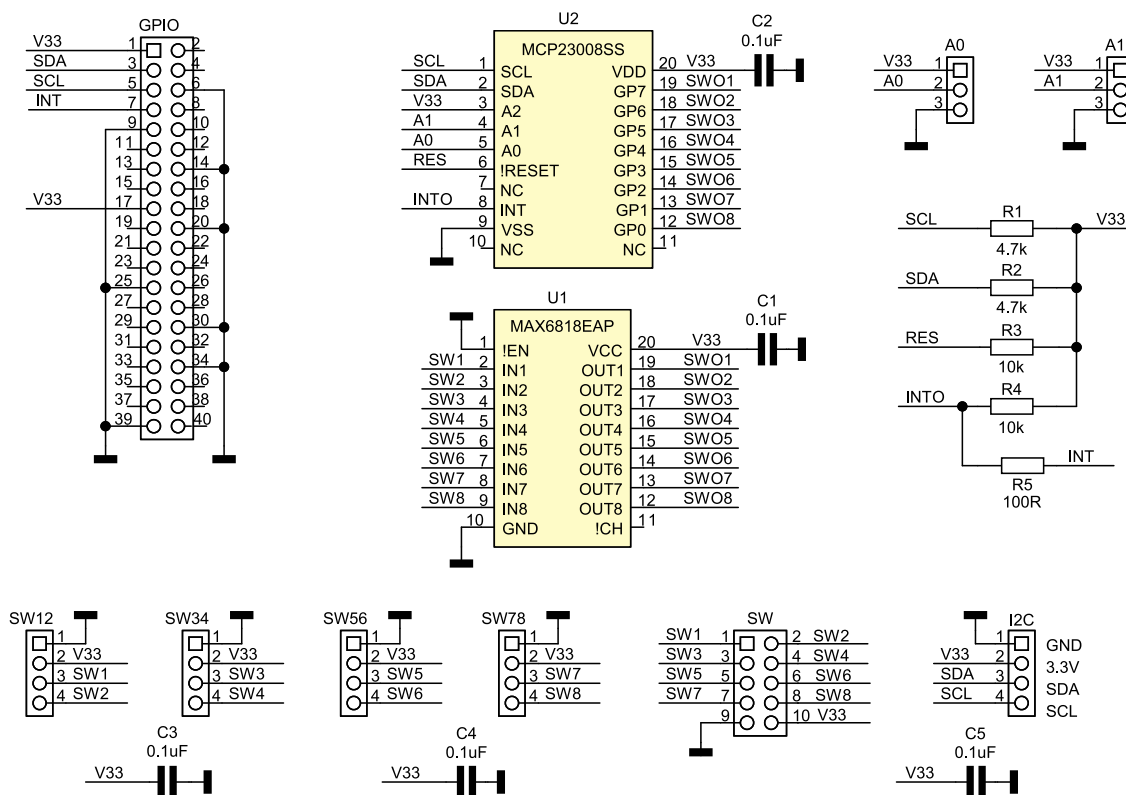
Moduł bazuje na układzie MAX6818, którego budowę wewnętrzną jednego z ośmiu kanałów pokazano na **rysunku 1**. Wejście zawiera elementy zabezpieczające, a dzięki zastosowanej logice wyjście aktywowane jest po ustabilizowaniu się stanu wejścia na czas dłuższy niż 40 ms.

Budowa i działanie

Schemat minimodułu pokazano na **rysunku 2**. Dzięki zastosowaniu specjalizowanego układu aplikacja jest niewiarygodnie



Rysunek 1. Budowa wewnętrzna MAX6818 (za notą Maxim)



Rysunek 2. Schemat modułu

Wykaz elementów, kupuj na stronie sklep.avt.pl (Warszawa, ul. Leszczynowa 11, tel. +48222578451, e-mail: handlowy@avt.pl)

Półprzewodniki:

U1: MAX6818EAP (SSOP20)

U2: MCP23008SS (SSOP20)

Rezystory:

R1, R2: 4,7 kΩ 5% (SMD0603)

R3, R4: 10 kΩ 5% (SMD0603)

R5: 100 Ω 5% (SMD0603)

Kondensatory:

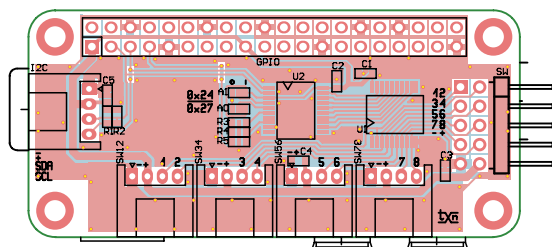
C1, C2, C3, C4, C5: 0,1 μF 16 V (SMD0603)

Pozostałe:

GPIO: złącze IDC40 żeńskie

SW12, SW34, SW56, SW78, I²C: złącze Grove kątowe (110990037)

SW: złącze IDC10 kątowe



Rysunek 3. Schemat płytki PCB

prosta i składa się z U1 typu MAX6818 oraz ekspandera GPIO magistrali I²C U2 typu MCP23008. Prostota okupiona jest niestety nieco wyższą ceną U1 w porównaniu do kilkunastu elementów niezbędnych dla realizacji dyskretnych np. na bramkach HC14.

Współpracujące elementy stykowe podłączone są do złączy SWxx zgodnych ze standardem Grove. Każde złącze umożliwia podłączenie pary styków. Element stykowy podłączony jest do wejść SWx i masy układu. Opcjonalnie płytka umożliwia wlotowanie złączy SW typu IDC10, na które wyprowadzono wszystkie sygnały SW1...8 i zasilanie 3,3 V. Wyjście SW0x układu U1 aktywowanie jest po ustabilizowaniu się stanu wejścia SWx na czas dłuższy niż 40 ms.

Wolne od zakłóceń sygnały SW01...8 podłączone są do ekspandera GPIO magistrali I²C. Dla uproszczenia obwodu drukowanego odwrócono podłączenia styków SW1...8 z układem U2. MCP23008 ma możliwość generowania przerwania po zmianie stanu wejścia,

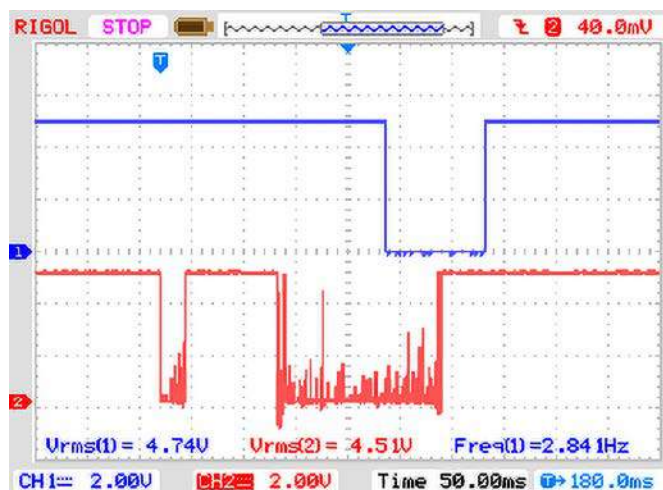
po wlotowaniu rezystora R5, przerwanie doprowadzone jest do GPIO4 Raspberry.

Moduł zasilany jest napięciem 3,3 V z Raspberry, pobór prądu nie przekracza kilku mA. Zwory adresowe A0, A1 umożliwiają ustawienie czterech adresów modułu (0x24...0x27). Dodatkowo wprowadzono magistralę I²C (złącze I²C, 3,3 V)

Montaż i uruchomienie

Moduł zmontowany jest na miniaturowej płytce dwustronnej, której schemat został pokazany na **rysunku 3**. Moduł nie wymaga uruchamiania, po ustaleniu adresu zworami A0, A1, podłączeniu do Raspberry Pi oraz elementów stykowych jest gotowy do pracy.

Dla szybkiego sprawdzenia można zastosować pakiet i2c_tools. Domyślnie po resecie układ U2 skonfigurowany jest jako



Rysunek 4. Przebiegi w układzie eliminatora drgań styków

brama wejściowa, której stan wyprowadzeń można odczytać poleceniem:

```
i2cget -y 1 0x24 0x09
```

Przy pomocy polecenia:

```
i2cset -y 1 0x24 0x01 0xFF
```

można zanegować logicznie stan wejść SW1...8. Szczegółowa konfiguracja opisana jest w notce katalogowej MCP23008.

Przykładowy oscylogram pokazujący skuteczność układu został pokazany na **rysunku 4**, (2-IN, 1-OUT). Jeżeli wszystko działa poprawnie, moduł można zastosować we własnych aplikacjach. Będzie przydatny w projektach automatyki domowej czy robotyce.

Adam Tatuś, EP

REKLAMA



KOMPUTERY RASPBERRY PI I MODUŁY ARDUINO



<http://sklep.avt.pl>

**Podstawowe parametry:**

- układ pracuje poprawnie przy napięciach zasilania 3,3...18 V przy obciążeniu do 3 A,
- napięcia wejściowe ze złącz INA i INB kluczowane są na wyjście OUT, przy czym aktywny jest zasilacz o wyższym napięciu,
- spadek napięcia wynosi mniej niż 50 mV.

* **Uwaga!** Elektroniczne zestawy do samodzielnego montażu. Wymagana umiejętność lutowania! Podstawową wersją zestawu jest wersja [B] nazywana potocznie KIT-em (z ang. zestaw). Zestaw w wersji [B] zawiera elementy elektroniczne (w tym [UK] – jeśli występuje w projekcie), które należy samodzielnie wstawić w dotychczasową płytkę drukowaną (PCB). Wykaz elementów znajduje się w dokumentacji, która jest podlinkowana w opisie kitu. Mając na uwadze różne potrzeby naszych klientów, oferujemy dodatkowe wersje:

- wersja [C] – zmontowany, uruchomiony i przetestowany zestaw [B] (elementy wlutowane w płytkę PCB),
 - wersja [A] – płytka drukowana bez elementów i dokumentacji.
- Kity, w których występuje układ scalony wymagający zaprogramowania, mają następujące dodatkowe wersje:
- wersja [A+] – płytka drukowana [A] + zaprogramowany układ [UK] i dokumentacja,
 - wersja [UK] – zaprogramowany układ.

Dodatkowe materiały do pobrania ze strony www.ulubionykiosk.pl/media

- Sterownik mikrosilnika krokowego do Pi Pico (EP 12/2022)
- Radiomodem ISM do Raspberry Pi Zero (EP 11/2022)
- Moduł LoRa do RPi Pico (EP 9/2022)
- Moduł z wyświetlaczami numitron (EP 8/2022)
- Interfejs aparatury kontrolnej i sygnalizacyjnej standardu M22 do Raspberry Pi (EP 7/2022)
- Sterownik mikrosilników prądu stałego do Rpi Pico (EP 7/2022)

Nie każdy zestaw AVT występuje we wszystkich wersjach! Każda wersja ma załączony ten sam plik PDF! Podczas składania zamówienia upewnij się, którą wersję zamawiasz! <http://sklep.avt.pl>

W przypadku braku dostępności na stronie sklepu osoby zainteresowane zakupem płytek drukowanych (PCB) prosimy o kontakt via e-mail: kity@avt.pl.

Moduł redundancji zasilania do komputerów SBC

Zapewnienie stabilnego zasilania jest podstawowym wymogiem koniecznym przy projektowaniu urządzeń elektronicznych. Problem jest szczególnie istotny dla komputerów jednopłytkowych SBC, które w dzisiejszych czasach realizują całkiem poważne zadania, a ewentualny zanik zasilania może uszkodzić system plików, naruszyć integralność danych, nie wspominając o oczywistej przerwie w pracy urządzenia. Dwa źródła zasilania współpracujące z modułem redundancji pozwalają ograniczyć skutki wielu awarii zasilania.

Najprostsze przełączanie źródeł w układach zasilania bezprzerwowego bazuje na diodach Schottkiego. Pomimo niewielkiego spadku napięcia w kierunku przewodzenia w zastosowaniach, w których zależy nam na najwyższej możliwej sprawności szczególnie w przypadku systemów zasilanych niskim napięciem 3,3...5 V (Raspberry Pi) oraz maksymalnym wykorzystaniu energii, konieczne jest bardziej złożone rozwiązanie.

W zaprezentowanym rozwiązaniu zastosowano układ kontrolera diody idealnej typu LTC4373 (Analog Devices), którego schemat wewnętrzny pokazano na **rysunku 1**. Gwarantuje on spełnienie wszystkich wymagań stawianych opisanej aplikacji.

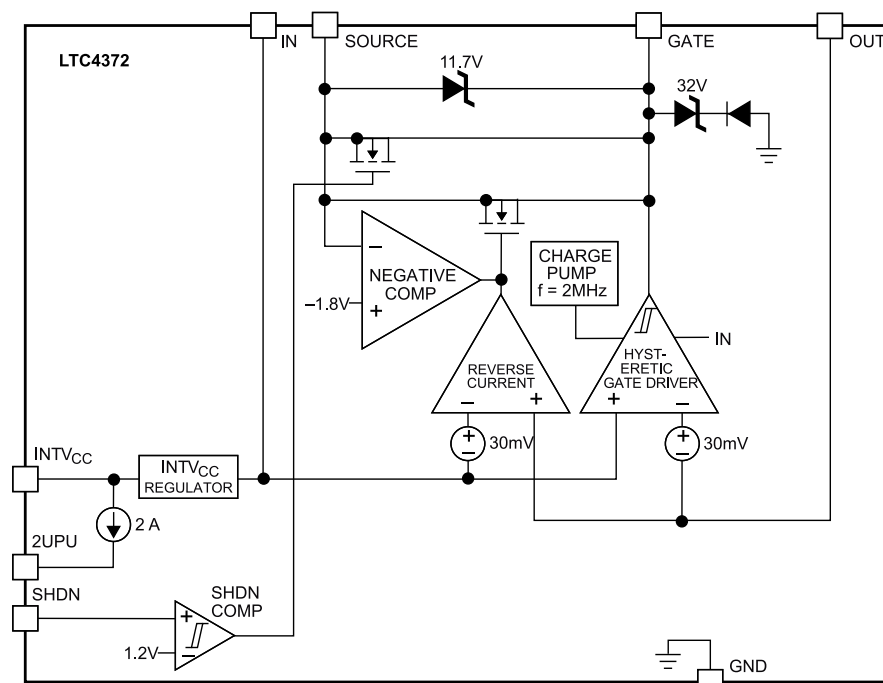
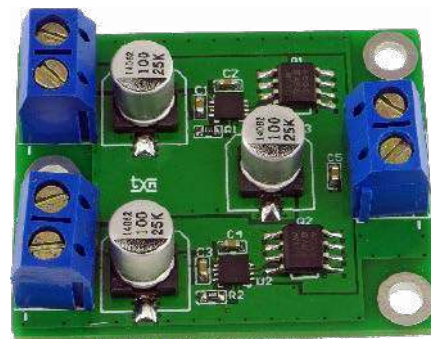
Budowa i działanie

Schemat modułu redundancji zasilania został pokazany na **rysunku 2**. Pracuje poprawnie przy napięciach zasilania 3,3...18 V przy obciążeniu do 3 A, co w większości wypadków wystarcza do zasilania klasycznych komputerów SBC. Składa się z dwóch identycznych bloków kontrolera diody idealnej U1/Q1 i U2/Q2. LTC4373 odpowiada za detekcję obecności napięcia wejściowego UV, wbudowany komparator steruje tranzystorem MOSFET zapewniając niski spadek napięcia

w kierunku przewodzenia oraz odcięcie i minimalny prąd wsteczny, gdy napięcie spadnie poniżej 1,191 V. Wbudowana pompa ładunkowa oraz driver zapewnia niezawodne

sterowanie klucza Q1 w szerokim zakresie napięcia wejściowego i obciążenia.

Napięcia wejściowe z zasilaczy doprowadzone są do złącz INA, INB, gdzie w zależności



Rysunek 1. Struktura wewnętrzna LTC4373 (za notą Analog Devices)

Wykaz elementów, kupuj na stronie sklep.avt.pl (Warszawa, ul. Leszczyńska 11, tel. +48222578451, e-mail: handlowy@avt.pl)

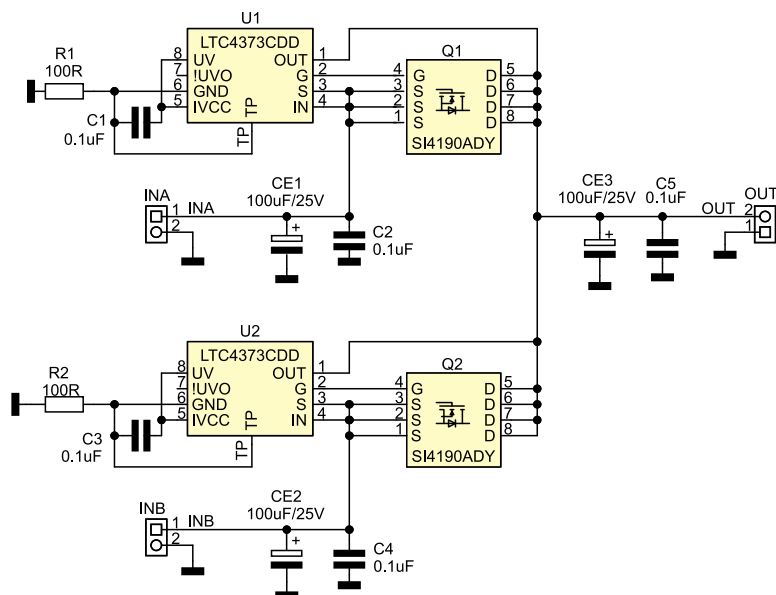
Rezystory:
R1, R2: 100 Ω (SMD0603)

Kondensatory:
C1, C2, C3, C4, C5: 0,1 μF ceramiczny 50 V (SMD0603)

CE1, CE2, CE3: 100 μF/25 V elektrolityczny low ESR 8 mm

Półprzewodniki:
Q1, Q2: SI4190ADY tranzystor MOSFET (SO8)
U1, U2: LTC4373CDD (DFN8)

Pozostałe:
INA, INB, OUT: złącze śrubowe DG126-5.0-2

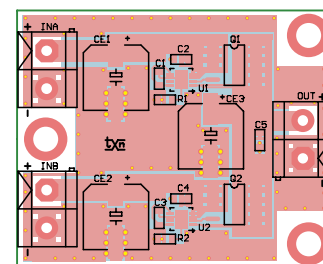


Rysunek 2. Schemat modułu

od ich poprawności kluczowane są na wyjście OUT (aktywny jest zasilacz o wyższym napięciu). Zanik lub zwarcie jednego z napięć zasilania INA/B powoduje bezprzerwowe przełączenia na sprawne źródło.

Montaż i uruchomienie

Układ zmontowany jest na niewielkiej płytce drukowanej, której schemat został pokazany na **rysunku 3**. Montaż nie wymaga opisu, należy tylko poprawnie przylutować



Rysunek 3. Schemat płytki PCB

pady termiczne układów. Moduł nie wymaga uruchamiania, warto jednak przy pomocy dwóch regulowanych zasilaczy i sztucznego obciążenia sprawdzić poprawność przełączania oraz spadek napięcia na kluczach w kierunku przewodzenia. W modelu spadek napięcia wynosił mniej niż 50 mV przy napięciach zasilania z zakresu 3,3...18 V i prądzie obciążenia 3 A. W zależności od doboru tranzystorów Q1, Q2 i napięć pracy kondensatorów układ można przystosować do pracy przy wyższym napięciu np. 24...36 V i z nieco większymi prądami obciążenia.

Adam Tatuś, EP

REKLAMA

Sięgnij po archiwalne wydania ELEKTRONIKI PRAKTYCZNEJ

Przesyłka
GRATIS

Zamów wygodnie na
www.UlubionyKiosk.pl



Historia rozwoju elektroniki implantowalnej

Niezależnie od tego, jakiej dziedzinie się przyjrzymy, zauważymy tę samą zasadę – postęp jest równoznaczny z przekraczaniem kolejnych granic, a nierzadko nawet łamaniem odwiecznych konwencji. W medycynie, a zwłaszcza technologii medycznej, widać to szczególnie wyraźnie, co pokażemy na przykładzie kilku wybranych grup implantowalnych urządzeń elektronicznych.

Kiedy w grudniu 1809 r. Ephraim McDowell postanowił otworzyć brzuch kobiety cierpiącej na gigantycznych rozmiarów, ponad 5-kilogramową torbiel jajnika, tłum zebrany pod domem lekarza o mało nie zlinczował go za – zdaniem zebranej ludności – próbę morderstwa. Na szczęście (dla wszystkich) operacja udała się – McDowell zapisał się w historii jako ojciec chirurgii ogólnej, a jego następcy z biegiem lat coraz rzadziej byli postrzegani jako bezduszni oprawcy, a coraz częściej – jako wybawcy ostatniej szansy.

Nietykalne pozostawało jednak serce – tego organu bali się



Fotografia 1. Ludwig Rehn – pionier kardiochirurgii, pierwszy lekarz, który z sukcesem zoperował serce człowieka (<https://t.ly/zCW5>)

dotknąć skalpelem nawet najlepsi operatorzy ówczesnej medycyny i na wyrost wieszczili, iż pierwszy, który ośmieli się zbeczczyć święty organ, zostanie z miejsca odrzucony przez całe środowisko lekarskie. A jednak, gdy jesienią 1896 roku do szpitala we Frankfurcie nad Menem trafił pchnięty nożem, 22-letni ogrodnik, chirurg Ludwig Rehn (**fotografia 1**) dokonał pierwszej udanej próby zszycia rany kłutej serca, ratując w ten sposób wykrwawiającego się pacjenta z objętej śmierci. Zamiast społecznej kary, na chirurga spłynęła sława pioniera i protoplasty kardiochirurgii, choć przed Rehmem byli już nieliczni śmiałkowie, którzy (bez sukcesu) próbowali dokonać tego samego.

Chirurgia rozwijała się coraz intensywniej, wtedy jednak jeszcze nikt (być może z wyjątkiem nielicznych wizjonerów-marzycieli) nie myślał poważnie o zostosowaniu jej do wszczepiania w ludzkie ciało jakichkolwiek bardziej złożonych urządzeń, mających na celu wspomaganie, regulację funkcji lub nawet zastępowanie narządów dotkniętych chorobą. Zamiast tego chirurdzy skupiali się wyłącznie na usuwaniu lub naprawianiu patologii anatomicznych, choć równie często wykorzystywali otwarcie jam ciała w celach diagnostycznych – ujrzenie wnętrza ciała pacjenta było jedyną metodą „medycznego rekonesansu” przed erą współczesnych technik obrazowania.

Wreszcie jednak, w 1958 roku, pod skórę klatki piersiowej pacjenta ze Szwecji, Arne Larssona (**fotografia 2**) trafił pierwszy w historii rozrusznik serca zastosowany u człowieka, co zapoczątkowało całkowicie nową erę w leczeniu najbardziej wymagających przypadków, dotąd skazanych na niechybną śmierć. Początki historii elektroniki implantowalnej są tożsame właśnie z tym wydarzeniem, choć w rzeczywistości powstanie pierwszego na świecie aktywnego implantu poprzedzało kilka innych wynalazków i pewne tragiczne wydarzenia, będące impulsem do szybkiego wdrożenia nowatorskiej technologii.



Fotografia 2. Arne Larsson – pierwszy pacjent, który otrzymał implantowalny rozrusznik serca (i dwadzieścia pięć, ulepszonych konstrukcji w ciągu pozostałych 43 lat swojego życia).
Źródło: <https://t.ly/XAZH>

Na początku był... metronom

Pierwsze szeroko stosowane klinicznie rozruszniki wykorzystywane przy ciężkich zaburzeniach rytmu, zwłaszcza u pacjentów z tzw. blokiem serca, opracował Paul Zoll, bostoński kardiolog zainspirowany wcześniejszymi urządzeniami badanymi przez inne zespoły naukowców. PM-65 – bo tak brzmiała handlowa nazwa modelu – generował impulsy o napięciu 50...150 V, które były następnie dostarczane do serca przez niewielkie, metalowe elektrody, umieszczone na skórze pacjenta, a dodatkowo wyposażony został we własny elektrokardiograf. Zdecydowanie nie była to konstrukcja idealna – pomijając duże rozmiary i zasilanie sieciowe (które ograniczało zasięg mobilności systemu do długości kabla zasilającego), bodaj najistotniejszym problemem zauważonym podczas użytkowania rozrusznika okazały oparzenia skóry, spowodowane bolesnymi impulsami wysokiego napięcia. Co jednak najważniejsze – rozrusznik działał, czego dowodem byli kolejni uratowani pacjenci (w tym widoczny na **fotografii 3**



Fotografia 3. 76-letni Pincus Shapiro (po prawej) – pierwszy pacjent, który otrzymał 3-miesięczną terapię rozrusznikiem zewnętrznym PM-65 (dolny, większy moduł generatora impulsów stymulujących był wyposażony w dodatkową „nadstawkę” EKG z niewielką lampą oscyloskopową, widoczną na górze urządzenia). W środku – żona pacjenta, po lewej – Seymour Furman, kardiolog (<https://t.ly/9zB0>)



Fotografia 4. Earl Bakken w swoim „garażowym” laboratorium (rok 1955). Źródło: <https://t.ly/S2CS>

Pincus Shapiro – 76-letni przedsiębiorca, którego serce, dzięki trzy-miesięcznej terapii zewnętrzną elektrostymulacją, udało się wrócić do „stanu używalności”).

Podobnie jak każde inne urządzenie nie wyposażone w rezerwowe źródło energii, rozrusznik PM-65 działał tylko tak długo, jak pozwalała na to dostępność napięcia sieciowego. 31. października 1957 r. potężna awaria miejskiej sieci energetycznej w Minneapolis wyłączyła sprzęt obecny w salach pacjentów – niepodtrzymywany przez generatory rezerwowe, obsługujące jedynie sale operacyjne i pooperacyjne. Trzygodzinna przerwa w dostawie energii spowodowała śmierć dziecka, którego serce było podtrzymywane przez jeden z rozruszników PM-65. To tragiczne wydarzenie skłoniło jednego z pionierów kardiologii – Clarence’a Waltona Lillehei’a – do skontaktowania się z inżynierem Earlem Bakkenem (**fotografia 4**), który 8 lat wcześniej założył w niewielkim garażu firmę... Medtronic (**fotografia 5**) – początkowo stanowiącą zaplecze techniczno-serwisowe dla pobliskiego szpitala. Obydwaj panowie postanowili nie dopuścić do kolejnych tragedii i uniezależnić działanie rozrusznika od zaniku napięcia sieciowego. Bakken przypomniał sobie, że rok wcześniej w czasopiśmie Popular Electronics widział prosty projekt elektronicznego metronomu, oparte go na oscylatorze blokującym, zbudowanym – co niezwykle ważne – na tranzystorach, a nie na prądożernych lampach, będących podstawowym budulcem rozruszników pierwszej generacji. Inżynier nieznacznie zmodyfikował układ, zamknął go w 10-centymetrowym, aluminiowym pudełku i wyposażył w banalnie prosty panel sterujący (**fotografia 6**). Po krótkich testach, wykonanych na psie w szpitalnym laboratorium, urządzenie trafiło z powrotem do Lillehei’a, który – ku zaskoczeniu Bakkena – już następnego dnia



Fotografia 5. Domowy garaż, który posłużył za pierwszą siedzibę założonej w 1949 roku firmy Medtronic (<https://t.ly/3RUZ>)

zastosował je do podtrzymywania życia dziewczynki przebywającej na oddziale pooperacyjnym. Jak widać, sześć dekad temu ścieżka od prototypu do wdrożenia technologii medycznej była nieporównanie krótsza, niż dziś... Firma Medtronic rozrosła się natomiast do postaci jednego z największych koncernów na rynku technologii medycznych, stając się potentatem w zakresie zaawansowanej diagnostyki i terapii w różnych obszarach medycyny. Dziś koncern zatrudnia 90 tysięcy pracowników i uzyskuje roczny przychód rzędu 31 miliardów dolarów (dane za rok 2022).

Dwa tranzystory i pasta do butów

Zastosowanie tranzystorów oraz zasilania baterijnego umożliwiło wielokrotną miniaturyzację stymulatorów względem konstrukcji „wózkowych”, zaś przejście ze stymulacji przezskórnej na nasierdziową (tj. z elektrodami mocowanymi bezpośrednio na sercu) pozwoliło zniwelować problem oparzeń i bólu. Niestety w branży medycznej każda decyzja ma swoje konsekwencje – w tym przypadku wyprowadzenie elektrod z klatki piersiowej poza obszar ciała pacjenta wiązało się z ryzykiem infekcji w miejscach wyjścia kabli z otworów w skórze. Z tego powodu – a także z uwagi na ergonomię – zaczęto poszukiwania jeszcze lepszej metody stabilizacji rytmu serca. Stąd już prosta droga prowadziła do kolejnego przełomu, którego dokonali: chirurg Ake Senning oraz wynalazca Rune Elmqvist. Ósmego października 1958 roku odbyła się pierwsza w historii implantacja wszczepialnego rozrusznika (fotografia 7).

Konstrukcja bazowała na dwóch tranzystorach PNP (rysunek 1) – pierwszy z nich pracował w roli generatora blokującego, drugi – stopnia wyjściowego, sprzężonego zmiennoprądowo z elektrodami nasierdziowymi. Co ciekawe, układ był wyposażony w bardzo prosty obwód... ładowania indukcyjnego – cewka widoczna po lewej stronie schematu okalała elektronikę i współpracowała z jednopółprzewodnikowym



Fotografia 6. Pierwszy, zewnętrzny rozrusznik baterijny Bakkena (<https://t.ly/K-8t>)



Fotografia 7. Replika pierwszego rozrusznika wszczepialnego (Rune Elmqvist, Ake Senning).
Źródło: <https://t.ly/L0o->



Fotografia 8. Arne Larsson z żoną, Else-Marie Larsson. Naciskając na chirurga, by ten ratował jej męża, kobieta prawdopodobnie nie spodziewała się, że jej determinacja w walce o życie ukochanego przyspieszy rozwój elektroniki medycznej (<https://t.ly/9cUZQ>)

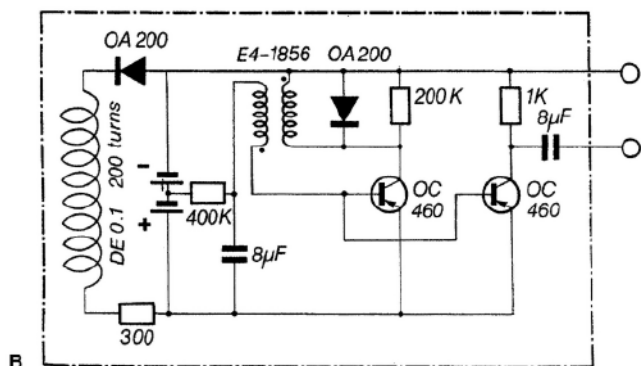
prostownikiem, umożliwiając doładowywanie dwóch połączonych szeregowo akumulatorów Ni-Cd (o pojemności 60 mAh każdy) za pomocą zewnętrznego pola magnetycznego o częstotliwości 150 kHz. Urządzenie musiało być doładowywane co 12 godzin za pomocą 25-centymetrowej cewki, mocowanej na klatce piersiowej pacjenta i zasilanej generatorem lampowym.

Całość zmontowanej metodą na pająkę elektroniki zalano biogodną żywicą epoksydową (Araldite), w formie wykonanej z... puszek po pastcie do czyszczenia butów marki Kiwi. Średnica tak wytworzonej obudowy wynosiła około 55 mm, a grubość 16 mm. Urządzenie generowało impulsy o amplitudzie rzędu 2 V i fizjologicznej częstotliwości 70...80 uderzeń/min. Pierwsza wersja stymulatora, wszczepiona z zastosowaniem torakotomii lewostronnej (nacięcia w ścianie klatki piersiowej), pracowała zaledwie 8 godzin, zaś kolejna już cały tydzień. Za rekrutację pierwszego pacjenta odpowiadała jego żona (fotografia 8), która – dowiedziawszy się o trwających badaniach nad implementacją rozrusznika – postanowiła przekonać chirurga do wszczepienia urządzenia jej mężowi. Nie zraziła jej nawet odpowiedź lekarza, który poinformował zdesperowaną kobietę, że jego niewielki zespół nie dysponuje jeszcze aparatem gotowym do implantacji. Else Marie Larsson skwitowała brak finalnej wersji rozrusznika bardzo krótko: „w więc zróbcie taki!”. Trudno dziwić się jej desperacji, gdyż ukochany mąż każdego dnia zmagał się z 20...30 omdleniami pochodzenia kardiogenego. Po zabiegu mężczyzna został pierwszym na świecie człowiekiem, który powrócił do normalnej aktywności, będąc całkowicie zależnym od stymulatora i przeżył jeszcze 43 lata, podczas których urządzenie wymieniano jeszcze 25 razy. Pacjent zmarł niemal 1,5 roku po Senningu, który – nawiasem mówiąc – był jego rówieśnikiem.

Warto dodać, że projekt pierwszego rozrusznika wszczepialnego od razu spotkał się ze znacznym zainteresowaniem środowiska techniki medycznej. Firma Elema Schonander, zatrudniająca Rune Elmqvista, po kilku akwizycjach i wynikających stąd zmianach nazwy (Siemes-Elema, Siemens-Pacesetter) została wykupiona przez St. Jude Medical w 1994 roku, po czym... St. Jude sam został w 2017 roku wchłonięty przez holding Abbott, będący dziś jednym z najsilniejszych graczy na rynku zaawansowanych technologii z zakresu kardiologii, neuromodulacji, diagnostyki molekularnej i diabetologicznej.

Efekt motyla, czyli jak pomylenie kodu paskowego rezystorów wpłynęło na rozwój kardiologii

Jak to zwykle bywa w przypadku przełomowych odkryć i wynalazków, tak i rozrusznik serca został opracowany równolegle



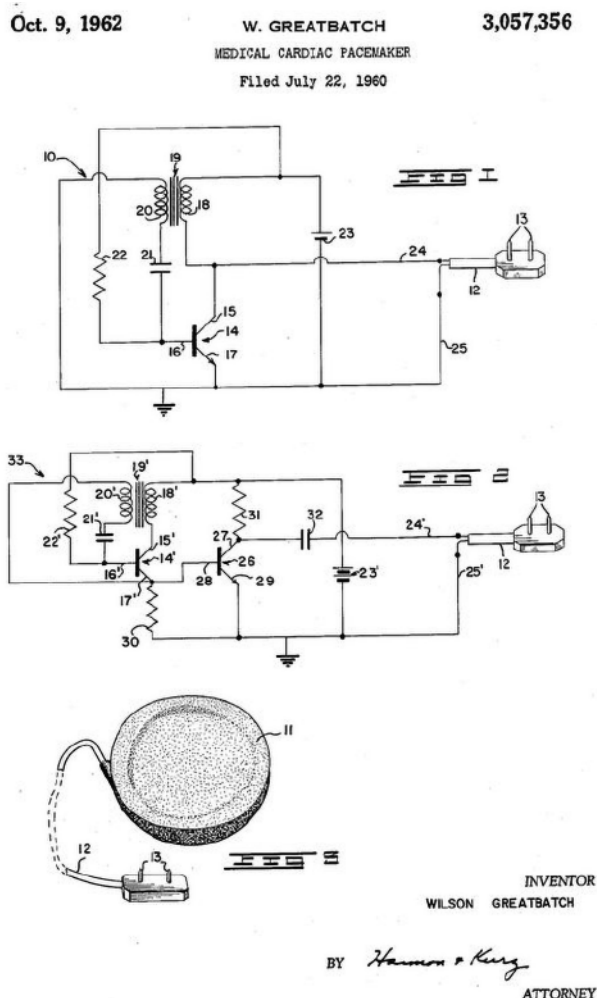
Rysunek 1. Schemat ideowy rozrusznika Elmqvista (<https://t.ly/q0lnx>)

w innym zespole. Tym razem jednak innowacja powstała poniekąd przez przypadek – inżynier Wilson Greatbatch (**fotografia 9**), pracując nad nowym typem generatora elektronicznego, błędnie odczytał wartość rezystora, który zamierzał włączyć w obwodzie bazy jednego z tranzystorów. Mylnie identyfikując kod paskowy elementu, zamiast planowanego podzespołu o rezystancji 10 k Ω , z pudełka wziął rezystor 1 M Ω . Okazało się, że taka zmiana pozwoliła zbudować generator, który 1,8-milisekundowe impulsy wytwarzał z częstotliwością około 1 Hz, a co więcej – przez nieaktywną część okresu drgań układ pobierał pomijalny wręcz prąd zasilania.

Greatbatch błyskawicznie skojarzył taką charakterystykę pracy z generatorem idealnym do użycia w roli stymulatora wszczepialnego, którego budowę rozważał już kilka lat wcześniej, zaś całą sytuację związaną z pomyłką wyjaśnił (jako człowiek głęboko wierzący) działaniem siły wyższej. Rozpoczął poszukiwania chirurga chętnego do współpracy w celu rozwinięcia tej koncepcji. Dr William Chardack ze szpitala w Buffalo zgodził się wejść w kooperację z inżynierem, wskutek czego – w maju 1957 – wystartowały testy na zwierzętach. W 1960 roku Greatbatch zgłosił konstrukcję do amerykańskiego urzędu patentowego.



Fotografia 9. Wilson Greatbatch prezentujący skonstruowany przez siebie rozrusznik (<https://t.ly/Ljss>)



Rysunek 2. Oryginalne rysunki z opisu patentowego US3057356A autorstwa W. Greatbatcha (https://t.ly/R_bk)

Projekt rozrusznika występował w dwóch wersjach – jednorozrusznikowej oraz dwutorowej, w której (podobnie, jak w konstrukcji Elmqvista) drugi tranzystor pracował w roli stopnia wyjściowego, zwiększającego moc impulsu dostarczanego do mięśnia sercowego (**rysunek 2**). Pierwsza implantacja omawianej konstrukcji miała miejsce w 1960 roku, a niedługo później Greatbatch wraz ze współtwórcą rozrusznika – drem Williamem Chardackiem – sprzedał wyłączną licencję na konstrukcję urządzenia firmie Medtronic. Dziś głównym spadkobiercą dorobku inżyniera, który rozstał się później z korporacją z Minnesoty, jest firma Integer Holdings Corporation – łączny dorobek inżyniera przekracza 100 patentów i zgłoszeń patentowych, zaś Integer pozostał w obszarze elektroniki implantowalnej, dla której dziś oferuje szerokie portfolio komponentów i technologii.

Psy pod napięciem, czyli o trudnej drodze do wdrożenia ICD

Michel Mirowski (**fotografia 10**) był lekarzem polskiego pochodzenia – wykształcony w Lyonie, odbył rezydenturę w Izraelu oraz staże w słynnym szpitalu Johna Hopkinsa w Baltimore, a także w Instytucie Kardiologii w Mexico City. W 1967 roku dowiedział się o śmierci swojego przyjaciela, spowodowanej nawracającym częstoskurczem komorowym. Dla Mirowskiego, dotkniętego już wieloma tragicznymi wydarzeniami (podczas II Wojny Światowej naziści w ramach Holokaustu wymordowali całą jego rodzinę, co zmusiło go do emigracji), ta strata stała się impulsem do rozwoju całkowicie nowego typu urządzenia implantowalnego. Wszczepienia pierwszego kardiovertera-defibrylatora (ICD) dokonał dr Levi Watkins w 1980 roku u 57-letniej pacjentki z wieloma epizodami częstoskurczu komorowego w historii leczenia.

W przeciwieństwie do rozruszników, początki historii rozwoju ICD były jednak wielokrotnie dłuższe i zdecydowanie nie usłane różami. Prace Mirowskiego zaczęły się w 1968 roku przy udziale kardiologa Mortona Mowera² i pomimo niebywalej innowacyjności, spotkały się z silnym oporem środowiska medycznego. Po niemal 100 latach od słynnej operacji Ephraima McDowella – jak na złość – znów wybawcy tysięcy ludzkich istnień zostali posądzeni o działanie nieetyczne. Po pierwsze – aby sprawdzić działanie wszczepionego urządzenia, trzeba było sztucznie wywołać epizod migotania komór (ciężkie zaburzenie rytmu serca), co samo w sobie stanowi zagrożenie dla pacjenta (pomimo ściśle kontrolowanych warunków i doskonałego zaplecza aparaturowego oraz farmakologicznego podczas zabiegu). Po drugie – w owych czasach detekcja częstoskurczu oraz migotania komór była znacznie trudniejsza, niż dziś, w dobie niskomocowych procesorów DSP. Po trzecie – przerwanie epizodu tachykardii wydawało się mało realne przy użyciu niewielkiego, zasilanego baterijnie urządzenia. Po czwarte – idea wszczepienia wysokonapięciowego generatora pod skórę pacjenta wydawała się podówczas bardzo kontrowersyjna, pomimo oczywistych zalet, jakie technologia ICD mogła przynieść kardiologii.

Nieufność środowiska lekarskiego była na tyle zaawansowana, że nawet nagrania wideo z przebiegu testów na 25 psach, u których urządzenie było w stanie nie tylko na żądanie wywołać, ale przede wszystkim szybko i skutecznie zakończyć epizod zaburzeń rytmu serca, nie zdołały przekonać kolegów po fachu, a zespół spotkał się nawet z absurdalnymi zarzutami, że psy były... wytrenowanymi aktorami, zaś cały eksperyment został sfałszowany.

Opór świata medycyny sprawił, iż coraz silniejszy wówczas Medtronic – zatrudniający w owym czasie prawie 5000 pracowników i generujący roczny przychód na poziomie 250 mln. dolarów (!) – odrzucił koncepcję ICD, pomimo wstępnego zainteresowania wspomnianego już wcześniej Earla Bakkena. Zdanie decydentów było jednoznaczne – środowisko medyczne trzeba będzie przekonywać przez kolejne 20 lat do praktycznego stosowania ICD. I choć pomylili się oni zaledwie o kilka lat (FDA zaaprobowала użycie ICD już w 1985 roku, czyli 17 lat po rozpoczęciu prac), to i tak Mirowskiemu oraz

jego zespołowi należy się wielkie uznanie za doprowadzenie tak ryzykownego projektu do samego końca. Tym bardziej, że pierwsze badania kliniczne były obciążone sporą śmiertelnością (spowodowaną co prawda nie przez samo urządzenie, ale przez konieczność implantacji elektrod nasierdżiowych w procedurze torakotomii), co przy braku silnej determinacji zespołu zapewne doprowadziłoby do przerwania prac.

Pierwsze konstrukcje ICD „cierpiały” z powodu dość długiego czasu niezbędnego na ponowną aktywację po wyładowaniu, co sprawiło, że znaczne wysiłki zostały włożone w przyspieszenie ładowania urządzenia. Po pionierskim kardiowerterze-defibrylatorze o oznaczeniu AID stworzona została wersja AID-B (fotografie 11 i 12), częściowo rozwiązująca ten problem, choć – podobnie, jak jej poprzedniczka – także produkowana na indywidualne zamówienie. Warto bowiem dodać, że w owym czasie urządzenia nie były jeszcze wyposażone w funkcję programowania do określonych potrzeb danego pacjenta. Wbudowane baterie litowe zapewniały 3-letnią pracę urządzenia, w trakcie której mogło ono dostarczyć do 100 „strzałów” ratujących życie, przy czym każdy impuls miał energię rzędu 25...30 J.

Rozwoju ciąg dalszy – więcej, dłużej i lepiej

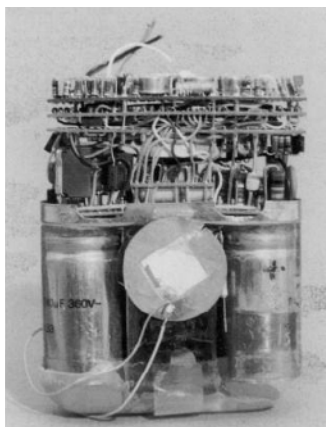
Dalsze, przełomowe pod względem funkcjonalnym, osiągnięcia w dziedzinie elektrostymulacji serca były przejawem rozwiązywania kolejnych problemów i niedogodności ówczesnych projektów, jednak przy wyraźnej kontynuacji wypracowanej dotąd ogólnej formy urządzeń. Powróćmy zatem



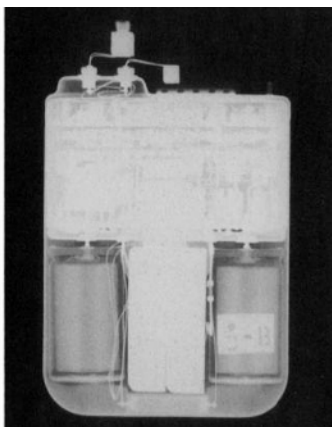
Fotografia 10 Michel Mirowski – twórca pierwszego wszczepialnego kardiowertera-defibrylatora (<https://t.ly/7nJO>)



Fotografia 11. Jeden z pierwszych kardiowerterów-defibrylatorów (ICD) z wczesnych lat 80. XX w. – model AID-B (<https://t.ly/5qlB>)

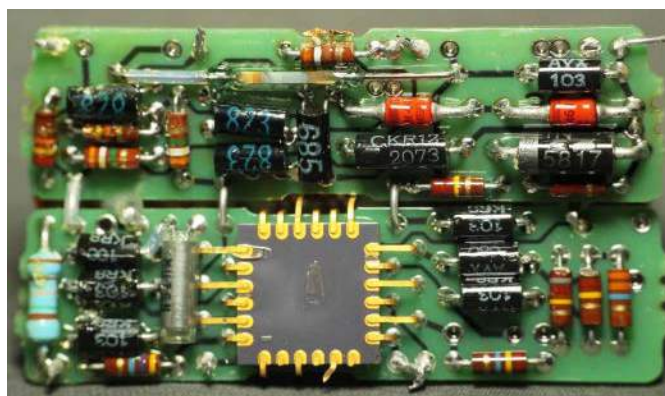


Fotografia 12. Wnętrze ICD typu AID-B. Widoczna elektronika urządzenia oraz dwa duże kondensatory 140 μ F/360V, współpracujące z generatorem impulsów HV (<https://t.ly/5qlB>)

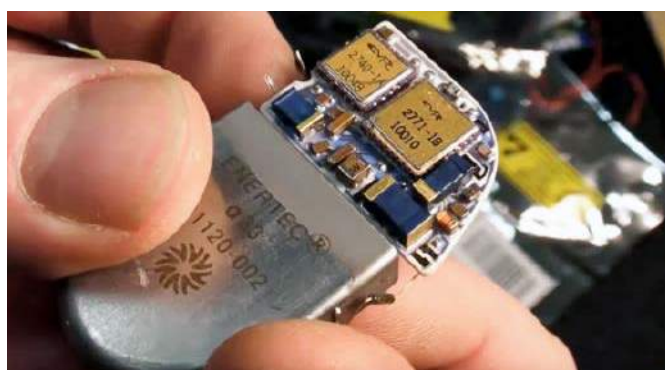


Fotografia 13. Rozwój rozruszników serca oraz ICD wiązał się z nieuchronną miniaturyzacją urządzeń we wszystkich trzech osiach, choć największa różnica między starszymi, a nowszymi konstrukcjami, jest widoczna na grubości obudowy (<https://t.ly/feng>)

jeszcze na chwilę do lat 60. XX w. W kolejnych wersjach rozruszników żywiczne obudowy zostały wyparte przez szczelne puszkę tytanowe (fotografia 13), zaś elektrody nasierdżiowe – wymagające implantacji za pomocą torakotomii – zostały w 1962 roku zastąpione przez łatwiejsze i mniej inwazyjne elektrody wprowadzane za pomocą procedur przeznaczeniowych pod wizyjną kontrolą z użyciem fluoroskopii (ciągłego obrazowania klatki piersiowej pacjenta za pomocą promieniowania rentgenowskiego), co pozwoliło na wykonywanie procedury implantacji lekarzom nie będącym chirurgami. Zmiana konstrukcji i sposobu mocowania elektrod była konieczna, gdyż twórcy kolejnych generacji urządzeń elektroterapeutycznych wciąż zmagali się po drodze z szeregiem problemów, niezwiązanych już bezpośrednio z samą elektroniką – elektrody nasierdżiowe szybko traciły swoją funkcjonalność przez pogarszające się z dnia na dzień parametry elektryczne interfejsu metal-tkanka, zaś przewody nie wytrzymywały kolejnych cykli obciążeń mechanicznych i pękały tak, jak każdy inny rodzaj kabla.



Fotografia 14. Elektronika rozrusznika wykonana w technologii THT (źródło: archiwum własne Autora)



Fotografia 15. Wnętrze rozrusznika po usunięciu tytanowej obudowy. Elektronika wykonana w technologii SMD (<https://t.ly/VZA->)



Fotografia 16. Rozrusznik Chardack-Greatbatch wykonany w obudowie żywicznej, w technologii pośredniej – „ścieżki” wykonano w postaci wyprofilowanych blaszek, łączących poszczególne komponenty (<https://t.ly/ykKh>)

Coraz rzadziej wewnątrz obudowy można było spotkać płataninę przewodów, która ustępowała miejsca płytkom drukowanym – początkowo w technologii przewlekanej (**fotografia 14**), choć tak szybko, jak było to możliwe, producenci przeszli na montaż powierzchniowy (**fotografia 15**). Co ciekawe, pomiędzy erą metody na pająka, a PCB, można było także spotkać konstrukcje pośrednie, czego dowodem jest artystyczna wręcz forma jednego z wczesnych modeli firmy Medtronic, pokazana na **fotografii 16**. Warto dodać, że ludzko podobne rozwiązanie jest stosowane do dziś, jednak nie do łączenia elektroniki wewnątrz obudowy, ale do wykonywania połączeń pomiędzy stykami złączy elektrod a portem wyprowadzonym z hermetycznej obudowy rozrusznika/ICD (**fotografia 17**).

W latach 70. XX wieku zbudowano i wszczepiono pierwszy rozrusznik programowalny za pomocą pola magnetycznego, ta sama dekada przyniosła także implantację pierwszego rozrusznika z zasilaniem nuklearnym (Medtronic ogłosił powstanie takiej konstrukcji dwa lata wcześniej). Radioizotopowe generatory termoelektryczne (**fotografia 18**), pomimo wielokrotnie dłuższego czasu eksploatacji w porównaniu do chemicznych źródeł energii, nie zostały jednak spopularyzowane z uwagi na szereg problemów natury formalnej – zastosowany w urządzeniach Pluton-238 był silnie toksyczny, co w okresie, gdy rozruszniki miały jeszcze nierazdo problemy ze szczelnością obudowy, mogło skończyć się tragicznie dla pacjenta. Poza tym radioaktywny charakter zasilania utrudniał

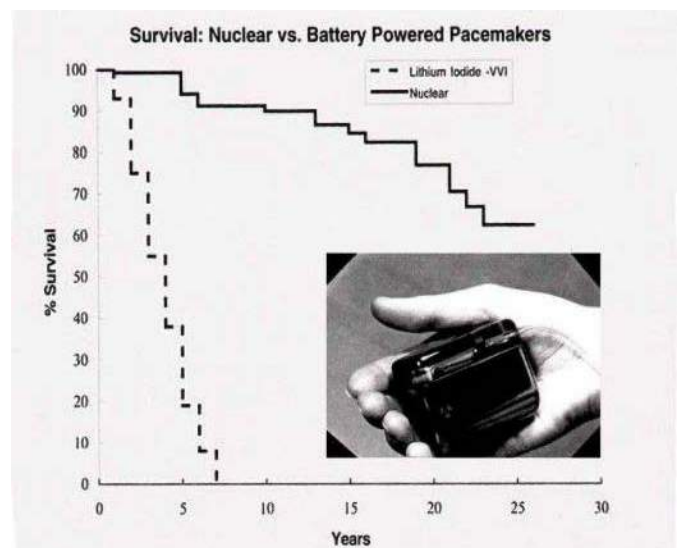


Fotografia 17. Blok gniazd do podłączenia elektrod we współczesnym urządzeniu do terapii CRT (<https://t.ly/xdwl>)



Fotografia 18. Wnętrze rozrusznika marki Medtronic z zasilaniem nuklearnym. Widoczny (zamknięty w osobnej obudowie) radioizotopowy generator termoelektryczny (https://t.ly/eX_5)

pacjentom swobodne podróżowanie, zaś na wszczepiających takie rozruszniki lekarzach prawo wymuszało konieczność usunięcia urządzeń po śmierci pacjentów oraz odpowiedniego zutylizowania szkodliwych odpadów. Nie zawsze było to jednak możliwe, gdyż... pacjenci umierali długo po przejściu ich lekarzy na emeryturę (choć



Rysunek 3. Porównanie żywotności baterii litowo-jodowej oraz źródła nuklearnego opartego na Plutonie-238 (<https://t.ly/jof6>)



Fotografia 19. Pierwszy rozrusznik dostosowujący częstotliwość generatora do aktywności ruchowej pacjenta – model Activitrac II 8412 firmy Medtronic. Uwagę zwracają niewielkie wymiary urządzenia, pomimo zaimplementowania w nim dodatkowej funkcjonalności (60×43×12 mm). Źródło: https://t.ly/_BeE



Fotografia 20. Porównanie wielkości klasycznego, współczesnego rozrusznika serca oraz ultranowoczesnego modelu Micra TPS marki Medtronic (<https://t.ly/jdyB>)

same rozruszniki działały nadal i być może pracują do dziś, gdyż żywotność zasilania nuklearnego ocenia się na ponad 80 lat! – rysunek 3).

Sytuację odmieniło wprowadzenie nowego typu baterii – najpierw cynkowo-rtęciowych, a potem litowo-jodowych (opracowanych przez Greatbatcha i wydłużających czas pracy urządzeń do około 10 lat). Koniec lat 70. przyniósł rozruszniki z komunikacją dwukierunkową, rozpoczęto też stosowanie stymulatorów dwukomorowych (1978). W 1985 r. pojawiły się pierwsze urządzenia z korekcją szybkości rytmu serca w zależności od aktywności fizycznej użytkownika, co umożliwiło dostrojenie rytmu serca do stale zmieniających się potrzeb fizjologicznych pacjenta. Rozrusznik Activitrac™ firmy Medtronic (fotografia 19) został bowiem wyposażony w piezoelektryczny czujnik drgań, który wykrywał zwiększoną aktywność pacjenta i na tej podstawie podkręcał częstotliwość generatora. Echo tej przełomowej technologii – udoskonalonej przez zastosowanie bardziej zaawansowanych algorytmów i zmianę z prostych czujników piezo na niskomocowe akcelerometry MEMS – można znaleźć, także we współczesnych konstrukcjach.

Co ciekawe, technologia elektrostymulacji rozwija się nadal, a najnowsze zdobycze technologii nie wiążą się tylko z wprowadzaniem innowacji w postaci funkcji telemetrycznych, ulepszonych algorytmów detekcji zaburzeń rytmu serca, czy też wbudowywaniem dodatkowych funkcji pomiarowych. I choć miniaturyzacja jest doprawdy spektakularna – co widać najlepiej na przykładzie miniaturowego rozrusznika Micra TPS, zaimplantowanego po raz pierwszy w 2013 roku (fotografia 20) – to całkiem niedawno byliśmy świadkami jeszcze jednego przełomu w tej dziedzinie medycyny. Dopiero na przełomie XX i XXI wieku powstała bowiem terapia resynchronizująca (CRT), w ramach której stymulacji podlegają nie tylko prawe jamy serca (przedsionek i komora), lecz także lewa komora. Dziś rozruszniki (CRT-P) oraz defibrylatory (CRT-D – fotografia 21) wyposażone w dodatkową elektrodę tworzą jedno z podstawowych narzędzi terapeutycznych w niewydolności serca, spowodowanej niektórymi zaburzeniami rytmu (powodującymi rozszynchronizowanie



Fotografia 21. Współczesne urządzenie do terapii resynchronizującej CRT-D – Rivacor 7 HF-T QP marki Biotronik (<https://t.ly/qPSj>)

skurczów obu komór serca). W tym przypadku jednak mamy do czynienia – od strony technicznej – z rozszerzonymi wersjami opisanych wcześniej urządzeń (rozrusznika oraz ICD), nie zaś z całkowicie nową technologią, budowaną od podstaw.

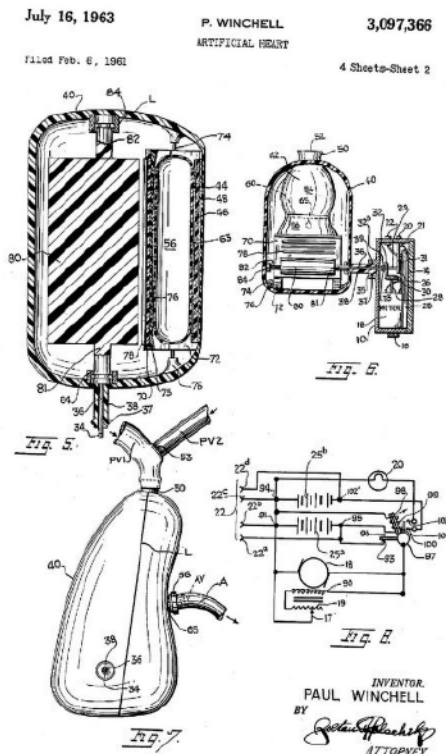
Stymulatory serca stworzyły podwaliny innych obszarów elektroniki implantowalnej. Doświadczenia zdobyte na przestrzeni kilku dekad pozwoliły opracować rozmaite rodzaje neurostymulatorów mózgu i rdzenia kręgowego, implantów regulujących pracę żołądka, pęcherza moczowego, a nawet baroreceptorów, odpowiedzialnych za autonomiczną regulację ciśnienia krwi. W każdym z wymienionych przypadków można już na pierwszy rzut oka zauważyć niebywałe podobieństwa konstrukcyjne – choć poszczególne implanty istotnie różnią się między sobą pod względem zastosowanych metod stymulacji, algorytmów sprzężenia zwrotnego, liczby elektrod, czy też dodatkowych funkcjonalności, jedno jest pewne – opisane powyżej, doniosłe wynalazki techniki medycznej uitorowały drogę do eksploatacji kolejnych branż medycyny.

Długotępe psy i patenty brzuchomówcy

Elektryczne wspomaganie układu krążenia rozwijało się niezwykle szybko już od połowy XX wieku, nadal jednak pozostawały problemy niezwiązane bezpośrednio z zaburzeniami rytmu serca. Wszelkie poważniejsze schorzenia natury hemodynamicznej – czyli, najprościej rzecz ujmując, wynikające z patologii w zakresie mechanicznej akcji serca, wciąż okazywały się śmiertelnym zagrożeniem dla pacjentów. Wraz z rozwojem kardiologii udawało się ratować pacjentów w coraz gorszym stanie zdrowia, choć w przypadkach schyłkowej niewydolności serca jedynym wyjściem wydawała się jego transplantacja. Kardiochirurgia poradziła sobie także z tym problemem, zaś rozwój immunosupresji oraz zaawansowanych technik diagnostycznych umożliwił lepsze prowadzenie najbardziej wymagających pacjentów. Rozwiązanie zagadnienia przeszczepu rzuciło jednak nowe światło na kolejne problemy – liczba serc kompatybilnych biologicznie z dawcami jest stosunkowo mała, a skrajnie niewydolni krążeniowo pacjenci nie mogą zbyt długo czekać na odpowiedniego dawcę. Medycyna znów zwróciła się w stronę inżynierii po pomoc – niezbędne okazało się bowiem opracowanie całkowicie sztucznego serca, które albo odciąży naturalne serce pacjenta, albo w pełni zastąpi je w jego klatce piersiowej.

Pierwsze próby zastąpienia biologicznego serca maszyną były wykonywane jeszcze przed spektakularnym rozwojem kardiochirurgii, a wynikały często z zupełnie innych pobudek. Pod koniec lat 30. XX wieku rosyjski lekarz Vladimir Demikhov (ten sam, którego nazwisko obiegło świat po udanym przeszczepie dodatkowej psiej głowy do ciała innego, żyjącego czworonoga) przeprowadził próby sztucznego serca na zwierzętach. Historia zna szereg tego typu eksperymentów, wśród których da się zauważyć wyraźną tendencję stopniowego wydłużania czasu przeżycia stworzeń poddawanych badaniom, liczonego w godzinach, dniach, a następnie tygodniach, a nawet miesiącach.

Pierwsze sztuczne serce przeznaczone do wszczepienia do ciała człowieka zostało opracowane przez duet w składzie: Paul Winchell (wzięty hollywoodzki aktor o polskich korzeniach, brzuchomówca i wynalazca, który dzięki medycznemu wykształceniu był w stanie opracować szereg patentów z pogranicza technologii i medycyny) oraz Henry Heimlich (amerykański torakochirurg i wynalazca, który do historii medycyny przeszedł głównie za sprawą manewru Heimlicha – do dziś pozostającego podstawową metodą pierwszej pomocy przy zadławieniach). Patent na rozwiązanie został przyznany w roku 1963 (rysunek 4). Zaledwie 3 lata później jeden z najsłynniejszych kardiochirurgów XX wieku – Adrian Kantrowitz, który zasłynął z pierwszego w USA przeszczepu serca (u niego odbywał jeden ze staży prof. Zbigniew Religa) – dokonał pierwszego wszczepienia urządzenia do wspomaganie lewej komory serca. Po kolejnych trzech latach pierwsze w historii, całkowicie sztuczne serce trafiło do klatki



Rysunek 4. Część rysunków z pierwszego w historii patentu dotyczącego sztucznego serca (US3097366)

piersiowej pacjenta oczekującego na przeszczep i z sukcesem pracowało tam przez 64 godziny (fotografia 22).

Początek lat 80. XX wieku przyniósł pierwsze działające egzemplarze sztucznego serca, przeznaczonego do długotrwałej eksploatacji. Zostały opracowane przez Roberta Jarvika – wynalazcę i lekarza, którego choroba ojca zainspirowała do zajęcia się budową mechanicznego zamiennika skrajnie niewydolnego mięśnia sercowego. Serce Jarvik-7, wszczepione w 1982 emerytowanemu dentyście przez Williama DeVriesa, pracowało w ciele pacjenta przez 112 dni, zaś kolejny pacjent przeżył z urządzeniem aż 620 dni. Projekt Jarvik-7 (fotografia 23) udało się doprowadzić do etapu stosowalności klinicznej, choć z uwagi na problemy z utrzymaniem jakości produkcji oraz serwisu, FDA wycofała w 1990 roku pozwolenie na korzystanie ze sztucznego serca Jarvika.



Fotografia 22. Oryginalny egzemplarz pierwszego w historii, wszczepionego pacjentowi sztucznego serca Liotta-Coolley, będącego pomostem pozwalającym bezpiecznie doczekać transplantacji żywego organu. Ekspонат jest wystawiony w muzeum w Waszyngtonie (https://t.ly/BurR)



Fotografia 23. Jarvik-7 Pierwsze, permanentne sztuczne serce (https://t.ly/Rpmm)

Po kilku dekadach okazało się, że w wielu przypadkach zamiast podłączania mechanicznych komór lub nawet pełnego zamiennika tego organu wystarczy... niewielka pompka, umieszczona w nacięciu w dolnej części mięśnia sercowego i podłączona do aorty wstępującej. Dziś takie produkty oferuje szereg firm, w tym także założone przez Jarvika w 1987 roku przedsiębiorstwo Jarvik Research, Inc. (fotografia 24). Znow to, co kiedyś wydawało się niemożliwe, okazało się niemalże rutynową procedurą, stosowaną w szpitalach na całym świecie. Nadal konstruktorzy i lekarze zmagają się jednak z szeregiem problemów – podobnie jak w przypadku pierwszych rozruszników z elektrodami nasierdziowymi, tak i teraz, w przypadku przewodów pneumatycznych zasilających wszczepione do klatki piersiowej sztuczne serce, dają o sobie znać infekcje wokół miejsca wyprowadzenia przewodów. I to właśnie one okazują się jedną z najczęstszych komplikacji podczas terapii LVAD.



Fotografia 24. Jarvik 2015 – jeden ze współcześnie stosowanych modeli pomp do wspomagania lewej komory serca (LVAD). Niewielkie urządzenie jest przeznaczone dla małych dzieci, stąd tak kompaktowe wymiary całości (porównanie z baterią typu AA). Źródło: https://t.ly/gylr



Rysunek 5. Szkic wczesnej wersji butelki ledejskiej (<https://t.ly/ORLl>)

„Butelka z prądem”, zupa i koci telefon

Elektryczność fascynowała człowieka już od wielu stuleci, jednak dopiero w XVIII wieku zaczęło coraz bardziej świadomie wykorzystywać to zjawisko w celach eksperymentalnych, a nieco później – także w zastosowaniach praktycznych. Benjamin Wilson już w 1748 roku wykorzystał butelkę ledejską (rysunek 5) – archetyp kondensatora do... eksperymentalnej stymulacji głowy niedosłyszącej kobiety. Przykładając jedną z elektrod do jej skroni obserwował reakcję pacjentki na wysokonapięciowe wyładowanie. O dziwo, powtórzenie eksperymentu w ciągu kilku dni spowodowało poprawę słuchu kobiety, co rzecz jasna zachęciło Wilsona i kolejnych badaczy do podążania tą ścieżką. Pomysł z Anglii dotarł nawet do Włoch, gdzie Alessandro Volta postanowił sprawdzić go na własnej skórze (dosłownie), opisując

w 1800 roku, że nieprzyjemne odczucie wyładowania elektrycznego, które poczuł – jak to opisuje – w mózgu, uznał za niebezpieczne i zaprzestał powtarzania tego wątpliwego etycznie eksperymentu. Co jednak ważniejsze: zwrócił uwagę na szum, słyszalny podczas stymulacji i porównał ów dźwięk (a dokładniej – wrażenie dźwiękowe) do odgłosu... gotującej się zupy.

Dopiero w 1930 roku dwóch badaczy z Princeton – Charles William Bray i Ernest Glenn Weaver (fotografia 25) – przeprowadziło interesujący, choć dość makabryczny eksperyment z wykorzystaniem (na szczęście nieprzytomnego) kota. Po otwarciu czaszki zwierzęcia podłączyli bowiem do nerwu słuchowego... kabel, biegnący do pobliskiego telefonu. Podczas, gdy jeden z naukowców mówił do „podłączonego” ucha, drugi nasłuchiwał przebiegów wytworzonych przez nerw słuchowy w oddalonym o 15 metrów, dźwiękoszczelnym pokoju. Eksperyment miał niezwykle istotne znaczenie dla późniejszej protetyki słuchu, pokazał bowiem, że nerw słuchowy – w odróżnieniu od innych nerwów – odpowiada na pobudzenie w zupełnie inny sposób. Podczas gdy pozostałe nerwy reagują zwiększeniem częstotliwości impulsów na wzrost amplitudy stymulacji, w tym przypadku sygnał elektryczny jest silnie skorelowany zarówno z amplitudą, jak i częstotliwością stymulacji. Innymi słowy, ucho zaczęło być postrzegane jako biologiczny mikrofon, przetwarzający odebrany dźwięk do postaci elektrycznej kopii oryginalnej fali akustycznej.

W połowie XX w. neurofizjolog André Djourno oraz chirurg Charles Eyriès wykorzystali ponowną operację jednego z pacjentów (przeprowadzaną z powodu komplikacji po innym zabiegu ucha) do bezpośredniej stymulacji nerwu słuchowego z użyciem niewielkiej cewki indukcyjnej. Pomimo wstępnych sukcesów – pacjent odzyskał częściowo zdolność słyszenia prostych bodźców – pierwszy oraz drugi egzemplarz urządzenia przestały po pewnym czasie zdawać egzamin, co zniechęciło naukowców do kontynuowania badań.

Wyniki prac opublikowała prasa, a wycinek jednej z gazet, przyniesiony przez pacjenta, zainspirował Williama House’a i Johna Doyle’a do rozpoczęcia własnych prac nad protezowaniem słuchu. Panowie współpracowali przy tym z inżynierem elektronikiem, bratem Doyle’a o imieniu Jim. Na początku stycznia 1961 roku badacze zaimplantowali pojedynczą elektrodę, a niecały miesiąc później zastąpili ją już czterokanałową sondą, którą jednak po niedługim czasie musieli usunąć. Dość szybko współpraca



Fotografia 25. Charles William Bray i Ernest Glenn Weaver (<https://t.ly/kj8E>)



Fotografia 26. Pierwszy implant ślimakowy dopuszczony przez FDA w 1984 do użytku u dorosłych (<https://t.ly/74L0>)

Doyle'ów z Housem zakończyła się, zarówno z powodu zalewu telefonów od potencjalnych pacjentów (wieści zbyt szybko rozniosły się przez rozrochoną prasę), jak i z przyczyn osobistych – bracia zabronili lekarzowi dostępu do dokumentacji elektronicznej, pomimo iż to przecież House zainicjował całą współpracę.

House nie poddał się jednak i już w 1972 roku wszczepiono pierwszy na świecie, noszony przez pacjenta (dziś określilibyśmy go mianem wearable) system transmisji indukcyjnej z mocowaniem magnetycznym, opracowany przez jego zespół. Dwa lata później odbyła się pierwsza międzynarodowa konferencja poświęcona elektrostymulacji nerwu słuchowego, a sam temat stał się coraz bardziej gorący z naukowego, jak i użytkowego punktu widzenia. Ostatecznie, w 1984 roku FDA dopuściła do użytku pierwszy, komercyjny jednokanałowy implant ślimakowy dla dorosłych (**fotografia 26**). Choć już 3 lata przed tym wydarzeniem miała miejsce pierwsza implantacja tego typu urządzenia u dziecka (**fotografia 27**), wielokanałowe implanty pediatryczne dla dzieci w wieku 2+ zostały dopuszczone dopiero w roku 1990, zaś dla niemowląt w pierwszym roku życia – dopiero w 2000 r.

Podsumowanie

Od wielu lat media i niektórzy naukowcy mają nas wizję rychłego połączenia technologii z mózgiem człowieka poprzez realizację idei interfejsu mózg-maszyna. Nie brak też doniesień o innych metodach spektakularnej transformacji człowieka w bionicznego cyborga, naszpikowanego elektroniką od stóp do głów. Inni znów panikują przed „zaczepianiem” społeczeństw i globalną kontrolą. Po raz kolejny okazuje się jednak, że prawdziwych innowacji należy szukać zupełnie gdzie indziej. Nie na pierwszych stronach gazet i magazynów popularnonaukowych, ale na salach operacyjnych i w laboratoriach badawczo-rozwojowych wielkich koncernów, które rozwinęły się z małych, garażowych firm lub ciasnych warsztatów, założonych przez wybitnych inżynierów przed kilkoma dekadami.

Historia pokazuje, że droga do osiągnięcia obecnego poziomu techniki implantowalnej była długa i wyboista, a jak to zwykle bywa w medycynie – pochłonęła też wiele ofiar, zarówno wśród zwierząt, na których prowadzono (z obecnego punktu widzenia niezbyt etyczne) badania podstawowe, jak i wśród pierwszych pacjentów, zmagających się z poważnymi powikłaniami po wszczepieniu prototypowych urządzeń. Na najwyższy podziw i uznanie zasługują wszyscy ci, którzy – wykazując się najwyższą odwagą – świadomie położyli się na stołach operacyjnych i pozwolili innowatorom na wdrożenie najświeższych, pionierskich wyników prac, prowadzonych w pocie czoła



Fotografia 27. Dr William House z Tracy Husted – pierwszą pacjentką pediatryczną z wszczepionym implantem ślimakowym (1981). Źródło: <https://t.ly/6fXU>

przez najbardziej otwarte umysły ubiegłego wieku. To właśnie im Autor pragnie zadedykować niniejszy artykuł.

inż. Przemysław Musz

Główne źródła informacji wykorzystanych w artykule:

1. Thorwald J., Stulecie chirurgów, Wyd. Literackie, 1987
2. Aquilina O. A brief history of cardiac pacing. Images Paediatr Cardiol. 2006 Apr; 8(2):17-81
3. Ramsdale, D. r., Rao, A., Cardiac Pacing and Device Therapy, wyd. Springer London, 2012
4. Efimov I.R., Kroll M.W., Tchou P.J., Cardiac Bioelectric Therapy. Mechanisms and Practical Implications, wyd. Springer Science+Business Media, 2009
5. A Legacy of Innovation. The Medtronic Story, dost. online: <https://t.ly/vxC4>
6. Mudry A, Mills M., The Early History of the Cochlear Implant: A Retrospective. JAMA Otolaryngol Head Neck Surg. 2013;139(5):446–453
7. <https://t.ly/9Adw>
8. <https://t.ly/OBJpe>
9. <https://t.ly/PPWY>
10. https://t.ly/56_WX

REKLAMA

**Czytaj artykuły zanim
zostaną wydane
w formie papierowej
na
[www.ep.com.pl/
EPwtoku](http://www.ep.com.pl/EPwtoku)**

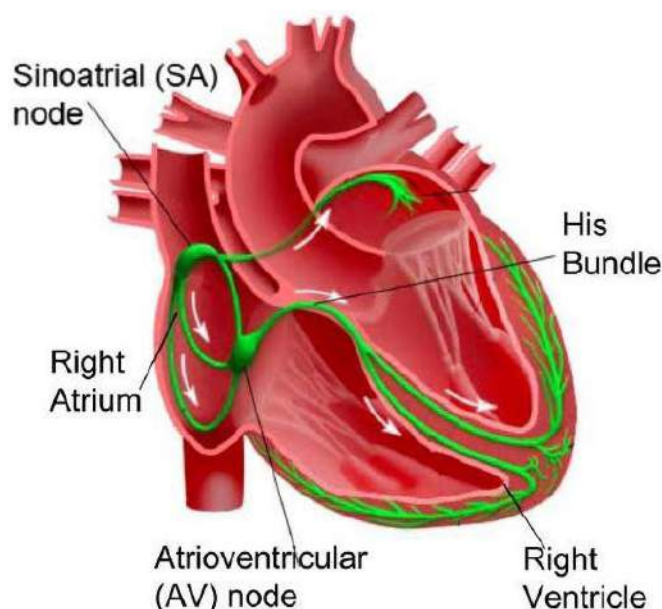


Aktywne implanty okiem inżyniera medycznego

Wprowadzenie urządzenia elektronicznego do ciała człowieka oraz pozostawienie go w nim na długi czas (lub nawet dożywotnio) stawia konstruktorom wymogi dalece bardziej rygorystyczne, niż w większości obszarów elektroniki. Choć w tym przypadku nie mamy już do czynienia z silnymi wibracjami, spotykany mi choćby w przemyśle czy motoryzacji, zaś zakres temperatur pracy jest bodaj największy z możliwych, to na urządzenia wszczepialne czyha szereg innych zagrożeń, co więcej – one same mogą spowodować szereg niebezpieczeństw dla zdrowia i życia pacjenta. Nic więc dziwnego, że w praktyce niemal wszystkie implanty aktywne „wpadają” w III, najwyższą klasę urządzeń medycznych, co oznacza, że (potencjalnie) stwarzają największe ryzyko w przypadku jakiegokolwiek awarii lub błędu. W artykule zaprezentujemy najważniejsze informacje dotyczące konstrukcji urządzeń wszczepialnych oraz przeznaczonych dla nich akcesoriów, omówimy także kluczowe parametry wybranych implantów elektronicznych.

Rynek aktywnych implantów – choć obejmuje wyroby o znaczeniu kluczowym dla współczesnej terapii i diagnostyki – bardzo wolno rozszerza się o nowe rodzaje produktów. Wynika to z szeregu czynników, związanych zarówno z ograniczonym zapotrzebowaniem rynkowym w tym zakresie, jak i z wyśrubowanymi wymogami formalnymi (certyfikacja, dopuszczenie do użytku, badania kliniczne), limitami technicznymi (moc i pojemność zasilania, objętość obudowy, wymogi szczelności), etc. Dobrym przykładem będą tutaj sztuczne serca – choć ich początki sięgają kilku dekad wstecz (zainteresowanym fascynującą historią techniki urządzeń wszczepialnych polecamy artykuł *Historia rozwoju elektroniki implantowalnej* w tym numerze *Elektroniki Praktycznej*), to z wielu względów okazało się, że często znacznie lepszym rozwiązaniem jest albo zastosowanie prostszych (i tańszych) urządzeń do tymczasowego wspomagania lewej komory serca, albo możliwie szybkie wykonanie klasycznej transplantacji biologicznego organu. Rynek bezlitośnie weryfikuje początkowe założenia – niektóre technologie (np. rozruszniki serca czy też implantowalne kardiowertery-defibrylatory) z powodzeniem przeżywają próbę czasu, podczas gdy inne stopniowo zanikają we mgle technologicznej historii.

Z powyższych względów zdecydowaliśmy się omówić zagadnienia konstrukcyjne i projektowe implantów aktywnych na przykładzie dwóch zągębiających się pod wieloma względami grup urządzeń: rozruszników serca (kardiostymulatorów) oraz implantowalnych kardiowerterów-defibrylatorów (ICD). Znaczną część podanych w artykule informacji można jednak ekstrapolować na inne grupy urządzeń, zwłaszcza rozmaitych neurostymulatorów – te same warunki panujące wewnątrz ciała ludzkiego, wymuszają bowiem stosowanie sprawdzonych rozwiązań m.in. w zakresie doboru materiałów, źródeł zasilania, sposobów produkcji obudowy i złączy elektrod, etc.



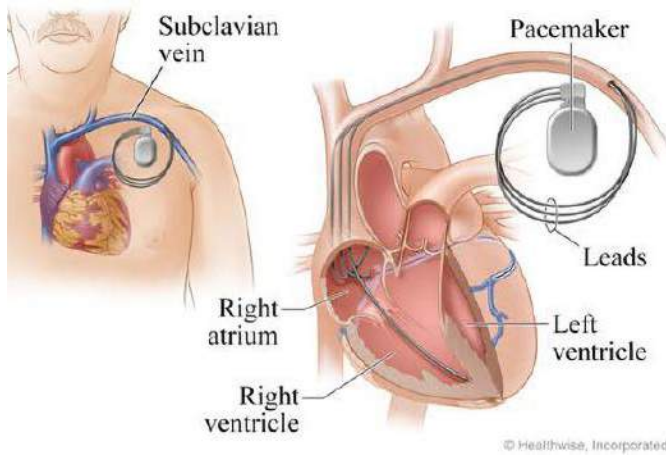
Rysunek 1. Uproszczone zobrazowanie układu bodźcotwórczo-przewodzącego serca. Oznaczenia: *sinoatrial node* – węzeł zatokowo-przedsionkowy, *right atrium* – prawy przedsionek, *right ventricle* – prawa komora, *atrioventricular node* – węzeł przedsionkowo-komorowy, *His bundle* – pęczek Hisa (https://t.ly/_dG2)

Podstawowe informacje o układzie bodźcotwórczo-przewodzącym serca

W zdrowym sercu tempo i kolejność kurczenia się przedsionków oraz komór są nadawane przez tzw. układ bodźcotwórczo-przewodzący, będący grupą wyspecjalizowanych kardiomiocytów (komórek mięśnia sercowego). Źródłem rytmicznych pobudzeń o częstotliwości regulowanej w pewnym zakresie przez inne, pozasercowe czynniki (autonomiczny układ nerwowy oraz hormony) jest tzw. węzeł zatokowo-przedsionkowy, a taki „zdrowy”, naturalny, rytm jest z tegoż powodu nazywany rytmem zatokowym. Układ bodźcotwórczo-przewodzący można wyobrazić sobie jako swego rodzaju „wewnętrzne okablowanie” mięśnia sercowego (**rysunek 1**), które jednak nie tylko przewodzi impulsy elektryczne wywołujące skurcze, ale także opóźnia je, przy czym czasy propagacji w poszczególnych odcinkach znacząco się różnią, dzięki czemu akcja serca może być hemodynamicznie skuteczna (tj. odpowiednia ilość krwi, spływającej pod niskim ciśnieniem do przedsionków, jest następnie pompowana do płuc oraz aorty).

Tryby pracy rozruszników

Na rodzaj zaburzeń rytmu serca ma wpływ szereg czynników – niektóre z nich są powiązane z nieprawidłowościami funkcjonowania samego węzła zatokowo-przedsionkowego, inne zaś wynikają z zakłócenego przewodzenia w dalszych częściach układu bodźcotwórczo-przewodzącego (np. w wyniku częściowej martwicy w obrębie blizny pozawałowej). Zaburzenia te w wielu przypadkach wymagają korekcji poprzez narzucenie rytmu za pomocą sztucznego rozrusznika lub jedynie wspomaganie lub koordynację naturalnych pobudzeń tylko wtedy, gdy jest to naprawdę potrzebne.



Rysunek 2. Typowe ułożenie elektrod w przypadku urządzenia do terapii resynchronizującej serca. Oprócz elektrod w prawym przedsionku i prawej komorze (stosowanych w klasycznej elektrostymulacji dwujamowej) wykorzystywana jest także dodatkowa, trzecia elektroda, umieszczona w (lub na) ścianie lewej komory (<https://t.ly/3L2R>)

Z powyższego opisu wynika kilka istotnych założeń koncepcyjnych w kwestii budowy i funkcjonowania rozrusznika:

- **Miejsce stymulacji** – w zależności od tego, na jakim etapie układu bodźcotwórczo-przewodzącego występuje anomalia, niezbędne może być stymulowanie na poziomie obydwu prawych jam serca, prawej komory lub (obecnie znacznie rzadziej) tylko prawego przedsionka. Stąd też pojawia się konieczność zastosowania jednego lub dwóch kanałów wyjściowych, choć w przypadku tzw. terapii resynchronizującej serca liczba kanałów zwiększa się do trzech – patrz **rysunek 2** (więcej informacji na ten temat podamy w dalszej części artykułu).
- **Tryb stymulacji** – pierwsze rozruszniki miały postać prostych, jednokanałowych (dziś nazwalibyśmy je jednojamowymi) stymulatorów o stałej częstotliwości wyjściowej i działały całkowicie asynchronicznie względem biologicznego rytmu serca. Takie rozwiązanie miało szereg wad i choć pozwalało leczyć niektóre zaburzenia rytmu, generowało także pewne zagrożenia dla pacjenta. Elektrofizjologia serca jest zagadnieniem bardzo złożonym i nieprawidłowe zastosowanie stymulacji (np. wygenerowanie impulsu w określonej części cyklu pracy serca) może zapoczątkować epizod groźnych dla życia pacjenta arytmii komorowych. Takie przypadki miały miejsce po rozpoczęciu stosowania rozruszników pierwszej generacji, stąd już w latach 60. XX w. do użytku wprowadzono urządzenia pracujące w trybie stymulacji „na żądanie” (*demand pacing*), co sprowadzało się do detekcji spontanicznej aktywności serca, a następnie wyzwalań (lub wstrzymywania) impulsu w zależności od niej. Tak więc do kanałów wyjściowych należało dołożyć także obwody wzmacniaczy, komparatory oraz układy czasowe, które umożliwiały uzależnienie działania generatora od aktualnego etapu cyklu pracy serca.

Opisane powyżej parametry stymulacji mają absolutnie krytyczne znaczenie dla zastosowania urządzenia u konkretnego pacjenta, stąd też w praktyce klinicznej stosuje się międzynarodowy kod NASPE/BPEG Generic, nazywany w skrócie NBG (North American Society of Pacing and Electrophysiology, British Pacing and Electrophysiology Group). Powszechnie używane oznaczenia kodują:

- Miejsce stymulacji (*pacing*):
 - O – brak,
 - A – przedsionek (*Atrium*),
 - V – komora (*Ventricle*),
 - D – obie jamy ($D = A + V$).
- Miejsce detekcji spontanicznej aktywności serca (*sensing*):
 - O – brak,



Fotografia 1. Nowoczesny rozrusznik dwujamowy z automatyczną regulacją częstości impulsów stymulujących (DDDR) – model Effecta DR marki Biotronik (<https://t.ly/0bnz>)

- A – przedsionek,
- V – komora,
- D – obie jamy ($D = A + V$).
- Tryb pracy:
 - O – brak,
 - T – wyzwolenie impulsu stymulującego po detekcji sygnału (*Triggered*),
 - I – hamowanie stymulacji po wykryciu aktywności własnej serca (*Inhibited*),
 - D – praca w obu trybach (tryb „mieszany” w zależności od częstości rytmu spontanicznego).
- Modulacja częstotliwości generatora impulsów:
 - O – brak,
 - R – zmienna częstotliwość impulsów stymulujących,
- Stymulacja wielopunktowa:
 - O – brak,
 - A – przedsionek,
 - V – komora,
 - D – obie jamy ($D = A + V$).

Kod NBG jest – obok rysunku prezentującego układ wprowadzeń (A – wyprowadzenia przedsionkowe, V – komorowe) – najważniejszym oznaczeniem technicznym, spotykanym na obudowie każdego stymulatora. Rzadko jednak można spotkać symbole dłuższe, niż 4 znaki – obecnie najczęściej można znaleźć urządzenia typu DDD czy DDDR (**fotografia 1**). Co ciekawe, pomimo kilku rewizji kodu (pierwsza wersja, opracowana przez ICHD, powstała w już 1974 roku), niektórzy



Fotografia 2. Chińskie rozruszniki dwujamowe typu DDDR oraz DDDC (<https://t.ly/4Ghf>)

producenci nadal wykorzystują dodatkowe oznaczenia, które zniknęły z „oficjalnej” wersji zrewidowanej na przełomie XX i XXI wieku. Jako przykład można tutaj wskazać urządzenia chińskiej marki Qinming typu DDDC (litera C we wcześniejszych wersjach kodu oznaczała funkcje komunikacyjne/telemetryczne, które dziś są już standardem w każdym rozruszniku – patrz **fotografia 2**). Warto też dodać, że niektórzy producenci oznaczają rozruszniki jednojamowe z użyciem litery S (single), np. SSI – w tym przypadku elektroda stymulatora może być wszczepiona w komorze lub w przedsionku, zaś po wykryciu pobudzeń spontanicznych o odpowiednio wysokiej amplitudzie, generowanie impulsów zostaje wstrzymane.

Kardiowertery-defibrylatory (ICD)

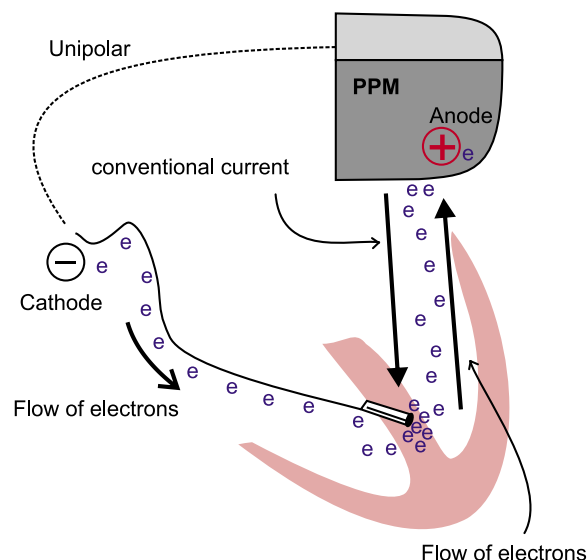
Podczas, gdy sztuczne rozruszniki serca (kardiostymulatory) są przeznaczone do utrzymywania lub regulowania rytmu serca, kardiowertery-defibrylatory mają za zadanie przerwać groźne dla życia zaburzenia rytmu (takie, jak częstoskurcz komorowy czy też migotanie komór), niejako „resetując” serce i przywracając je w ten sposób do normalnej pracy. Istotna różnica pomiędzy kardiowersją, a defibracją, wynika z momentu, w którym dostarczany jest do serca wysokoenergetyczny impuls. Defibracja stanowi procedurę asynchroniczną – wysokoenergetyczne wyładowanie następuje niezależnie od aktualnej fazy cyklu pracy serca. Kardiowersja zaś polega na dostarczeniu impulsu o mniejszej energii, ale w ściśle określonym momencie aktywności bioelektrycznej serca i jest stosowana w tych przypadkach, w których asynchroniczne wyładowanie mogłoby pogorszyć stan pacjenta poprzez wywołanie jeszcze groźniejszych zaburzeń rytmu. Warto dodać, że kardiowertery-defibrylatory mogą posiadać także funkcję stymulacji, w ramach której pełnią rolę klasycznego rozrusznika serca.

Dla defibrylatorów także powstał kod, będący swego rodzaju adaptacją oznakowania NBG. W tym przypadku mamy do czynienia z czterema literami, kolejno:

- Jama w której następuje wyładowanie:
 - O – brak,
 - A – przedsionek (**A**trium),
 - V – komora (**V**entricle),
 - D – obie jamy (D = A + V).
- Jama serca stymulowana antytachyarytmicznie (tj. w celu przeciwdziałania tachyarytmii, czyli zaburzeniom rytmu przebiegającym ze zbyt wysoką częstotliwością skurczów):
 - O – brak,
 - A – przedsionek,
 - V – komora,
 - D – obie jamy (D = A + V).
- Metoda detekcji częstoskurczu:
 - E – elektrogram (zapis aktywności elektrycznej wewnątrz serca),
 - H – hemodynamika (zapis parametrów związanych z mechaniką pracy serca – rozwiązania oparte np. na czujnikach ciśnienia krwi są jednak niespotykane w praktyce).
- Jama serca stymulowana antybradyarytmicznie (tj. w celu przeciwdziałania bradyarytmii, czyli zaburzeniom rytmu przebiegającym ze zbyt niską częstotliwością skurczów):
 - O – brak,
 - A – przedsionek,
 - V – komora,
 - D – obie jamy (D = A + V).

CRT

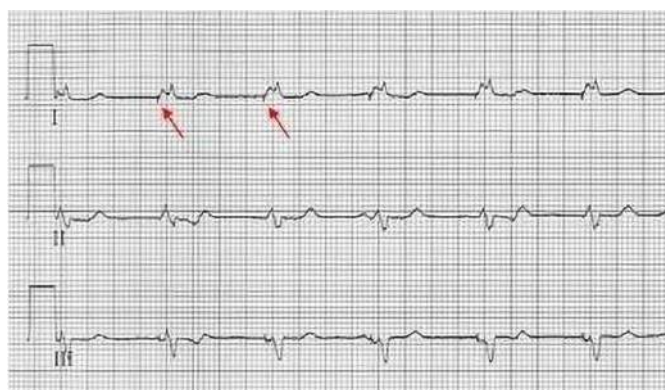
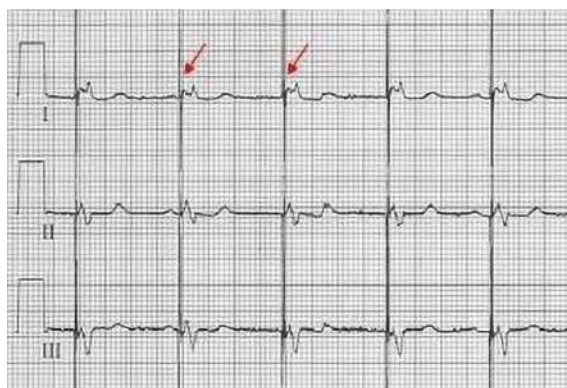
Terapia resynchronizująca serca jest rozszerzeniem klasycznych terapii z użyciem rozruszników lub kardiowerterów-defibrylatorów. Poprzez zastosowanie dodatkowej, trzeciej elektrody, umieszczonej w rejonie lewej komory serca, możliwe jest zmniejszenie tzw. dyssynchronii. Zaburzenie to polega na patologicznym zwiększeniu opóźnienia pomiędzy elektrycznym pobudzeniem komór, a ich



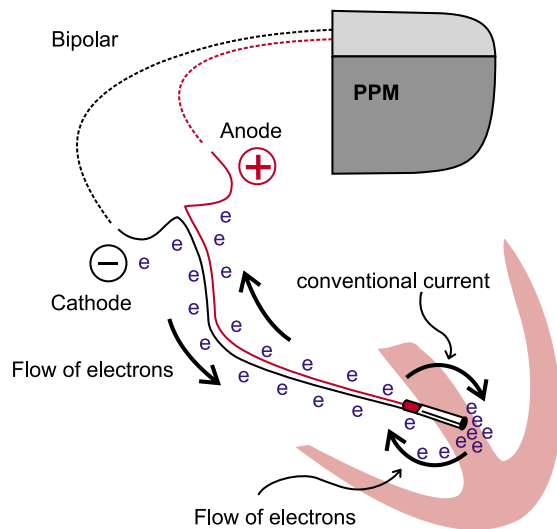
Rysunek 3. Schemat stymulacji unipolarnej (<https://t.ly/ACZO>)

wynikowym skurczem. Jeżeli różnica czasu jest znaczna, serce nie jest w stanie skutecznie pompować właściwej ilości – stan taki określamy jako redukcja rzutu minutowego serca. Terapia CRT zyskała szczególne zainteresowanie klinicystów w kwestii leczenia pacjentów z niewydolnością serca (*heart failure*, HF) – czyli tych, których układ krążenia nie jest w stanie sprostać fizjologicznym potrzebom organizmu w zakresie ukrwienia narządów (najprościej mówiąc, serce pacjenta z HF jest za słabe, by zapewnić organizmowi właściwe funkcjonowanie, co powoduje m.in. niezwykle szybkie męczenie się przy jakimkolwiek wysiłku, a w skrajnych przypadkach nawet podczas spoczynku).

Urządzenia CRT to w rzeczywistości dwie podgrupy urządzeń – CRT-P (rozruszniki z dodatkowym kanałem dla lewej komory) oraz CRT-D (rozszerzenie klasycznego ICD o kanał lewokomorowy).



Rysunek 4. Przykładowe zapisy EKG z widocznymi impulsami pochodzącymi od stymulacji unipolarnej (po lewej) i bipolarnej (po prawej). Źródło: <https://t.ly/wwwUn>



Rysunek 5. Schemat stymulacji bipolarnej (<https://t.ly/ACZO>)

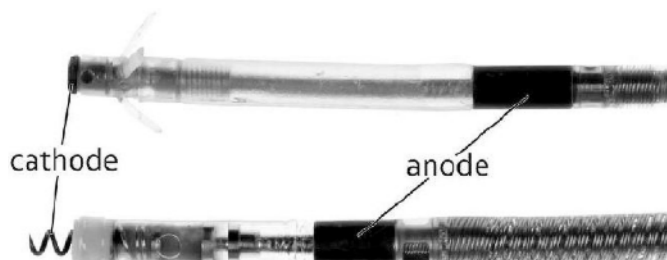
Polaryzacja i układ elektrod

Pierwsze rozruszniki serca korzystały z elektrod mocowanych na zewnętrznej powierzchni mięśnia sercowego (tzw. elektrody nasierdziowe), co umożliwiało chirurgom umiejscowienie wyprowadzeń wyjściowych generatora dokładnie tam, gdzie było to konieczne z medycznego punktu widzenia. Przejście na elektrody implantowane z dostępu naczyniowego spowodowało jednak, iż konstruktorzy rozruszników zaczęli wykorzystywać tytanową obudowę urządzenia jako anodę, zaś do wnętrza serca wszczepiana była tylko sama katoda (rysunek 3). Takie rozwiązanie – określane mianem stymulacji unipolarnej – miało pewne zalety, do których należy zaliczyć przede wszystkim:

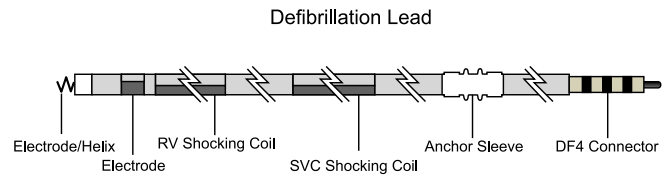
- prostszą konstrukcję (a zatem niższy koszt) elektrody,
- mniejszą średnicę zewnętrzną elektrody i – co za tym idzie – większą elastyczność.

Z uwagi na szeroki obszar tkanek pacjenta objęty polem elektrycznym (wytworzanym podczas impulsu pomiędzy katodą, a obudową stymulatora), w zapisach EKG u pacjentów z wszczepionymi elektrodami unipolarnymi można zauważyć bardzo wyraźne „szpilki” odpowiadające impulsom stymulatora (rysunek 4), choć – z drugiej strony – w niektórych przypadkach zbyt silne artefakty mogą utrudniać właściwą obserwację i interpretację sygnału elektrokardiograficznego, zwłaszcza w uproszczonych konfiguracjach odprowadzeń o zmniejszonej liczbie kanałów pomiarowych. Stymulacja unipolarna ma jednak szereg znacznie istotniejszych wad, do których zaliczyć należy m.in. większą podatność na artefakty pochodzenia pozasercowego (co nie dziwi z uwagi na fakt, iż wzmacniacz wejściowy stymulatora zbiera sygnał z większego obszaru ciała, w dodatku o dość wysokiej impedancji) oraz ryzyko przypadkowego pobudzenia pobliskich mięśni lub nerwów obwodowych.

Rozwój technologii produkcji elektrod kardiostymulatorów doprowadził do opracowania (w latach 80. XX wieku) elektrod bipolarnych,



Fotografia 3. Przykładowe elektrody bipolarne (pokazano końcówki po stronie implantacji). Dolny odcinek oznacza skalę 3 mm (źródło: <https://t.ly/aWPU>)



Rysunek 6. Budowa 4-stykowej elektrody kardiowertera-defibrylatora (<https://t.ly/cr1F>)

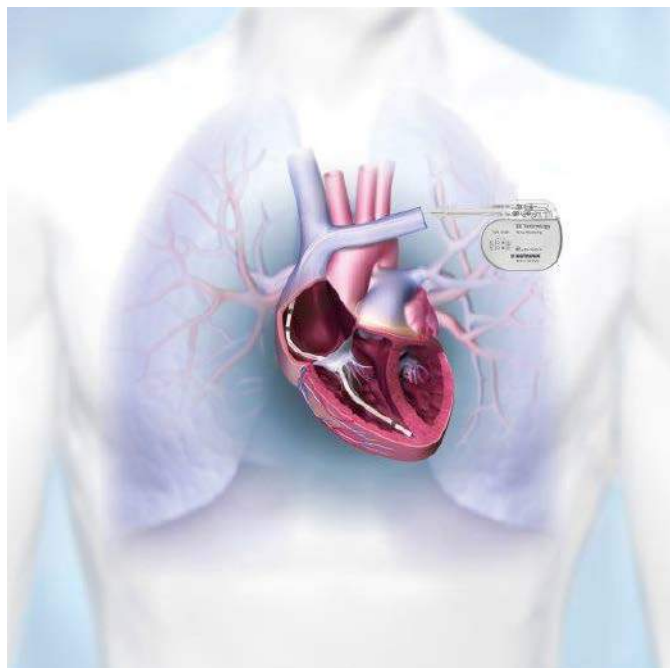
w których anoda ma postać pierścienia, umieszczonego proksymalnie (tj. – w tym przypadku – bliżej stymulatora) w stosunku do końcówek (katody). Takie rozwiązanie sprawia, że obieg prądów stymulujących jest wielokrotnie krótszy, niż w przypadku stymulacji unipolarnej (zamyka się bowiem na obszarze rzędu kilku... kilkunastu milimetrów – rysunek 5, fotografia 3), co prowadzi m.in. do znacznego osłabienia artefaktów w zapisach EKG oraz redukcji ryzyka błędnej aktywacji rozrusznika (spowodowanej np. zakłóceniami EMI czy też aktywnością mięśni szkieletowych). Warto dodać, że – z technicznego punktu widzenia – zastosowanie elektrod bipolarnych nie uniemożliwia przeprogramowania go w razie potrzeby na tryb unipolarny. W takim przypadku anoda jest po prostu elektronicznie przełączana z pierścienia, znajdującego się we wszczepionej elektrodzie, na metalową puszkę obudowy elektrostymulatora (rzecz jasna, dostępność takiej funkcji zależy od modelu stymulatora).

Wszczepialne kardiowertery-defibrylatory korzystają z bardziej rozbudowanych elektrod, posiadających dodatkowe powierzchnie stykowe, przeznaczone do dostarczania wysokonapięciowych impulsów w prawej komorze serca oraz żyły głównej górnej (SVC – superior vena cava). Przykładowe rozwiązanie pokazano na rysunku 6.

Gwoli ścisłości warto dodać, iż na rynku pojawiły się także rozwiązania o nieco zbliżonej konstrukcji, ale przeznaczone do zgoła innego celu. Urządzenia bazujące na technologii DX marki Biotronik (fotografia 4) umożliwiają „zdalną” detekcję aktywności prawego przedsionka, jednak bez konieczności implantacji dedykowanej dlań elektrody – zamiast niej wykorzystywane są dodatkowe pierścienie „dipolowe”, umieszczone w tak dobranej odległości od końcówki części komorowej, że po wszczepieniu znajdują się one na wysokości przedsionka (rysunek 7). Umożliwia to znaczące skrócenie procedury implantacji ICD (co przekłada się na mniejsze obciążenie pacjenta oraz oszczędność czasu i pieniędzy szpitala) u pacjentów, którzy wprawdzie nie wymagają stymulacji przedsionków, jednak są zagrożeni występowaniem epizodów migotania przedsionków, które pozostałyby niewykrywalne przy użyciu samej tylko elektrody komorowej. Co jednak ważniejsze, zastosowanie „zdalnej” detekcji aktywności przedsionków pozwala zredukować ryzyko „pomylenia” arytmii nadkomorowych z komorowymi, co ma niebagatelne znaczenie dla bezpieczeństwa terapii.



Fotografia 4. Kardiowerter-defibrylator z serii DX marki Biotronik (<https://t.ly/Tpfi>)

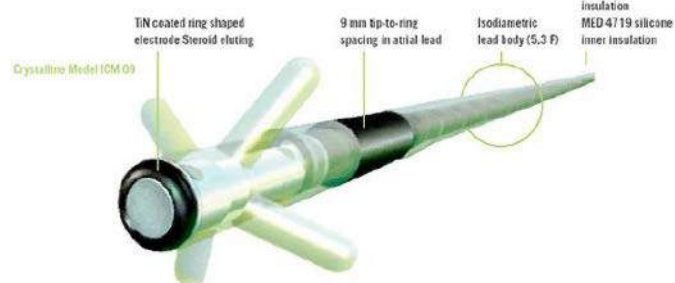


Rysunek 7. Schematyczne zobrazowanie ułożenia elektrody, biegnącej od urządzenia ICD z technologią Biotronik DX do prawej komory serca. Widoczne dwa pierścienie na wysokości prawego przedsionka umożliwiają zdalny sensing aktywności w tym obszarze (<https://t.ly/Tpfi>)

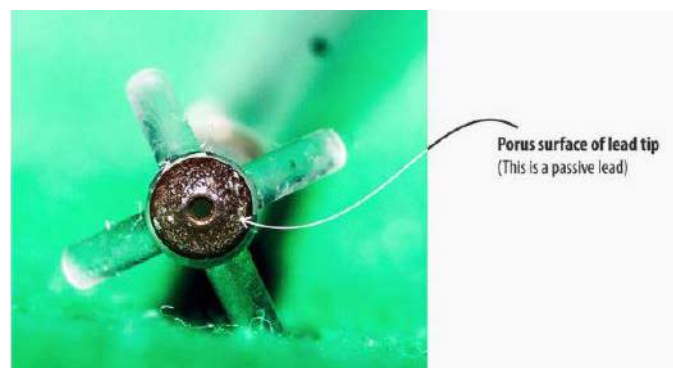
Budowa końcówek elektrod

Konstrukcja elektrod musi spełniać szereg istotnych wymogów technicznych, do których należą przede wszystkim:

- możliwość niezawodnego zamocowania końcówki w docelowym punkcie mięśnia sercowego,
- odpowiednia elastyczność i wysoka odporność na zginanie,
- dobry kontakt elektryczny z tkanką mięśnia sercowego,



Rysunek 8. Budowa elektrody z mocowaniem pasywnym (<https://t.ly/Aima>)



Fotografia 5. Końcówka elektrody pasywnej. Widoczna silnie chropowata powierzchnia katody (zaznaczona strzałką) i cztery haczyki umożliwiające zakleszczenie w tkance miokardium (<https://t.ly/v cXr>)



Fotografia 6. Akcesoria z serii TightRail Guardian marki Philips, przeznaczone do usuwania elektrod z serca pacjenta (<https://t.ly/mYy8>)

- wtyk umożliwiający hermetyczne i niezawodne zamocowanie w gnieździe stymulatora.

Wszystkie elektrody kardiostymulatorów oraz ICD można podzielić ze względu na sposób ich kotwiczenia w tkance miokardium.

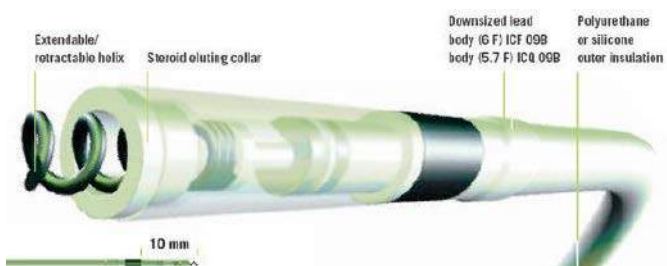
- **Elektrody z mocowaniem pasywnym** (rysunek 8) mają obłą (zwykle sferyczną lub spłaszczoną) końcówkę o silnie chropowatej powierzchni (fotografia 5), ułatwiającej wrastanie w tkankę i zwiększającej efektywną powierzchnię kontaktu elektrycznego. Implantacja odbywa się poprzez dociśnięcie elektrody do mięśnia sercowego, najczęściej w rejonie koniuszka prawej komory – miejsce to, z uwagi na duże zagęszczenie tzw. beleczek mięśniowych, doskonale „współpracuje” z niewielkimi haczykami, które zakleszczając się w rozbudowanej strukturze miokardium powodują zablokowanie elektrody w pożądanej pozycji (gwoździści należy wspomnieć, że istnieje też możliwość fiksacji elektrod pasywnych w prawym przedsionku). Do zalet tego rozwiązania należy zaliczyć mniejszą inwazyjność implantacji, która nie wiąże się z bezpośrednim uszkodzeniem mięśnia sercowego.

Pewną wadą elektrod pasywnych jest trudność z ich usuwaniem po długotrwałej eksploatacji przez pacjenta – z tego też względu powstał szereg urządzeń, wprowadzanych przez dostęp naczyniowy (przy wykorzystaniu elektrody jako przewodnika) i służących do bezpiecznego wycięcia elektrody z przyrośniętej tkanki. W zależności od modelu, urządzenia te bazują na specjalnych, mechanicznych nożach obrotowych (fotografia 6) lub cewnikach z wbudowanymi na obwodzie światłowodami, transmitującymi promieniowanie lasera ultrafioletowego o długości fali równej 308 nm, generowanego impulsowo z częstotliwością powtarzania od 25 do 80 Hz (fotografia 7).

- **Elektrody z mocowaniem aktywnym** są pozbawione haczyków, zaś końcówka (katoda) ma postać cienkiej, ostro zakończonej spiralki (rysunek 9), którą lekarz samodzielnie wkręca bezpośrednio w tkankę mięśnia



Fotografia 7. Końcówka robocza cewnika GlideLight marki Philips, przeznaczonego do usuwania elektrod metodą laserową (LLD). Źródło: <https://t.ly/PqYa>



Rysunek 9. Budowa elektrody do mocowania aktywnego (https://t.ly/_TZz)

sercowego (np. ścianę komory lub przegrodę międzykomorową), obracając w tym celu specjalne narzędzie, umieszczone na zewnętrznym końcu elektrody (przed jej podłączeniem do gniazda stymulatora). W czasie wszczepiania z końcówki zsuwana jest osłona, chroniąca układ naczyniowy i struktury serca przed uszkodzeniem (w tym szczególnie perforacją) podczas wprowadzania.

Istotną wadą elektrod tego typu jest znacznie większa inwazyjność, wiążąca się z koniecznością wkręcenia gwintu bezpośrednio w mięsień sercowy. Z drugiej strony, mocowanie aktywne jest uważane za bardziej niezawodne i uniwersalne, w porównaniu do fiksacji pasywnej, nieco łatwiej przebiega też usuwanie elektrod aktywnych po dłuższym czasie od implantacji.

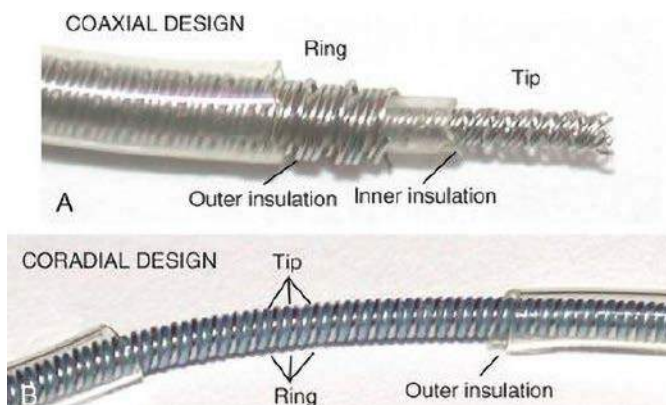
Aspekty materiałowe i konstrukcyjne w produkcji przewodów elektrod implantowalnych

Środowisko panujące wewnątrz ciała człowieka jest – wbrew pozorom – bardzo agresywne dla wprowadzonych doń urządzeń i akcesoriów, głównie z uwagi na chemiczne oddziaływanie płynów ustrojowych na stosowane w inżynierii biomedycznej materiały. Co więcej, duże znaczenie dla zmęczenia mechanicznego ma także istotna ruchomość – nie należy zapominać, że elektrody wszczepiane do ludzkiego serca są poddawane nie tylko zginaniu w wyniku ruchu mięśni w rejonie obręczy barkowej, ale także – a raczej przede wszystkim – powtarzalnym cyklom pracy serca. Przyjmując, że tętno pacjenta wynosi 70 uderzeń/minutę, otrzymujemy zawrotną liczbę prawie 37 milionów (!) cykli zginania w ciągu roku – a należałoby przecież uwzględnić wzrost częstości rytmu serca podczas zwiększonej aktywności fizycznej pacjenta.

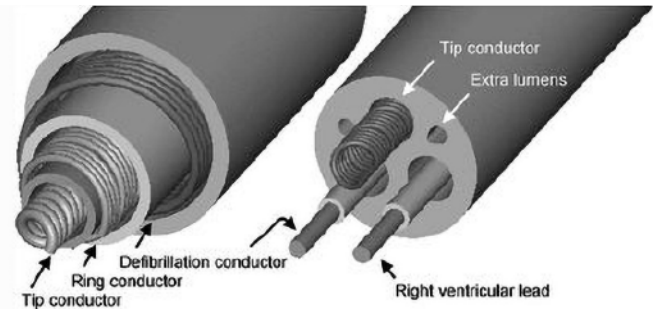
Proste (i bardzo zawodne) przewody stosowane w pierwszych generacjach rozruszników serca zostały dość szybko zastąpione znacznie bardziej wytrzymałymi mechanicznie przewodami spiralnymi, które – z uwagi na doskonałe osiągi – są stosowane do dziś. Istnieją jednak dwie główne odmiany, różniące się sposobem nawinięcia przewodów (w przypadku elektrod bipolarnych):

- **konstrukcja koncentryczna (coaxial)** – końcówka elektrody (katoda) jest podłączona za pomocą środkowego przewodu spiralnego, osłoniętego cienką warstwą izolacji od zewnętrznego oplotu, łączącego wtyk stymulatorowy z pierścieniem anody (fotografia 8). Całość jest pokryta zewnętrzną izolacją, bezpośrednio kontaktującą się z ciałem pacjenta.
- **konstrukcja typu coradial** – w tym przypadku obydwa przewody są nawinięte w sposób, przypominający nieco uzwojenie bifilarne: cienka warstwa izolacji drutu nawojowego umożliwia wspólne nawinięcie obu żył, prowadzących do poszczególnych styków elektrody (patrz fotografia 7). Niewątpliwą zaletą takiego rozwiązania jest zmniejszenie liczby wewnętrznych warstw elektrody, co redukuje jej wynikową średnicę.

W przypadku wszczepialnych kardiowerterów-defibrylatorów (ICD) istnieje konieczność poprowadzenia – oprócz przewodów służących



Fotografia 8. Porównanie budowy elektrod typu coaxial i coradial (<https://t.ly/Z2RS>)



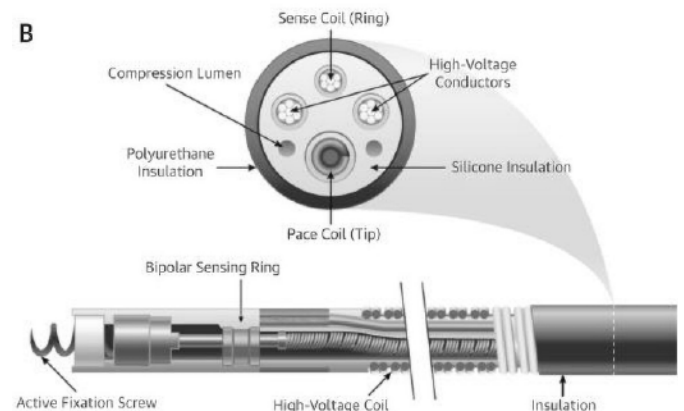
Rysunek 10. Porównanie konstrukcji dwóch przewodów elektrod ICD: koncentrycznego (po lewej) i wielożyłowego (po prawej). Źródło: <https://t.ly/AHnn>

do detekcji spontanicznej aktywności serca oraz ewentualnego przesyłu słabych impulsów stymulujących – także dodatkowego okablowania, umożliwiającego dostarczanie do miokardium wysokonapięciowych wyładowań podczas procedury kardiowersji lub defibrylacji. Choć istnieje możliwość wyprodukowania przewodów koncentrycznych o liczbie przewodników przekraczającej 2 (rysunek 10), to w praktyce stosuje się podejście określane mianem *multilumen design*. Co ciekawe, choć tego typu okablowanie przypomina (pod względem koncepcyjnym) przemysłowe kable hybrydowe, to producenci zdecydowali się pozostać przy stosowaniu żył o konstrukcji spiralnej, wspomaganych jedynie przez „zwykłe” przewodniki linkowe. Przykłady pokazano na rysunkach 11 i 12. Dodatkowe, puste kanały, oznaczane na rysunkach mianem *compression lumen*, zwiększają odporność całości na zewnętrzne siły ściskające.

Od strony materiałowej, elektrody stawiają inżynierom szczególnie restrykcyjne wymagania. Wszczepione przez naczyniowo przewody muszą być bowiem wysoce biokompatybilne, odporne mechanicznie, stabilne pod względem chemicznym oraz możliwie elastyczne i jednocześnie stosunkowo wygodne do prowadzenia w świetle naczynia krwionośnego. Dodatkowo, istotnym ograniczeniem są wymiary całości – średnica elektrod mieści się przeważnie w zakresie od około 1,1 mm do około 2,2 mm. Z tego też względu producenci elektrod do rozruszników i ICD stosują najwyższej klasy, dopuszczone do użytku medycznego materiały izolacyjne, w tym najczęściej:



Rysunek 11. Różne odmiany przewodów wielożyłowych stosowanych w elektrodach ICD marki Medtronic (<https://t.ly/aUJE>)



Rysunek 12. Schemat konstrukcji nowoczesnej elektrody ICD (<https://t.ly/-n0G>)

- **poliuretan** – biogodny, odporny na rozciąganie, szczególnie chętnie wykorzystywany z uwagi na bardzo niski współczynnik tarcia, znacząco ułatwiający przesuwanie kabla wewnątrz naczyń podczas implantacji. Do wad poliuretanu należy tendencja do tworzenia mikropęknięć oraz reaktywność z jonami metali. Przykładowy poliuretan termoplastyczny (TPU), wykorzystywany w urządzeniach implantowalnych, to Biomedical Elasthane™ 55D marki DSM Biomedical B.V.
- **silikon** – obojętny, biogodny i biostabilny materiał o wysokiej elastyczności, charakteryzujący się jednak dość dużym współczynnikiem tarcia (w porównaniu do poliuretanu oraz PTFE) oraz mniejszą odpornością na uszkodzenia mechaniczne podczas implantacji (np. w wyniku kontaktu z ostrym narzędziem). Pewnym problemem jest także tendencja silikonu do pęcznienia w kontakcie z płynami, choć i ta cecha materiału została wykorzystana praktycznie. Niektórzy producenci oferują bowiem specjalne... płyny spęczniające, ułatwiające pracę z silikonowymi rurkami o małym przekroju (poprzez chwilowe rozszerzenie światła rurki znacznie łatwiej jest umieścić w jego wnętrzu montowane elementy). Jako przykład wysoce biogodnego silikonu medycznego można wskazać dwuskładnikowy MED-4719 marki NuSil.
- **fluoropolimery (PTFE oraz ETFE)** – oprócz bardzo wysokiej biogodności i minimalnego współczynnika tarcia, zapewniają także możliwość zmniejszenia wymiarów (przekroju) elektrody dzięki zastosowaniu cienkich warstw materiału izolacyjnego. Do wad tej grupy materiałów należy stosunkowo wysoka sztywność oraz trudności stwarzane przez nie w procesie produkcyjnym. Warto dodać, że polimer GORE™ ePTFE (extended PTFE) był nawet stosowany do dodatkowego pokrywania elektrod ICD, co znacząco ułatwiało ich usuwanie z ciała pacjenta – dzięki śliskiej powierzchni możliwe jest bowiem zredukowanie ryzyka obrastania elektrody przez otaczające tkanki [1].

Do produkcji przewodzących części elektrod stosuje się szereg specjalistycznych stopów, do których zaliczyć należy przede wszystkim niklowo-kobaltowy stop MP35N (NiCoCrMo) o doskonałej odporności na korozję i wysokiej wytrzymałości mechanicznej, jednak – niestety – także dość wysokiej oporności właściwej. Z tego też względu typowa rezystancja (zaledwie półmetrowych) przewodów elektrod wynosi aż kilkadziesiąt omów. Elementy pozostające w bezpośrednim kontakcie z tkanką miokardium są natomiast wykonywane na bazie stopów platyny (np. platyna-iryd), dodatkowo platynowanych lub pokrywanych warstwą azotku tytanu. Wyjątkiem są wysokonapięciowe przewody elektrod, stosowanych w kardiowerterach-defibrylatorach – w tym przypadku, w celu obniżenia strat transmisyjnych, stosowane są stopy o niskiej rezystancji, bazujące m.in. na platynie i irydzie bądź srebrze.

Zarówno elektrody z mocowaniem pasywnym, jak i modele wyposażone w końcówkę do fiksacji aktywnej, są dziś powszechnie wyposażane w nośniki leków steroidowych (gł. octanu deksametazonu – DXA), którego celem jest zapobieganie ostremu lub przewlekłemu zapaleniu tkanki wokół elektrody. Stany zapalne w dłuższym okresie eksploatacji rozrusznika powodują bowiem powstawanie zwłóknień otaczających punkt implantacji elektrody i prowadzących do stopniowego pogarszania kontaktu elektrycznego katody z mięśniem sercowym, co w efekcie podwyższa wymagany próg czułości stymulatora oraz zwiększa poziom energii, niezbędnej do efektywnego wywołania skurczu. Nietrudno domyślić się zatem, że konsekwencją stosowania aktywnych farmakologicznie elektrod jest... oszczędność energii i wydłużenie czasu pracy stymulatora. Pierścienie zawierające nośniki leków – zarówno w elektrodach do mocowania aktywnego, jak i pasywnego – zostały odpowiednio oznaczone na rysunkach 8 i 9.

Rozruszniki bezelektrodowe

W ostatnich latach rynek aparatury medycznej obiegły informacje o wprowadzeniu do użycia klinicznego całkowicie nowego typu rozrusznika, określanego mianem bezelektrodowego. Pojęcie to może być



Fotografia 9. Porównanie rozmiarów dwóch rozruszników bezelektrodowych: Nanostim (St. Jude) i Micra TPS (Medtronic). Źródło: <https://t.ly/gjdm>

nieco mylące z technicznego punktu widzenia, gdyż urządzenie istotnie posiada elektrody, jednak są one na stałe przytwierdzone do obudowy i – zamiast długich przewodów wprowadzanych do serca przez naczyniowo – mają postać specjalnych haczyków, mocujących rozrusznik... bezpośrednio do ściany prawej komory serca, gdyż niewielkie wymiary urządzenia pozwalają wprowadzić go przez żyłę udową, bez konieczności wykonywania kieszonki dla implantu pod skórą klatki piersiowej.

Porównanie rozmiarów dwóch wprowadzonych do klinicznego użytku konstrukcji rozruszników bezelektrodowych – Micra TPS (wyprodukowanego przez firmę Medtronic) oraz Nanostim (marki St. Jude Medical) zostały pokazane na **fotografii 9**. Ten pierwszy – w chwili pisania niniejszego artykułu najmniejszy elektrostymulator na świecie – ma wymiary 25,9 mm (długość) × 6,7 mm (średnica obudowy) i masę zaledwie 2 gramów (!), zaś do zamocowania w tkance miokardium służą „wąsy” przypominające nieco system stosowany w konwencjonalnych elektrodach pasywnych. Dla porównania – Nanostim ma długość aż 42 mm i średnicę niewiele mniejszą, bo równą 5,99 mm, zaś jego mocowanie opiera się na wkręcającej końcówce, działającej w sposób zbliżony do elektrod przeznaczonych do fiksacji aktywnej. Co ciekawe, górna granica żywotności (czasu pracy) zaledwie nieznacznie różni się na korzyść urządzenia Nanostim (9,8 roku w porównaniu do 9,6 roku dla Micra TPS).

Obecnie zdania odnośnie użycia rozruszników bezelektrodowych są podzielone – pomimo oczywistych zalet (związanych m.in. z redukcją ilości dostępów chirurgicznych czy też doskonałym efektem estetycznym, ograniczającym ranę pozabiegową tylko do pachwiny), wątpliwości budzą trudności w eksplantacji (usunięciu) takiego aparatu z wnętrza serca, czy też ryzyko perforacji mięśnia sercowego podczas implantacji rozrusznika. Firma Medtronic, w ostrzeżeniach dotyczących potencjalnych komplikacji podczas procedury wszczepiania urządzenia Micra TPS, zwraca uwagę, iż usunięcie urządzenia może być utrudnione z uwagi na otaczające je zwłóknienia w tkance mięśnia sercowego. Obejściem problemu jest zatem... całkowite wyłączenie stymulatora i pozostawienie go w sercu do końca życia pacjenta – trudno byłoby bowiem



Fotografia 10. Szczegóły „anatomiczne” rozrusznika Micra TPS (<https://t.ly/-x5h>)



CAUTION: The Aveir™ DR dual-chamber leadless pacemaker is an investigational device limited by Federal (U.S.) local law and Medical Device Regulation for investigational use only.

Fotografia 11. Dwojamowy system kardiostymulacji oparty na synchronizowanych bezprzewodowo rozrusznikach bezelektrodowych (<https://t.ly/h4h->)

wyobrazić sobie otwieranie klatki piersiowej (przez sternotomię lub przynajmniej torakotomię) i samego serca jedynie w celu usunięcia starego rozrusznika.

Ciekawe rozwiązanie w zakresie kardiostymulatorów bezelektrodowych opracowała firma Abbott. Dostrzegłszy problem ograniczenia zastosowań tego rodzaju urządzeń jedynie do sytuacji klinicznych, w których wystarczająca jest stymulacja jednojamowa, inżynierowie stworzyli system złożony z... dwóch rozruszników (niewiele różniących się pomiędzy sobą rozmiarami), z których jeden jest przeznaczony do implantacji w komorze, zaś drugi – w przedsionku (**fotografia 11**). Obydwa urządzenia komunikują się pomiędzy sobą bezprzewodowo, co pozwala na w pełni synchroniczną pracę pomimo braku połączenia galwanicznego. Ten interesujący system został po raz pierwszy wszczepiony u człowieka w klinice w Cleveland (USA) na początku ubiegłego roku, zaś eksperymentalny zabieg był częścią ogólnoustawowego badania klinicznego.

Kolejną innowacją, o której zdecydowanie warto wspomnieć, jest druga generacja rozruszników bezelektrodowych firmy Medtronic – Micra™ AV. W tym przypadku producent zdecydował się na wykorzystanie... akcelerometru w celu wykrywania aktywności przedsionków. Takie rozwiązanie umożliwia pracę w trybie VDD (stymulacja komory, detekcja w obu jamach), nawet pomimo braku fizycznego (galwanicznego) połączenia z prawym przedsionkiem – sygnałem informującym o skurczu przedsionków są zafalowania w odczycie z czujnika przyspieszenia. Micra™ AV jest zatem pierwszym kardiostymulatorem na świecie, który do sensingu wykorzystuje metodę inną, niż elektrogram.

Złącza elektrod

Obszar technologii skupionych wokół terapii bioelektrycznej serca – podobnie, jak niemal wszystkie inne dziedziny techniki medycznej – podlega silnej standaryzacji, z korzyścią zarówno dla lekarzy (uproszczenie procedur implantacji urządzeń i większa elastyczność w kompletowaniu zaopatrzenia), jak i pacjentów (możliwość dobrania optymalnych akcesoriów dla potrzeb danego pacjenta). Nie inaczej jest w przypadku gniazd i wtyków, stosowanych do łączenia elektrod z urządzeniami (rozrusznikami oraz ICD).

Współczesne kardiostymulatory wszczepialne korzystają niemal wyłącznie ze standardu określanego jako IS-1 lub IS-4. Złącza IS-1 (**rysunek 13**) występują w dwóch wersjach: IS-1 Uni (unipolarnej) oraz IS-1 Bi (bipolarnej) – choć wymiarowo są takie same, różnią

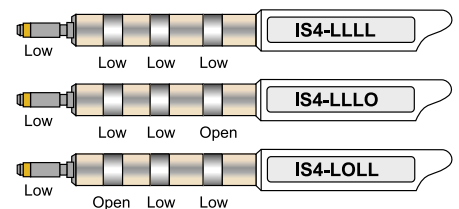


Rysunek 13. Porównanie konstrukcji złączy IS-1 Bi (na górze) i IS-1 Uni (na dole). Czerwone elementy oznaczają pola stykowe, niebieskie – uszczelnienia, pozostała część to izolacja przewodu oraz złącza (<https://t.ly/nwUq>)



Fotografia 12. Stymulator resynchronizujący (CRT-P) marki Boston Scientific. Widoczne oznaczenia gniazd elektrod: dwa główne złącza IS-1 (dla prawego przedsionka i prawej komory) oraz jedno IS-4 w konfiguracji LLLL (<https://t.ly/4Gb5>)

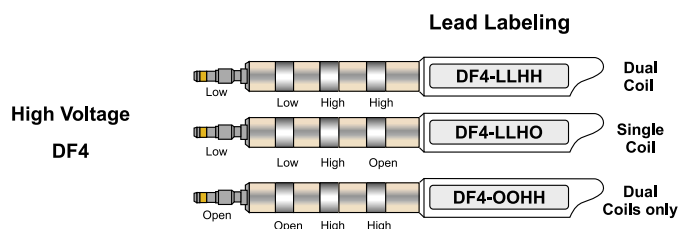
Low Voltage
IS4



Rysunek 14. Przykładowe oznaczenia układu połączeń wtyków elektrodowych typu IS-4. Low oznacza niskie napięcie, open – brak połączenia (<https://t.ly/MHw->, <https://t.ly/D4Q0>)

się obecnością dodatkowego pierścienia (w wersji bipolarnej), służącego do podłączenia anody. Podobną koncepcję konstrukcyjną zastosowano w przypadku czteroelektrodowego złącza IS-4, wykorzystywanego w bardziej rozbudowanych stymulatorach (głównie CRT – patrz **fotografia 12**). Oprócz oznaczenia typu złącza, na obudowach kardiostymulatorów często można też spotkać dodatkowy kod, określający poziom potencjałów oraz – w ogóle – obecność połączenia elektrycznego z danym stykiem gniazda. Przykładowe konfiguracje zostały pokazane na **rysunku 14**.

Złącza IS-1 i IS-4 są przeznaczone do pracy z niskimi napięciami (zwykle rzędu kilku woltów), stosowanymi w kardiostymulatorach. Wszczepialne kardiowertery-defibrylatory wymagają jednak rozwiązań znacznie bardziej zaawansowanych pod względem wytrzymałości dielektrycznej. Dlatego też urządzenia ICD oraz CRT-D korzystają z dedykowanych złączy typu DF-1 oraz DF-4 – ich budowa jest bardzo zbliżona do IS-1/IS-4, nawet gołym okiem można jednak dostrzec pewną różnicę – wtyki DF-4 mają końcowy pin o zmniejszającej się skokowo średnicy, podczas gdy złączach IS-4 stopniowanie średnicy nie występuje (pin ma kształt walcowy). Na **rysunku 15** pokazano spotykane w praktyce konfiguracje połączeń DF-4, w tym wersje mieszane (wyprowadzenia niskonapięciowe i defibrylacyjne w ramach tej samej elektrody).



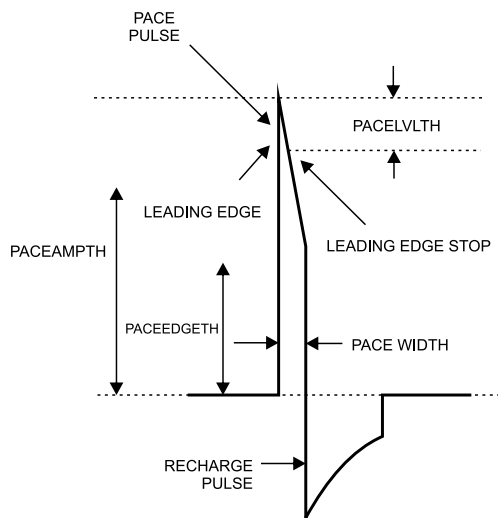
Rysunek 15. Przykładowe oznaczenia układu połączeń wtyków elektrodowych typu DF-4. Low oznacza niskie napięcie, high – wysokie napięcie, open – brak połączenia (<https://t.ly/MHw->, <https://t.ly/D4Q0>)

Technical drawing of a shaft assembly with dimensions in millimeters. The drawing shows a shaft with a central section (3) and end sections (5, 4, 8). Keyways are shown in sections 5, 4, and 8. A table of dimensions is provided above the shaft. The dimensions are: 18.7, 17.7, 16.04, 13.43, 11.53, 9.92, 7.02, 4.41, 2.51. The shaft is labeled with A, B, C, and D. The dimensions are in millimeters.

Section	Dimension 1	Dimension 2	Dimension 3	Dimension 4	Dimension 5	Dimension 6	Dimension 7	Dimension 8	Dimension 9
5	18.7	17.7	16.04	13.43	11.53	9.92	7.02	4.41	2.51
4	18.7	17.7	16.04	13.43	11.53	9.92	7.02	4.41	2.51
8	18.7	17.7	16.04	13.43	11.53	9.92	7.02	4.41	2.51

[illegible]

eprasa.pl f462ec3a6c



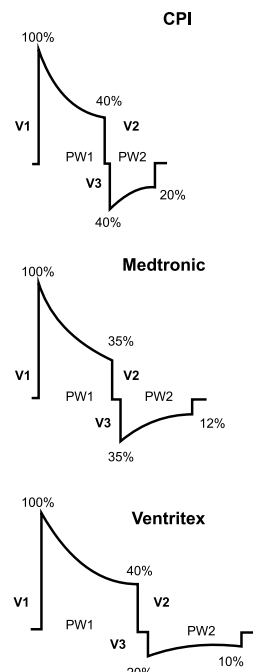
Rysunek 20. Typowy, bifazowy impuls kardiostymulatora
(<https://t.ly/5ZWH>)

jakie muszą spełniać urządzenia (oraz ich kluczowe komponenty), aby mogły one skutecznie realizować powierzone zadanie kliniczne.

W przypadku sztucznych rozruszników serca typowy impuls ma kształt dwufazowy (**rysunek 20**) – „główna” część przebiegu wywołuje pożądany skurcz mięśnia sercowego, zaś następujący po niej, dłuższy impuls o znacznie niższej amplitudzie ma za zadanie (mówiąc w ogromnym uproszczeniu) „rozładować” pojemności tkanek serca w celu uniknięcia zjawiska polaryzacji DC. Typowa szerokość impulsu głównego to 400...600 μ s, zaś amplituda to około 1,5...2,5 V. Warto jednak zwrócić uwagę, że rzeczywiste parametry u danego pacjenta są zależne od szeregu czynników, w tym przede wszystkim od miejsca implantacji i związanej z tym pobudliwości mięśnia na stymulację w danym punkcie. Podczas programowania rozrusznika ustala się zatem optymalne nastawy czasu i amplitudy, stąd też producenci kardiostymulatorów oferują dość szerokie zakresy tych parametrów – przykładowo, rozruszniki z serii Effecta marki Biotronik mogą pracować z amplitudami ustawianymi w zakresie od 200 μ V do 7,5 V i szerokościami impulsów od 100 μ s do 1,5 ms. Warto dodać, że współczesne kardiostymulatory umożliwiają automatyczne ustawianie, a także bieżącą korekcję amplitudy, co pozwala na kompensację zmian impedancji, wynikających z progresji narostu tkanek wokół katody.

Osobnym zagadnieniem jest czułość, wyrażana w miliwoltach – oznacza ona nic innego, jak tylko próg napięcia o określonej polaryzacji, którego przekroczenie (przez potencjały wewnątrzsercowe) zostanie zinterpretowane przez rozrusznik jako załamek P (depolaryzacja przedsionka) lub R (depolaryzacja komory). „Książkowa” wartość czułości przy zastosowaniu konfiguracji unipolarnej wynosi 2,5 mV (wg normy EN 45502-2-1) – zbyt niski próg prowadzi do zjawiska określanego mianem *oversensing* (stymulator reaguje na artefakty, interpretując je jako oczekiwaną aktywność danej jamy), zaś zbyt wysoki – *undersensing* (stymulator pozostaje nieczuły na aktywność m. sercowego, na którą powinien zareagować). W przypadku wspomnianej serii Effecta (Biotronik) czułość można ustawić w zakresie od 100 μ V do nawet 7,5 mV.

Oprócz amplitudy i szerokości impulsów, klinicznie istotne są także liczne parametry czasowe, związane z synchronizacją poszczególnych impulsów. Programowalne są zatem takie parametry, jak maksymalna częstość akcji serca w poszczególnych scenariuszach pracy (np. według zegara czasu rzeczywistego – rozróżnienie dzień/noc), opóźnienie AV (maksymalny czas pomiędzy wykryciem aktywności lub stymulacją przedsionka, a wyzwoleniem impulsu stymulującego komorę), czy też tzw. okresy refrakcji i wygaszania (ang. *blanking*) dla poszczególnych jam serca, określające (mówiąc w dużym uproszczeniu) czasową dezaktywację sensingu w poszczególnych kanałach po wystąpieniu w nich określonego „zdarzenia”.



Rysunek 21. Bifazowe przebiegi impulsów defibrylacji, stosowane przez trzech różnych producentów ICD, obecnych na rynku w latach 90. ub. wieku. Warto dodać, że podczas gdy urządzenia marek CPI i Medtronic korzystały z pojedynczych kondensatorów wysokonapięciowych, kardiowerter-defibrylator Ventritex łąadował dwa szeregowo kondensatory (po 300 µF każdy) w pierwszej fazie (V1) i tylko jeden z nich w fazie drugiej (V2). Źródło: <https://t.ly/S68d>

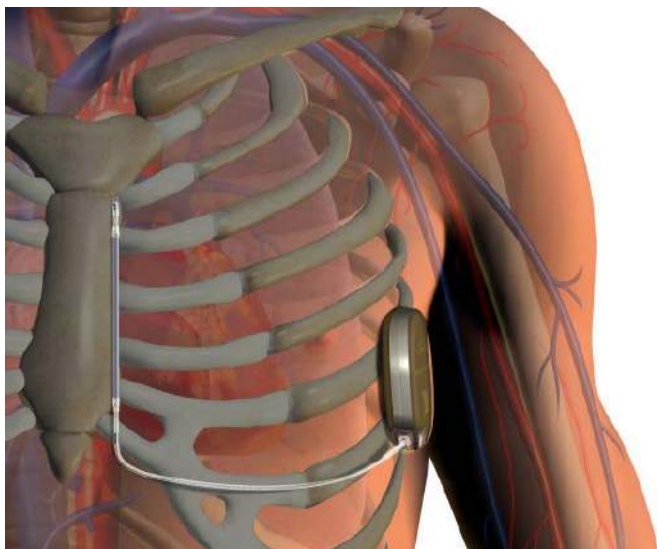
Ponieważ jednak szczegółowe wyjaśnienie znaczenia tych parametrów wymagałoby zagłębienia się w elektrofizjologię mięśnia sercowego oraz dokładnego omówienia poszczególnych trybów stymulacji, dlatego wspomnianych zagadnień nie będziemy szerzej omawiać na łamach niniejszego artykułu – zainteresowanych odsyłamy do bibliografii, znajdującej się na końcu tekstu.

W przypadku wszczepialnych kardiowerterów-defibrylatorów podstawowym parametrem jest energia wstrząsu, dostarczanego przez wysokonapięciowe elektrody do mięśnia sercowego, wyrażana w dżulach [J]. Impulsy stosowane we współczesnych ICD mają najczęściej postać dwufazową (**rysunek 21**) i w typowych warunkach dostarczają energię na poziomie od kilku do kilkunastu J. Większość kardiowerterów-defibrylatorów oferuje energię maksymalną na poziomie do 35 J, zaś w przypadku modeli określanых jako *ultra-high energy* górną granicą może być nawet wartość 42 J. Uzyskanie tak dużej ilości energii jest możliwe dzięki zastosowaniu przetwornicy typu flyback, podwyższającej napięcie baterii (rzędu 2,8...3,2 V) do kilkuset woltów, przy czym warto wiedzieć, że narastające zbrocze impulsu terapeutycznego może wiązać się z przepływem prądu na poziomie nawet kilkunastu amperów (!). Wartości energii odpowiadającej określonym napięciom szczytowym można znaleźć w materiałach producentów – przykładowe dane dla ICD z serii Inventa 7 marki Biotronik podano w **tabeli 1**.

Warto dodać, że od niedawna w użyciu są także defibrylatory typu S-ICD (*subcutaneous ICD*). Od klasycznych

Tabela 1. Napięcia szczytowe i energie wstrząsu dla trzech różnych poziomów docelowych – dane z instrukcji ICD Inventra 7 (Biotronik). Źródło: <https://t.ly/VaK1>

Energia wstrząsu (docelowa)	Napięcie szczytowe (zakres tolerancji)	Wartość energii wstrząsu (zmierzona)	Wartość napięcia szczytowego (zmierzona)
1 J	90...120 V	0,82 J	92 V
22 J	440...480 V	19,8 J	458 V
45 J	620...690 V	41,3 J	667 V



Rysunek 22. Ilustracja miejsca implantacji kardiowertera-defibrylatora typu S-ICD oraz współpracującej z nim elektrody podskórnej (<https://t.ly/8CQt>)

kardiowerterów-defibrylatorów różnią się one przede wszystkim rodzajem elektrody, która zamiast śródsierdziowo, implantowana jest podskórnie, wzdłuż lewego brzegu mostka (**rysunek 22**). Z uwagi na znacznie wyższą impedancję, wprowadzaną przez warstwy mięśni, osierdziej, etc. (w porównaniu do impedancji widzianej przez elektrody ICD wszczepione przeznaczyńowo), konieczne jest użycie wyższych poziomów energii defibrylacji na poziomie do 80 J – choć, z uwagi na brak bariery skórnej – jest to i tak 4...5-krotnie niższa energia w por. do stosowanych w medycynie ratunkowej i intensywnej terapii defibrylatorów zewnętrznych. Urządzenia S-ICD są stosowane u pacjentów, którzy z pewnych przyczyn (np. w związku z anatomicznymi patologiami naczyniowymi lub ogólnym stanem zdrowia) nie mogą być poddani implantacji klasycznego ICD.

Dodatkowe funkcje stymulatorów i kardiowerterów-defibrylatorów

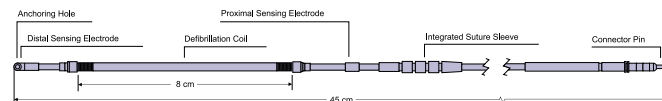
Producenci implantów aktywnych dokładają wszelkich starań, by – już po zakończeniu procedury implantacji – możliwe było dokonywanie niezbędnych korekt parametrów terapii, okresowych przeglądów i diagnostyki ewentualnych problemów. Wymiana lub usunięcie urządzenia wiąże się bowiem nieodłącznie z koniecznością reoperacji, co sprawia, że wszystkie możliwości oferowane przez zewnętrzne programatory są od zawsze w cenie. Mało tego – miniaturyzacja pamięci półprzewodnikowych i opracowanie kolejnych generacji energooszczędnych układów sterujących umożliwiło wprowadzenie do implantów rozbudowanych funkcji rejestracji epizodów zaburzeń rytmu serca, które w doskonałej jakości mogą być odtworzone przez lekarza i zbadane w celu szczegółowej oceny. Zapis krótkich, zwykle około 10-sekundowych fragmentów elektrogramu, umożliwia wprowadzenie niezbędnych zmian w parametrach elektroterapii, czy też towarzyszącej jej farmakoterapii.

Programowanie

Programowanie parametrów urządzenia – dawniej przeprowadzane za pośrednictwem cewki indukcyjnej przykładowej w okolicę miejsca implantacji – dziś zostało zdominowane przez krótkozasięgowe łącza radiowe. Ustawianie i kontrola

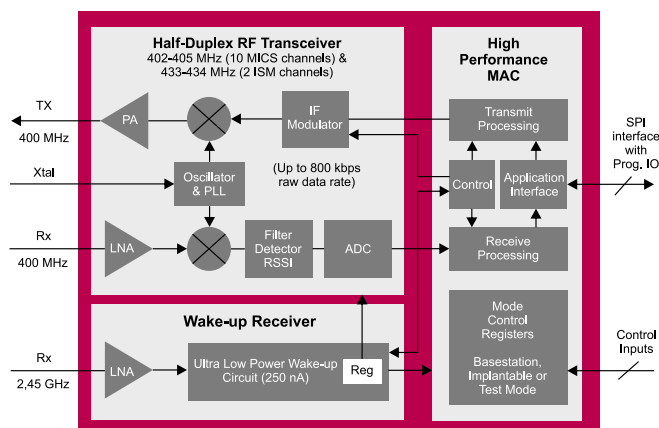


Fotografia 13. Kardiowerter-defibrylator S-ICD z serii EMBLEM™ marki Boston Scientific (<https://t.ly/mGd->)



Rysunek 23. Elektroda EMBLEM™ S-ICD typu 3501 (Boston Scientific). Źródło: <https://t.ly/Eamo>

ZL70101 Simplified Block Diagram



Rysunek 24. Schemat blokowy transceiwera ZL70101, przeznaczego do urządzeń implantowalnych (<https://t.ly/uRWM>)

parametrów stymulatora oraz odczyt rozbudowanego pakietu danych diagnostycznych (dotyczących tak samego urządzenia, jak i fizjologii pacjenta) jest podstawową funkcjonalnością wszystkich współczesnych implantów elektrokardiologicznych. Standardem komunikacyjnym, który ukształtował dzisiejszy świat elektroniki implantowalnej, jest MICS – *Medical Implant Communication System* – opracowany przez FCC w 1999 roku i „wchłonięty” przez nowszy, udostępniony 10 lat później standard MedRadio (*Medical Device Radiocommunications Service*). Pierwszy z nich opierał się na paśmie 402–405 MHz, drugi rozszerzył je o 1 MHz z każdej strony. Kolejne edycje uzupełniły całość o podzakres pasma S, wykorzystywany w sieciach MBAN (*Medical Body Area Network*) – 2360–2400 MHz, a także o kolejne 24 megaherce w podprzedziałach 413–419 MHz, 426–432 MHz, 438–444 MHz oraz 451–457 MHz. Z biegiem czasu na rynku pojawiły się nawet dedykowane do celów medycznych transceiwery scalone, np. układ ZL70101 (**rysunek 24**), intensywnie wykorzystywany m.in. przez firmę Biotronik w zastosowaniach telemetrycznych. Transceiver zapewnia dwukierunkową komunikację w paśmie 400 MHz oraz dodatkowy front-end 2,45 GHz, umożliwiający wybudzenie układu ze stanu uśpienia przez zewnętrzny kontroler.

Telemetria

Telemetria umożliwia zdalną kontrolę parametrów pracy stymulatora lub kardiowertera-defibrylatora i coraz częściej stanowi ważny element telemedycznej opieki nad pacjentem kardiologicznym. Dawniej wprowadzane były rozwiązania stacjonarne, takie jak Merlin@Home marki St. Jude Medical (przejętej przez koncern Abbott) – patrz **fotografia 14**. Urządzenie – do dziś dostępne w ofercie producenta – umożliwia zdalny odczyt danych przez lekarza, przy użyciu utrzymywanej przez Abbott platformy Merlin.net™ Patient Care Network. Informacje zebrane z implantu są przesyłane poprzez łącze telefoniczne, co z jednej strony zapewnia łatwość instalacji i zwiększa dostępność dla osób starszych, z drugiej – ogranicza możliwość używania transmittera przez pacjenta jedynie do czasu, w którym przebywa on w domu.

Nowsze rozwiązania korzystają z podejścia, które dziś – w dobie IoT, czy nawet IoE – chyba nikogo już nie dziwi: łącze Bluetooth Low Energy umożliwia komunikację z rozrusznikiem lub ICD przy użyciu dowolnego urządzenia mobilnego, rzecz jasna w powiązaniu z przeznaczoną do tego celu aplikacją mobilną. Jako przykład należy tutaj wskazać technologię BlueSync, opracowaną przez markę Medtronic – przesył informacji odbywa się z użyciem silnego szyfrowania



Fotografia 14. Merlin@Home™ – urządzenie telemetryczne St. Jude Medical (obecnie Abbott). Źródło: <https://t.ly/uf1x>

w standardzie NIST, zaś w roli bramy sieciowej może wystąpić nie tylko smartfon lub tablet pacjenta, lecz także specjalne urządzenie – MyCareLink Relay™ (fotografia 15), które pod odebraniem danych z implantu może przesyłać je za pomocą sieci WiFi albo połączenia LTE.

Kontrola impedancji

Automatyczna kontrola impedancji pozwala na zaawansowaną diagnostykę jakości kontaktu elektrod z tkanką miokardium i stanowi jedną z podstawowych funkcji współczesnych implantów kardiologicznych. W tym celu generator jest przestawiany w tryb niskiej amplitudy, wytwarzając impulsy podprogowe, których poziom jest zbyt niski, aby wywołać skutki fizjologiczne (skurcz mięśnia sercowego), ale wystarczający do sprawdzenia, czy impedancja tkanek widziana od strony wyjścia urządzenia mieści się w dopuszczalnym przedziale. Dane na temat impedancji są nie tylko wykorzystywane przez procesor urządzenia – mogą być także odczytywane przez lekarza w trakcie rutynowej kontroli działania implantu, a co więcej – w niektórych przypadkach nawet wysyłane za pomocą łącza telemetrycznego.

IEGM

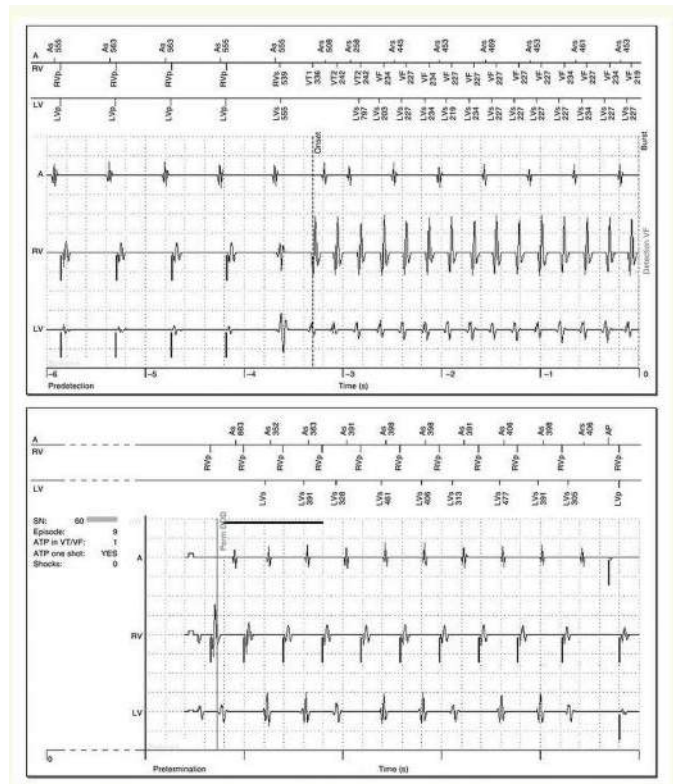
IEGM (ang. intracardiac electrogram) to inaczej zapis wewnątrzsercowego elektrogramu, czyli spontanicznej aktywności serca rejestrowanej przez front-end wejściowy (kanały sensingu). Funkcja odczytu IEGM jest niezwykle użyteczna, pozwala bowiem lekarzowi zobaczyć dokładnie to, co „widzi” rozrusznik za pomocą wszczepionych elektrod. Współczesne urządzenia są w stanie rejestrować zapisy IEGM z wysoką rozdzielczością, co daje niespotykane wcześniej możliwości zaawansowanego „debugowania” urządzeń po wszczępieniu w ciało pacjenta oraz szczegółowego określania charakteru wykrytych zaburzeń serca (np. arytmii) – patrz rysunek 25.

Autoregulacja rytmu względem aktywności fizycznej

Fizjologiczną reakcją organizmu na wysiłek fizyczny lub stany emocjonalne jest przyspieszenie częstości akcji serca. W przypadku „awarii” naturalnego rozrusznika (węzła zatokowo-predsionkowego) odpowiedzialność za ten parametr spoczywa na kardiostymulatorze – generowanie impulsów ze stałą częstotliwością uniemożliwiłoby bowiem właściwe ukrwienie organów wewnętrznych i mięśni szkieletowych (co było zresztą jedną z głównych wad



Fotografia 15. Przekaznik telemetryczny MyCareLink Relay™ marki Medtronic (<https://t.ly/ToIs>)



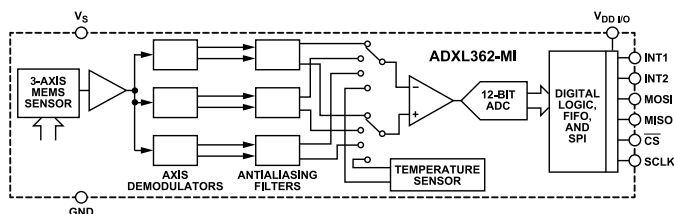
Rysunek 25. Przykładowy elektrogram wewnątrzsercowy, zarejestrowany przez kardiowerterdefibrylator Lumax HF-T marki Biotronik (<https://t.ly/fnjb>)

stymulatorów w pierwszych dwóch dekadach ich rozwoju). Dziś regulacja rytmu sterowana przez sygnał z wbudowanego akcelerometru jest standardem we współczesnych rozrusznikach. Rzecz jasna, koniecznością jest tutaj zastosowanie niskomocowego czujnika, który nie nadwerży budżetu energetycznego urządzenia. Firma Analog Devices opracowała czujnik spełniający to wymaganie – układ ADXL362-MI (rysunek 26) pobiera zaledwie 1,8 μ A prądu zasilania podczas pomiaru z częstotliwością próbkowania równą 100 Hz, zaś w trybie aktywacji ruchem – zaledwie 270 nA. Jest to zresztą jeden ze stosunkowo niewielu układów, przeznaczonych explicite do urządzeń implantowalnych, czy w ogóle do wyrobów medycznych klasy III (tematyka metod doboru komponentów do takich aplikacji nadawałaby się zresztą co najmniej na osobny artykuł).

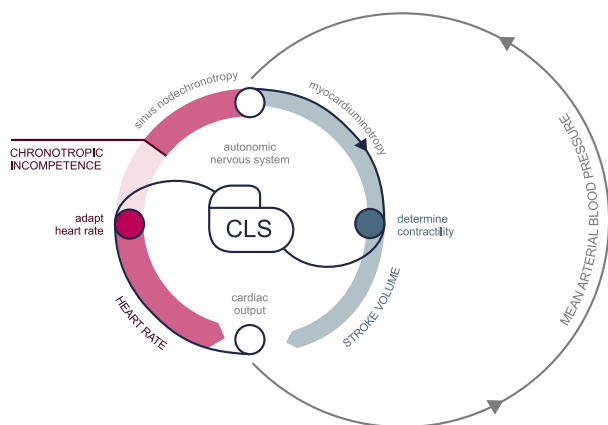
Warto dodać, że interesującą technologię – wykraczającą daleko poza (dość proste) dopasowywanie szybkości stymulacji do zapisów z akcelerometru – opracowała firma Biotronik w ramach funkcjonalności, określanej jako CLS – Closed Loop Stimulation. Jak sama nazwa wskazuje, regulacja rytmu odbywa się w zamkniętej pętli sprzężenia zwrotnego, zaś sygnałem wejściowym dla algorytmu jest impedancja, mierzona w czasie rzeczywistym, w każdym cyklu pracy serca (rysunek 27). Technologia opiera się na wykorzystaniu tzw. regulacji inotropowej – kurczliwość mięśnia sercowego (której zmiany mają swoje odzwierciedlenie także w jego impedancji, mierzonej unipolarnie przez elektrodę komorową) – jest regulowana przez centralny układ nerwowy, w odpowiedzi na czynniki, które w prawidłowej fizjologii prowadzą do przyspieszenia akcji serca. Z tego też względu technologia CLS działa nie tylko podczas zwiększonego wysiłku fizycznego, ale także pobudzenia emocjonalnego, dzięki czemu stymulacja ma charakter znacznie bardziej zbliżony do naturalnego rytmu zatokowego.

Tryb magnesu

Magnet mode – czyli tryb magnetyczny lub tryb magnesu – jest niezmiennie stosowany od kilku dekad zarówno w rozrusznikach serca, jak i kardiowerterach-defibrylatorach. Zbliżenie magnesu



Rysunek 26. Schemat blokowy akcelerometru ADXL362-MI marki Analog Devices (<https://t.ly/ZOfA>)



Rysunek 27. Stymulacja w pętli zamkniętej (Closed Loop Stimulation) opiera się na „wpięciu” rozrusznika w pętle autoregulacyjne organizmu, przy czym sygnałem wejściowym jest impedancja mierzona przez elektrodę komorową, zaś wyjściowym – częstość stymulacji. Źródło: <https://t.ly/U3vW>

stałego (fotografia 16) do wszczepionego urządzenia powoduje przejście w bezpieczny tryb terapii, nieczuły na zakłócenia w kanałach sensingu (np. na artefakty wynikające z użycia noży elektrochirurgicznych w trakcie operacji). Choć dokładne funkcjonowanie urządzenia zależy od jego modelu oraz nastaw wgranych przez lekarza, to ogólna zasada jest następująca: w rozrusznikach dezaktywowane są kanały wejściowe służące do detekcji spontanicznej aktywności serca (co oznacza włączenie trybu asynchronicznego, np. AOO, VOO czy DOO), zaś w urządzeniach typu ICD dezaktywowana jest funkcja wstrząsów terapeutycznych (przy pozostawieniu działania ewentualnych funkcji rozrusznika).

Warto w tym miejscu dodać, że tryb magnetyczny (aktywowany automatycznie po wykryciu pola magnetycznego przekraczającego określoną wartość indukcji progowej, np. 1,5 mT) nie stanowi wystarczającego zabezpieczenia przed wpływem pól obecnych podczas badania pacjenta metodą rezonansu magnetycznego (MRI). Przed rozpoczęciem badań obrazowych urządzenie powinno być bowiem przestawione przez kardiologa w specjalny tryb pracy za pomocą zewnętrznego programatora. Kwestia kompatybilności implantów z MRI jest ponadto zagadnieniem znacznie szerszym – urządzenia dopuszczone warunkowo do pracy w środowisku pól magnetycznych obecnych w gantrze tomografu MRI (określane jako *MRI-conditional*) muszą być bowiem od początku do końca zaprojektowane i przebadane pod kątem bezpieczeństwa. Co ciekawe, wbrew pozorom to nie silne, stałe pole magnetyczne nadprzewodzącego elektromagnesu jest tutaj czynnikiem krytycznym – znacznie bardziej groźne dla implantu (a więc także samego pacjenta) są bowiem szybkozmienne pola cewek gradientowych oraz promieniowanie elektromagnetyczne (RF), które nie tylko mogłyby z łatwością zakłócić pracę obwodów urządzenia, ale także doprowadzić do nieakceptowalnego przegrzania obudowy lub wywołać mechaniczne drgania komponentów wewnętrznych.

Konstrukcja obudowy

Hermetyczna, biogodna, niemagnetyczna obudowa implantu jest jednym z kluczowych aspektów w kwestii bezpieczeństwa oraz

niezawodności urządzenia. Od wielu lat standardem jest umieszczanie urządzeń implantowalnych w lekkich, biokompatybilnych obudowach tytanowych. Ich budowa opiera się zwykle na konstrukcji dwuczęściowej – w jednej z precyzyjnie wyprofilowanych pokryw umieszczana jest płytka drukowana (zwykle wykonywana w formie sztywno-giętkiej) wraz z zasilaniem (baterią i – w przypadku ICD – kondensatorem wysokonapięciowym) oraz specjalnym, hermetycznym przepustem (fotografia 17), służącym do połączenia ze stykami, znajdującymi się w bloku złączy elektrod (określanym często mianem *header*). Po dociśnięciu pokryw są one następnie spawane laserowo (rysunek 28), zaś na przedniej części wykonuje się oznakowania za pomocą laserowej grawerki CNC (fotografia 18). Do odizolowania wewnętrznej powierzchni obudowy od znajdującej się w niej elektroniki producenci powszechnie używają osłony kaptonowej, co zresztą nie dziwi z uwagi na doskonałe właściwości dielektryczne tego materiału. Typową konstrukcję przykładowego stymulatora wszczepialnego pokazano na rysunku 29.

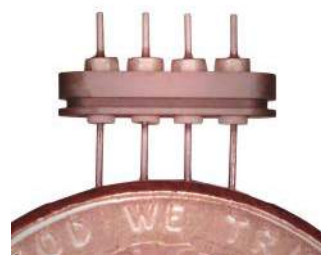
Osobnym zagadnieniem jest produkcja bloku złączy. W tym przypadku istnieje kilka możliwych technik wykonania, zaś sam proces stanowi przedmiot szeregu międzynarodowych patentów. Jedną z możliwości jest umieszczenie elementów przewodzących (styków i wyprowadzeń) na specjalnych rdzeniach w formie wtryskowej – po wykonaniu wtrysku rdzenie są usuwane, pozostawiając puste kanały, do których wsuwane będą złącza elektrod. Blok złączy, wykonany z tworzywa sztucznego, okazuje się ponadto doskonałym miejscem do zainstalowania anteny, co zobrazowano na rysunku 30.

Baterie

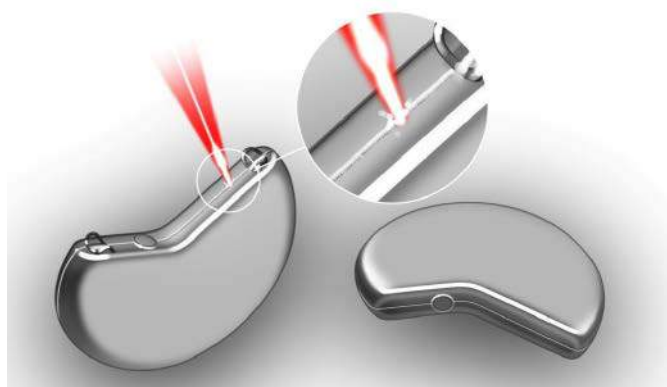
Pierwsze rozruszniki z zasilaniem akumulatorowym stwarzały nieakceptowalne ryzyko w przypadku, gdyby pacjent zapomniął (lub z jakichś względów nie mógł) doładować urządzenia – u osób „zależnych” od rozrusznika (*pacemaker-dependent*) rozładowanie źródła zasilania stanowiłoby jednak śmiertelne zagrożenie. Między innymi z tego względu szybko rozpoczęto poszukiwania



Fotografia 16. Magnes pierścieniowy marki Medtronic, przeznaczony do pracy z urządzeniami wszczepialnymi (https://t.ly/P_f1)



Fotografia 17. Przepust hermetyczny do aplikacji implantowalnych (<https://t.ly/KjOr>)



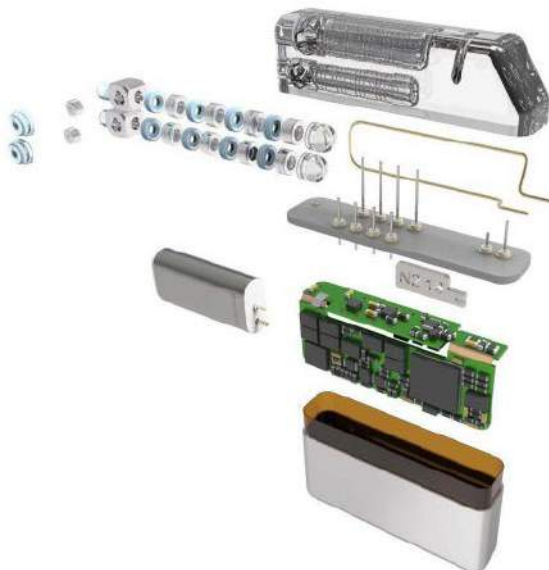
Rysunek 28. Ilustracja procesu spawania laserowego tytanowej obudowy implantu (<https://t.ly/CNGE>)

baterii, które byłyby w stanie sprostać potrzebom, stawianym przed inżynierami przez technologię urządzeń implantowalnych. Ze względów bezpieczeństwa współczesne baterie są wykonywane w hermetycznych obudowach (fotografia 19), przy czym metalowa puszką jest zwykle połączona galwanicznie z jedną z elektrod (często z anodą, choć nie jest to uniwersalna reguła). Podsumowanie najważniejszych parametrów pięciu typów baterii stosowanych w urządzeniach wszczepialnych zestawiono w tabeli 2.

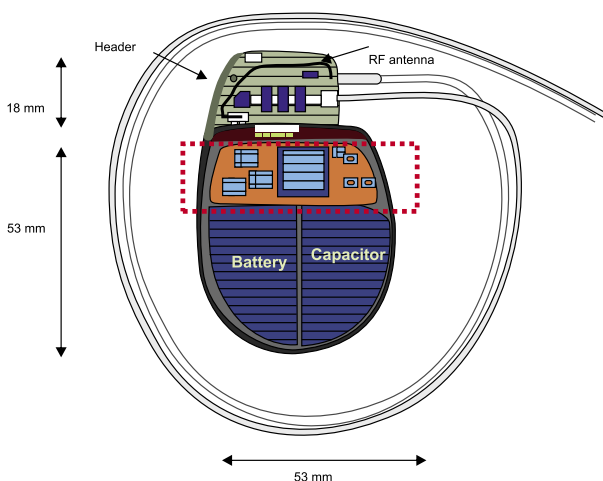
Dość szybko po stworzeniu pierwszych kardiostymulatorów opracowane zostały baterie litowo-jodowe, określane skrótowo Li/I_2 , a ściślej



Fotografia 18. Znakowanie laserowe obudowy implantu (<https://t.ly/oaf7>)



Rysunek 29. Widok rozstrzelony przykładowego stymulatora wszczepialnego (<https://t.ly/Z7ZD>)



Rysunek 30. Widok 3D kardiowertera-defibrylatora. W górnej części można zauważyć antenę prętową, umieszczoną nad złączami elektrod (<https://t.ly/IEwC>)

rzecz ujmując – Li/I_2 -PVP, w których katoda jest wykonana z kompleksu jodu i poli-2-winylopirydyny. Do głównych zalet tego rodzaju ogniwi należy bardzo wysoki poziom bezpieczeństwa (potwierdzony przez bezproblemową eksploatację w milionach urządzeń przez kilka ostatnich dekad) oraz doskonała wolumetryczna gęstość energii na poziomie $1,0 \text{ Wh/cm}^3$. Baterie Li/I_2 -PVP nie nadają się jednak do urządzeń o zwiększonym poborze mocy, niezależnie od tego, czy ma on miejsce tylko okazjonalnie (np. w kardiowerterach-defibrylatorach), czy też fazy wyższego zapotrzebowania na moc są dłuższe (np. we wszczepialnych pompach infuzyjnych). Warstwa jodku litu (LiI), pełniąc rolę separatora oraz stałego elektrolitu, ma bowiem silną tendencję do zwiększania grubości w miarę rozładowywania ogniwa, co prowadzi do znacznego wzrostu impedancji wewnętrznej (choć istnieje możliwość ograniczenia tego efektu przez pokrycie litowej anody warstwą poliwinylpirydyny – amerykański patent W. Greatbatcha nr 3957533).

Z uwagi na tradycyjny już zakres zastosowań baterii litowo-jodowych – zdecydowanie najczęściej można je znaleźć w rozrusznikach – producenci podają zwykle charakterystyki rozładowania dla stałego obciążenia rezystancją $100 \text{ k}\Omega$. Nie ma w tym jednak nic dziwnego – większość kardiostymulatorów operuje właśnie w zakresie prądu (uśrednionego) w zakresie kilku... kilkunastu μA , zaś zwiększanie obciążenia bardzo szybko obniża napięcie ogniwa – patrz rysunek 31.

Przykładową baterię „katalogową” – LiS 2650 marki Litronik – pokazano na rysunku 32. Metalowa puszką, z jednej strony tradycyjnie zaokrąglona, ma wymiary całkowite $27,33 \times 25,58 \times 5,1 \text{ mm}$, co daje łączną objętość $2,7 \text{ cm}^3$. Nominalne napięcie ogniwa to $2,8 \text{ V}$, pojemność dla rozładowania rezystancją $100 \text{ k}\Omega$ wynosi $1,21 \text{ Ah}$, zaś poziom samorozładowania nie przekracza 7% dla 10-letniego okresu użytkowania.

Urządzenia o większym zapotrzebowaniu energetycznym często korzystają z baterii litowych wykonanych na bazie tlenku manganu (IV) – Li/MnO_2 . Oferują nieco wyższe napięcie nominalne ($3,0 \dots 3,2 \text{ V}$) i o połowę wyższą pojemność wolumetryczną, choć gęstość energii jest bardzo zbliżona do baterii litowo-jodowych. Zaletą ogniwi Li/MnO_2 , która dyskwalifikuje źródła na bazie jodu w aplikacjach takich jak ICD, jest zdolność do pracy z impulsami prądowymi rzędu kilku amperów, co ma niebywałe znaczenie dla procesu ładowania kondensatorów za pomocą przetwornic podwyższających napięcie. Baterie Li/MnO_2 występują w wersjach zoptymalizowanych zarówno dla urządzeń o średnim, jak i dużym poborze mocy, a przewidywana żywotność wynosi zwykle do 10 lat.

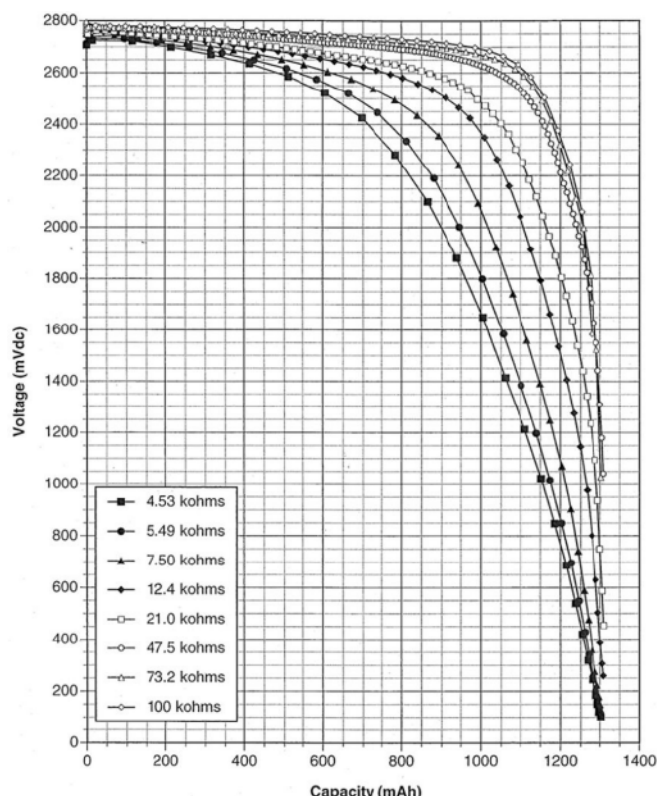
Baterie litowe mogą także mieć katodę wykonaną w oparciu o wysokostabilną strukturę monofluorku węgla (Li/Cf_x). Takie



Fotografia 19. Baterie marki Integer przeznaczone do aplikacji w urządzeniach implantowalnych (<https://t.ly/bjEo>)

Tabela 2. Porównanie podstawowych parametrów pięciu podstawowych typów baterii stosowanych w urządzeniach implantowalnych (<https://t.ly/kpFgu>)

Rodzaj baterii	SEM [V]	Napięcie nominalne [V]	Pojemność wolumetryczna [mAh/cm³]	Gęstość energii [mWh/g]
Li/I_2	2,8	2,8	1041	210...270
Li/MnO_2	3,3	3,0	1540	230...270
Li/Cf_x	3,1	3,0	2335	440
Li/SVO	3,2	3,2	1510	270
C/LiCoO_2	4,2	3,9	783	155



Rysunek 31. Charakterystyki rozładowania ogniw Li/I₂-PVP przy obciążeniu o charakterze rezystancyjnym (<https://t.ly/qgx0>)

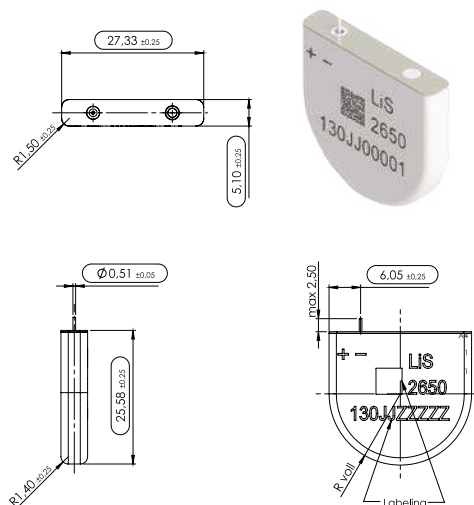
rozwiązanie oferuje niską impedancję wewnętrzną, przewidywalną charakterystykę rozładowania (pozwalającą w prosty sposób określać pozostałą pojemność ogniwa na podstawie jego napięcia) oraz bardzo wysoką pojemność i gęstość energii. Źródła Li/Cf_x są wykorzystywane m.in. w zaawansowanych kardio- i neurostymulatorach oraz wszczepialnych pompach infuzyjnych.

W zakresie urządzeń o dużym zapotrzebowaniu energetycznym i wysokiej mocy chwilowej (szczególnie ICD) sporą popularność zyskały baterie z katodą wykonaną na bazie związku srebra, wanadu i tlenu (Ag2V4O11). Ogniwa Li/SVO oferują parametry nieco zbliżone do Li/MnO₂ – pojemność wolumetryczna wynosi około 1300...1500 mAh/cm³, zaś samorozładowanie przebiega z szybkością około 2%/rok. Napięcie znamionowe jest równe 3,2 V. Warto dodać, że w ramach poszukiwania jeszcze wydajniejszych źródeł energii dla urządzeń implantowalnych, liczne badania zostały skierowane na tory źródeł hybrydowych, łączących monofluorek węgla z SVO, co doprowadziło do stworzenia ogniw określanych jako Li/CF_x – SVO – chętnie wykorzystywanych w konstrukcjach ICD, produkowanych m.in. przez markę Biotronik.

Nie sposób pominąć jednak udziału bardziej konwencjonalnych technologii elektrochemicznych, w tym baterii i akumulatorów litowo-jonowych. Dość powiedzieć, że wspomniane w artykule rozruszniki z serii Micra marki Medtronic bazują właśnie na takich źródłach zasilania, choć z powodu ograniczeń wymiarowych nie mamy tutaj do czynienia z bateriami stanowiącymi osobne komponenty – ogniwo zostało zintegrowane bezpośrednio z obudową urządzenia. Akumulatory, choć ze względów bezpieczeństwa nie są wykorzystywane w urządzeniach podtrzymujących i ratujących życie (rozrusznikach i kardiowerterach-defibrylatorach), można jednak spotkać w niektórych implantach – takich, jak neurostymulatory, których przerwa w funkcjonowaniu nie stanowi dla pacjenta bezpośredniego, śmiertelnego ryzyka.

Podsumowanie

Technologie wykorzystywane w urządzeniach implantowalnych stanowią mariaż dorobku kilku pokoleń inżynierów elektroników, specjalistów w zakresie materiałoznawstwa, mechaniki, elektrochemii, informatyki i oczywiście lekarzy, którzy nadzorowali (a najczęściej



Rysunek 32. Wymiary przykładowej baterii litowo-jodowej – LiS 2650 marki Litronik (<https://t.ly/c1n0>)

też inicjowali) prace nad kolejnymi przełomowymi rozwiązaniami. Wymogi bezpieczeństwa i niezawodności są tutaj wyśrubowane do maksimum, gdyż pozornie niewielkie usterki mogą z łatwością doprowadzić do śmierci pacjentów albo przynajmniej poważnego uszczerbku na zdrowiu.

Przez wiele lat wyścig na rynku implantów kardiologicznych odbywał się na polach wydłużania czasu pracy urządzeń, dopracowywania algorytmów sterujących stymulacją lub wstrząsami terapeutycznymi, wprowadzania coraz bardziej „fizjologicznych” metod stymulacji oraz – rzecz jasna – miniaturyzacji formy urządzeń w celu zwiększenia wygody pacjentów i poprawienia efektu estetycznego. Ostatnie lata przyniosły jednak szereg innowacji, które rzucają nowe światło na rynek implantów aktywnych – rozruszniki implantowane w całości wewnątrz jam serca i wykorzystujące detekcję drgań mechanicznych zamiast elektrod przedsionkowych to rozwiązania całkowicie nowe, znacznie podwyższające poprzeczkę tak z rynkowego, jak i technologicznego punktu widzenia.

Ograniczenia są jednak nieubłagane – inżynierowie wciąż walczą z granicami żywotności baterii, co od czasu do czasu budzi dyskusję na temat przyszłościowych metod zasilania implantów. Uwaga naukowców w tym zakresie jest skierowana przede wszystkim na technologie odzyskiwania energii (*energy harvesting*) – z przemian biochemicznych, gradientów temperatury, czy nawet drgań generowanych przez serce oraz mięśnie szkieletowe pacjenta. Optymizm, jak zawsze, należy jednak trzymać w ryzach – doświadczenie pokazuje bowiem, że w praktyce najlepiej bronią się rozwiązania sprawdzone „w boju”, choć często znacznie mniej spektakularne, niż szumne doniesienia prasy popularnonaukowej.

inż. Przemysław Musz, EP

Wybrane źródła:

- [1] DOI: 10.1111/pace.12074
- [2] Wytyczne ESC 2021 dotyczące stymulacji serca i terapii resynchronizującej serca, Vol 79, Supp. VI (2021): Zeszyty Edukacyjne 6/2021
- [3] Aquilina O. *A brief history of cardiac pacing*. Images Paediatr Cardiol. 2006 Apr; 8(2):17–81
- [4] Ramsdale, D. R., Rao, A., *Cardiac Pacing and Device Therapy*, wyd. Springer London, 2012
- [5] Efimov I.R., Kroll M.W., Tchou P.J., *Cardiac Bioelectric Therapy. Mechanisms and Practical Implications*, wyd. Springer Science+Business Media, 2009
- [6] Bock DC, Marschik AC, Takeuchi KJ, Takeuchi ES. *Batteries used to Power Implantable Biomedical Devices*. Electrochim Acta. 2012 Dec 1;84:10.1016

Znowu móc chodzić

– technologia urzeczywistnia marzenia

Osoby cierpiące na zaburzenia nerwowo-mięśniowe marzą o możliwości chodzenia bez kul. Egzoszkielet Autonomy może być spełnieniem tego marzenia. System aktywnego wspomagania chodzenia wspiera osłabione mięśnie i umożliwia intuicyjną sekwencję ruchów, która naśladuje sekwencję naturalną. Dodatkową siłę dostarcza sześć mikrosilników. Aby usprawnić harmonijną interakcję między egzoszkieletem a użytkownikiem, firma FAULHABER opracowała innowacyjny, wszechstronny silnik z czujnikiem momentu obrotowego.

W medycynie rozróżnia się ponad 800 zaburzeń nerwowo-mięśniowych. Jak wynika z nazwy, dotyczą one zarówno nerwów, jak i mięśni. Niektóre oddziałują na cały organizm, inne skupiają się w określonych obszarach. Na szczęście większość z nich występuje stosunkowo rzadko, ale i tak wielu chorych pacjentów cierpi z powodu poważnych ograniczeń mobilności. Dzieje się tak, ponieważ mimo różnych przyczyn i sposobów rozwoju, zaburzenia te mają jedną wspólną cechę: osłabienie mięśni (dystrofia mięśniowa), która w wielu przypadkach ma przebieg postępujący.

„Jeżeli osłabione są mięśnie nóg, chodzenie staje się coraz trudniejsze, aż do momentu, w którym bez podparcia jest ono niemożliwe” – wyjaśnia Mohamed Bouri, lider grupy badawczej REHA Assist (Rehabilitation and Assistive Robotics, robotyka na rzecz rehabilitacji i wsparcia) z uniwersytetu technicznego w szwajcarskiej Lozannie (EPFL). „Mięśnie nadal pracują, ale nie są na tyle silne, aby zapewnić stabilne stanie lub niezależne poruszanie nogami. Jak łatwo się domyślić, ma to ogromny wpływ na zakres ruchów i poziom życia pacjentów. Skutki przypominają następstwa porażenia połowicznego po udarze.

Częściowe wspomaganie w lekkiej wersji

Mohamed Bouri opowiada o będących już w użyciu tradycyjnych egzoszkieletach bazujących na technologii inspirowanej humano-idami. Urządzenia te sprawiają, że osoby sparaliżowane mogą poruszać się bez kul, jednak ich masa przekracza 40 kg. Autonomy – opracowany przez REHA Assist waży zaledwie 25 kg i współpracuje z osłabionym, jednak wciąż funkcjonującym układem kostno-szkieletowym pacjenta.

Urządzenie składa się z gorsetu umieszczanego wokół tułowia oraz mankietów mocowanych do kończyn dolnych. Po każdej stronie znajdują się trzy silniki dostarczające siłę, której brakuje mięśniom. Pierwszy z silników odpowiada za zginanie i prostowanie bioder, a drugi za te same ruchy kolan. Rolą trzeciego silnika jest wspieranie odwodzenia i przywodzenia nóg w stawach biodrowych – innymi słowy ruchu nóg w bok od linii środkowej ciała.

Wszystkie silniki pomagają pacjentom w utrzymaniu równowagi i chodzeniu w pozycji wyprostowanej. Przeprowadzone ostatnio badanie kliniczne obejmujące osoby doświadczające problemów z chodzeniem wykazały, że egzoszkielet Autonomy spełnia swoją rolę,

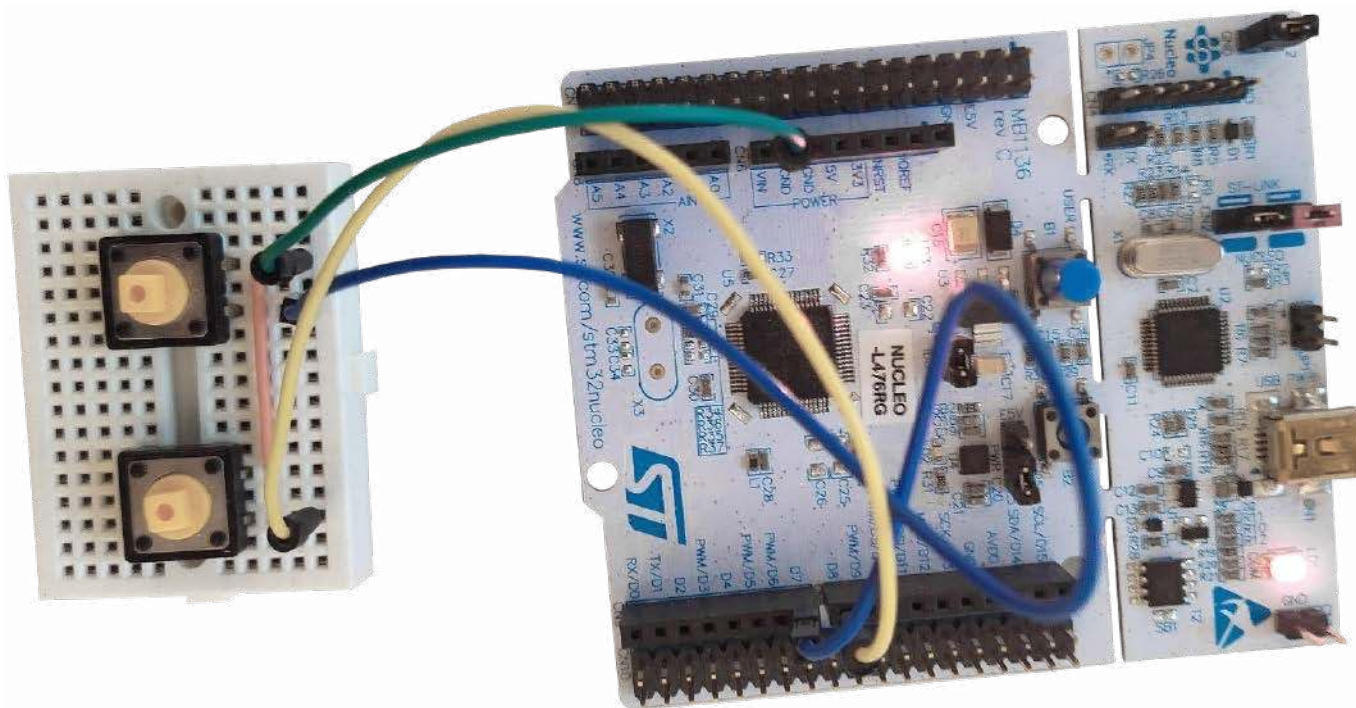


zapewniając wsparcie, ale nie ograniczając swobody ruchów, zgodnie z intencjami użytkowników. Zakres ruchów stawów i kadencja chodu pozostały bez zmian.

Moc i potencjał rozwojowy napędów

FAULHABER jest dostawcą sześciu jednostek napędowych montowanych w urządzeniach. Ich najważniejszym elementem jest bezszczotkowy silnik 3274 BP4 o średnicy 32 milimetrów. Oferuje on największą moc spośród wszystkich dostępnych na rynku silników w swojej klasie wielkości. Moc przenoszona jest przez przekładnię planetarną 42 GPT z wałem opracowanym specjalnie pod kątem tego zastosowania. Magnetyczny enkoder IE3 przekazuje dane o położeniu do sterownika. Czujniki momentu obrotowego wbudowane są w przekładnię czterech silników sterujących zginaniem i prostowaniem. Wymagania dotyczące jednostek napędowych są typowe dla mikrosilników najwyższej klasy. Do najważniejszych właściwości w tym zastosowaniu należą wysoka moc przy najmniejszej możliwej pojemności i masie, a także precyzja, niezawodność i długi okres eksploatacji.

Faulhaber
www.faulhaber.com



Wbudowane sieci neuronowe w STM32 (1)

Sieci neuronowe są obecnie najbardziej rozpowszechnionym algorytmem sztucznej inteligencji. W ostatnim czasie powstały także wersje przeznaczone dla mikrokontrolerów, gdzie znajdują zastosowanie w wykrywaniu wzorców w zbieranych danych. W artykule zaprezentuję, jak można wytrenować sieć w bibliotece TensorFlow, a następnie, jak ją uruchomić na mikrokontrolerze używając dostępnych bibliotek. Eksperymenty zostały wykonane na płycie Nucleo-L476RG, ale powinny zadziałać, także na innych układach firmy ST dysponującymi odpowiednią ilością pamięci.

Pierwszym przykładem będzie mało przydatna w praktyce, ale idealna, jako przykład startowy sieć realizująca działanie funkcji XOR. Na początku przygotowujemy dane szkoleniowe i utworzymy model sieci oraz wytrenujemy jej wagi w pakiecie TensorFlow. Następnie wygenerujemy model TensorFlow lite i na jego podstawie otrzymamy kod dla mikrokontrolera STM32.

Uczymy sieć

Do trenowania sieci zastosujemy środowisko Google Colab [1]. Możemy tam tworzyć notatniki i uruchamiać je w chmurze. Dzięki temu mamy już zainstalowane potrzebne biblioteki. Cały kod użyty w projekcie znajduje się w [2]. Na **listingu 1** pokazano kod odpowiedzialny za importowanie bibliotek, z których będziemy korzystać. Są to moduły:

- **TensorFlow**, który dostarcza funkcji dla sieci neuronowych,
- **NumPy** odpowiadający za efektywne obliczenia numeryczne.

Dla sprawdzenia wyświetlamy wersję TensorFlow. Ja używałem wersji 2.8.2.

Listing 1. Dołączanie bibliotek

```
import tensorflow as tf
import numpy as np
print("TensorFlow version:", tf.__version__)
```

Kolejnym krokiem jest przygotowanie treningowego zbioru danych (**listing 2**). Ponieważ w naszym przypadku mamy tylko 4 punkty, dla których będziemy używać naszej sieci, to będą one zarówno zbiorem treningowym jak i użytym do oceny modelu. Zarówno wejścia jak i wyjścia naszej sieci będą liczbami zmiennoprzecinkowymi typu `float32`. Przyjmijmy więc, że stan 0 zakodujemy jako `-127`, a 1 jako `+127`. Natomiast stan wyjścia będziemy oceniali na podstawie znaku. Zero i więcej będą odpowiadać logicznej 1, natomiast liczby ujemne będą interpretowane jako 0.

Listing 2. Zbiór uczący

```
data_x = np.array([
    [-127.0, -127.0],
    [-127.0, 127.0],
    [127.0, -127.0],
    [127.0, 127.0]
], dtype=np.float32)

data_y = np.array([
    -127.0,
    127.0,
    127.0,
    -127.0
], dtype=np.float32)
```

Dla naszego zadania wystarczyłaby mniejsza sieć, ale wtedy algorytm uczenia miałby problem ze znalezieniem satysfakcjonujących nas wag. Moglibyśmy je dobrać ręcznie (co dla naszego małego przypadku byłoby proste). Ale my po prostu zwiększymy rozmiar modelu. Został on pokazany na **rysunku 1** oraz **listingu 3**. Stosujemy trzy warstwy. W dwóch pierwszych są po cztery neurony z nieliniowością typu *ReLU* (*Rectified linear unit linear* – poprawiona jednostka liniowa). Kształt tej funkcji zaprezentowano na **rysunku 2**. Można ją zapisać jako:

$$y = \max(0, x)$$

Na końcu mamy pojedynczą warstwę z wyjściem liniowym. Nie możemy wykorzystać tu nieliniowości *ReLU*, ponieważ na wyjściu

chcemy otrzymywać, także wartości ujemne symbolizujące logiczne 0.

Na początku parametry sieci są wypełniane wartościami losowymi. Możemy sprawdzić jaki uzyskamy wynik dla naszych 4 punktów uruchamiając następujące działanie:

```
predictions = model(data_x).numpy()
```

Ja uzyskałem wynik:

```
array([[ -206.96419 ],
       [ -132.3428  ],
       [  -82.97228 ],
       [ -21.927877]], dtype=float32)
```

Odpowiada on zwróceniu logicznego 0 dla każdej pary wejść. Jest on dość daleki od oczekiwanych przez nas wartości. Aby to poprawić musimy przeprowadzić trening sieci. Konfigurujemy więc sposób obliczania błędu pomiędzy wartością oczekiwaną, a zwracaną oraz wybieramy metodę optymalizacji – wpisujemy:

```
model.compile(optimizer='rmsprop',
              loss='mae', metrics=['mae'])
```

Następnie uruchamiamy szkolenie poleceniem:

```
model.fit(data_x, data_y,
          epochs=1000)
```

W moim przypadku wystarczyło 1000 iteracji (parametr *epochs*) dla uzyskania satysfakcjonującego wyniku. Aby ocenić, czy efekt szkolenia jest dla nas akceptowalny obliczamy wynik dla potrzebnych nam 4 punktów i sprawdzamy, czy znak odpowiada oczekiwanej przez nas wartości. Ja otrzymałem:

```
array([[ -127.08677 ],
       [  0.72788846],
       [  0.72788846],
       [ -127.08033 ]],
      dtype=float32)
```

Jeżeli wynik jest satysfakcjonujący możemy wygenerować model TensorFlow lite i zapisać go do pliku. W tym celu uruchamiamy kod:

```
converter = tf.lite.TFLiteConverter.from_keras_model(model)
tflite_model = converter.convert()
open('xor.tflite', 'wb').write(tflite_model)
```

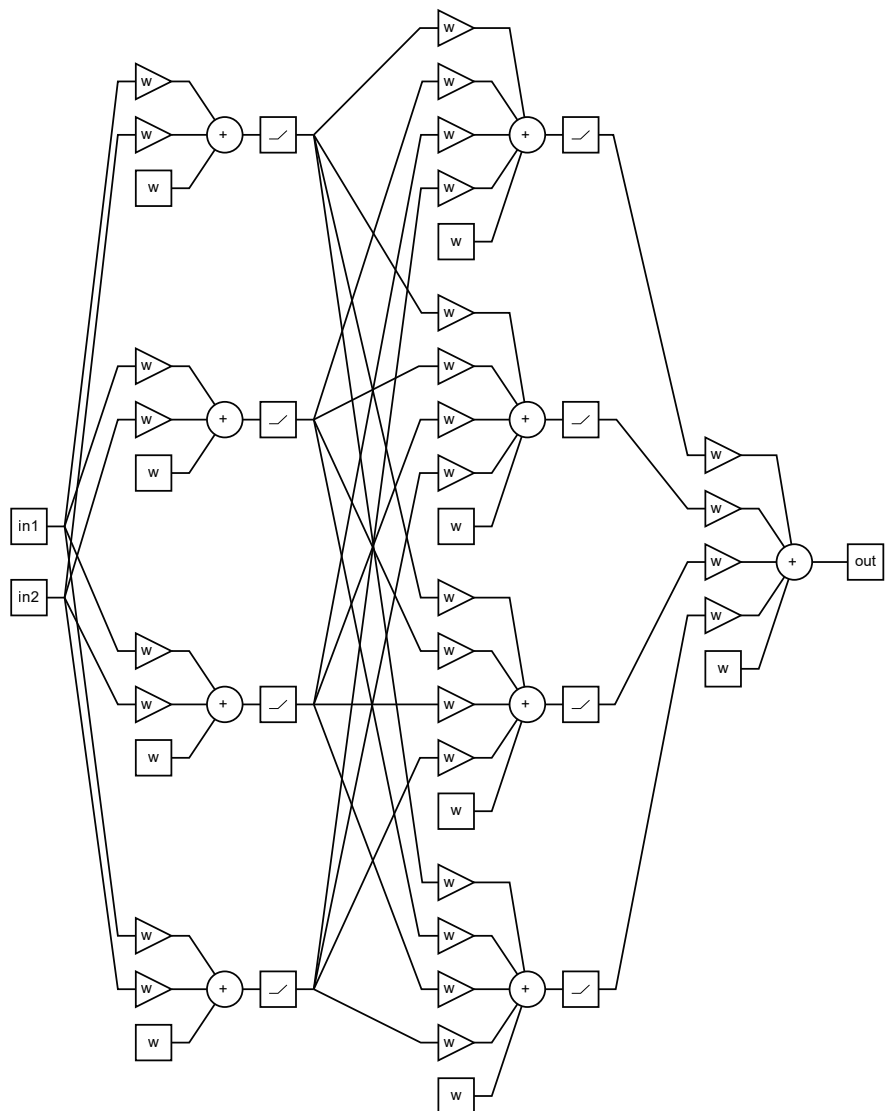
Teraz pobieramy wygenerowany plik. Klikamy ikonę folderu w lewej części okna, co powoduje rozwinięcie panelu *Pliki* pokazanego na **rysunku 3** i pobieramy plik *xor.tflite*. Jeżeli nie wiemy, w którym miejscu w drzewie katalogów został on zapisany, możemy wywołać w nowym polu typu *Kod* polecenie *pwd*. Zawartość pobranej sieci możemy podejrzeć na przykład za pomocą programu Neutron [3], albo bezpośrednio w środowisku STM32CubeIDE.

Mikrokontroler

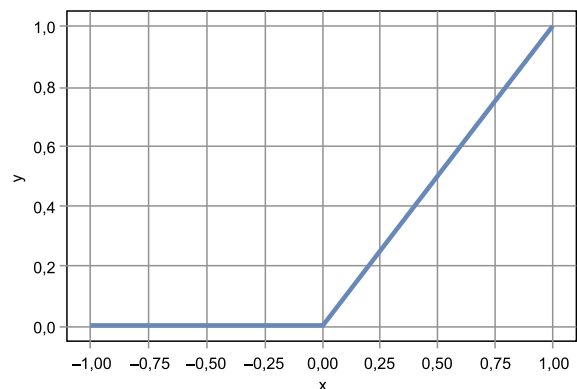
Sieć jest już zaprojektowana. Musimy teraz przygotować program dla mikrokontrolera. Utworzymy go w środowisku STM32CubeIDE, które można pobrać z [4]. Jest ono zbudowane na bazie edytora Eclipse z dołączonym generatorem kodu Stm32Cube oraz kompilatorem.

Listing 3. Opis kolejnych warstw sieci

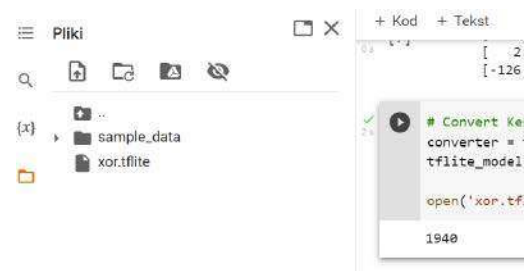
```
model = tf.keras.models.Sequential([
    tf.keras.layers.Dense(4, activation='relu'),
    tf.keras.layers.Dense(4, activation='relu'),
    tf.keras.layers.Dense(1, activation='linear'),
])
```



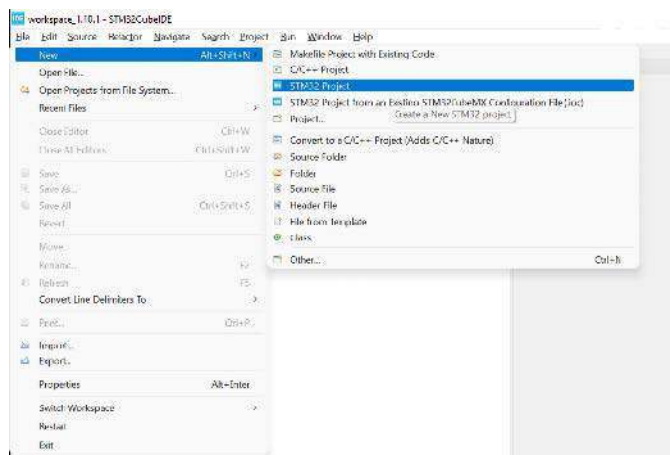
Rysunek 1. Model sieci



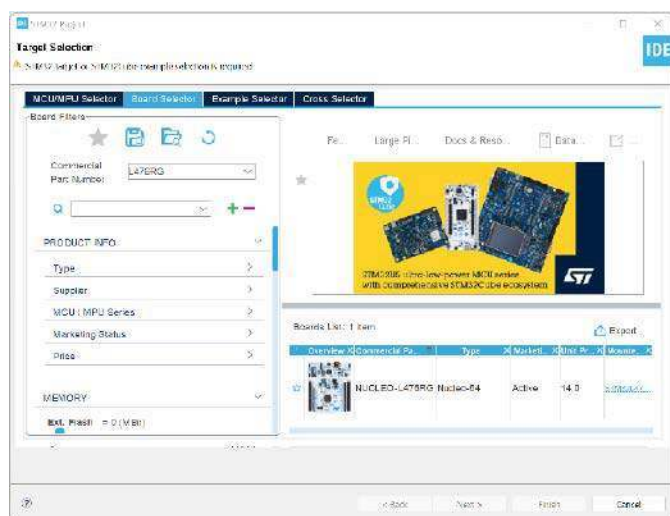
Rysunek 2. Poprawiona jednostka liniowa – ReLU



Rysunek 3. Pobieramy wygenerowany plik z wytrenowanym modelem sieci



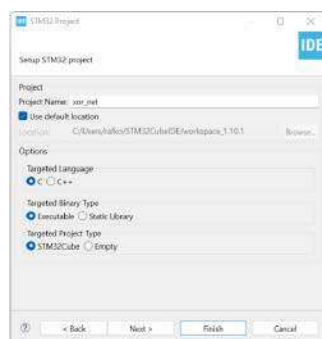
Rysunek 4. Tworzymy nowy projekt



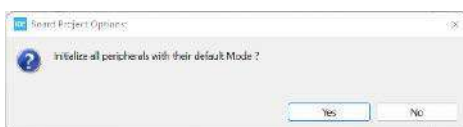
Rysunek 5. Wybór zestawu startowego

Zacznijmy od pobrania biblioteki do sztucznej inteligencji. Z menu wybieramy *Help* i klikamy *Manage embedded software packages*. W oknie wybieramy zakładkę *STMicroelectronics*, a z listy wybieramy *X-CUBE-AI* w najnowszej wersji i naciskamy *Install Now*.

Następnie tworzymy nowy projekt, co pokazano na **rysunku 4**. W kolejnym oknie (**rysunek 5**) przechodzimy do zakładki *Boards Select* (wybór płytek) i znajdujemy płytkę *Nucleo*, której chcemy użyć. W moim przypadku jest to *L476RG*. W następnym oknie (**rysunek 6**)

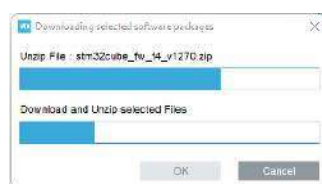


Rysunek 6. Wybór nazwy projektu

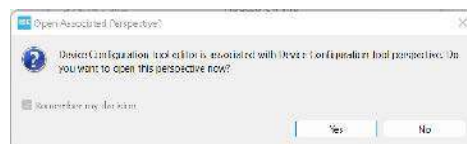


Rysunek 7. Pytanie o inicjalizację peryferiów

wyberamy nazwę projektu i klikamy *Finish*, aby zakończyć. Następnie pojawi się pytanie o zainicjalizowanie peryferiów dostępnych na płytce (**rysunek 7**). Potwierdzamy klikając *OK*. Jeżeli nie mamy wymaganych bibliotek na dysku, to zostaną one automatycznie pobrane o czym



Rysunek 8. Pobieranie brakujących bibliotek



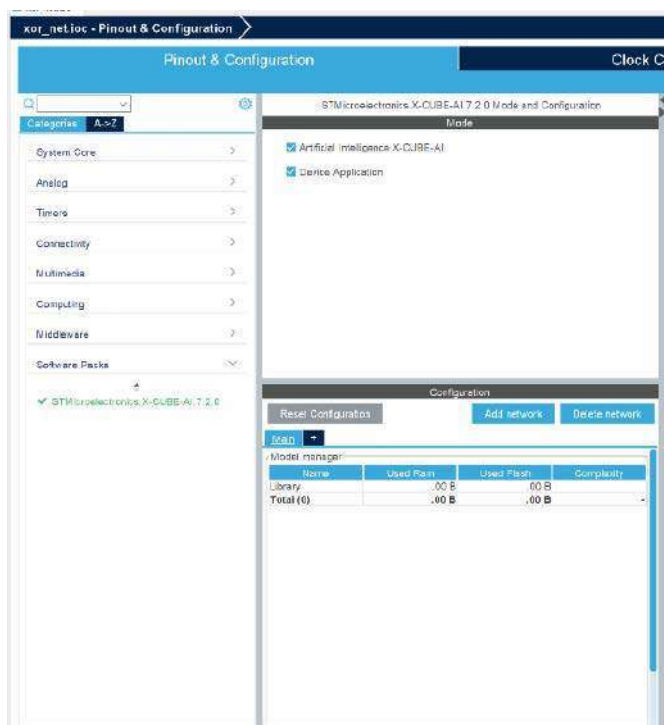
Rysunek 9. Pytanie o otwarcie widoku konfiguracji



Rysunek 10. Dodajemy bibliotekę CubeAI



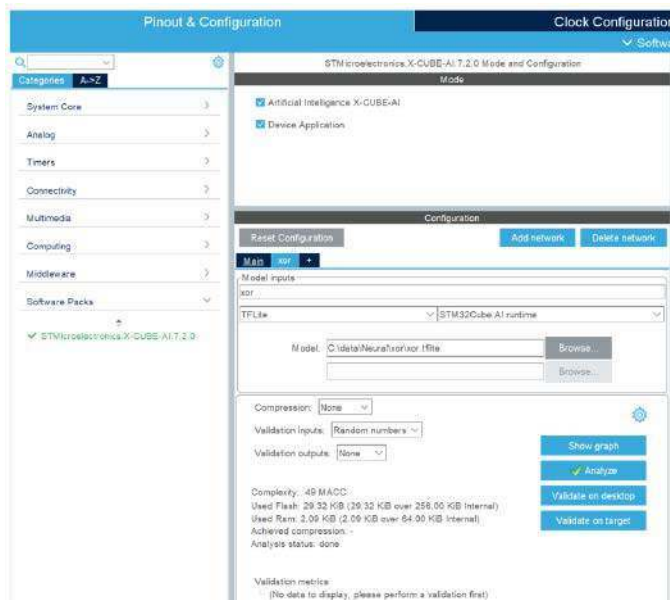
Rysunek 11. Lista bibliotek



Rysunek 12. Widok konfiguracji biblioteki CubeAI

powiadamia nas ekran podobny do pokazanego na **rysunku 8**. Gdy projekt zostanie utworzony, zostaniemy zapytani, czy otworzyć ekran konfiguracji (**rysunek 9**) – klikamy *Yes*.

Zobaczymy otwarte okno pozwalające na konfigurację projektu. Zacznijmy od skonfigurowania sieci neuronowej. W tym celu na górnym pasku klikamy *Software package* i z rozwiniętego menu wybieramy *Select Components* (**rysunek 10**). Ukaze się lista dostępnych paczek (**rysunek 11**), rozwijamy *STMicroelectronics X-CUBE-AI*. Zaznaczamy *tic* w polu *X-CUBE-AI Core*. Natomiast w zakładce *Application* wybieramy z listy *Application Template* (szablon aplikacji). Po zatwierdzeniu możemy dostać pytanie, czy chcemy zmienić



Rysunek 13. Konfiguracja nowej sieci



Rysunek 14. Złożoność sieci oraz ilość wymaganej pamięci

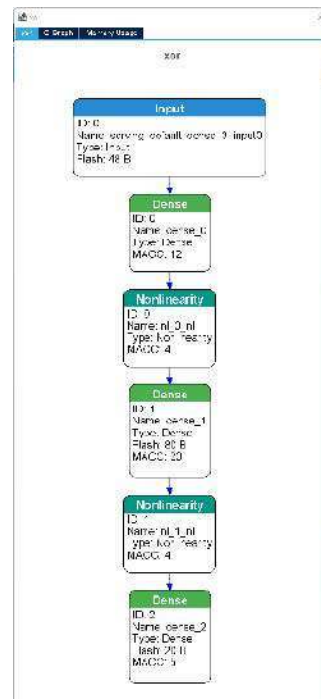
ustawienia zegara, w celu zwiększenia wydajności. Zgadamy się klikając Yes.

Wracamy do głównego okna konfiguratora. W prawym panelu rozwijamy *Software Package* i wybieramy *STMicroelectronics.X-CUBE-AI*. Pojawi się ekran jak na **rysunku 12**. Widzimy, że w polu *Mode* zaznaczone są obie opcje – zarówno biblioteka jak i przykładowa aplikacja. W polu *configure* klikamy przycisk *Add network*. Zobaczymy okno jak na **rysunku 13**, w polu tekstowym *Model Inputs* wpisujemy nazwę naszej sieci. Ja wybrałem *xor*. Wybieramy rodzaj sieci na *TFLite* oraz *STM32Cube.AI runtime*.

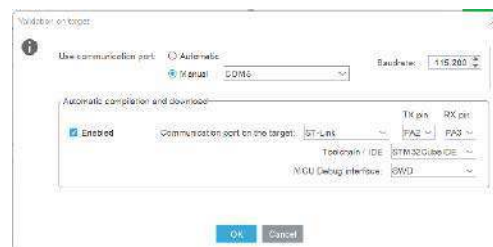
Klikamy *Browser* i wybieramy pobrany wcześniej plik *xor.tflite*. Znajdziemy go, także w repozytorium [5] i naciskamy przycisk *Analyze*. Otworzy się nowe okno, w którym będą pojawiały się informacje o postępach. Gdy się zakończy klikamy OK. Po chwili, w głównej aplikacji zobaczymy parametry naszej sieci oraz ilość potrzebnej pamięci RAM i Flash (**rysunek 14**). Klikając *Show Graph* możemy wyświetlić schemat załadowanej sieci (**rysunek 15**). Ciekawą funkcją jest *Validate on Target*. Pozwala on na wygenerowanie testowej aplikacji, która zostanie uruchomiona bezpośrednio na naszej płytce. Najpierw jednak musimy wygenerować kod. W tym celu po prostu zapisujemy nasz projekt. Wtedy powinno się pojawić okno dialogowe z pytaniem, czy wygenerować pliki. Potwierdzamy.

Jeżeli zostaniemy przeniesieni do widoku edycji kodu musimy z powrotem wrócić do widoku konfiguracji. W tym celu w znajdujący się w lewej części okna panel z listą plików wybieramy i otwieramy plik o rozszerzeniu *.ino*. Teraz możemy już nacisnąć przycisk *Validate on target*. Po kliknięciu zobaczymy okno konfiguracji (**rysunek 16**). Przy

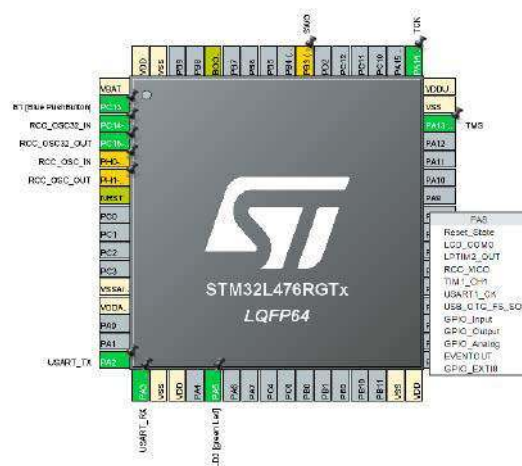
ustawieniach portu szeregowego wybieramy *Manual* i zaznaczamy port szeregowy pod którym system operacyjny widzi naszą płytkę. W drugiej części zaznaczamy *Enabled* przy automatycznym tworzeniu i wgraniu projektu. Pozostałe opcje pozostawiamy bez zmian. Zatwierdzamy klikając OK. Otworzy się okno, w którym zobaczymy informację o postępach oraz wyniki. Możemy dowiedzieć się o ile różnią się wartości uzyskane na komputerze oraz na mikrokontrolerze. Znajdziemy także czas wykonania się poszczególnych warstw (**listing 4**). Widzimy, że pojedyncze wywołanie naszej sieci zajmuje 36 μ s, czyli 2892 cykle zegara. Znajdziemy także rozpiškę jaki procent czasu zajmuje która warstwa. Gdy zapoznamy się z wynikami zamykamy okno naciskając OK.



Rysunek 15. Graf przedstawiający budowę załadowanej sieci



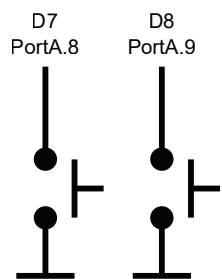
Rysunek 16. Konfiguracja testu w sprzęcie



Rysunek 17. Konfiguracja portów mikrokontrolera

Listing 4. Podsumowanie wygenerowanej aplikacji

```
Results for 10 inference(s) - average per inference
device       : 0x415 - STM32L4x6xx @80/80MHz fpu,art_lat=4,art_prefetch,art_icache,art_dcach
duration     : 0.036ms
CPU cycles   : 2892
cycles/MACC  : 64.28
c_nodes      : 5
c_id  m_id  desc      output      ms      %
-----
0      0      Dense (0x104) (1,1,1,4)/float32/16B 0.009  25.4%
1      0      NL (0x107)    (1,1,1,4)/float32/16B 0.004  12.3%
2      1      Dense (0x104) (1,1,1,4)/float32/16B 0.010  28.1%
3      1      NL (0x107)    (1,1,1,4)/float32/16B 0.004  12.1%
4      2      Dense (0x104) (1,1,1,1)/float32/4B  0.008  22.1%
-----
                                0.036 ms
```



Rysunek 18. Podłączenie przycisków

Wróćmy jednak do naszego projektu. Zostało nam jeszcze skonfigurowanie portów mikrokontrolera. Przechodzimy do głównego ekranu konfiguratora, tego z rysunkiem układu scalonego (rysunek 17). Widzimy, że PA5, do którego jest podłączona dioda LED został już skonfigurowany. Pozostaje nam jeszcze ustawienie wejść dla przycisków. Podłączymy je na płytce stykowej zgodnie z rysunkiem 18. Klikamy więc lewym przyciskiem myszy na *PortA.8*, a następnie *PortA.9* i dla każdego z nich wybierzmy *GPIO_Input*. Następnie w prawym panelu rozwijamy wybieramy *System Core* i klikamy *GPIO* (rysunek 19). W środkowym panelu wybieramy najpierw *PortA.8*, a następnie *PortA.9* i dla każdego włączamy podciąganie do plusa (*pull-up*). W tym samym panelu widzimy, że *PortA.5*, do którego podłączona jest dioda LED znajdująca się w zestawie, został już skonfigurowany jako wyjście. Możemy także sprawdzić, że konfiguracja portu szeregowego, podłączonego do USB, także jest gotowa. Gdy zapiszemy projekt, zostaniemy zapytani, czy wygenerować kod, potwierdzamy. Kolejne okienko pyta, czy przełączyć w tryb edycji kodu. Tu także potwierdzamy.

Printf

Aby uprościć debugowanie przekierujemy standardową funkcję z biblioteki C: *printf* na port szeregowy. W tym celu w pliku *Core/Src/main.c* w bloku */* USER CODE BEGIN PV */* dodajemy definicję dwóch funkcji (listing 5). Funkcja *_write* służy do wysłania ciągu znaków na standardowe wyjście. W niej każdy znak jest przekazywany do funkcji *_io_putchar*. Druga z nich wywołuje funkcję z biblioteki HAL, która odpowiada za obsługę portu szeregowego. Sam UART został automatycznie skonfigurowany przy wyborze płytki na prędkość 115200.

Drugą częścią jest włączenie wypisywania liczb zmiennoprzecinkowych. Aby ograniczyć rozmiar kodu wynikowego ta opcja jest domyślnie wyłączona. Na prawym pasku z listą plików znajdujemy nasz projekt i klikamy go prawym przyciskiem myszy. Z listy wybieramy *Properties*. W nowym oknie (rysunek 20) przechodzimy do zakładki *C/C++ Build* → *Settings* → *MCU Settings* i zaznaczamy *tic* przy opcji *Use float with printf*.

Uruchamiamy sieć

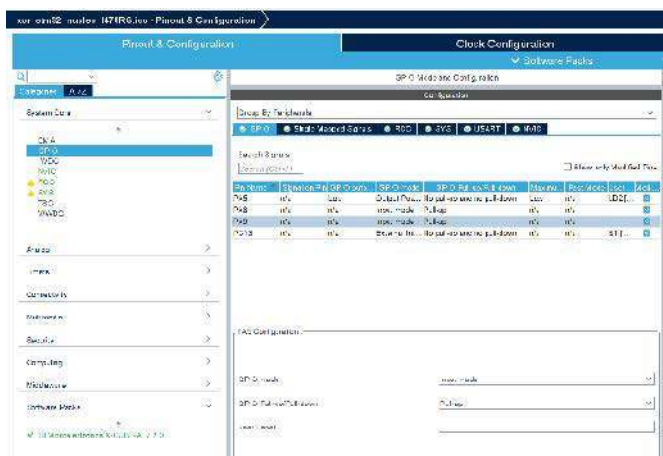
Czytając dalej plik *main.c* znajdziemy, że przy konfiguracji wywoływana jest funkcja *MX_X_CUBE_AI_Init()*, a w głównej pętli programu *MX_X_CUBE_AI_Process()*. Ich implementację znajdziemy w *X-CUBE-AI/App/app_x-cube-ai.c*. Musimy zmodyfikować jedynie drugą z nich. Jej zawartość po zmianach pokazuje listing 6. Na początku definiujemy zmienne: *nn_input* jest tablicą z wejściami dla sieci, a do *nn_output* zostaną wpisane wyniki. Rozmiar wejścia i wyjścia jest zdefiniowany przez stałe i wynosi odpowiednio 2 i 1, co jest zgodne z tym co skonfigurowaliśmy w TensorFlow. Wejścia sieci są odczytywane z przycisków. Jeśli przycisk jest naciśnięty wpisujemy 127, a w przeciwnym razie -127.

Następnie przypisujemy wskaźniki do tablic z danymi do pól data struktur z wejściami i wyjściami sieci. Wywołanie obliczeń następuje w funkcji *ai_xor_run*. Jeżeli otrzymany wynik jest dodatni, dioda led jest zaświecana, a gdy ujemny gaszona. Na końcu wypisujemy wartość zwróconą przez sieć na port szeregowy. Cały program można pobrać z repozytorium [5]. Podłączamy przyciski do pinów D7 i D8 zgodnie ze schematem z rysunku 18. Kompilujemy projekt i programujemy płytkę za pomocą przycisku *Run()*. Gotowy model został pokazany na fotografii tytułowej. Stan diody LD2 powinien być funkcją xor przycisków. Gdy otworzymy monitor portu szeregowego zobaczymy dokładny wynik uzyskany z sieci neuronowej.

Listing 5. Obsługa funkcji printf

```
int _io_putchar(int ch){
    HAL_UART_Transmit(&huart2, (uint8_t *)&ch, 1, HAL_MAX_DELAY);
    return ch;
}

int _write(int file, char *ptr, int len){
    int DataIdx;
    for (DataIdx = 0; DataIdx < len; DataIdx++){
        _io_putchar(*ptr++);
    }
    return len;
}
```



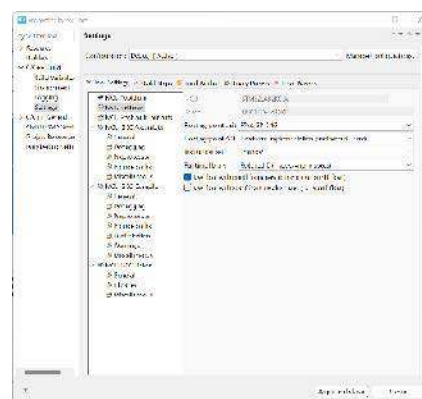
Rysunek 19. Konfiguracja wejść/wyjść

Uruchomiliśmy pierwszą sieć neuronową na mikrokontrolerze. Jednak obliczanie funkcji xor nie jest zbyt intrygującym zadaniem. Dlatego w drugiej części uruchomimy rozpoznawanie kształtów.

Rafał Kozik
rafkozik@gmail.com

Bibliografia:

- [1] <https://bit.ly/3Q9v7b5>
- [2] <https://bit.ly/3lgeYyR>
- [3] <https://bit.ly/3Grh2Ct>
- [4] <https://bit.ly/2XsPWIH>
- [5] <https://bit.ly/3Q3weJx>



Rysunek 20. Włączamy obsługę zmiennych float w funkcji printf

Listing 6. Funkcja MX_X_CUBE_AI_Process

```
void MX_X_CUBE_AI_Process(void){
    /* USER CODE BEGIN 6 */

    ai_i32 batch;

    float nn_input[AI_XOR_IN_1_SIZE];
    float nn_output[AI_XOR_OUT_1_SIZE];

    nn_input[0] = HAL_GPIO_ReadPin(GPIOA, GPIO_PIN_8) ? 127.0 : -127.0;
    nn_input[1] = HAL_GPIO_ReadPin(GPIOA, GPIO_PIN_9) ? 127.0 : -127.0;

    ai_input->data = nn_input;
    ai_output->data = nn_output;

    batch = ai_xor_run(xor, ai_input, ai_output);

    HAL_GPIO_WritePin(GPIOA, GPIO_PIN_5, nn_output[0]>=0 ? 1 : 0);
    printf("%f\r\n", nn_output[0]);

    if (batch != 1) {
        ai_log_err(ai_xor_get_error(xor), "ai_xor_run");
    }

    /* USER CODE END 6 */
}
```


Ulubiony Kiosk

Pobierz bezpłatnie multimedialne dodatki do tego wydania Elektroniki Praktycznej

**Projekty, miniprojekty, materiały do
artykułów i kursów oraz wiele innych!**



*** Kupiłeś magazyn
w Ulubionym
Kiosku lub masz
prenumeratę?
Multimedialne dodatki
będą odblokowane
automatycznie!**

*** Zakupiłeś czasopismo
u zewnętrznego
dystybutora?
Odblokuj bibliotekę
multimediów
samodzielnie.**

Szczegóły na UlubionyKiosk.pl/media



Zastosowanie komponentów rad-hard

Wpływ promieniowania na układy scalone i inne półprzewodniki

Wszyscy rozumieją ogólne niebezpieczeństwo związane z promieniowaniem (samo słowo w naszych głowach brzmi dość groźnie), ale skupiają się głównie na wpływie promieniowania jonizującego na organizmy żywe. Czy podobne zagrożenie występuje w odniesieniu do urządzeń elektronicznych? Okazuje się, że promieniowanie może mieć istotny wpływ na struktury półprzewodnikowe. W zaprezentowanym artykule omówimy efekty, jakie może wywołać promieniowanie w układach scalonych i gdzie występuje największa potrzeba ochrony wrażliwej elektroniki przed wpływem wysokoenergetycznych cząstek. Finalnie podany zostanie też opis pewnych metryk i definicji terminów, stosowanych w opisie tego rodzaju podatności, aby łatwiejsze było zrozumienie literatury technicznej dotyczącej tego zagadnienia.

W 2003 roku w Belgii odbywały się wybory. Maszyny do głosowania przyjmowały głosy na dwa sposoby: poprzez wpisanie ich do komputera lub poprzez zapisanie swojego głosu na karcie magnetycznej, którą wrzucało się następnie do urny, wyposażonej w magnetyczny czytnik rzeczonych kart. To typowa implementacja systemów do elektronicznego zliczania głosów w wyborach, jakie coraz częściej stosowane są na świecie. Podczas zliczenia głosów z komputera wykryto błąd polegający na tym, że jeden kandydat miał dokładnie 4096 głosów więcej, niż było to możliwe przy danej liczbie oddanych kart do głosowania. Tę dało się sprawdzić, zliczając fizyczne karty magnetyczne, znajdujące się w urnie. Liczba ta nie jest przypadkowa: $2^{12} = 4096$. W dalszych badaniach odkryto, że tranzystor odpowiedzialny za sygnalizację trzynastego bitu w jakimś miejscu całego toru cyfrowego (W pamięci? W procesorze? Tego nie wiadomo) został uderzony przez foton promieniowania kosmicznego. Jego energia była wystarczająca do zmiany stanu tego bitu z 0 na 1. Spowodowało to dodanie do zmiennej przechowującej ilość głosów na danego kandydata wartość dokładnie 2^{12} .

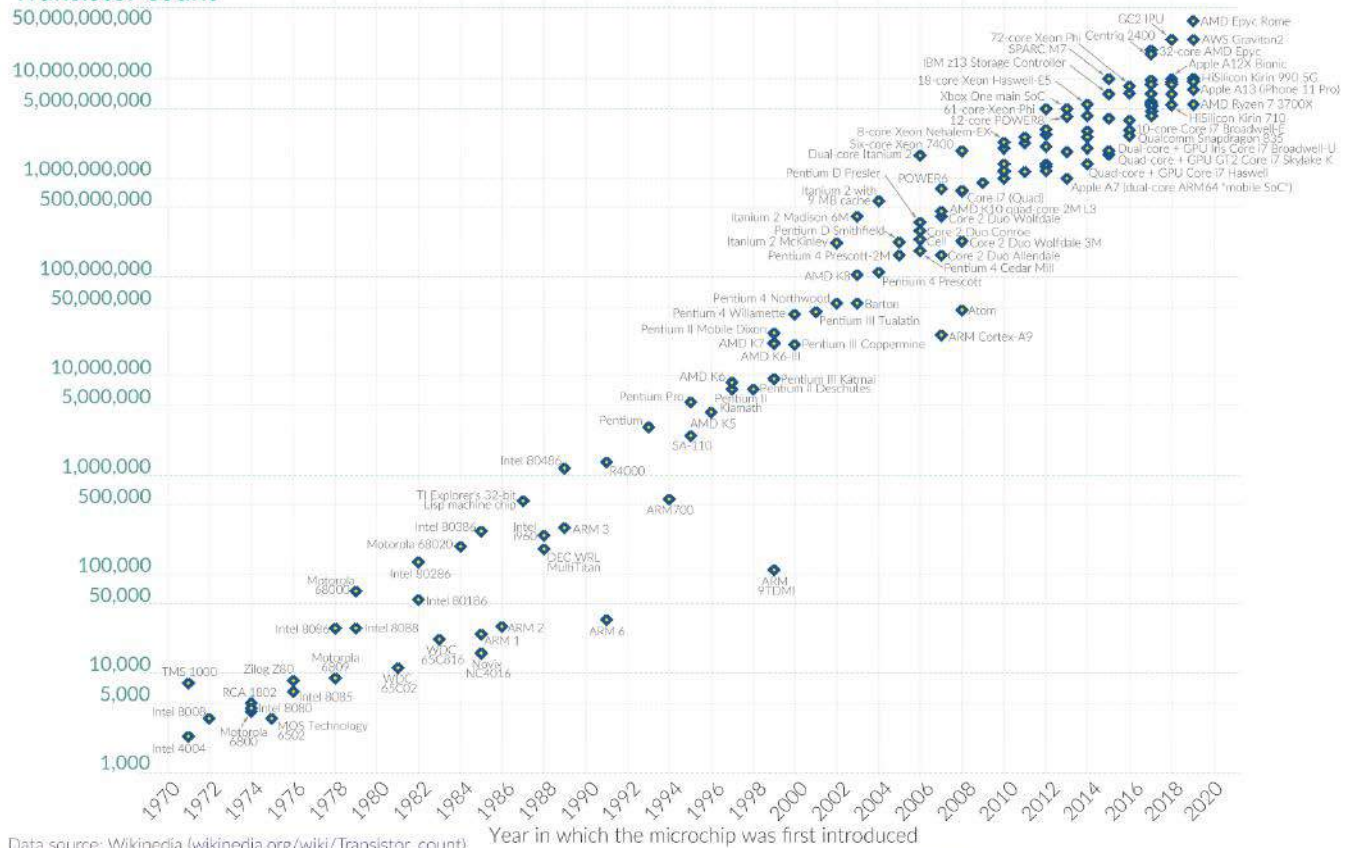
Nie była to pierwsza taka sytuacja. Wcześniej, w 1996 roku, firma IBM badała podobny problem. Okazało się, że układy scalone, pochodzące

Moore's Law: The number of transistors on microchips doubles every two years

Moore's law describes the empirical regularity that the number of transistors on integrated circuits doubles approximately every two years. This advancement is important for other aspects of technological progress in computing – such as processing speed or the price of computers.

Our World
in Data

Transistor count



Data source: Wikipedia (wikipedia.org/wiki/Transistor_count)

OurWorldinData.org – Research and data to make progress against the world's largest problems.

Licensed under CC-BY by the authors Hannah Ritchie and Max Roser.

Rysunek 1. Prawo Moora – liczba tranzystorów w układzie scalonym podwaja się, co dwa lata (Zaadaptowano z OurWorldInData.org)

z jednego z zakładów produkcyjnych firmy, mają podwyższoną skłonność do występowania tzw. pojedynczych zdarzeń wzbudzenia (*Single Event Upsets* – SEU). Są one rodzajem miękkiego błędu, polegającego na zmianie wartości bitu pod wpływem zewnętrznego oddziaływania, np. kwantu promieniowania czy zewnętrznego pola. W przypadku wspomnianych układów IBM problemem okazał się lekko podwyższony poziom związków uranu w ceramice obudowy. Zwiększony poziom promieniowania przy samej strukturze krzemowej spowodował zwiększenie poziomu SEU.

Podczas analizowania przyczyn i skutków problemu SEU, badacze z IBM obliczyli także, że „komputer zwykle doświadcza jednego błędu wywołanego promieniowaniem kosmicznym na 256 megabajtów pamięci RAM, na miesiąc”. Wyliczenie to ma prawie 30 lat. W tym czasie ilość tranzystorów w układach scalonych wzrosła 1000-krotnie (**rysunek 1**), a napięcia sterowania układami spadły prawie o rząd wielkości, co sprawia, że układy te są jeszcze bardziej podatne na zdarzenia, takie jak w Belgii.

Rodzaje promieniowania i miejsce jego typowego występowania

Kosmos to wysoce radioaktywne środowisko. Skutki ekspozycji na to promieniowanie skutecznie niszczą funkcjonalność większości komercyjnych układów scalonych. Istnieją cztery główne przyczyny tego promieniowania:

- uwiecznione elektrony,
- uwiecznione protony,
- protony słoneczne,
- promienie kosmiczne.

Przyjrzyjmy się wszystkim tym źródłom promieniowania i ich charakterystyce. Różnią się one istotnie parametrami, a co za tym

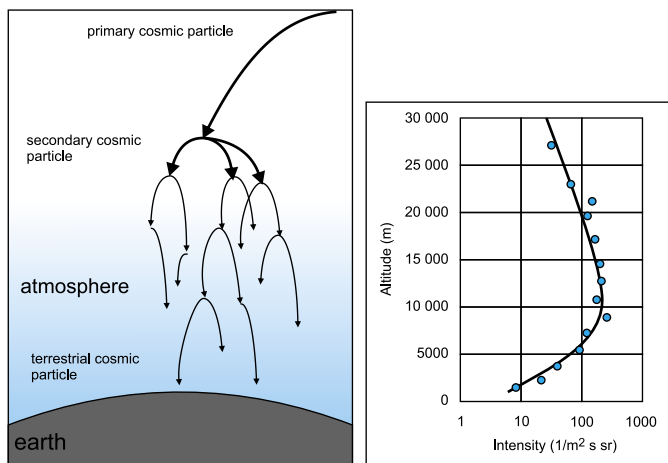
idzie mechanizmami i intensywnością wpływu na urządzenia półprzewodnikowe.

Uwięzione elektrony to ujemnie naładowane cząstki, które mają stosunkowo niewielką masę, ale są również niezwykle energetyczne. Ze względu na ich małą masę zwykle znajdują się na bardzo wysokich orbitach, takich jak orbity geosynchroniczne (lub orbity GEO), które znajdują się około 36 000 km nad ziemią.

Uwięzione protony to dodatnio naładowane cząstki uwięzione przez pole grawitacyjne planety. Mają mniej energii niż elektrony, ale są około 2000 razy masywniejsze niż elektrony. Występują w dużych stężeniach na niskich wysokościach. Na przykład, na niskiej orbicie okołoziemskiej (LEO), czyli orbitach znajdujących się na wysokości ok. 1400...2000 km nad powierzchnią Ziemi, źródła promieniowania są zdominowane przez uwięzione protony.

Protony słoneczne są podobne do protonów uwięzionych, z wyjątkiem tego, że pochodzą ze słońca, głównie z wiatru słonecznego, wyładowań koronowych i innych zjawisk na powierzchni naszej gwiazdy. Z uwagi na to nie są one uwięzione w polu grawitacyjnym naszej planety.

Ostatnim, ze źródeł promieniowania w kosmosie są tzw. promienie kosmiczne. To zbiorcza nazwa na promieniowanie, składające się m.in. z cząstek alfa, ciężkich jonów i protonów. Pochodzą one z wielu źródeł – emitowane są przez słońce, ale pochodzą także spoza naszego układu słonecznego, głównie z innych gwiazd, w szczególności wybuchów supernowych, a także z blazarów – aktywnych galaktyk, które emitują szerokie widmo promieniowania. Energia cząstek promieniowania kosmicznego może być niezwykle wysoka – o kilka rzędów wielkości wyższa niż np. sztucznie przyspieszane cząstki w najpotężniejszych akceleratorach, takich jak ten w ośrodku badawczym CERN. Ale to nie te cząstki bezpośrednio powodują awarię urządzeń na ziemi.



Rysunek 2. Kaskada cząstek generowanych w ziemskiej atmosferze na skutek „ostrzału” promieniowaniem kosmicznym (po lewej) oraz natężenie promieniowania w funkcji wysokości (po prawej)

Podczas podróży w kierunku Ziemi wiele cząstek jest odchylanych przez Słońce i ziemskie pole magnetyczne. Generowane wtedy tzw. promieniowanie hamowania – deszcz nowych, wtórnych, trzeciorzędowych cząstek (protonów, neutronów, elektronów, mionów itd.). Zjawisko to zachodzi także w naszej atmosferze, w związku, z czym natężenie promieniowania kosmicznego zmienia się wraz z wysokością. Do wysokości około 10...15 tysięcy metrów nad poziomem morza dominuje zjawisko generacji nowych cząstek, podczas gdy na wyższych pułapach dominuje ich absorpcja. Na powierzchni Ziemi dociera któraś z kolei generacja początkowego „pierwotnego kosmosu”. Zobrazowano to na **rysunku 2a**. Typowy strumień wynosi 20 neutronów na cm^2 na godzinę (na poziomie morza), jak pokazano na **rysunku 2b**. Z tego opisu można wywnioskować, że strumień ziemskich cząstek kosmicznych zależy od wysokości. Zależności od szerokości geograficznej, wynikające z wpływu ziemskiego pola magnetycznego i rzeczywistej aktywności Słońca, można pominąć przy oszacowaniu pierwszego rzędu.

Na ziemi także mamy źródła promieniowania, są to głównie instalacje atomowe, zwłaszcza siłownie jądrowe. W nich również musi pracować elektronika, często w ich radioaktywnym wnętrzu.

Tabela 1. Typowe poziomy dawek promieniowania w różnych środowiskach pracy urządzeń elektronicznych

Rodzaj promieniowania	Środowisko		
	Reaktor atomowy (normalna praca przez 40 lat)	Magazyn broni atomowej (praca przez 20...25 lat)	Przestrzeń kosmiczna (misja kosmiczna trwająca 5 lat)
Promieniowanie gamma	$10^3 \dots 10^8$ rad	2×10^3 rad	–
Neutrony	$10^9 \dots 10^{14}$ n/ cm^2	–	–
Elektrony i protony	–	–	$10^5 \dots 10^6$ rad

Tabela 2. Wpływ promieniowania gamma i rozprzeczonych neutronów na urządzenia krzemowe.

Rodzaj promieniowania	Poziom energii	Podstawowy mechanizm interakcji	Pierwotne efekty w krzemie i SiO_2	Wtórne efekty w krzemie i SiO_2
Fotony	Niska energia	Efekt fotoelektryczny	Zjawiska jonizacyjne	Przemieszczenie cząstek
	Średnia energia	Efekt Comptona		
	Wysoka energia	Kreacja par		
Neutrony	Niska energia (neutrony termiczne)	Wychwyt neutronu i dalsze reakcje jądrowe	Przemieszczenie cząstek	Zjawiska jonizacyjne
	Wysoka energia (neutrony szybkie)	Rozpraszanie elastyczne		

Elektronika pracuje także np. w magazynach broni atomowej i może z naszego punktu widzenia wydawać się to nieistotne, to jest wręcz przeciwnie – jej poprawne działanie jest gwarantem bezpiecznego składowania tego rodzaju uzbrojenia.

Charakter promieniowania, jaki spotyka układy elektroniczne na ziemi, jest inny, niż w przypadku promieniowania, na jakie wystawione są urządzenia w przestrzeni kosmicznej. W najnowszych generacjach reaktorów to głównie wysokie temperatury robocze (uzyskiwane w celu zwiększenia sprawności cieplnej) są problemem, ponieważ ograniczają żywotność czujników i materiałów elektronicznych. Ponadto reaktory te wykorzystują zarówno strumień neutronów szybkich, jak i strumień neutronów termicznych. Neutrony termiczne mają znacznie niższą energię, równą około 0,025 eV. Z kolei szybkie neutrony osiągają energie w zakresie 1...10 MeV.

Każdy rodzaj neutronów w inny sposób oddziałuje z materią i ma odmienne mechanizmy degradacji półprzewodników, które omówione będą w dalszej części. Na przykład powtarzająca się ekspozycja na strumień neutronów termicznych powoduje transmutację materiałów poprzez absorpcję tychże neutronów. Nie dzieje się tak z neutronami szybkimi, ponieważ mają one znacznie niższy przekrój czynny na absorpcję (prawdopodobieństwo, że zostaną zaabsorbowane przez inną cząstkę). W **tabeli 1** podsumowano typowe poziomy dawek różnego rodzaju promieniowania w różnych scenariuszach.

Wpływ promieniowania na urządzenia półprzewodnikowe

Głównym źródłem problemów, powodowanych przez promieniowanie kosmiczne w urządzeniach półprzewodnikowych są ciężkie jony. Te masywne, naładowane cząstki niosą ze sobą wysoką energię, a co za tym idzie mogą powodować poważne uszkodzenia urządzeń. Jednak mniejsze i lżejsze kwanty energii także mogą oddziaływać z materią elementów elektronicznych. Mówiąc o wpływie promieniowania na elektronikę, rozważyć musimy dwie klasy interakcji – chwilowe zdarzenia, nie pozostawiające po sobie śladów w strukturze półprzewodnika oraz obserwowany w czasie, destrukcyjny wpływ promieniowania na strukturę elementu (który na ogół też jest obserwowany dopiero po czasie, gdy układ wykazuje makroskopowe efekty degradacji materiałów).

To, jaki charakter, ma dane oddziaływanie zależne jest głównie od energii i rodzaju promieniowania, które bombarduje dany układ. Oczywiście, dawka także nie jest bez znaczenia, ale kluczowym aspektem jest energia kwantów promieniowania, „atakujących” półprzewodnik. Mechanizm typowej interakcji pomiędzy półprzewodnikiem, a kwantem promieniowania jest relatywnie prosty. Rozważmy typowy tranzystor polowy. Głównymi materiałami, z jakich składa się taki element są półprzewodnik (krzem) i izolator kanału (typowo dwutlenek krzemu). Rozważmy również dwa rodzaje promieniowania – jeden, pod postacią wysokoenergetycznych fotonów (np. promieniowanie gamma), a drugi pod postacią rozprzeczonych neutronów. W **tabeli 2** podsumowano różne mechanizmy interakcji. W ogólności, zjawiska związane z jonizacją w materiale są bardziej związane z zjawiskami chwilowymi (których, jak opisano powyżej, efekty mogą być

jednak długotrwałe dla układu), a te związane z przemieszczeniem atomów w strukturze krystalicznej są raczej długotrwałe i przekładają się na postępującą degradację materiału.

W ogólności obserwowane są trzy zjawiska, podczas interakcji cząstek wysokoenergetycznych z materiałem:

- Jonizacja materiału, poprzez interakcję z elektronami powłok atomów,
- Przemieszczenie atomów w materiale,
- Reakcje jądrowe atomów materiału.

Wszystkie te mechanizmy mogą koegzystować. Co więcej, efekty niektórych z nich, mogą wywoływać kolejne zdarzenia. Na przykład neutron może najpierw oddziaływać z jądrem, powodując uszkodzenie struktury krystalicznej (powstanie defektu) poprzez przemieszczenie atomu, a następnie generować wtórne, naładowane cząstki, które jonizują materiał, jeśli mają wystarczającą energię. Wszystko uzależnione jest od energii cząstki.

W przypadku zderzenia wysokoenergetycznych naładowanych cząstek przeważa efekt jonizujący. Tylko niewielki ułamek ich energii jest wykorzystywany do przemieszczenia atomów. Cząstki neutralne są odpowiedzialne z kolei raczej za uszkodzenia spowodowane przemieszczeniem lub nawet kaskady przemieszczeń, w przypadku cząstek o naprawdę wysokiej energii. Dalej, można by przystąpić do opisu fizyki, stojącej za tymi zjawiskami, jednak ramy czasopisma stanowiąco nie pomieściłyby takiego, nawet przyspieszonego kursu fizyki... Literatura wskazana się na końcu artykułu z pewnością może być dobrym punktem wyjścia do tego (karkołomnego, jeśli nie mamy już jakiegось wykształcenia fizycznego za sobą) zadania.

Chwilowy wpływ i zjawiska w półprzewodniku

Na powierzchni ziemi najczęściej będziemy obserwować pojedyncze zdarzenia – SEU, wynikające ze zderzenia cząstki promieniowania kosmicznego ze złączem PN półprzewodnika. Najczęściej dotyczy to zderzenia z układami pamięci, jak opisano dokładnie poniżej. Zakłócenia SEU są spowodowane uderzeniami promieniowania jonizującego w elementy pamięci, takie jak komórki pamięci konfiguracji, pamięć użytkownika i rejestry. W zastosowaniach naziemnych głównymi źródłami promieniowania jonizującego budzącymi obawy są cząstki alfa emitowane z radioaktywnych zanieczyszczeń w materiałach, wysokoenergetyczne neutrony wytwarzane w wyniku oddziaływania promieni kosmicznych z atmosferą ziemską oraz neutrony termiczne, które w większości przypadków są termalizowanymi neutronami wysokoenergetycznymi, ale mogą być również wytwarzane w sprzęcie stworzonym przez człowieka.

Badania prowadzone w ciągu ostatnich 20 lat doprowadziły do opracowania materiałów o wysokiej czystości do konstrukcji układów scalonych (szczególnie ich obudów), minimalizujących ilość występujących SEU, powodowane przez promieniowanie alfa. Nieuniknione neutrony atmosferyczne pozostają dziś główną przyczyną efektów SEU. Miękkie błędy są losowe i zdarzają się zgodnie z prawdopodobieństwem związanym z poziomami energii, strumieniem i podatnością komórek konkretnej pamięci półprzewodnikowej.

Z definicji SEU nie niszczą obwodów, ale mogą powodować błędy danych. W systemach komputerowych jedną z najbardziej wrażliwych części są na ogół pamięci podręczne 1 i 2 poziomu, ponieważ muszą być one bardzo małe i bardzo szybkie, co oznacza, że nie utrzymują dużego ładunku. Oznacza to, że nawet niewielkie wyładowanie wywołane zderzeniem z relatywnie niskoenergetycznym kwantem promieniowania kosmicznego może zmienić stan komórki pamięci. Podobna sytuacja ma miejsce w przypadku pamięci RAM, jednak układy pamięci operacyjnej systemów komputerowych są nieco mniej podatne na SEU, niż pamięć podręczna. Dokładnie tego rodzaju zdarzenie mogło mieć miejsce podczas wyborów w Belgii, co opisywano we wstępie do artykułu.

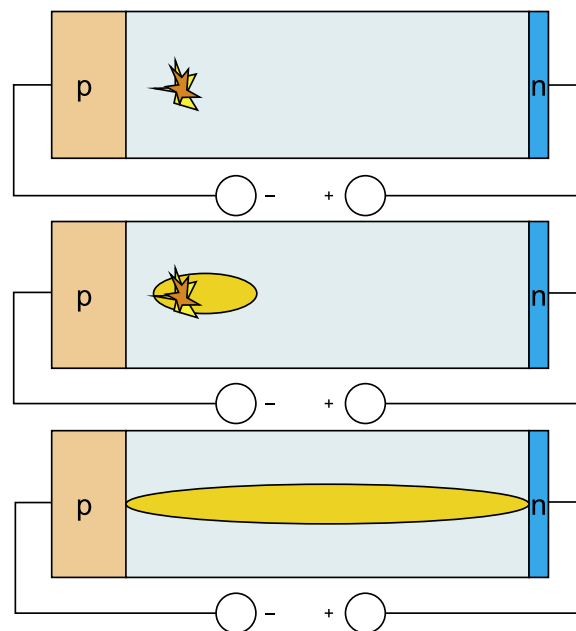
Kolejnym słabym punktem jest maszyna stanów w sterowaniu mikroprocesorem, ze względu na ryzyko wejścia w stany martwe (bez wyjść), jeśli zmieni się w nieprzewidywany sposób zmienna licznika

stanu. Taki efekt oczywiście mogą wywoływać też bardziej klasyczne błędy, takie jak naruszenia ochrony pamięci itd. W przypadku maszyny stanów, która steruje procesorem, ryzyko wystąpienia SEU jest relatywnie mniejsze, z uwagi na to, że obwody te muszą sterować całym procesorem, więc mają stosunkowo duże tranzystory, aby zapewnić stosunkowo duży prąd i nie są tak wrażliwe, jak mogłoby się wstępnie wydawać.

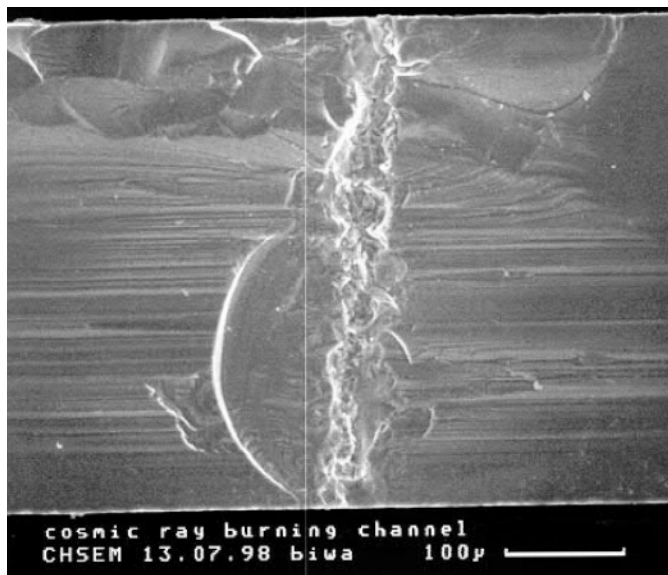
W obwodach cyfrowych i analogowych pojedyncze zdarzenie może spowodować propagację jednego lub więcej impulsów napięciowych (zakłóceń) w obwodzie, co określa się, jako występowanie tzw. stanów przejściowych pojedynczego zdarzenia (*Single Event Transient* – SET). Ponieważ propagujący się impuls nie jest technicznie zmianą stanu, jak w SEU w pamięci, należy rozróżnić SET i SEU. Jeśli SET propaguje się przez obwody cyfrowe, to może spowodować odczyt nieprawidłowej wartości w sekwencyjnej jednostce logicznej, wtedy sytuację taką uznaje się za SEU. Typowe SET to chwilowe szpilki napięcia, obserwowane w sygnale, które mogą np. powodować chwilowe zakłócenie pomiarów lub też doprowadzić nawet do uszkodzenia układu (patrz dalsza część artykułu).

Stały wpływ i uszkodzenia materiałów i układów

Większość cząstek kosmicznych przechodzi przez urządzenia półprzewodnikowe bez żadnej interakcji. Jeśli już dojdzie do interakcji, istnieje kilka jej mechanizmów, opisanych powyżej. Część takich interakcji może mieć katastrofalne skutki dla układów scalonych. Wyróżniamy dwa główne mechanizmy, opisane powyżej – jonizacyjny i defektowy. W przypadku jonizacji atomów, na ogół zjawisko takie powoduje niegroźne zdarzenia, takie jak zmiana wartości bitów itp. jednak, w przypadku układów, w których mamy do czynienia z dużym polem elektrycznym zdeponowana energia cząstki kosmicznej może prowadzić do zjonizowania lokalnej chmury ładunku, która jest wzmacniana przez pole elektryczne w elemencie. Na obciążonym urządzeniu, chmura może być na tyle duża, że powoduje zwarcie elementu i może zostać zaobserwowany krótki impuls wysokiego prądu. Efekt ten jest wykorzystywany w detektorach cząstek do eksperymentów fizycznych w celu identyfikacji i zliczania cząstek o wysokiej energii. W większości urządzeń półprzewodnikowych, które pracują z dużym polem elektrycznym (np. tranzystory mocy) zdeponowana energia może doprowadzić do powstania tzw. streamera, cienkiego, przewodzącego obszaru w blokującym przyrządzie



Rysunek 3. Mechanizm powstawania przewodzącego streamera w elemencie pracującym zaporowo



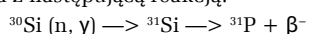
Rysunek 4. Efekt przepływu wysokiego prądu przez stramer, widziany w obrazie mikroskopowym przekroju diody półprzewodnikowej

półprzewodnikowym, jak pokazano na **rysunku 3**. W takim przypadku urządzenie może ulec zniszczeniu, jak pokazano na **rysunku 4**.

Z pewnym prawdopodobieństwem cząstki promieniowania o odpowiednio dużej energii, mogą po zderzeniu – elastycznym bądź nieelastycznym – z atomem, wybić go z jego miejsca, generując zaburzenie w strukturze krystalicznej półprzewodnika. Takie jednoatomowe zaburzenie sieci krystalicznej nazywamy defektem punktowym. Półprzewodniki pierwiastkowe grupy IV mają sieć krystaliczną w strukturze diamentu. Stała sieci dla krzemu wynosi 0,543 nm – o połowę tej odległości musi przesunąć się atom, aby wypadł z sieci krystalicznej. W ten sposób powstaje tzw. wakans (defekt w postaci pustego miejsca w strukturze krystalicznej) oraz atom pozasieciowy (defekt w postaci atomu, który znajduje się poza strukturą krystaliczną). Te pierwotne defekty są wysoce mobilne w temperaturze pokojowej i dlatego będą migrować na duże odległości. Mogą one albo zniknąć z materiału, rekombinując w materiale objętościowym, albo zostaną uwieszone przez inne atomy zanieczyszczeń, powodując powstanie bardziej stabilnych defektów wtórnych lub kompleksów defektów.

Wpływ takich defektów na działanie układu może być bardzo różny, ale niezmiennie negatywny. Pojedynczy defekt tego rodzaju nie jest zagrożeniem dla układu, ale już nagromadzenie tych defektów spowodować może np. zauważalne obniżenie prądu maksymalnego elementu, przewodności itp. Na szczęście większość cząstek promieniowania ma małą szansę na wywołanie takiego efektu. Wysokoenergetyczne fotony (np. promieniowanie gamma) niemalże nigdy nie oddziałują pierwotnie z atomami w ten sposób, jednak wygenerowane przez nich elektrony w zjawisku Comptona, podczas rozpraszania wstecznego, mogą już w ten sposób oddziaływać. Z drugiej strony, ciężkie, wysokoenergetyczne jony (takie jak promieniowanie alfa) tylko około 0,1% swojej energii przekazują w postaci przemieszczeń atomów struktury krystalicznej.

Trzecim efektem, pozostawiającym stałe uszkodzenia w krzemie, są reakcje jądrowe podczas napromieniowywania neutronami. Jeśli dochodzi do nieelastycznej absorpcji neutronów, przez atomy, dochodzi do tak zwanej transmutacji neutronowej. W przypadku krzemu mamy do czynienia z następującą reakcją:



Po zaabsorbowaniu neutronu termicznego krzem zmienia się w niestabilny izotop krzemu, a następnie transmutuje w atom fosforu przy towarzyszącej emisji promieniowania beta (elektronu). Zatem pod wpływem promieniowania krzem powoli zamienia się w fosfor, który stosowany jest, jako domieszka typu N w półprzewodnikach.

www.ti.com

SLVSA211 – JUNE 2011 – REVISED JANUARY 2014

16-Mb RADIATION-HARDENED SRAM

Check for Samples: SMV512K32-SP

FEATURES

- 20-ns Read, 13.8-ns Write Through Maximum Access Time
- Functionally Compatible With Commercial 512K x 32 SRAM Devices
- Built-In EDAC (Error Detection and Correction) to Mitigate Soft Errors
- Built-In Scrub Engine for Autonomous Correction
- CMOS Compatible Input and Output Level, Three State Bidirectional Data Bus
 - 3.3 ± 0.3-V I/O, 1.8 ± 0.15-V CORE
- Radiation Performance ⁽¹⁾
 - Uses Both Substrate Engineering and Radiation Hardened by Design (HBD) ⁽²⁾
 - TID Immunity > 3e5 rad (Si)
 - SER < 5e-17 Upsets/Bit-Day (Core Using EDAC and Scrub) ⁽³⁾
 - Latch up immunity > LET = 110 MeV (T = 398K)
- Available in a 76-Lead Ceramic Quad Flatpack
- Engineering Evaluation (EM) Samples are Available ⁽⁴⁾

(1) Radiation tolerance is a typical value based upon initial device qualification. Radiation Data and Lot Acceptance Testing is available – contact factory for details.

(2) HardSiT™ technology and memory design under a license agreement with Silicon Space Technology (SST).

(3) SER calculated using CREME96 for geosynchronous orbit, solar minimum.

(4) These units are intended for engineering evaluation only. They are processed to a non-compliant flow (e.g. no burn-in, etc.) and are tested to temperature rating of 25°C only. These units are not suitable for qualification, production, radiation testing or flight use. Parts are not warranted for performance on full MIL specified temperature range of -55°C to 125°C or operating life.

Rysunek 5. Fragment karty katalogowej pamięci SRAM o zwiększonej odporności na promieniowanie

O efektach powolnego zmieniania się krzemu typu P w typ N w urządzeniu półprzewodnikowym pisać nie trzeba – są oczywiste. Podobne efekty występować mogą w innych półprzewodnikach, np. german pod wpływem neutronów transmutuje głównie w gal, arsen lub selen, zależnie od tego, z jakim izotopem germanu mamy do czynienia.

Statystyki opisujące podatność elementów

W przypadku chwilowych awarii układu, powodowanych przez promieniowanie, określa się maksymalną ilość takich awarii w przeliczeniu na jednostkę czasu w założonych warunkach ekspozycji na promieniowanie. Im mniejsza jest ta wartość, tym lepiej.

Centralnym parametrem używanym do opisywania odporności elementów na trwałe uszkodzenia powodowane przez promieniowanie jest całkowita pochłonięta dawka (*Total Integrated Dose* – TID). Określa się dla elementów pewną wartość TID, dla której nie występują w układzie żadne negatywne efekty (a ściślej mówiąc – występujące efekty nie wychodzą poza określoną tolerancję). Im wyższa dawka, tym układ jest bardziej odporny.

Na **rysunku 5** pokazano fragment przykładowej karty katalogowej układu scalonego (w tym wypadku statycznej pamięci RAM) o zwiększonej odporności na promieniowanie. W karcie katalogowej podano TID – nie mniejsze niż 3×10^5 rad oraz częstotliwość występowania SER – poniżej 5×10^{-17} na bit, na dzień w warunkach orbity geostacjonarnej, przy minimum aktywności słonecznej (warunki te są znormalizowane i zgodne z standardem, w tym przypadku CREME96, czyli narzędziem stworzonym przez NASA do estymacji poziomu promieniowania).

Podsumowanie

Układy o zwiększonej odporności na promieniowanie są klasycznie potrzebne w szeregu dziedzin, takich jak elektronika kosmiczna, awionika i systemy pracujące w elektrowniach atomowych. Z uwagi na rosnącą podatność na wpływ promieniowania, wraz ze zmniejszaniem się wymiaru charakterystycznego urządzeń półprzewodnikowych, odporność na promieniowanie kosmiczne jest coraz bardziej istotna w zwykłej, ziemskiej elektronice, przynajmniej w systemach o wymaganej wysokiej dostępności/bezawaryjności, jak elektronika medyczna czy przemysłowe systemy bezpieczeństwa funkcjonalnego.

Badania nad elektroniką tolerującą promieniowanie gwałtownie wzrosły w ciągu ostatnich kilku lat, co zaowocowało wieloma interesującymi podejściami do modelowania efektów promieniowania i projektowania odpornych na promieniowanie układów scalonych i systemów wbudowanych. Badania te są silnie napędzane rosnącym zapotrzebowaniem na elektronikę odporną na promieniowanie do zastosowań kosmicznych, systemów eksperymentalnych fizyki wysokich energii, takich jak Wielki Zderzacz Hadronów w CERN, oraz wieloma ziemskimi systemami. Chociaż efekty całkowitej dawki

w masowych układach CMOS są dobrze znane, wciąż niewiele wiadomo na temat wpływu promieniowania na układy produkowane w zaawansowanych technologiach, takich jak (FD-)SOI czy FinFET.

Marzenie o umożliwieniu wysokowydajnych obliczeń, przetwarzania sygnałów i komunikacji w najtrudniejszych i najbardziej zróżnicowanych środowiskach z występowaniem promieniowania stawia elektronice wiele wyzwań. Nieuchronnie łączy badaczy z kilku dyscyplin, począwszy od fizyki jądrowej i fizyki ciała stałego, poprzez zaawansowane metody modelowania i kreatywne techniki projektowania obwodów, aż po zastosowanie progresywnych algorytmów i głębokiego uczenia maszynowego w celu optymalizacji wydajności systemu dla najróżniejszych zastosowań w najtrudniejszych warunkach.

Nikodem Czechowski, EP

Źródła:

1. C. Baraniuk „The computer errors from outer space”, BBC Future 12 października 2022.
2. J.M.G. Strachan-Deol, K.L. Knowles, L.L. Thompson, A.L. Kelly, „P6 3 Error 404”, Journal of Physics Special Topics 20 (2021).
3. P. Murley, G. Srinivasan, „Soft-error Monte Carlo modeling program, SEMM”, IBM Journal of Research and Development 40 (1996).
4. M. Roser, H. Ritchie, E. Mathieu „Technological Change” OurWorldInData.org (2013).
5. K.E. Holbert, „Space Radiation Environmental Effects. Courses in Electrical Engineering”, Arizona State University. 2007.
6. R. Sayyah, T.C. Macleod, F.D. Ho „Radiation-hardened electronics and ferroelectric memory for spaceflight systems”, Ferroelectrics 413 (2011).
7. A. Belousov, „Radiation Effects on Semiconductor Devices in High Energy Heavy Ion Accelerators” praca doktorska, Darmstadt 2014.
8. M. Moll „Radiation Damage in Silicon Particle Detectors”, praca doktorska, Hamburg 1999.
9. „Cosmic ray failures of power semiconductor devices” ABB Semiconductors Native Post – Cosmic Ray 06/2019.
10. C. Findeisen, E. Herr, M. Schenkel, R. Schlegel H. Zeller, „Extrapolation of cosmic ray induced failures from test to field conditions for IGBT modules” Microelectronics Reliability 38 (1998).
11. C. Claeys, E. Simoen, „Basic Radiation Damage Mechanisms in Semiconductor Materials and Devices. In: Radiation Effects in Advanced Semiconductor Materials and Devices”. Springer Series in Materials Science 57, Berlin, Heidelberg. 2002.
12. https://ecss.nl/item/?glossary_id=1628
13. <https://www.ti.com/product/SMV512K32-SP>
14. (red.) P. Leroux, „Radiation Tolerant Electronics”, MDPI 2019.

REKLAMA

Miernik uniwersalny 5999 [V, A, Ω, F, Hz, DutyC, temp] True RMS, MT-1707 Pro'sKit

Pomiary, zakresy:

napięcie DC [V]:	600 m/6/60/600/1000 ±(0,5%+3)
napięcie AC [V]:	6/60 ±(0,8%+3); 600/750 ±(1%+10)
prąd DC [A]:	600 μ/60 m/600 m ±(0,8%+3); 10
±(1,5%+10)	
prąd AC [A]:	60 m/600 m ±(1%+3); 10 ±(2%+10)
rezystancja [Ω]:	600/6 k/60 k/600 k ±(0,8%+3); 60 M/600
M ±(1,2%+30)	
pojemność [F]:	1 n~9.999 n ±(4,0%+30); 10 n~1 m
	±(2,5%+10); 1 m~60 m ±(5,0%+30)
częstotliwość [Hz]:	9.999 Hz~9.999 MHz ±(1%+3)
współczynnik wypełnienia [Duty Cycle]:	0,1%...99,9%
temperatura [°C/°F]:	-20°C do 1000°C ±(1%+3)

Funkcje, cechy:

wyświetlacz LCD 5999 podświetlany, podświetlenie miejsca pomiaru, True RMS, test diody, test ciągłości obwodu, Data Hold, wybór zakresu: ręczny, impedancja wejściowa do pomiaru napięcia DC ok. 10 MΩ, Auto Power Off, wskaźnik polaryzacji, wskaźnik przekroczenia zakresu, wskaźnik niskiego napięcia baterii, NCV – wbudowany bezkontaktowy detektor napięcia AC, zabezpieczenia: bezpiecznik 0,5 A/1000 V i 10 A/1000 V, normy: CE, CAT. III – 1000 V, CAT. IV – 600 V, zasilanie 1× bateria 9 V (np. 6F22), wymiary: 190×89×53 mm, waga netto: 315 g, waga brutto: 510 g

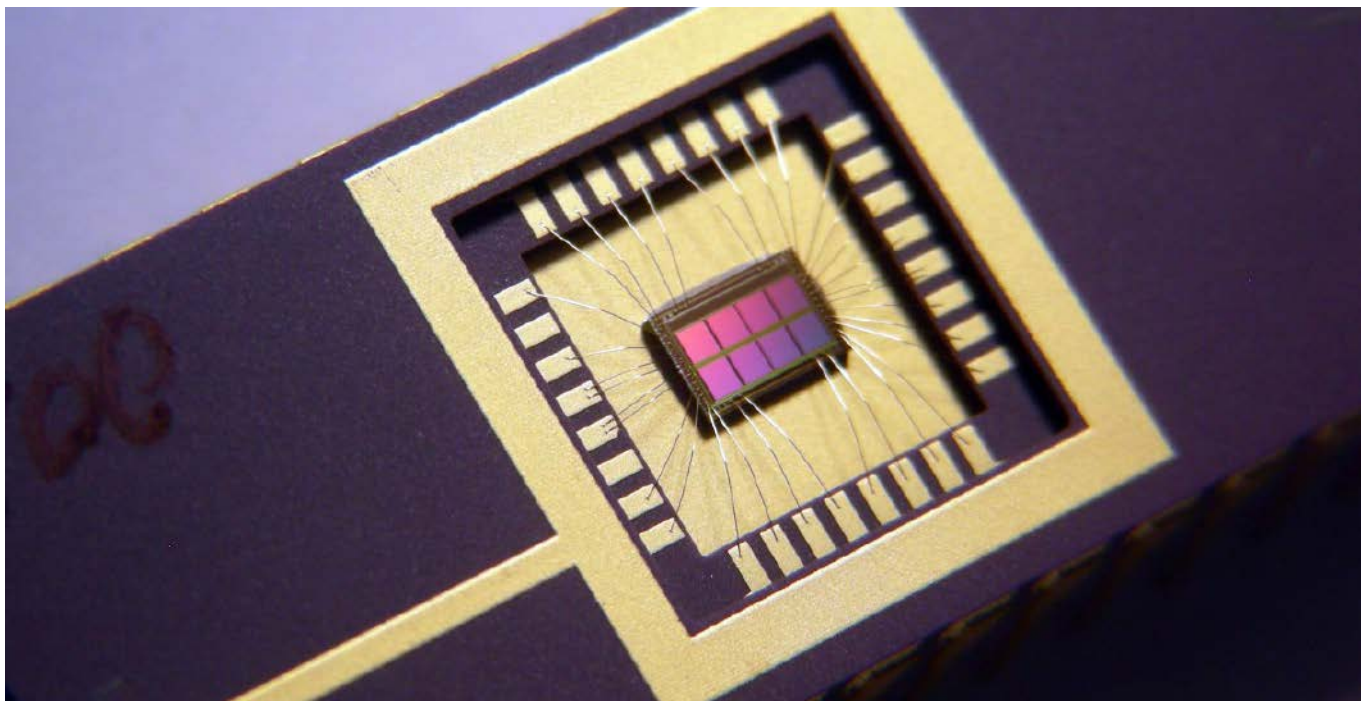
159,00 zł



AVT SPV Sp. z o.o.
03-197 Warszawa, ul. Leszczyńska 11
tel. +48 22 257 84 49, handlowy@avt.pl
sklep.avt.pl



Przedstawiona oferta cenowa ma charakter informacyjny i nie stanowi oferty handlowej w rozumieniu Art.66 par.1 Kodeksu Cywilnego



Strategie ochrony elektroniki przed promieniowaniem

Promieniowanie jonizujące, rozpędzone cząstki... istnieje wiele atomowych zagrożeń dla elektroniki pracującej w środowiskach kosmicznych, awionice czy reaktorach atomowych. W związku, z tym trzeba zabezpieczać urządzenia, aby mogły one (w miarę) niezawodnie pracować. W zaprezentowanym artykule omówimy kilka strategii, jakie stosuje się w urządzeniach elektronicznych, aby skompensować ich podatność na promieniowanie.

Wszechświat jest przepełniony różnymi rodzajami promieniowania. W większości środowisk, tutaj na ziemi, jesteśmy odizolowani od niego, ale wraz z coraz intensywniejszą eksploracją przestrzeni kosmicznej, przez coraz precyzyjniejszą elektronikę, problem ten staje się coraz wyraźniejszy. W ciągu ostatnich dekad na całym świecie wysłano wiele bezzałogowych i załogowych misji w przestrzeń kosmiczną. Te naszpikowane elektroniką pojazdy muszą być ekstremalnie niezawodne z uwagi na koszty misji kosmicznych, jak i potencjalne zagrożenie życia, w przypadku misji załogowych.

Urządzenia elektroniczne byłyby nieprzydatne na orbicie, gdyby nie tolerowały promieniowania kosmicznego i regularnie ulegały awariom. Z tego powodu promieniowanie, jego wpływ na elektronikę i różne metody zapobiegania lub korygowania niekorzystnego wpływu promieniowania na obwody elektroniczne stały się dużym obszarem badań w inżynierii lotniczej. Projektowanie i wykonywanie obwodów lub urządzeń odpornych na promieniowanie nazywa się często hartowaniem radiacyjnym obwodu lub urządzenia. Oznacza to, że urządzenia lub obwody są utwardzone na działanie promieniowania i stąd nazwa – rad-hard.

Spośród wielu technologii, które są stosowane do implementacji struktur rad-hard, takich jak elementy CMOS, BJT i FPGA, układy CMOS zyskały ogromną popularność, jako korzystna opcja dla

elektroniki kosmicznej. Zalety, jakie mają – niewielki rozmiar, duża szybkość działania, a co najważniejsze, ich zdolność do pracy przy niskim zużyciu energii, są niesamowicie istotne w przypadku urządzeń pracujących w kosmosie przy bardzo ograniczonym budżecie mocy.

Istnieje kilka aspektów, jakie są optymalizowane w układach, aby utwardzić je na wpływ promieniowania. Zabiegi te skupiają się głównie w dwóch obszarach – konstrukcji samych urządzeń półprzewodnikowych oraz architekturze układu, która jest istotna zwłaszcza w przypadku systemów cyfrowych, gdzie rozwiązania programowe mogą kompensować pewną część problemów wywoływanych przez promieniowanie. W dalszej części artykułu przyjrzymy się obu obszarom, które służą do utwardzania systemów elektronicznych.

Mechanizmy oddziaływania promieniowania

Istnieje kilka mechanizmów opisujących, jak różnego rodzaju promieniowanie oddziałuje z urządzeniami półprzewodnikowymi. Sposób i siła tego oddziaływania zależne są głównie od rodzaju cząstki i jej energii. Wysokoenergetyczne cząstki, które silnie oddziałują z jądrami powodują uszkodzenia półprzewodnika, zmianę domieszkania itd., przyczyniając się do degradacji elementu i finalnie jego uszkodzeniu. Rozpędzone elektrony z kolei, uderzając w półprzewodnik, mogą powodować jego jonizację, co ma wpływ na działanie i może indukować zakłócenia w postaci szpilek napięcia. To może doprowadzić do uszkodzenia układu. Dokładniejszy opis tych mechanizmów można znaleźć w drugim artykule poświęconym tematyce rad-hard, jaki prezentujemy w tym wydaniu „Elektroniki Praktycznej”.

Awarie

Kiedy ciężki jon uderza w obwód elektroniczny i przechodzi przez jego czuły punkt, generuje przejściowy impuls napięcia w tym punkcie. Nazywa się to zdarzeniem przejściowym (SET). Kiedy taki jon uderza np. obszar dyfuzji w urządzeniu MOS, SET indukuje

jonizację w tym obszarze, co często powoduje zmianę stanu logicznego. Tego rodzaju błędy stanów są określane, jako zakłócenia pojedynczego zdarzenia (SEU).

Impuls jonizacji indukowany przez ciężkie jony może spowodować włączenie struktury tyrystorowej p-n-p-n (obecnej w objętościowych obwodach CMOS). Powoduje to utrzymanie struktury w niskiej impedancji między katodą a anodą, a pasożytniczy tyrystor pozostaje włączony przez cały czas. Nazywa się to zatrzaśnięciem układu może być bardzo destrukcyjnym efektem promieniowania jonizującego, ponieważ powoduje, że urządzenie nie reaguje na żadne sygnały sterujące.

Uszkodzenia

Do uszkodzeń urządzenia dochodzi na skutek kilku mechanizmów. Wszystkie one prowadzą finalnie do awarii układu, jednak mogą się różnić tempem, czy też dynamiką pojawiania się. Pewne uszkodzenia będą nagłe, inne będą powodowały powolne pogarszanie parametrów układu.

Jonizacja nośników w półprzewodniku może doprowadzić nie tylko do przełączenia się stanu logicznego w układzie, ale także do powstania lokalnego zwarcia, co oczywiście prowadzi do przepływu sporego prądu, który może uszkodzić dany fragment struktury półprzewodnikowej.

W przypadku innych zjawisk obserwuje się dwa mechanizmy – kinetyczne uszkodzenia półprzewodników, na skutek zderzeń pomiędzy atomem półprzewodnika i cząstką promieniowania (najczęściej zjonizowanym jonem, protonem itp.) oraz efekty atomowe, np. transmutacja, na skutek zderzenia neutronu z jądrem atomowym.

Pierwszy z efektów powoduje powstanie defektu, a nawet dwóch, w strukturze półprzewodnikowej. W momencie, gdy atom zostaje wybity z sieci krystalicznej, pozostaje po nim puste miejsce, tzw. wakans, który jest rodzajem defektu struktury krystalicznej półprzewodnika. Wybity atom lokuje się poza siecią, formując tzw. atom pozasieciowy – inny rodzaj defektu. Tego rodzaju defekty półprzewodnika powodują, generalnie, pogorszenie jego rozmaitych parametrów, związanych z przewodnictwem itp.

Drugi z efektów skutkuje zmianą atomów krzemu czy domieszek w inne atomy na skutek reakcji jądrowych. Dokładny opis tych procesów wykracza poza ramy tego artykułu, ale wystarczy wspomnieć, że może prowadzić do zmiany domieszkowania półprzewodnika z N na P i odwrotnie. Trochę więcej informacji na temat tych zjawisk znaleźć można w sąsiednim artykule poświęconym wpływowi promieniowania na półprzewodniki. Oczywiście, w przypadku degradacji na skutek pojawiania się defektów czy transmutacji atomów i zmian domieszkowania półprzewodnika, kluczem jest skala. Pojedyncze zdarzenie tego rodzaju nie ma istotnego wpływu na większość urządzeń półprzewodnikowych. Jednak uszkodzenia te, z biegiem czasu w promieniotwórczym środowisku, kumulują się, w końcu doprowadzając do istotnego pogorszenia parametrów układu, czyniąc go, de facto, nieużywalnym – niesprawnym.

Rozwiązania sprzętowe

Podstawowym krokiem, jaki można przedsięwziąć, jest utwardzenie samych urządzeń półprzewodnikowych. Robi się to na ogół na etapie projektowania i wytwarzania struktur półprzewodnikowych tych urządzeń, dlatego też w tym artykule omówimy jedynie główne techniki, stosowane w tych procesach, nie zagłębiając się zbyt w ich mechanizmy działania. Tematyka projektowania i wytwarzania struktur półprzewodnikowych, jakkolwiek szalenie ciekawa, leży poza zakresem tematycznym EP.

Podłoża izolacyjne

Utwardzone chipy są często produkowane na podłożach izolacyjnych zamiast zwykłych substratów półprzewodnikowych. Powszechnie stosowane są technologie, takie jak krzem na izolatorze (SOI) czy nawet krzem na szafirze (SOS). SOI wykorzystuje warstwę izolującą na podłożu krzemowym zamiast konwencjonalnych

substratów krzemowych w celu zmniejszenia pojemności pasożytniczych. Różnica między urządzeniami SOI a konwencjonalnymi polega głównie na tym, że złącze krzemowe znajduje się nad izolatorem elektrycznym. Pierwsze przemysłowe wdrożenie SOI miało miejsce w sierpniu 1998 roku. Korzyści płynące z techniki SOI w porównaniu z konwencjonalnymi technologiami produkcji to:

- niższa pojemność pasożytnicza dzięki izolacji od krzemu objętościowego w substracie,
- poprawiona odporność na zatrzaśnięcie dzięki całkowitej izolacji struktur domieszkowanych N oraz P.

Z punktu widzenia produkcji podłoża SOI są kompatybilne z większością konwencjonalnych procesów. Podstawową barierą we wdrażaniu SOI jest drastyczny wzrost kosztów substratów, który przyczynia się do wzrostu całkowitego kosztu produkcji o około 10...15%.

SOS to heteroepitaksjalny proces produkcji układów scalonych, składający się z cienkiej warstwy krzemu wyhodowanej na płycie szafirowej. Proces ten jest częścią rodziny technologii CMOS SOI. Jest używany głównie w zastosowaniach lotniczych i wojskowych ze względu na swoją naturalną odporność na promieniowanie. Pierwsza zaleta szafiru polega na tym, że jest on doskonałym izolatorem elektrycznym, zapobiegającym rozprzestrzenianiu się prądów błądzących spowodowanych promieniowaniem. Drugą zaletą jest to, że może być on wytwarzany w tych samych fabrykach, które produkują zwykłe układy na zwykłych substratach krzemowych. Wadami SOS jest głównie bardziej złożony proces produkcji i wysoka cena podłoży szafirowych. Dodatkowo, podłoża te są cięższe niż układy krzemowe, co w pewnych przypadkach może być problemem.

Jeśli chodzi o tolerancję na promieniowanie, podczas gdy zwykłe chipy komercyjne mogą wytrzymać od 5 do 10 krad, chipy klasy kosmicznej produkowane w technologii SOI lub SOS mogą przetrwać dawki, o co najmniej kilka rzędów wielkości wyższe. Kiedyś wiele chipów z serii 4000 było dostępnych w wersjach utwardzanych, obecnie układy te są mniej popularne.

Bipolarne układy scalone

Bipolarne układy scalone zawierają tranzystory bipolarne (BJT) jako główne elementy. Tranzystory tego rodzaju do działania potrzebują zarówno elektronów, jak i dziur. Przepływ ładunku w BJT wynika z dwukierunkowej dyfuzji nośników ładunku przez złącze między dwoma obszarami o różnych stężeniach nośników. Ten sposób działania kontrastuje np. z tranzystorami polowymi, w których tylko jeden typ nośnika jest zaangażowany w działanie.

Nowa koncepcja Super Junction (SJ) dla energoelektroniki może pomóc w podtrzymaniu trendu zmniejszania strat przełączania przyszłej elektroniki. Tranzystory bipolarne z SJ (SJBT) wykazują wiele podobieństw z superzłączowym tranzystorem MOSFET. Po kilku dekadach rozwoju, pseudomorficzny tranzystor o dużej ruchliwości elektronów wykonany z GaAs (PHEMT) okazał się być urządzeniem o wysokiej sprawności i niskim koszcie, który można komercyjnie produkować. Tego rodzaju elementy mogą skutecznie konkurować z układami polowymi, co jest istotne, bo jeśli chodzi o tolerancję na promieniowanie, bipolarne układy scalone mają generalnie wyższą tolerancję promieniowania niż obwody CMOS. Seria LS5400 może wytrzymać 1000 krad, a wiele urządzeń ECL może wytrzymać do 10 000 krad.

SRAM odporny na promieniowanie

Aby tolerować promieniowanie, pamięć DRAM bazująca na kondensatorach jest często zastępowana przez bardziej wytrzymałą (ale większą i droższą) pamięć SRAM. Urządzenie SRAM używane są głównie do odczytu, w aplikacjach takich, jak pamięć konfiguracyjna dla układów FPGA, można uodpornić na efekty promieniowania do bardzo wysokiego poziomu, dodając rezystor o dużej wartości do komórek pamięci.

Urządzenie SRAM używane głównie w stanie odczytu są zwykle zapisywane tylko raz po włączeniu zasilania, aby zdefiniować funkcję

układu scalonego, a w większości aplikacji nigdy nie jest zmieniana zawartość zaprogramowanej po włączeniu zasilania pamięci. Możliwa jest również produkcja SRAM na podłożach SOI, co oferuje wartościową charakterystykę utwardzania układu pamięci na wynik pojedynczego zdarzenia (SEU) czy zwiększenia całkowitej dopuszczalnej dawki promieniowania dla takiego układu.

Podłoże z materiału o szerokim pasmie wzbudzenia

Wyższą tolerancję na defekty można uzyskać stosując podłoże o szerokiej przerwie energetycznej, gdyż materiały te na ogół mają małe atomy i silnie elektrycznie wiązania atomowe, co zmniejsza prawdopodobieństwo wytwarzania defektów przez promieniowanie.

Półprzewodniki z szeroką przerwą energetyczną, takie jak azotek galu (GaN) czy węgiel krzemu (SiC), okazały się jednymi z najbardziej obiecujących materiałów na przyszłe komponenty elektroniczne wysokiej mocy. Oferują one ogromne korzyści w zakresie sprawności energetycznej, ale także wysoką niewrażliwość na promieniowanie, możliwość działania w wysokich temperaturach itp. Nadal w tym zakresie prowadzi się wiele prac badawczo-rozwojowych. W szczególności konieczna jest poprawa jakości substratów krystalicznych, aby zwiększać wydajność i niezawodność tych elementów. Konieczne są także dalsze prace badawcze, aby lepiej zrozumieć fizykę tych półprzewodników, poprawić wzrost materiałów i dalej zoptymalizować sprawność tych urządzeń.

Ochrona układu przed promieniotwórczością

Jest to intuicyjne działanie, mające na celu zmniejszenie ekspozycji struktury krzemowej na promieniowanie. Naukowcy badają różne materiały, nawet tak egzotyczne jak gleba księżycowa, jako osłonę przed promieniowaniem kosmicznym. Badania te opierają się na pomiarach, jak i modelowaniu zjawisk fizycznych. Zastosowanie odpowiednich materiałów zapewnia znaczną ochronę przed pierwotnym promieniowaniem kosmicznym i zdarzeniami związanymi z cząstkami słonecznymi. Głównym czynnikiem, jaki mówi o efektywności w ekranowaniu danego pierwiastka jest jego liczba atomowa Z – im wyższa, tym generalnie dany ekran jest skuteczniejszy. Badane są metale takie jak aluminium reprezentujące materiał o niskim/średnim Z czy wolfram reprezentujący materiał o wysokim Z . Wyniki obliczeń wskazują, że dla tłumienia promieniowania wymaganego dla typowej elektroniki używanej w misji Jowisza, materiał o niskim Z w połączeniu z materiałem o wysokim Z są gorszą osłoną na tę samą masę powierzchniową materiału o wysokim Z . Jednakże, gdy do ochrony elektroniki bardzo wrażliwej na promieniowanie wymagane jest masowe ekranowanie o masie powierzchniowej powyżej 10 g/cm², wówczas połączenie warstw z niską i wysoką liczbą atomową zapewnia optymalniejszy i co istotne lżejszy ekran radiacyjny.

Wytwarzanie chipów za pomocą zubożonego boru

Zubożony bor składa się tylko z izotopu boru-11. Jest to istotne, ponieważ promieniowanie kosmiczne wytwarza wtórne neutrony, jeśli uderzy w struktury statku kosmicznego, a neutrony powodują rozszczepienie boru-10, jeśli jest on obecny w półprzewodnikach statku kosmicznego. Podczas takiego rozszczepiania wytwarza się promieniowanie gamma, cząstka alfa i jon litu. Powstałe produkty rozszczepienia mogą przekazać swój ładunek do pobliskich struktur półprzewodnikowych, powodując różnego rodzaju problemy.

W konstrukcjach półprzewodników utwardzanych na promieniowanie jednym ze środków zaradczych jest użycie zubożonego boru do domieszkowania. Prawie nie zawiera on boru-10. Bor-11 jest w dużej mierze odporny na promieniowanie. Ogólnie rzecz biorąc, zubożony bor jest stosowany w warstwie pasywacyjnej szkła borokrzemianowego w celu ochrony układów scalonych. Szkło borofosfokrzemianowe (BPSG), jest rodzajem szkła krzemianowego, które zawiera dodatki zarówno w postaci boru, jak i fosforu.

Rozwiązania architektoniczne i programowe

Inne techniki hartowania systemów obejmują stosowanie środków logicznych, programowych i architektonicznych, aby zwiększyć odporność całego systemu na promieniowanie. Wymagają one oceny niezawodności systemu, aby wykryć miejsca kluczowej podatności i móc zastosować odpowiednie środki zaradcze.

Korekcja błędów pamięci

Pamięć DRAM zapewnia zwiększoną ochronę przed błędami w oparciu o kody korekcji błędów. Pamięć z korekcją błędów, znana, jako pamięć ECC (*Error Correcting Codes*) jest szczególnie odpowiednia dla aplikacji o wysokiej odporności na błędy, takich jak serwery, a także aplikacji kosmicznych ze względu na promieniowanie. Zawiera dodatkowe bity parzystości do sprawdzania i ewentualnie poprawiania uszkodzonych danych.

Ponieważ efekty promieniowania mogą zmienić zawartość pamięci, nawet, jeśli system nie uzyskuje w danym momencie dostępu do pamięci RAM, do ciągłego czyszczenia pamięci RAM należy używać tak zwanego obwodu skrubera. Zwykle następuje to w trzech krokach podczas odświeżania pamięci:

- odczyt danych,
- sprawdzanie bitów parzystości pod kątem błędów danych,
- zapisywanie poprawek do pamięci RAM.

Tradycyjne kontrolery pamięci z korekcją błędów wykorzystują tzw. kody Hamminga, chociaż niektóre mogą korzystać z potrójnej redundancji modułowej (TMD). Przeplot pozwala rozłożyć efekt pojedynczego promienia kosmicznego, który potencjalnie zaburza wiele fizycznie sąsiadujących bitów w wielu słowach poprzez powiązanie sąsiednich bitów z różnymi słowami. Dopóki pojedynczy błąd zdarzenia (SEU) nie ma zasięgu większego niż pewien ustalony próg błędów w jakimkolwiek konkretnym słowie między dostęпами, można go poprawić i utrzymać wolne od błędów dane dla systemu.

Schematy korekcji błędów są szeroko stosowane zarówno w architekturach pamięci, jak i przy komunikacji. Jednak najnowocześniejsze technologie CMOS i dyskowe mają bardzo mały poziom błędów, który może sięgać zaledwie jednego miliarda, dlatego rygorystyczna korekcja błędów nie zawsze jest konieczna. Wraz z postępującą miniaturyzacją układy są jednak coraz bardziej narażone na błędy w sygnałach, przez co konieczne jest wprowadzanie nowych rozwiązań. W skali nano wymagane są silniejsze i wielokrotne kody korygujące błędy, takie jak kody BCH. Czy też nowe układy do realizacji procedur dekodowania i kodowania kodów Hamminga, w którym zwiększa się niezawodność pamięci kosztem tylko kilku nanosekund opóźnienia w czasie dostępu do pamięci.

Chociaż kody Hamminga są w stanie skorygować pojedynczy błąd w bloku fizycznych bitów, stają się mniej produktywnie w przypadku wysokiego wskaźnika błędów. W zastosowaniach praktycznych kod BCH (250, 32, 45) może zapewnić 99,9956% poprawności przy nawet 10% bitowej stopie błędów w pamięci, ale oczekuje się, że 1 bajt na każde 711 bajtów będzie wadliwy. Ogólnie rzecz biorąc, jeśli użyjemy tylko kodów korekcji błędów, konieczne będzie implementowanie bardzo silnych i złożonych kodów korekcji błędów, co spowoduje duży narzut w obszarze złożoności obliczeniowej czy opóźnień. Tym samym traci się też zalety, wynikające z stosowania pamięci produkowanych w technologiach z bardzo małym wymiarem charakterystycznym. Między innymi to jest przyczyną, że układy o zastosowaniach kosmicznych produkuje się raczej w starszych procesach technologicznych.

Redundancja podatnych elementów

W inżynierii redundancja to powielanie krytycznych komponentów systemu w celu zwiększenia jego niezawodności. W wielu systemach o znaczeniu krytycznym dla bezpieczeństwa, np. awionice, niektóre części systemu sterowania powinny być nawet potrójne. W potrójnie redundantnym systemie jego trzy komponenty podrzędne muszą ulec awarii, zanim nastąpi awaria całego systemu.

Ponieważ każdy z nich rzadko zawodzi i oczekuje się, że zawiedzie niezależnie. Prawdopodobieństwo, że wszystkie trzy zawiodą jednocześnie jest (dostatecznie) małe.

Nadmiarowe elementy mogą być stosowane na poziomie systemu lub na poziomie obwodu. Na poziomie systemu można zastosować np. trzy oddzielne płytki mikroprocesorowe, które mogą działać niezależnie. Można dodać logikę, aby wykrywać i wyłączyć płytke w przypadku powtarzających się błędów. Na poziomie obwodu z kolei pojedynczy bit można zastąpić trzema bitami i oddzielną logiką głosowania dla każdego bitu, aby w sposób ciągły określać jego wynik. Jednak ta strategia zwiększy obszar chipa pięciokrotnie, więc zarezerwowana jest tylko dla mniejszych projektów.

Podjęcie z zastosowaniem potrójnej redundancji ma tę zaletę, że systemy tego rodzaju w zasadzie nie mają negatywnego wpływu na wydajność działania systemu, w odróżnieniu od np. układów z zegarami kontrolnymi itp. opisanymi poniżej.

Systemy wielordzeniowe z głosowaniem

O szczególnym przypadku redundancji mówi się w przypadku systemów głosujących. Jest to architektura wykorzystująca kluczowe układy logiczne w postaci minimum trzech, równoległe pracujących systemów, połączonych na końcu układem „głosującym”. Obwód ten porównuje wyjścia tych identycznych subsystemów i wybiera odpowiedź większości. Na **rysunku 1** pokazano tego rodzaju architekturę w postaci schematu blokowego.

Jeśli jeden z obwodów doświadcza zdarzenia wywołanego promieniowaniem, które wpływa na wyjście, to jego wyjście będzie się różnić od pozostałych dwóch obwodów. Gdyby porównano tylko dwa identyczne obwody, różny stan wyjść pozwoli zidentyfikować, że nastąpiło zdarzenie, ale nie obwód, w którym wystąpiło, ani prawidłowy wynik działania. Przy trzech obwodach można określić prawidłowe wyjście (w oparciu o rozsądne założenie, że prawdopodobieństwo wystąpienia identycznych błędów w dwóch obwodach jest praktycznie zerowe). W zależności od oprogramowania, system może następnie zaakceptować dane wyjściowe większości modułów lub ponowić operację.

Na rynku dostępne są systemy tego rodzaju. W przypadku mikrokontrolerów/procesorów, umieszcza się w jednej obudowie trzy rdzenie. Często są one taktowane zegarem przesuniętym w fazie względem siebie i fizycznie rozsunięte lub obrócone, aby zminimalizować prawdopodobieństwo wystąpienia błędu w dwóch rdzeniach naraz.

Zegary kontrolne

W systemach krytycznych stosuje się również tzw. zegary nadzorujące, które mogą być wykorzystane do wykonania twardego resetu systemu, chyba, że zostanie wykonana jakaś sekwencja, która ogólnie wskazuje, że system działa, na przykład operacja zapisu w wbudowanego procesora. Przykładem takiego zegara może być Watchdog, implementowany w większości mikrokontrolerów.

Jeśli promieniowanie spowoduje nieprawidłową pracę procesora, jest mało prawdopodobne, że oprogramowanie będzie działać

wystarczająco poprawnie, aby np. wyczyścić licznik czasu zegara nadzorującego. Watchdog, gdy licznik zostanie przepełniony, wymusi reset systemu. Jest to uważane za ostatnią deskę ratunku w porównaniu do innych metod utwardzania radiacyjnego.

Ocena niezawodności

Należy zauważyć, że oprócz powyższych technik hartowania, bardzo ważnym tematem jest również sposób testowania niezawodności układu scalonego. Istnieje wiele technik i podejść statystycznych do testowania układów. Najnowsze z nich bazują na rozkładzie Weibulla. Kluczowe jest połączenie metod empirycznych i statystycznych bazujących na danych z testów żywotności układów, aby oszacować niezawodność komponentów systemu.

Kluczowe jest zbadanie wpływu czynników na przyspieszoną degradację układu, do czego często stosuje się sztuczne źródła promieniowania, takie jak np. akceleratory. Pozwala to na symulowanie określonych rodzajów promieniowania, a co za tym idzie, studiowanie jego wpływu na gotowy system, jego awaryjność, przyspieszone starzenie itd.

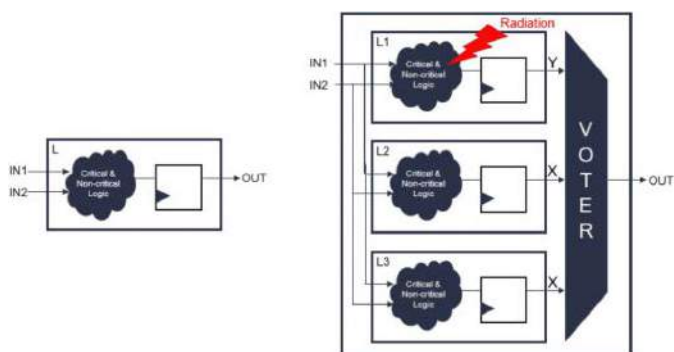
Podsumowanie

Zaprezentowane strategie zabezpieczania urządzeń półprzewodnikowych pozwalają zmniejszyć zawodność elektroniki narażonej na promieniowanie jonizujące, kosmiczne i inne. Tego rodzaju promieniowanie występuje w wielu środowiskach, ale kluczowe sektory, w jakich analizuje się podatność urządzeń półprzewodnikowych to systemy kosmiczne (satelity, pojazdy kosmiczne i inne), układy pracujące w reaktorach atomowych, systemy awioniki oraz, z uwagi na wymaganie wysokiej niezawodności, krytyczne elementy układów bezpieczeństwa itp. w systemach medycznych i bezpieczeństwa funkcjonalnego w przemyśle.

Poruszone w artykule aspekty sprzętowego zabezpieczania struktur półprzewodnikowych są, na ogół, poza zasięgiem elektroników, projektujących urządzenia elektroniczne. Nie oznacza to, że nie są one istotne. Dobierając układy do systemu narażonego na promieniowanie, należy starać się sięgać po układy rad-hard, które posiadają nie tylko zwiększoną odporność na promieniowanie, ale także są przebadane pod tym kątem, co pozwala na oszacowanie podatności i zawodności systemu. Szacowanie podatności systemu jest istotne, aby wprowadzić architektoniczne i programistyczne mechanizmy kompensowania podatności na promieniowanie. Istotna jest jednak wiedza, gdzie mogą występować problemy i jak często należy się ich spodziewać.

Zachowując odpowiednie podejście, modelując i testując zawodność układu, a także w wprowadzając odpowiednie środki zapobiegawcze na poziomie sprzętu i oprogramowania, można uzyskać systemy, które będą spełniały wymagania pracy w nawet najbardziej wymagających środowiskach, gdzie elementy elektroniczne narażone są na wysokie poziomy promieniowania.

Nikodem Czechowski, EP



Rysunek 1. Schemat blokowy układu z potrójną redundancją i systemem głosowania

Bibliografia:

1. Y. Fa-Xin, L. Jia-Rui, Zheng-Liang Huang, L. Hao, L. Zhe-Ming, „Overview of Radiation Hardening Techniques for IC Design”, Information Technology Journal 9 (2010).
2. R. Garg, N. Jayakumar, S. P. Khatri and G. S. Choi, „Circuit-Level Design Approaches for Radiation-Hard Digital Electronics” IEEE Transactions on Very Large Scale Integration (VLSI) Systems 17 (2009).
3. S. Nimmakayala, „A Survey of Radiation Hardened CMOS Techniques”, praca dyplomowa na Dept. of Electrical and Computer Engineering, Montana State University, 2009.
4. <https://bit.ly/3vDx7ic>
5. Y. Q. de Aguiar, „Predictive tools and Radiation-Hardening-by-Design (RHBD) techniques for SET and SEU in digital circuits”. Electronics. Université Montpellier, 2020

Stacja meteo, czyli czyszczenie szuflad (2)

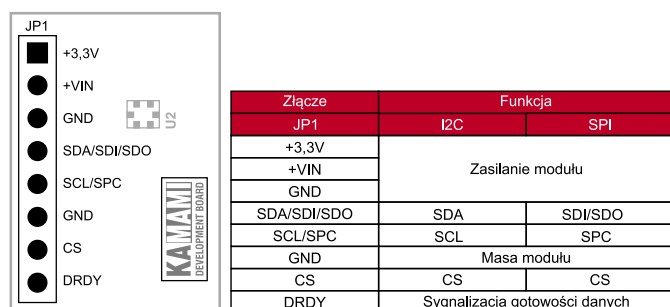
Jednym z etapów pracy nad projektem jest wybór i testowanie elementów składowych tak, aby spełniały założenia konstrukcyjne. Raz użyte moduły najczęściej lądują w szufladach czy regałach i po pewnym czasie może się uzbierać spora ilość dobrego sprzętu, z którym najczęściej nie wiadomo co zrobić. Kontynuujemy opis stacji meteo skonstruowanej zgodnie z założeniem, że użyte zostaną tylko części z zasobów nagromadzonych w szufladach.

Czujnik wilgotności KAmoHTS221

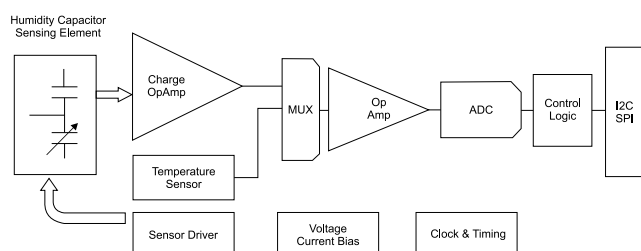
Moduł KAmoHTS221 ma wymiary 14×15×3 mm i podobnie, jak w przypadku KAmoLPS331, na dłuższej krawędzi umieszczono punkty lutownicze z wyprowadzeniami w rastrze 2,54 mm (**rysunek 13**). Podobnie, jak moduł barometru, moduł czujnika wilgotności ma możliwość zasilania napięciem w zakresie 2,5...5,5 V. Dwukierunkowe linie danych i zegara wyposażono w układ konwersji poziomów napięć.

Umieszczony na płytce modułu czujnik HTS221 mierzy wilgotność w zakresie od 20% do 80% RH z dokładnością $\pm 4,5\%$. Pomiar może być wykonywany z częstotliwością od 1 pomiaru na sekundę do 25 pomiarów na 2 sekundy. Oprócz pomiaru wilgotności czujnik mierzy temperaturę od -40°C do $+120^{\circ}\text{C}$ z dokładnością $\pm 0,5^{\circ}\text{C}$ w zakresie od $+15^{\circ}\text{C}$ do $+40^{\circ}\text{C}$, oraz $\pm 1^{\circ}\text{C}$ w zakresie 0°C do $+60^{\circ}\text{C}$.

Schemat blokowy układu został pokazany na **rysunku 14**. Podstawowym elementem toru pomiaru wilgotności jest sensor



Rysunek 13. Moduł czujnika wilgotności KAmoHTS221



Rysunek 14. Schemat blokowy czujnika HTS221



Rysunek 15. Sekwencja zapisania rejestru HTS221



Pierwsza część znajduje się pod adresem:
<https://ulubionykiosk.pl/media>

pojemnościowy. Jego pojemność zmienia się w zależności od wilgotności otoczenia. Konwersja zmiany pojemności na sygnał napięciowy jest wykonywana we wzmacniaczu Charge OpAmp. Sygnały napięciowe z tego wzmacniacza lub z czujnika napięciowego są konwertowane na postać cyfrową przez 16-bitowy przetwornik analogowo cyfrowy. Do komunikacji z mikrokontrolerem hostem może być zastosowany interfejs I2C lub SPI. My będziemy używać interfejsu I2C. Układ nie ma wyprowadzeń adresowych i 8-bitowy adres slave ma jedną wartość równą 0xBE dla zapisu danych i 0xBF dla odczytu danych.

Podobnie jak w przypadku czujnika ciśnienia konfiguracja parametrów pracy i odczytywanie mierzonych wartości odbywa się przez zapisywanie i odczytywanie wewnętrznych rejestrów. Standardowo będziemy potrzebować dwu funkcji: odczytania zawartości rejestru o podanym adresie i zapisania rejestru o podanym adresie. Zapisanie rejestru rozpoczyna się od wysłania przez mikrokontroler sekwencji START, a po niej adresu Slave 0xBE. Potwierdzenia adresu przez HTS221 pozwala na wysłanie przez mikrokontroler adresu rejestru, a po nim zapisywanej danej (**rysunek 15**).

Na **listingu 11** została pokazana procedura zapisywania rejestru z dwoma argumentami: **addr** – adres rejestru i **data** – dane do zapisania. **HTS_ADD** to adres slave HTS221. Pamiętajmy, że procedury obsługi magistrali I2C wygenerowane przez MCC wymagają podawania adresu 7-bitowego i HTS_ADD ma wartość 0x5F. Odczytanie zawartości rejestru rozpoczyna się od wysłania sekwencji START i adresu Slave z bitem R/W=0 (zapis) i bajt adresu rejestru. Potem jest wysyłana powtórna sekwencja START, adres Slave z bitem R/W=1, a po nim odczytywany jest jeden bajt zawartości zaadresowanego rejestru (**listing 12**, **rysunek 16**).

Listing 11. Zapisanie rejestru HTS221

```
#define HTS_ADD 0x5f
//zapisanie rejestru HTS221
void HTS221WriteReg(uint8_t addr, uint8_t data){
    I2C1_Write1ByteRegister(HTS_ADD, addr, data);
}
```

Listing 12. Odczytanie rejestru HTS221

```
uint8_t HTS221ReadReg(uint8_t addr){
    uint8_t data;
    data = I2C1_Read1ByteRegister(HTS_ADD, addr);
    return(data);
}
```


Zestawienie rejestrów sterujących i rejestrów wyników pomiarów zostało pokazane w **tabeli 4**. Pomiary wilgotności i temperatury mogą być wykonywane na żądanie, po wysłaniu polecenia (*one shot*), lub w sposób ciągły z określoną częstotliwością. Do programowania



Rysunek 16. Sekwencja odczytania rejestru HTS221

Tabela 4. Zestawienie rejestrów sterujących i rejestrów wyniku pomiaru układu HTS221

Rejestr	ADRES hex	POR	Funkcja
WHO_AM_I	0F	0xBC	ID układu
AV_CONF	10	0x7A	Rejestr konfiguracji rozdzielczości
CTRL_REG1	20	0	Rejestr kontrolny 1
CTRL_REG2	21	0	Rejestr kontrolny 2
CTRL_REG3	22	0	Rejestr kontrolny 3
STATUS_REG	27	0	Rejestr statusu
HUMIDITY_OUT_L	28	0	Młodsza część rejestru mierzonej wilgotności
HUMIDITY_OUT_H	29	0	Starsza część rejestru mierzonej wilgotności
TEMP_OUT_L	2A	0	Najstarsza część rejestru mierzonego ciśnienia
TEMP_OUT_H	2B	0	Najmłodsza część rejestru mierzonej temperatury

Tabela 5. Rejestr CTRL_REG1 układu HTS221

CTRL_REG1							
B7	B6	B5	B4	B3	B2	B1	B0
PD	res	res	res	res	BDU	ODR1	ODR2

PD – uruchamianie trybu obniżonego poboru mocy Power Down, PD = 0 układ uśpiony, PD = 1 układ aktywny.
BDU – odświeżanie rejestru wyjściowego, BDU = 0 dane odświeżane ciągle, BDU = 1 dane odświeżane po odczycie rejestrów ciśnienia i temperatury.
ODR1:ODR2 – częstotliwość odczytywania danych wyjściowych:
00 – pomiar na żądanie (One Shot),
01 – 1 Hz,
10 – 7 Hz,
11 – 12,5 Hz.

Tabela 6. Rejestr CTRL_REG2 układu HTS221

CTRL_REG2							
B7	B6	B5	B4	B3	B2	B1	B0
BOOT	res	res	res	res	res	HEATER	ONE_SHOT

BOOT – ustawienie tego bitu powoduje przepisanie ustawień kalibracyjnych z wewnętrznej pamięci FLASH do rejestrów kalibracyjnych.
HEATER – ustawienie tego bitu włącza wewnętrzne podgrzewanie w celu usunięcia wilgoci z kondensacji pary wodnej z czujnika wilgotności. Wyzerowanie bitu powoduje wyłączenie podgrzewania.
ONE_SHOT – ustawienie tego bitu inicjuje jednokrotny pomiar wilgotności i temperatury. Po wykonaniu pomiaru ONE_SHOT jest sprzętowo zerowany.

Listing 13. Inicjalizacja czujnika higrometru

```
void HTS221Init(void){
    HTS221WriteReg(0x20,0); //CTRL_REG1 Power Down
    HTS221Cal(); //odczytanie współczynników kalibracji
    HTS221WriteReg(0x10, 0x1b); //dokładność pomiarów
    //CTRL_REG1 Power On, One Shot, Update after read
    HTS221WriteReg(0x20, 0x84);
}
```

częstotliwości pomiarów jest przeznaczony rejestr CTRL_REG1 o adresie 0x20 opisany dokładnie w **tabeli 5**. Bit PD tego rejestru jest przeznaczony do wprowadzania układu w stan obniżonego poboru mocy, a bit BDU określa czy rejestr z danymi wyjściowymi wilgotności i temperatury ma być odświeżany automatycznie po wykonaniu pomiarów, czy też po odczycie przez mikrokontroler 8 starszych bitów wyniku poprzedniego pomiaru. Po włączeniu zasilania bit PD jest wyzerowany i żeby można było wykonywać pomiary należy do PD wpisać jedynkę.

Tabela 7. Zestawienie rejestrów kalibracyjnych układu HTS221

Rejestr	ADRES hex
H0_rH_x2	30
H1_rH_x2	31
T0_degC_x8	32
T1_degC_x8	33
T1/T0 msb	35
H0_T0_OUT	36
	37
H1_T1_OUT	3A
	3B
T0_OUT	3C
	3D
T1_OUT	3E
	3F

Listing 14. Odczytanie rejestrów kalibracyjnych

```
void HTS221Cal(void){
    int16_t T0,T1;
    struct {
        uint8_t H0_rH_x2; //addr 0x30
        uint8_t H1_rH_x2; //addr 0x31
        uint8_t T0_degC_x8; //addr 0x32
        uint8_t T1_degC_x8; //addr 0x33
        uint8_t T1_T0_msb; //addr 0x35
        int16_t H0_T0_OUT; //addr 36, 37
        int16_t H1_T0_OUT; //addr 3a, 3b
        int16_t T0_OUT; //addr 3c, 3d
        int16_t T1_OUT; //addr 3e, 3f
    }cal;

    //współczynniki kalibracji dla wilgotności
    cal.H0_rH_x2 = HTS221ReadReg(0x30);
    cal.H1_rH_x2 = HTS221ReadReg(0x31);
    H0_rh = (float)cal.H0_rH_x2/2;
    H1_rh = (float)cal.H1_rH_x2/2;

    cal.H1_T0_OUT = HTS221ReadReg(0x3b);
    cal.H1_T0_OUT <= 8;
    cal.H1_T0_OUT |= HTS221ReadReg(0x3a);

    cal.H0_T0_OUT = HTS221ReadReg(0x37);
    cal.H0_T0_OUT <= 8;
    cal.H0_T0_OUT |= HTS221ReadReg(0x36);

    H0_cal = cal.H0_T0_OUT;
    H1_cal = cal.H1_T0_OUT;

    //współczynniki kalibracji dla temperatury
    cal.T0_degC_x8 = HTS221ReadReg(0x32);
    cal.T1_T0_msb = HTS221ReadReg(0x35);
    T0 = cal.T1_T0_msb<3;
    T0 <= 8;
    T0 = T0|cal.T0_degC_x8;
    T0_degC = (float)T0/8;

    cal.T1_degC_x8 = HTS221ReadReg(0x33);
    cal.T1_T0_msb = HTS221ReadReg(0x35);
    T1 = (cal.T1_T0_msb & 0x0c);
    T1>= 2;
    T1 <= 8;
    T1 = T1 | cal.T1_degC_x8;
    T1_degC = (float)T1/8;

    cal.T0_OUT = HTS221ReadReg(0x3d);
    cal.T0_OUT <= 8;
    cal.T0_OUT |= HTS221ReadReg(0x3c);

    cal.T1_OUT = HTS221ReadReg(0x3f);
    cal.T1_OUT <= 8;
    cal.T1_OUT |= HTS221ReadReg(0x3e);

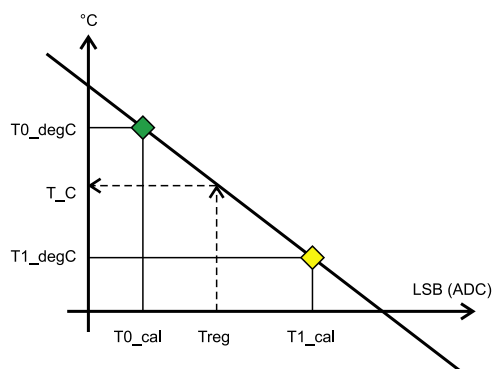
    T0_cal = cal.T0_OUT;
    T1_cal = cal.T1_OUT;
}
```

Wyzerowanie bitów ODR1 i ODR2 wprowadza układ w tryb pomiaru na żądanie. Żeby wyzwolić taki pomiar trzeba ustawić bit ONE_SHOT w rejestrze CTRL_REG2 (tabela 6). Po wyzwoleniu pomiaru, ale też w trybie pomiaru ciągłego trzeba testować, czy wynik pomiaru został zapisany do rejestrów wyjściowych. Żeby to zrobić należy odczytać zawartość rejestru statusowego STATUS_REG. Ustawienie bitu H_DA (bit b1 STATUS_REG) oznacza, że wynik pomiaru wilgotności został zapisany do rejestrów wyjściowych HUMIDITY_OUT_L i HUMIDITY_OUT_H, a ustawienie bitu T_DA (bit b0 STATUS_REG) oznacza, że wynik pomiaru temperatury został wpisany do rejestrów pomiaru temperatury. W czasie testów czujnika skonfigurowałem układ do pracy z wyzwaniem na żądanie i z ustawionym bitem BDU. Procedura inicjalizacyjna została pokazana na listingu 13.

W nieulotnej pamięci układu HTS221 umieszczono rejestry z zapisanymi w trakcie fabrycznej kalibracji wartościami kalibracji. W czasie sekwencji włączania zasilania układu dane kalibracyjne są przepisywane z pamięci Flash do rejestrów kalibracyjnych pokazanych w tabeli 7. W trakcie inicjalizacji układu mikrokontroler musi odczytać dane kalibracyjne po to, aby je potem zastosować przy każdorazowym odczycie wilgotności i temperatury. Procedura odczytu rejestrów kalibracyjnych została pokazana na listingu 14. W wyniku działania tej funkcji zapisywane są zmienne globalne pokazane na listingu 15. Te zmienne w połączeniu z 16-bitowymi danymi wyjściowymi będą służyły do wyliczenia mierzonej wartości. W przypadku temperatury wartości T0_cal, T1_cal, T0_degC i T1_degC są współzrędnymi wyznaczającymi prostą kalibracji – zostało to pokazane na rysunku 17. Do zmiennej Treg jest wpisana 16 bitowa wartość odczytana z rejestrów TEMP_OUT_L i TEMP_OUT_H. Mając wartości odczytane

Listing 15. Zmienne kalibracji

```
// Temperatura w stopniach C dla kalibracji
float T0_degC, T1_degC;
// Wyjściowa wartość temperatury dla kalibracji
int16_t T0_cal, T1_cal;
// Wilgotność dla kalibracji
float H0_rh, H1_rh;
// Wyjściowa wartość wilgotności dla kalibracji
int16_t H0_cal, H1_cal;
```



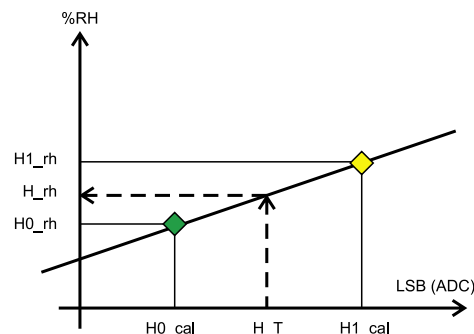
Rysunek 17. Wykres prostej kalibracji i wyliczanie temperatury w stopniach Celsjusza

Listing 16. Odczytanie, obliczenie i wyświetlenie temperatury

```
// odczytanie i wyświetlenie temperatury
void HTS221ReadTemp(char *str){
    int16_t Treg, tempt;
    float T_C, Temperature;

    // start pomiaru
    HTS221WriteReg(0x21, 1); // start One shot
    // czekaj na dane pomiaru
    while((HTS221ReadReg(0x27) & 1) == 0);

    Treg = HTS221TempRead(); // odczytaj rejestr temperatury
    // wylicz temperaturę na podstawie danych kalibracji
    T_C = ((float)(Treg - T0_cal)) / (T1_cal - T0_cal)
        * (T1_degC - T0_degC) + T0_degC;
    tempt = (int16_t)(T_C * 100);
    Temperature = ((float)tempt) / 100;
    // wyświetl temperaturę
    sprintf(str, "%4.1f °C", Temperature);
}
```



Rysunek 18. Wykres prostej kalibracji i wyliczanie wilgotności w %

z rejestrów kalibracji i wartość Treg wyliczamy wartość temperatury w Stopniach Celsjusza według zależności:

$$T_C = (Treg - T0_cal) / (T1_cal - T0_cal) * (T1_degC - T0_degC) + T0_degC$$

Na listingu 16 jest pokazana kompletna procedura wyzwolenia pojedynczego pomiaru, wyliczenia i wyświetlenia zmierzonej temperatury. Wartości kalibracyjne zostały jednokrotnie pobrane z rejestrów i wpisane do zmiennych w momencie inicjalizacji układu i nie ma potrzeby ich pobierać przy każdym pomiarze.

Podobnie wygląda sposób postępowania w przypadku odczytywania, przeliczania i wyświetlania wilgotności. Wilgotność jest wyliczana z zależności, której elementy obrazuje rysunek 18:

$$H_rh = (((H_T - H0_cal)) / (H1_cal - H0_cal)) * (H1_rh - H0_rh) + H0_rh$$

Procedura odczytywania przeliczania i wyświetlania wilgotności została pokazana na listingu 17. Jak widać do obliczania wyników niezbędne są wyliczenia z wykorzystaniem biblioteki operacji zmiennoprzecinkowej. Jeżeli dodamy do tego konwersję liczb zmiennoprzecinkowych na łańcuch ascii za pomocą funkcji *sprintf* i brak optymalizacji w bezpłatnej wersji kompilatora MPLAB XC8 to wiadomo, dlaczego kod wynikowy zajął ponad 50% całej pamięci flash o rozmiarze 16 k słów.

Czujnik temperatury DS18B20

Jak wspomniałem na wstępie, do pomiaru temperatury zastosowałem czujnik temperatury DS18B20 z magistralą 1-Wire. Jak wiemy sensory ciśnienia i wilgotności mają swoje własne sensory temperatury mierzące je z wystarczającą dokładnością. Jednak, żeby pomiar nie był zafałszowany czujnik temperatury musi być oddalony od wszelkich lokalnych źródeł ciepła. W praktyce sprowadza się to do wystawienia go do pomiaru na zewnątrz w pewnej odległości od obudowy urządzenia. Umieszczenie czujnika w obudowie urządzenia z pracującymi układami elektronicznymi zawyża pomiar przynajmniej o kilka stopni. Zależy to od wydajności źródeł ciepła w układzie na przykład zasilacza, czy pracującego mikrokontrolera, ale też od rodzaju obudowy. DS18B20 nadaje się idealnie do takiego zastosowania. Magistrala 1-Wire pracuje bez problemu z przewodami o długości 1...2 m, a sam czujnik jest hermetyczny i odporny na warunki atmosferyczne. Używane przeze mnie DS18B20 pracują bezawaryjnie wiele

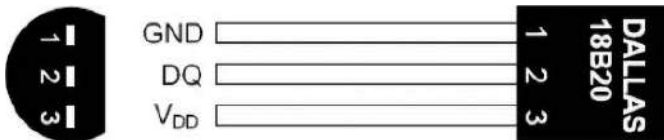
Listing 17. Odczytanie, obliczenie i wyświetlenie wilgotności

```
// odczytanie rejestrów i wyliczenie kompensacji
void HTS221ReadHum(char *str){
    int16_t H_T, humt;
    float H_rh, Hum;

    // start pomiaru
    HTS221WriteReg(0x21, 1); // start One shot
    while((HTS221ReadReg(0x27) & 2) == 0);

    H_T = HTS221HumRead(); // odczytaj rejestr wilgotności
    H_rh = ((float)(H_T - H0_cal)) / (H1_cal - H0_cal)
        * (H1_rh - H0_rh) + H0_rh;
    humt = (int16_t)(H_rh * 100);
    Hum = ((float)humt) / 100;

    sprintf(str + 8, "%2.1f", Hum);
    *(str+13) = '\0';
}
```

Rysunek 19. Obudowa czujnika DS18B20

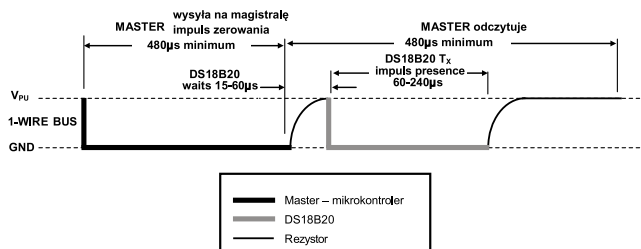
lat w normalnym środowisku wystawione na zewnątrz. Czujnik jest umieszczony w obudowie TO-92 został pokazany na **rysunku 19**.

Do komunikacji z mikrokontrolerem DS18B20 ma wbudowaną 1-przewodową magistralę 1-Wire składającą się z jednej dwukierunkowej linii DQ. Dodatkową zaletą 1-Wire jest programowy mechanizm identyfikacji układu dołączonego do magistrali na podstawie unikalnego kodu zapisanego w pamięci ROM. Ten mechanizm pozwala na dołączenie wielu układów do jednej magistrali. Z założenia wykorzystujemy jeden DS18B20 i oprogramowanie nie obsługuje odczytywania kodów identyfikacyjnych i identyfikacji na magistrali więcej niż jednego czujnika. Każde z urządzeń połączonych magistralą musi mieć wyjście typu otwarty dren, a sama linia sygnałowa DQ jest połączona z plusem zasilania przez rezystor 2...5 kΩ. Dwukierunkowa linia danych DQ może też spełniać rolę linii zasilającej układ. Wtedy wyprowadzenia GND i Vdd są zwarte i połączone z masą sterownika, a układ jest zasilany przez rezystor.

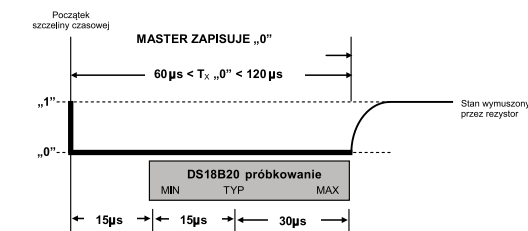
Magistrala 1-Wire jest zorganizowana według zasady master-slave. Układem master jest mikrokontroler, a układami slave dołączane do niego czujniki – w tym przypadku termometr DS18B20. Wymiana informacji poprzez magistralę 1-Wire składa się z czterech podstawowych sekwencji:

- sekwencji inicjalizacji (*reset and presence pulses*),
- zapisu zera,
- zapisu jedynki,
- odczytu bitu.

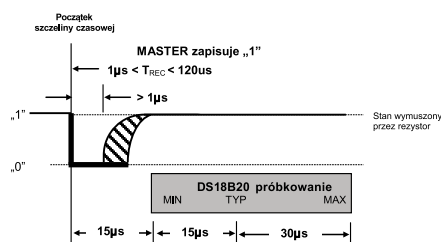
W czasie, kiedy nie jest wykonywana żadna z sekwencji linia DQ musi przejść w stan wysoki (stan *idle state*). Każda wymiana informacji z termometrem (zapisywanie i odczytywanie) DS18B20 zaczyna się sekwencją inicjalizacji wymuszeniem przez mikrokontroler na linii



Rysunek 20. Sekwencja inicjowania magistrali 1-Wire



Rysunek 21. Zapisywanie zera logicznego



Rysunek 22. Zapisywanie jedynki logicznej

Listing 18. Sekwencja inicjowania magistrali 1-wire

```
//linia DQ port RC4
#define DQ_IO PORTCbits.RC4
#define DQ_TRIS TRISCbits.TRISC4
uint8_t DS_Reset (void){
    uint8_t result;

    // ustaw na DQ stan niski
    //(stan linii 0 i linia wyjściowa)
    DQ_TRIS = 0;
    DQ_IO = 0;
    //odczekaj 480usek
    __delay_us(480);
    //linia DQ wejściowa
    //(stan wysoki ustala rezystor)
    DQ_TRIS = 1;
    __delay_us(70);

    if (DQ_IO)
        //stan niski na DQ - błąd
        result = 0;
    else
        //stan wysoki wykryto Presence pulse
        result = 1;
    //dokończenie sekwencji
    __delay_us(480);

    return(result);
}
```

DQ stanu niskiego (impulsu zerowania) . Czas trwania tego impulsu to minimum 480 µs (**rysunek 20**).

W naszym układzie linia DQ jest podłączona do linii portu RC4. Czujnik jest zasilany napięciem +3,3 V

Na **listingu 18** pokazano sekwencję *presence*. Procedura zwraca stan wysoki jeżeli czujnik został wykryty. Mikrokontroler wymusza stan niski na magistrali kiedy linia RC4 jest skonfigurowana jako wyjście (do rejestru portu PORTC4 jest programowo wpisane zero). Ustawienie RC4 jako wejście powoduje, że rezystor podłączony pomiędzy RC4 i plus zasilania wymusza na DQ stan wysoki.

Sekwencja zapisywania zera logicznego przez mikrokontroler do termometru została pokazana na **rysunku 21**. Mikrokontroler wymusza stan niski na magistrali przez co najmniej 60 µs. Jednak stan niski nie może trwać dłużej niż 120 µs, bo układ może to zinterpretować jako sekwencję zerowania.

Sekwencja zapisywania jedynki logicznej rozpoczyna się od wymuszenia stanu niskiego na magistrali. Stan ten nie może trwać dłużej

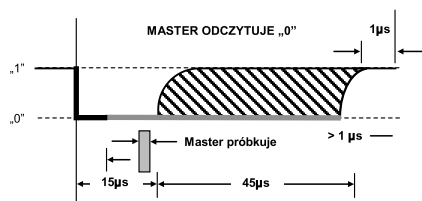
Listing 19. Zapis bitu na magistrali 1-wire

```
void DS_WriteBit(uint8_t val){
    // na magistrali zero logiczne
    //DQ linia wyjściowa
    DQ_TRIS = 0;
    DQ_IO = 0;

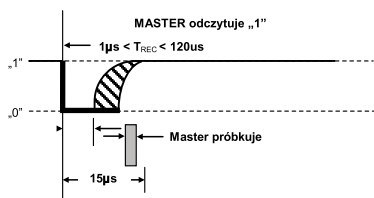
    //val = 1 zapisujemy jedynkę
    if (val){
        // czekaj kilka us
        asm ("nop");
        asm ("nop");
        asm ("nop");
        asm ("nop");
        //linia wejściowa
        //na magistrali jedynka logiczna
        DQ_TRIS = 1;
        //zakończ slot czasowy
        __delay_us(60);
    }else{
        //przez 60us zero logiczne
        __delay_us(60);
        // Ustaw jedynkę logiczną
        DQ_TRIS = 1;
        //zakończ slot czasowy
        __delay_us(10);
    }
}
```

Listing 20. Wysłanie danej 8-bitowej przez 1-wire

```
void DS_WriteByte(uint8_t val){
    uint8_t i;
    // pętla przesłania 8bitów LS-bit first
    for (i = 0; i < 8; i++){
        //wyślij bit na magistralę
        DS_WriteBit(val & 0x01);
        //przesuń o jedną pozycję
        //dla następnego notu
        val >>= 1;
    }
}
```



Rysunek 23. Odczytanie zera logicznego



Rysunek 24. Odczytanie jedynki logicznej

niż 15 µs. Po tym czasie na magistrali musi się pojawić stan wysoki przez czas co najmniej 45 µs (rysunek 22). Uniwersalna Procedura *DS_WriteBit* zapisuje zero logiczne lub jedynkę zależnie od wartości argumentu *val* (listing 19). Zapisanie dowolnej danej 8-bitowej jest już proste (listing 20). Odczytywanie danych rozpoczyna się od wymuszenia przez master na magistrali stanu niskiego przez minimum 1 µs. Dane wystawiane przez DS18B20 są prawidłowe przez

Listing 21. Odczytanie bitu za magistrali 1-wire

```
// odczytanie jednego bitu
uint8_t DS_ReadBit(void){
    uint8_t result;
    // na magistrali stan logiczny 0
    DQ_TRIS = 0;
    DQ_IO = 0;
    //czekaj kilka us
    asm ("nop");
    asm ("nop");
    asm ("nop");
    asm ("nop");
    // na magistrali stan logiczny 1
    DQ_TRIS = 1;
    asm ("nop");
    asm ("nop");
    asm ("nop");
    asm ("nop");
    //odczytaj stan linii DQ
    if (DQ_IO)
        result = 1;
    else
        result = 0;
    //zakończ szczelinę czasową
    _delay_us(60);
    return (result);
}
```

Listing 22. Zapisanie bajtu

```
//zapisanie bajtu z argumentu data
void DS_WriteByte(uint8_t data){
    uint8_t i;
    //pętla zapisania 8 bitów
    //pierwszy LSB (najmłodszy)
    for (i = 0; i < 8; i++){
        DS_WriteBit(data & 0x01);
        //przesunięcie danej
        //dla następnego bitu
        data >>= 1;
    }
}
```

Listing 23. Odczytanie bajtu

```
//odczytanie bajtu z czujnika przez magistrale 1-wire
uint8_t DS_ReadByte ( void ){
    uint8_t i, data;
    data = 0;
    //pętla odczytania 8 bitów
    for (i = 0; i < 8; i++){
        data >>= 1;
        // jeżeli odczytana jedynka
        if (DS_ReadBit() == 1)
            data |= 0x80;
    }
    return data;
}
```

maksymalnie 15 µs i przed upływem tego czasu mikrokontroler musi je przeczytać. Jeżeli na magistrali jest stan niski, to odczytane zostaje zero logiczne (rysunek 23), a jeżeli stan wysoki, to odczytana zostaje jedynka logiczna (rysunek 24).

Procedury komunikacji z termometrem używają programowych opóźnień wykorzystujących wykonywanie pustego rozkazu *NOP* rdzenia mikrokontrolera PIC16F. Ilość tych rozkazów może być inna dla mikrokontrolera taktowanego inną częstotliwością niż przyjęta w projekcie (16 MHz). Wyposażeni w procedury zapisania i odczytania bitu przez magistralę 1-Wire możemy zapisywać i odczytywać bajty (listing 22 i listing 23). Komunikacja z termometrem będzie polegała na wysyłaniu poleceń (komend) do układu i zapisywaniu lub odczytywaniu z układu zawartości rejestrów temperatury. Zastosujemy tu uproszczone sterowanie odczytywaniem temperatury. Jest to możliwe, bo do magistrali z założenia jest dołączony jeden czujnik i nie trzeba wykonywać sekwencji identyfikacji. Pełny opis komend i sposobu sterowania czujnikiem jest opisany w dokumentacji. My tutaj wykorzystamy tylko trzy komendy: *SKIP ROM DS*, *READ RAM* i *ConvertT*.

Procedura inicjowania pomiaru i odczytania wyniku została pokazana na listingu 24. Rozpoczynamy od wykonania sekwencji resetu *DS_Reset()*. Jeżeli procedura zwróci wartość zerową, to oznacza, że do magistrali nie jest podłączony sprawny czujnik (lub w ogóle go nie ma). Blokowane są wszystkie przerwania, żeby nie przeszkadzały

Tabela 8. Rejestr temperatury i tablica konwersji dla układu DS18B20

Temperatura (°C)	Wartość binarna (BIN)	Wartość wyjściowa (HEX)
+125	0000 0111 1101 0000	07D0
+85	0000 0101 0101 0000	0550
+25,0625	0000 0001 1001 0001	0191
+10,125	0000 0000 1010 0010	00A2
+0,5	0000 0000 0000 1000	0008
0	0000 0000 0000 0000	0000
-0,5	1111 1111 1111 1000	FFF8
-10,125	1111 1111 0101 1110	FF5E
-25,0625	1111 1110 0110 1111	FE6F
-55	1111 1100 1001 0000	FC90

Listing 24. Odczytanie rejestrów temperatury

```
uint8_t DS_GetTemperature(void){
    // zablokowanie wszystkich przerwań
    GIE = 0;

    if (DS_Reset() == 0){
        GIE = 1; // odblokowanie przerwań
        return 0; // nie wykryto sensora
    }

    //komenda SKIP ROM
    DS_WriteByte(DS_SKIP_ROM);
    // Start konwersji temperatury
    DS_WriteByte(DS_CONVERT_T);

    //Czekaj na zakończenie konwersji
    while (DS_ReadBit() == 0);
    // _delay_ms(1000);
    // czas konwersji jest zależny
    // od ustawionej rozdzielczości pomiaru
    if (DS_Reset() == 0){
        GIE = 1; // odblokowanie przerwań
        // nie wykryto sensora
        //po wykonaniu konwersji
        return 0;
    }

    //komenda SKIP ROM
    DS_WriteByte(DS_SKIP_ROM);
    // komenda zapisania wewnętrznych rejestrów wyniku
    DS_WriteByte(DS_READ_RAM);
    // odczytaj rejestry wyniku pomiaru
    TemperatureLow = DS_ReadByte();
    TemperatureHigh = DS_ReadByte();
    //koniec odczytywania rejestrów wyniku
    DS_Reset();
    // odblokuj przerwania
    GIE = 1;
    // dane odczytano poprawnie
    return 1;
}
```


Listing 25. Konwersja zawartości rejestrów temperatury na wartość dziesiętną ze znakiem

```
uint8_t ConvTemperature(char *buffer){
    uint8_t status, tenths, i;
    uint16_t temperature = 0;
    double temp;
    uint16_t tmp = 0;

    // status = 1 - wykryty sensor dane
    // w rejestrze temperatury są ważne

    // status = 0 - sensor nie został wykryty,
    // dane w rejestrze temperatury nie są ważne

    status = DS_GetTemperature();
    if (status == 0)
        return status;

    temperature = TemperatureHigh;
    temperature = (temperature << 8) | TemperatureLow;
    temp = (double)temperature / 16.0;
    sprintf (buffer+14, "%2.1f°C", temp);

    return(status);
}
```

w programowym odliczaniu dość krótkich opóźnień rzędu pojedynczych mikrosekund.

Pomiar temperatury jest wyzwalany komendą DS_CONVERT_T poprzedzona komendą SKIP_ROM. Czas trwania konwersji jest

zależny od rozdzielczości pomiaru: dla 9-bitowej to ok 94 ms, a dla 12-bitowej ok 750 ms. Domyślnie jest ustawiona konwersja 12-bitowa. W czasie pomiaru na linii DQ DS18B20 jest wymuszany stan niski. Testowanie linii DQ przez program pozwala na wykrycie zakończenia pomiaru – wtedy linia DQ przechodzi w stan wysoki. Wysłanie sekwencji komend DS_SKIP_ROM i DS_READ_RAM pozwala na odczytanie 2 bajtów wyniku i zapisanie ich do globalnych zmiennych `TemperatureLow` i `TemperatureHigh`.

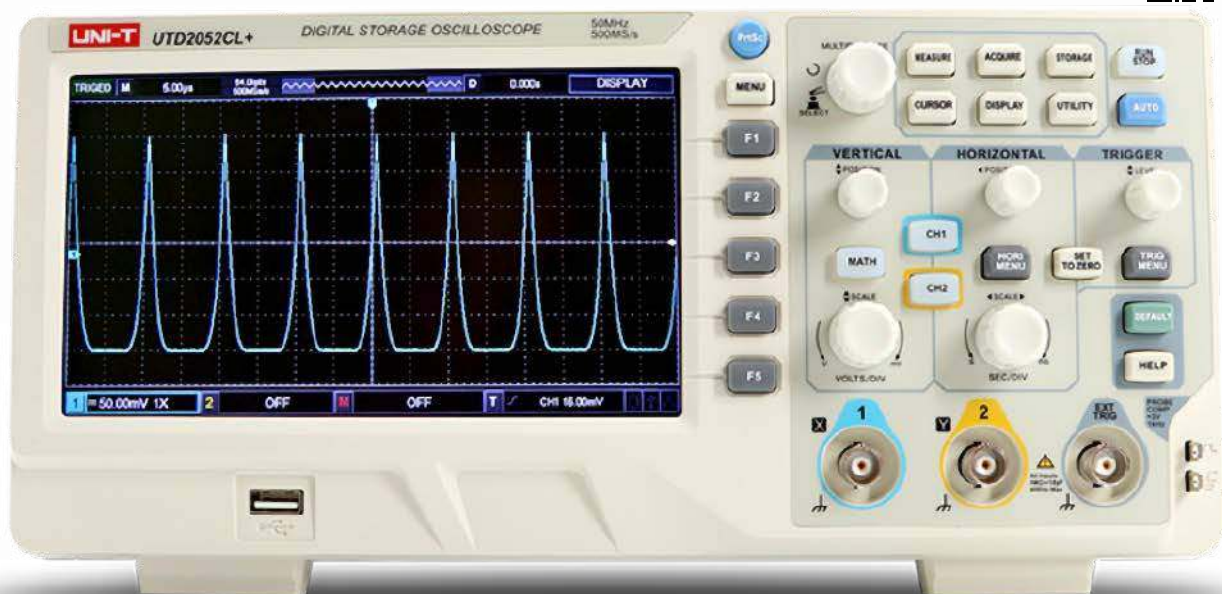
Kolejna procedura wykonuje konwersję odczytanych i zapisanych w kodzie U2 zawartości rejestrów temperatury na wartość dziesiętną ze znakiem (tabela 8). Dla rozdzielczości 12-bitowej, żeby wyliczyć wartość zmierzonej temperatury, trzeba zawartość rejestru podzielić przez 16. Tak uzyskaną wartość w formacie zmiennoprzecinkowym poddajemy konwersji na łańcuch znaków ascii przez funkcję `sprintf` (listing 25).

Ostatnim elementem naszej stacji meteo jest kanał transmisji radiowej pozwalający na przesyłanie danych pomiarowych. Jego konstrukcja i działanie zostaną opisane w ostatniej części artykułu, który ukaże się w kolejnym wydaniu EP.

Tomasz Jabłoński, EP

REKLAMA

Oscyloskop UNI-T UTD2052CL+ 1 169,00 zł



Szerokość pasma analogowego: 50 MHz
 Czas narastania: <7 ns
 Liczba kanałów: 2
 Podstawa czasu: 2 ns/div ~ 50 s/div
 A/D: 8 bit
 Czutość odchylenia pionowego: 1 mV/div ~ 20 V/div
 Sposób zapisu danych pomiarowych: setup, wave, bitmap
 Sposoby wyzwalania: zboczem, szerokością impulsu, sygnałem wideo, naprzemienne

Operacje matematyczne:

Interfejsy:

normy:
 zasilanie:
 wymiary:
 masa:

dodawanie, odejmowanie,
 mnożenie, dzielenie
 USB Host, USB Device, Pass/Fail
 CE, EN:61010-1
 230 V AC
 306×138×124 mm
 2,73 kg

W zestawie: oscyloskop, 2 sondy pomiarowe z dzielnikiem 1× i 10×, kabel zasilający, kabel USB, instrukcja obsługi



Przedstawiona oferta cenowa ma charakter informacyjny i nie stanowi oferty handlowej w rozumieniu Art.66 par.1 Kodeksu Cywilnego



AVT SPV Sp. z o.o.
 03-197 Warszawa, ul. Leszczynowa 11
 tel. +48 22 257 84 49, handlowy@avt.pl
sklep.avt.pl



Okres krótko po nowym roku to, dla wielu osób, okres wzmożonej aktywności fizycznej. Postanowienia noworoczne bardzo często obejmują chęć zrzućcenia kilku kilogramów lub poprawy naszej kondycji, a nawet, jeśli nie one, to ruch przyda się nam, aby spalić nadmiarowe, poświęteczne kalorie. Polecaną niemalże wszystkim aktywnością są spacer. Zrobienie 10 tysięcy kroków dziennie jest rekomendowane przez wielu specjalistów, jednak skąd wiedzieć, kiedy już pokonaliśmy tę barierę? Dobrze jest zaopatrzyć się w urządzenie, które pozwoli nam monitorować w jakiś sposób naszą aktywność – w prezentowanym rozwiązaniu będzie zliczać nasze kroki.

Pedometry

– kilka projektów do monitorowania naszej aktywności ruchowej

Spacer to dobry sposób na spalenie zbędnych kalorii ale mają też wiele innych zalet. Specjaliści wskazują, że regularne spacerowanie wpływa dobrze na nasze serce, układ mięśniowy i szkieletowy, a także nasz nastrój i samopoczucie. Zalety chodzenia można wymieniać w nieskończoność. Specjaliści klasyfikują chodzenie, jako umiarkowaną aktywność fizyczną, co w dzisiejszych czasach, gdy wiele osób pracuje siedząc za biurkiem, ma duże znaczenie. W zależności od naszego wieku powinniśmy robić dziennie 6...8 tysięcy kroków (osoby powyżej 60 roku życia) lub 8...10 tysięcy kroków (osoby młodsze), jeśli nasza sprawność fizyczna i inne uwarunkowania zdrowotne na to pozwalają.

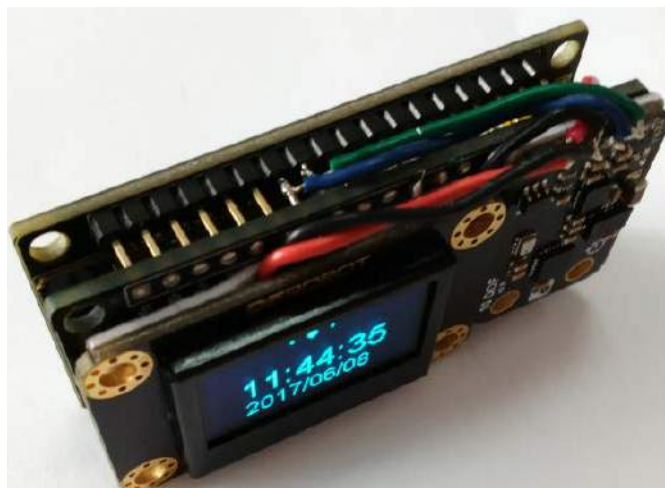
Zliczanie kroków jest podstawową metodą monitorowania naszego wysiłku podczas chodzenia. Pedometr to urządzenie, którego funkcją jest właśnie zliczanie kroków. Istnieje wiele metod takiego zliczania – część z nich przedstawiona jest w poniższym artykule, gdzie pokazujemy nie jedną, a trzy tego rodzaju konstrukcje. Opis każdej konstrukcji jest ograniczony, ale szczególny nacisk położono na sposób pomiaru/wykrywania kroków.

Akcelerometr

Pierwszy moduł bazuje na chyba najbardziej klasycznym rozwiązaniu – module inercyjnym, który monitoruje ruch noszącej urządzenie osoby i zlicza ruchy, które uznawane są za kroki. Takie podejście

jest najczęstsze, z takiego algorytmu korzystają również komercyjne pedometry, takie, jak spotykane w opaskach fitness czy telefonach komórkowych.

Zaprezentowana konstrukcja bazuje na module FireBeetle Board z mikrokontrolerem ESP32. Oprócz wspomnianej płytki, układ ma



Fotografia 1. Licznik kroków z akcelerometrem



Fotografia 2. Oporniki podciągające na płycie (uwaga – oporniki powinny podciągać do pinu 3V3, a nie AREF, jak na zdjęciu)

również ekran OLED 128×64. Do pomiaru kroków zastosowano 3-osiowy akcelerometr MEMS ADXL345. Gotowy moduł pedometru pokazano na **fotografii 1**. Wyświetlacz i akcelerometr podłączone są do mikrokontrolera przez interfejs I²C, dzięki czemu wystarczy połączyć dwie linie sygnałowe – danych i zegara, aby nawiązać komunikację. Obie linie muszą być podciągnięte do zasilania, jak pokazano na **fotografii 2** (**Uwaga! Na zdjęciu oporniki podciągnięte są do AREF, a nie znajdującego się obok pinu 3V3 na skutek pomyłki autora**). W projekcie podciągnięto je do zasilania opornikami 51 kΩ, jednak każda wartość pomiędzy 1 kΩ a ok. 50 kΩ powinna pracować w takim, prostym układzie o krótkich połączeniach poprawnie. Pozostało jeszcze podłączenie przycisków. Te podłączono pomiędzy linie D2, D3 i D4 i zasilanie (3V3).

Teraz pozostaje tylko załadowanie do układu firmware. Kod programu dla ESP32 jest dosyć złożony, gdyż sam układ oprócz zliczania ilości kroków liczy także czas, prezentuje go w postaci zegarka analogowego, zlicza spalone podczas chodzenia kalorie itd. Kluczowe elementy szkicu Arduino wykrywającego i zliczającego kroki pokazano na **listingu 1**. Funkcja `updateAdxl345` aktualizuje (inkrementuje) liczbę kroków, zapisaną w zmiennej globalnej `stepSum`. W pierwszej kolejności program pobiera dane z akcelerometru ADXL345 i przetwarza je do danych dla trzech osi, zapisanych, odpowiednio, w zmiennych `xx`, `yy` oraz `zz`. Sam algorytm używa tylko składowej `X` pomiaru przyspieszenia. Jeśli jest ono poniżej 100 (jednostek umownych, LSB) to nic się nie dzieje – funkcja drukuje przez UART aktualną liczbę kroków i wychodzi. Jeśli wartość bezwzględna różnicy przyspieszenia w tej iteracji pętli i poprzedniej iteracji (zmienna globalna `currentValue`) jest większa od 80 i obecna wartość jest mniejsza od poprzedniej, to licznik kroków jest inkrementowany. Następnie aktualizowana jest wartość przyspieszenia w zmiennej `currentValue`. Jeśli `xx` nie jest mniejsze od `currentValue`, to następuje tylko aktualizacja zapisanej wartości przyspieszenia. Finalnie, drukowana jest aktualna wartość kroków i program wychodzi z funkcji `updateAdxl345`.

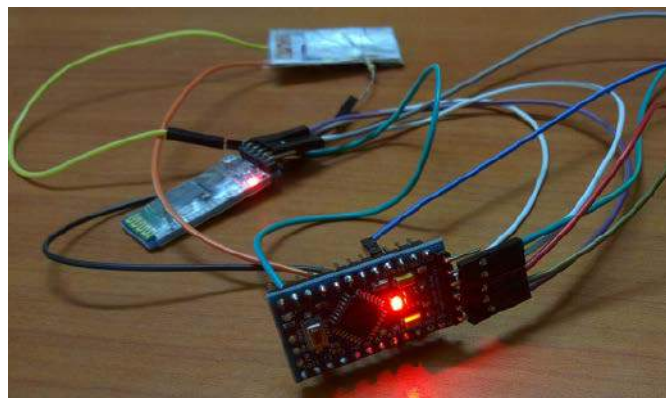
Funkcja `updateAdxl345` musi być okresowo wywoływana przez program (najlepiej z częstotliwością przygotowywania próbek danych

Listing 1. Listing programu dla ESP32 z akcelerometrem do zliczania kroków (fragment)

```
void updateAdxl345(void){
  // Odczyt danych z ADXL345
  readFrom(DEVICE, regAddress, TO_READ, buff);
  xx = (((int)buff[1]) << 8) | buff[0];
  yy = (((int)buff[3]) << 8) | buff[2];
  zz = (((int)buff[5]) << 8) | buff[4];

  if(xx < 100){
    sprintf(str, "%d", stepsSum);
    return;
  }

  if(fabs(xx - currentValue) > 80){
    if(xx < currentValue){
      stepsSum++;
    }
    currentValue = xx;
  }
  sprintf(str, "%d", stepsSum);
}
```



Fotografia 3. Pedometr z sensorem nacisku DIY

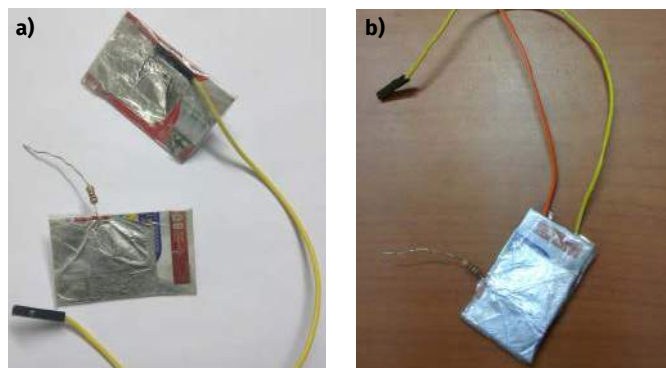
przez akcelerometr). Aktualizuje ona wartość zmiennej globalnej `stepSum`, z której korzystać mogą w dowolny sposób inne komponenty programowe.

Sensor naciskowy

Zamiast akcelerometru zastosować można sensor innego rodzaju. Autor drugiego zaprezentowanego tutaj projektu zaproponował prosty sensor naciskowy, który umieszcza się w butcie, pod stopą, co pozwala mu w łatwy sposób monitorować nacisk stopy na ziemię. Zmienia się on w rytm kroków, co pozwala na zliczanie ich przez moduł Arduino. Podobnie, jak w poprzednim przypadku, mikrokontroler zlicza nie tylko liczbę kroków, ale także oblicza prędkość chodu, ilość spalonych kalorii itd. Na **fotografii 3** pokazano gotowe urządzenie.

Podstawą projektu jest wykonanie sensora wykrywającego kroki. W tym celu należy wziąć dwa kawałki kartonu o wymiarach 40×20 mm. Następnie wycinamy dwa kawałki folii aluminiowej każdy o wymiarach 60×20 mm. Jedna strona każdego kartonu powinna być w całości przykryta folią aluminiową a druga strona mniej więcej do połowy długości. Następnie należy każdą z elektrod okleić taśmą klejącą tak, aby niewielka część folii aluminiowej po jednej stronie kartonu pozostała odsłonięta. W ten sposób wykonujemy dwie elektrody. Należy upewnić się, że pomiędzy folią a kartonem nie ma żadnej szczeliny. Jedyne miejsce, w którym powinna znajdować się szczelina, znajduje się pomiędzy odkrytymi fragmentami folii aluminiowej na powierzchni dwóch kartoników, gdy są one ze sobą złożone „w kanapkę”.

Do jednego kawałka folii dołączamy opornik 10 kΩ tak, aby stykał się z folią znajdującą się po tej stronie kartonu, która nie jest w pełni pokryta folią aluminiową. Należy przykleić go tak, aby opornik się nie ruszał. Do drugiego kawałka folii dołączyć należy dwa przewody (po tej samej stronie kartonu). Je również należy zakleić taśmą tak, aby przewody się nie poruszały. Istotne jest, aby przykleić oba przewody naraz, gdyż ciężko jest potem dokleić kolejny. Finalnie należy skleić obie elektrody ze sobą, tak, aby odizolowane powierzchnie folii aluminiowej skierowane były do siebie. Całość należy zabezpieczyć



Fotografia 4. Sensor nacisku; a) elektrody przed sklejeniem; b) gotowy sensor

Listing 2. Oprogramowanie dla Arduino dla pedometru z sensorem nacisku (fragment)

```

void loop(){
  // Poprzedni stan
  lastState = state;
  // Odczyt aktualnego stanu - sensor na pinie 12
  state = digitalRead(12);
  if (state == HIGH && lastState == LOW){
    // Inkrementacja licznika kroków
    count++;
  }
  // Pozostała logika programu
}

```

i zamocować do siebie folią aluminiową, pamiętając, aby nie ścisnąć elektrod ze sobą za mocno. Nie powinno być żadnego kontaktu między foliami, gdy układ jest rozprasowany.

Na koniec należy sprawdzić sensor za pomocą multimetru. Po naciśnięciu rezystancja sensora powinna się zmniejszyć. Gotowy sensor pokazany jest na **fotografii 4**. Sensor podłączony jest pomiędzy masę i zasilanie. Masa podłączana jest do jednego z przewodów, doklejonych do elektrody, a zasilanie do opornika 10 kΩ przy drugiej elektrodzie. Trzecie wyjście sensora – drugi z przewodów – podłączony jest do wejścia cyfrowego modułu Arduino Pro Mini. Działanie jest bardzo proste, układ sprawdza stan logiczny na wejściu podłączonym do sensora i monitoruje zbocza narastające w głównej pętli programu – *loop()*, jak pokazano na **listingu 2**. Jeśli stan pinu, zapisany w zmiennej *state* jest wysoki, a *lastState* – stan pinu w poprzedniej iteracji pętli programu – jest niski, oznacza to, że sensor naciśnięto, więc licznik kroków *count* jest inkrementowany. Dalsza logika programu wykorzystuje wartość licznika, do prezentacji ilości zrobionych kroków, spalonych kalorii itd.

Bez Arduino, bez programowania

Niewątpliwie Arduino zrewolucjonizowało elektronikę, w szczególności systemy amatorskie, jednak przed jego pojawieniem się na rynku, hobbyści również konstruowali różne interesujące układy. Przykładem projektu „z tamtych czasów” może być pedometr pokazany na **fotografii 5**. Zaprezentowany układ zawiera tzw. wyłącznik przechyłny (*tilt switch*). Jest to prosty element mechaniczny, który łączy obwód, w momencie, gdy zostanie pochylony. Elementy te często stosuje się, jako detektory wibracji itp. W tym układzie przełączenie tego sensora powoduje inkrementację licznika binarnego CD4024, który odpowiedzialny jest za zliczanie kroków. Przycisk S4 pozwala na zresetowanie tego licznika.

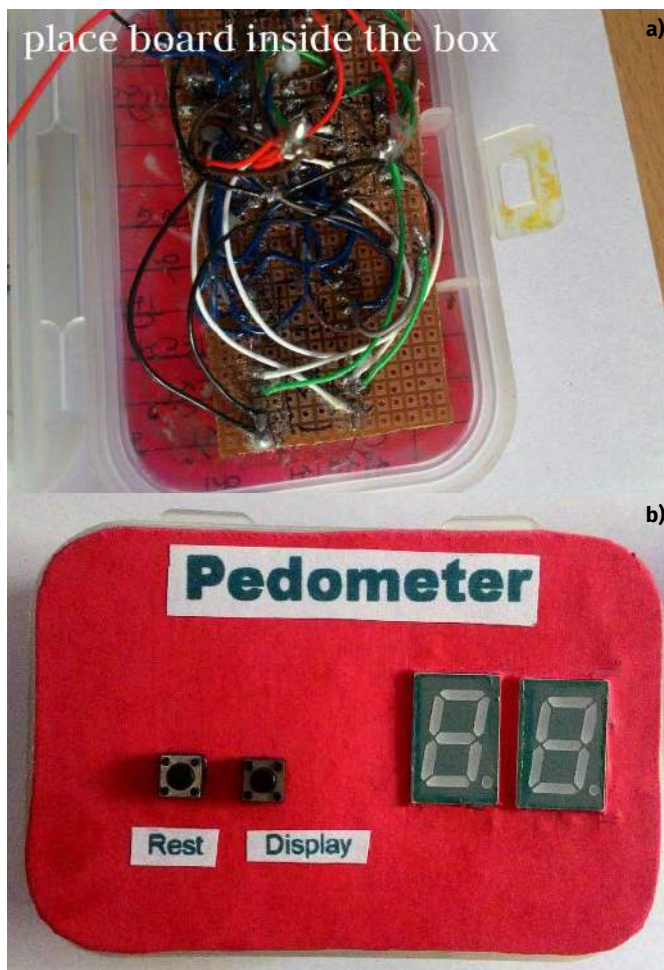
Do wyświetlania ilości kroków zastosowano dwa wyświetlacze siedmiosegmentowe. Za ich sterowanie odpowiedzialne są dwa liczniki dekadowe CD4026, które mają wyjście kompatybilne z tego rodzaju wyświetlaczem, dzięki czemu mogą one bezpośrednioysterować poszczególne diody LED w wyświetlaczu.

Całość uzupełnia buzzer, który piszczy w momencie, gdy zlicza krok oraz prosty układ zasilania z wyłącznikiem. Całe urządzenie może być zmontowane na kawałku płytki uniwersalnej i umieszczone w kompaktowej obudowie.

Alternatywne pomysły

Oprócz omówionych realizacji, zbierając materiały do tego artykułu, w sieci znaleziono szereg innych konstrukcji, na które warto zwrócić uwagę, ze względu na ciekawe lub bardzo uproszczone podejście. W zakresie tego ostatniego, na pewno na uwagę wskazuje możliwość zmodyfikowania komercyjnego pedometru z pomocą modułu Arduino. Tego rodzaju proste urządzenia, które zliczają kroki i wyświetlają je na ekranie mogą kosztować zaledwie kilka złotych. Ich funkcjonalności są jednak, lekko mówiąc, ubogie. Podłączenie Arduino do takiego układu pozwala na istotne rozbudowanie jego możliwości, szacowanie prędkości chodu, obliczanie spalonych kalorii itp.

Ciekawym miejscem, w jakim można umieścić moduł Arduino i licznik kroków są... buty (jak pokazano na fotografii tytułowej).



Fotografia 5. Pedometr na licznikach CD4026; a) uniwersalna płytka drukowana z elementami; b) układ w obudowie

Podświetlenie butów, które może być sterowane za pomocą Arduino, może indykować ilość kroków, tempo itd. Nie jest to może najbardziej ergonomiczna lokalizacja dla informacji o liczbie kroków, ale z pewnością rzuca się w oczy.

Finalnie, nie sposób wspomnieć o użyciu algorytmów uczenia maszynowego do analizy zbieranych danych. Nawet niewielkie mikrokontrolery obecnie dostarczają na tyle dużą moc obliczeniową, iż mogą być zaprzęgnięte do zadań związanych z inferencją i AI. Przykładowe algorytmy mogą obejmować wsparte AI zliczanie kroków, szacowanie prędkości ruchu, ilości spalonych kalorii, przebytej odległości czy też rozróżnianie chodu od biegu.

Podsumowanie

Nawet najbardziej zaawansowany gadżet nie spali za nas nadmiaru kalorii. Jednak urządzenia, takie jak licznik kroków, istotnie pomagają w monitorowaniu naszych ćwiczeń, co zapewnia optymalne obciążenie, dobrane do naszych potrzeb. Dodatkowo, samodzielnie wykonany gadżet, z pewnością zmotywuje nas mocniej do uprawiania sportów, chociażby takich, jak bieganie czy chodzenie, a to na pewno nam nie zaszkodzi.

Nikodem Czechowski, EP

Źródła

1. <https://www.instructables.com/How-to-Make-a-Step-Counter/>
2. <https://github.com/Chocho2017/pedometer-ESP32>
3. <https://bit.ly/3VK5gI8>
4. <https://www.instructables.com/Pedometer-2/>
5. <https://www.instructables.com/Arduino-Pedometer/>
6. <https://bit.ly/3VDQl1M>



Haptyczne wąsy czuciowe

Karnawał jest doskonałym okresem na przebieranki, a dzięki temu, także na eksperymenty z nietypowymi projektami elektronicznymi. Zaprezentowany projekt w pewnym sensie pozwala zamienić się w kota. Elektroniczne wibrysy, czyli włosy czuciowe, występujące u ssaków, m.in. kotów, psów i gryzoni, mogą być oryginalnym uzupełnieniem przebrania oraz interesującym eksperymentem naukowym.

Zaprezentowany projekt jest kolejnym z serii sensorów haptycznych, opracowanych przez Nicholas Gonyea. Eksperymentowanie z urządzeniami tego rodzaju ma kilka celów. Z jednej strony system taki może zapewniać zupełnie nowe zmysły, nieznane dotychczas ludziom, i może być początkiem nowego wynalazku. Z drugiej strony, jak wskazuje autor, takie rozwiązanie ma posłużyć do lepszego zrozumienia zachowań kotów, czy innych zwierząt z wibrysami. Finalnie, taki sensor może pomóc np. ludziom niewidomym, którzy muszą zastępować wzrok innymi sensorami, najczęściej dotykiem.

Układ działa, jako wzmacniacz sensoryczny, w tym przypadku dotykowo-dotykowy, tj. wzmacniający bodziec dotykowy i przekazujący go do użytkownika także w postaci zmysłu dotyku. Tego rodzaju sprzężenie nazywa się haptycznym (ze starogreckiego ΑΠΙΚΟΣ haptikós, czyli „dotykalny” lub „dotykowy”). Przykładami interfejsów haptycznych, jakie nas otaczają, są np. telefony komórkowe, które mają funkcję wibrowania.

Opis układu

Układ składa się z dwóch zestawów czujników ugięcia (w sumie jest ich osiem, po cztery na stronę). Odbierają one informacje dotykowe wynikające ze zgięcia sensora, znajdującego się u podstawy naszego elektronicznego wibrysa, przy kontakcie z obiektami w bezpośrednim otoczeniu użytkownika systemu. Początkowa informacja o napięciu/rezystancji sensora mierzona przez każdy czujnik jest następnie przetwarzana na informację o kącie zgięcia sensorów. Te informacje są następnie konwertowane na wyjście proporcjonalne w postaci impulsów prostokątnych o modulowanej szerokości (PWM), które sterują natężeniem wibracji elementów haptycznych, umieszczonych na twarzy użytkownika. Elementy te tworzą tak zwany wyświetlacz haptyczny. Składa się on z macierzy silników prądu stałego z niecentrycznymi masami (niesymetryczne dyski na osiach), które wibrują, gdy silnik pracuje. Im kręci się szybciej, tym wibracje są bardziej intensywne.

Każdy czujnik ma własny mikrokontroler w postaci SparkFun ProMini 3,3 V/8 MHz, który wykonuje konwersję wartości i steruje efektem haptycznym. W każdym ProMini działają dwie pętle (takie podejście było konieczne, aby złagodzić problemy z zakłóceniami EMI powodowanymi przez silnik). Jedna funkcja służy do pozyskiwania danych z wąsów, podczas gdy druga wysyła sygnał PWM do wyświetlacza haptycznego.

Dwa wyświetlacze wibracyjne dostarczają bodźców dotykowych na czoło użytkownika (lokalizacja efektorów może być inna,

zależnie od preferencji). Każdy silnik umieszczony w każdym z dwóch wyświetlaczy jest połączony z własnym wężem (ustawione są w tej samej orientacji, co węży na twarzy), i wibruje w taki sposób, aby odzwierciedlać kąt zgięcia węża. Intensywność wibracji jest proporcjonalna do kąta zgięcia węża. Dzięki zjawiskom substytucji/wzmacniania sensorycznego bardzo prawdopodobne jest to, że po pewnym okresie treningu każdy może rozszerzyć swój wcześniej istniejący aparat somatosensoryczny o sferę tych nowych wyrostków przypominających węży, tak jak robią to osoby niewidome wykorzystując substytucję sensoryczną dotykowo-wzrokową (np. przy użyciu laski). Używając tego modułu odpowiednio często jest nawet prawdopodobne, że kora somatosensoryczna – SS (pośrednicząca przy stymulacji receptorów skóry na czole) może z czasem rozwinąć nową „reprezentację wężów” mieszczącą się gdzieś w tym obszarze mózgu. Ten rodzaj poszerzenia sensorycznego został już zaprezentowany w wielu badaniach (np. wzbogacanie sensoryczne dla niewidomych – wzmocnienie czuciowe poprzez dotyk za pomocą wyświetlacza wibracyjnego noszonego wokół talii).



**Fotografia 1. 100-milimetro-
wy jednokierunkowy sensor
zginania firmy Flexpoint Sensor
Systems**



Fotografia 2. Wtyczka JST z krótkimi przewodami przylutowana do sensora zgięcia



Budowa wężów

Opracowanie odpowiednich sensorów do wężów, było kluczowe dla powodzenia całego projektu. Muszą one być wystarczająco elastyczne, aby naśladować prawdziwe węży, a jednocześnie wystarczająco sztywne, aby sensor powracał do prostej, niewygiętej pozycji. W tym celu zastosowano 100-milimetrowy jednokierunkowy sensor zginania firmy Flexpoint Sensor Systems (**fotografia 1**). Do nóżek sensora przylutowana jest następnie wtyczka JST, która pozwala na podłączenie sensora do dalszej części układu. Następnie, sensor zginania łączony jest z paskiem polistyrenu o długości ok 355 mm, szerokości 2 mm i grubości ok 0,75 mm za pomocą kleju silikonowego. Następnie na całość nałożona jest rurka termokurczliwa. Podstawa uformowana jest z masy plastycznej – kleju formownalnego Sugru.

Przejdźmy przez szczegółowy opis kroków budowy tych sensorów:

Lutujemy 3-wyprowadzeniową wtyczkę JST z krótkimi przewodami do sensora zgięcia. Usuujemy środkowy pin złącza JST (**fotografia 2**). Potrzebne jest około 15 mm przewodu pomiędzy wtyczką. Lutowane połączenia zabezpieczamy koszulkami termokurczliwymi.

Do sensora przyklejamy, za pomocą elastycznego kleju silikonowego, pasek z elastycznego polistyrenu. Pasek musi być mocno przytwierdzony do całej powierzchni sensora.



Fotografia 3. Za pomocą masy Sugru uformowano wokół podstawy czujnika, paska i wtyczki wzmacniający element



Fotografia 4. Element, do którego przymocowane są sensory

Za pomocą masy Sugru (lub podobnej) należy uformować wokół podstawy czujnika, paska i wtyczki wzmacniający element, upewniając się, że obejmuje on wszystkie te elementy. Należy pamiętać również, aby umieścić masę na tyle wysoko, aby w pełni zabezpieczyć pasek, ale nie za wysoko, aby nie ograniczać swobodę ruchu/zginania czujnika. Z masą, taką jak Sugru, nie należy się śpieszyć. Po nałożeniu nie ma się, co śpieszyć – typowo mamy około 30...45 minut, aż masa zacznie twardnieć. Zanim w pełni stwardnieje należy upewnić się, że wtyczka jest dobrze dopasowana do strony gniazda złącza JST (**fotografia 3**).

Powyższe operacje powtarzamy jeszcze siedem razy (lub dowolną inną liczbę razy, w zależności od tego, ile chcemy mieć wężów w systemie). Należy przyłożyć dużą wagę do staranności i powtarzalności budowy każdego czujnika, gdyż znacznie uprości to potem kalibrację sensorów.

Okablowanie efektorów haptycznych

Kiedy węży są gotowe, należy zamontować je w jakiś sposób na specjalnym hełmie, który służy do ich noszenia, wraz z elektroniką. Autor zaprojektował specjalne, zakrzywione ramie z punktem mocowania sensorów na końcu. Komponent zaprojektowany został w programie Adobe Illustrator i wykonany z przezroczystego akrylu o grubości ok. 1,6 mm za pomocą plotera laserowego. Po wycięciu węży zostały osadzone w części uchwytu za pomocą masy w Sugru.

Aby sprawić, że sposób mocowania będzie pozwalał na większą elastyczność, w ramieniu dodano wycięcie, w części, którą montuje się do hełmu. Przez wycięcie przekłada się dwie śruby, które służą do montażu do hełmu. Element, do którego przymocowane są sensory, można w ten sposób przesuwac, względem hełmu, jak pokazano na **fotografii 4**. Aby zbudować ten element należy:

Do jednej strony płaskiej taśmy kablowej przylutować skręcone przewody, które następnie podłączone zostaną do modułu sterującego.



Fotografia 5. Akrylowy element, po wycięciu i uformowaniu z użyciem opalarki



Fotografia 6. Silnik wibracyjny przygotowany do zamontowania w kasku



Fotografia 7. Efekторы haptyczne zamontowane po wewnętrznej stronie kasku

Gniazda JST dla złącz, znajdujących się na końcach wąsów, zostają przylutowane na drugim końcu taśmy.

Wszystkie połączenia lutowane należy zabezpieczyć koszulkami termokurczliwymi.

Autor wykorzystał w swoim systemie 4 wąsy na stronę, co oznacza, że taśma musi mieć 8 przewodów. Jeśli chcemy skorzystać z innej liczby wąsów w systemie, należy dobrać inną taśmę.

Akrylowy element, po wycięciu, można wygiąć korzystając z ciepła opalarki, aby lepiej dostosować do kształtu twarzy. Należy wygiąć ten płaski element tak, aby miejsce montażu sensorów dotykało twarzy (**fotografia 5**).

Wiązka przewodów umieszczana jest na akrylowym ramieniu i tymczasowo zabezpieczana taśmą. Następnie za pomocą masy Sugru należy ułożyć złącza JST dla wibrysów na części ramienia w kształcie dysku. Instalując ten element należy upewnić się, że po podłączeniu, wąsy nie będą na siebie nachodzić.

Po wyschnięciu kleju Sugru, można przymocować całość do wybranego hełmu, jak opisano powyżej.

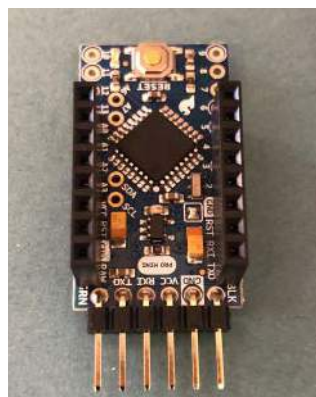
Elektronika i montaż układu

Konfiguracja wyświetlacza haptycznego w systemie jest dosyć prosta. Efekторы zainstalowane są w hełmie w taki sposób, że dotykają czoła w podobnym układzie, w jakim rozłożone są wąsy czuciowe. Na tym samym hełmie zainstalowana jest również sterująca systemem elektronika i źródło zasilania – bateria 5 V. Montaż całego układu jest relatywnie prosty. Wystarczy po kolei, realizować opisane poniżej kroki.

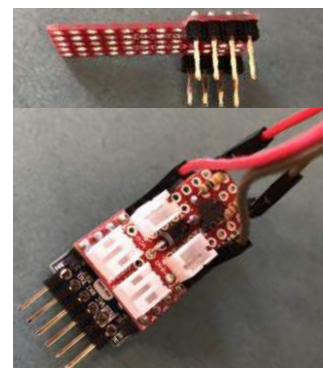
W pierwszej kolejności musimy przygotować efekторы haptyczne – silniki wibracyjne. Do silników tych lutujemy skręcone przewody, które następnie zakańczamy przewodami z 2-pinowym złączem JST.



Fotografia 8. Akumulator zamontowany w górnej części kasku



Fotografia 9. Dwa 8-pinowe złącza żeńskie dolutowane do modułu ProMini



Fotografia 10. Płytkę Pro Mini Protoshield rozbudowujemy o dwa rzędy goldpinów po cztery styki

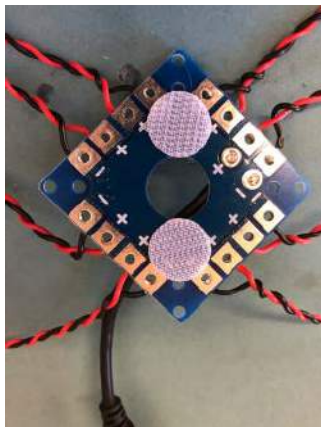
Wszystkie połączenia lutowane trzeba zabezpieczyć rurką termokurczliwą. Finalnie, na silniku wibracyjnym montowany jest samoprzylepny rzep (strona z haczykami), jak pokazano na **fotografii 6**. Te czynności powtarzamy tyle razy, ile jest sensorów i silników.

Druga strona naklejek z rzepem montowana jest na kasku (**fotografia 7**). W jego górnej części montowany jest akumulator, jak pokazano na **fotografii 8**. Obok baterii montowany jest rozdzielacz zasilania – prosta płytka, która rozdziela zasilanie z USB na wymaganą ilość linii (każdy z ośmiu modułów z mikrokontrolerem).

Każdy sensor wyposażony jest w osobny mikrokontroler. Moduł sterujący jednym wąsem składa się z modułu SparkFun Arduino Pro Mini, Pro Mini Protoshield, tranzystora BC337, diody prostowniczej 1N4001, rezystora 220 Ω , rezystora 10 k Ω i zestawu złącz, zgodnie z dalszym opisem w kolejnych krokach.

W pierwszej kolejności lutujemy dwa 8-pinowe złącza żeńskie do modułu ProMini, a następnie kątowe złącze męskie z sześcioma pinami, jak pokazano na **fotografii 9**. Po drugiej stronie modułu montujemy rzep, który pozwoli następnie zamontować moduł na hełmie.

Do płytki Pro Mini Protoshield lutujemy dwa rzędy goldpinów po cztery piny, jak pokazano na **fotografii 10a**. Na płytce tej instalowany jest także tranzystor do sterowania silnikiem wibracyjnym, dioda zabezpieczająca przed szpilkami napięcia z silnika (uzwojenie silnika generuje je w momencie zatrzymywania się) oraz konieczne złącza – dla silnika i dla modułu zasilania. Autor nie udostępnił schematu tego urządzenia, jednakże połączenia opisanych elementów powinny być oczywiste dla każdego, mającego do czynienia z elektroniką. Jedyne, co musimy wiedzieć, to to, że sensor podłączamy do pinu A0, a silnik do pinu 5 (przez tranzystor sterujący). Jeśli chcemy podłączyć te elementy do innych pinów mikrokontrolera, pamiętajmy o wprowadzeniu tych zmian w kodzie, pokazanym



Fotografia 11. Płytki rozdzielające zasilanie montowane na rzep

w dalszej części artykułu. Na **fotografii 10b** pokazano gotową płytkę.

Powtórz powyższy montaż kolejne siedem razy, aby umieścić łącznie osiem modułów z elektroniką, po jednym dla każdej płytki z mikrokontrolerem.

Do płytki dystrybucyjnej zasilania dołączamy dwa fragmenty odciętego przewodu USB z wtyczką USB-A, którą podłączymy następnie do powerbanka. Do wyjść z płytki dystrybucyjnej dołączamy z kolei osiem przewodów z złączami JST z dwoma pinami, które posłużą do zasilania poszczególnych płytek z mikrokontrolerem. Płytki rozdzielające dla zasilania montowane są finalnie na rzep na kasce, jak pokazano na **fotografii 11**.

Po obu stronach kasku, montowane są po cztery moduły z mikrokontrolerami, także na rzepach. Finalnie łączymy ze sobą wszystkie wejścia i wyjścia (sensory i silniki).

Na **fotografii 12** pokazana jest gotowa konstrukcja urządzenia. Teraz wystarczy do każdego z mikrokontrolerów załadować firmware, aby nasze haptyczne wąsy ożyły.

Oprogramowanie

Finalnym elementem systemu jest oprogramowanie. Przed wgraniem do układu, musimy zaopatrzyć się w płytkę/moduł konwertera USB-UART, np. taki z układem z rodziny FTDI. Arduino ProMini nie ma takiego konwertera na pokładzie. W samym kodzie także trzeba wprowadzić kilka zmian i wartości kalibracyjnych, osobno dla każdego wąsa. Opisane poniżej czynności powtarzamy dla każdego z mikrokontrolerów, kontrolujących pojedynczą parę sensor – efektor. Poszczególne wartości – unikalne dla każdego sensora, wprowadzane są do kodu, pokazanego na **listingu 1**.

W pierwszej kolejności, korzystając z multimetru, mierzymy napięcie zasilania na każdym z wąsów (VCC), jak i rezystancję opornika 10 kΩ w układzie. Wartości te wprowadzane są następnie w odpowiednie miejsca w kodzie. Następnie należy sprawdzić pozostałe zmienne, w szczególności te, które kontrolują wejścia i wyjścia z układu (np. *mtr*, *flexADC* itp.). Po przygotowaniu oprogramowania można podłączyć wszystkie obwody do zasilania, podłączyć moduł FTDI (podłączony do USB komputera) do pierwszego ProMini i przesłać do mikrokontrolera skompilowany kod.

Po tym przystępujemy do kalibracji układu. Gdy będzie on gotowy do pracy, zobaczymy na monitorze portu szeregowego, wydrukowane wartości: *Resistance* (opór), *Bend* (zgięcie) i *Forward Intensity* (intensywność). Każdy wąż jest unikalny i będzie miał nieco inny opór bazowy, dlatego wartości te będą różne. Ustawiamy

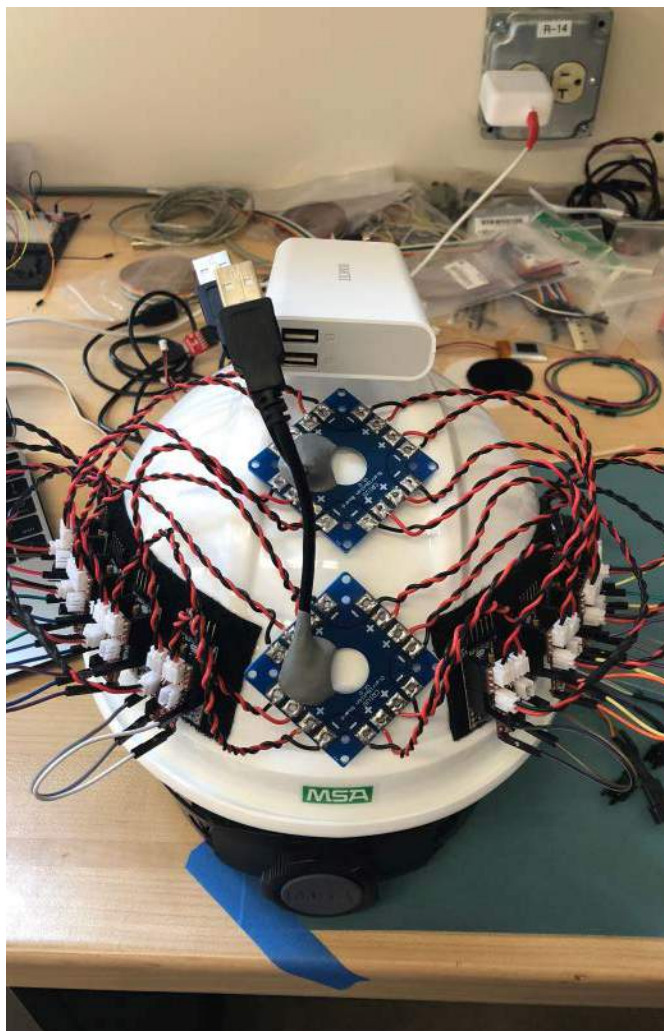
Listing 1. Firmware mikrokontrolera, znajdującego się w każdym wąsie

```
// Zmierzone napięcie zasilania - linia 3,3V
const float VCC = 3.33;
// Zmierzona rezystancja opornika 10k
const float R_DIV = 10000.0;

// Parametry kalibracyjne
// Zmierzona rezystancja prostego sensora
float STRAIGHT_RESISTANCE = 60000.0;
// Zmierzona rezystancja zgiętego sensora
float BEND_RESISTANCE = 90000.0;
// Kąt zgięcia sensora
float BEND_ANGLE = 90.0;
// Pin podłączenia sensora
int sen = A0;
// Pin podłączenia silnika
int mtr = 5;

void setup(){
  pinMode(sen, INPUT);
  pinMode(mtr, OUTPUT);
  Serial.begin(115200);
}

void loop(){
  // Odczyt z ADC
  int flexADC = analogRead(0);
  // Wyznaczenie napięcia
  float flexV = flexADC * VCC / 1023.0;
  // Wyznaczenie rezystancji
  float flexR = R_DIV * (VCC / flexV - 1.0);
  Serial.println("Resistance: " + String(flexR) + " ohms");
  // Mapowanie wartości rezystancji na kąt
  float angle = map(flexR, STRAIGHT_RESISTANCE, BEND_RESISTANCE, 0.0, BEND_ANGLE);
  Serial.println("Bend: " + String(angle) + " degrees");
  // Mapowanie wartości kąta na siłę wibracji
  float FWD_intensity = map(angle, 5.0, 80.0, 30.0, 130.0);
  Serial.println("Forward Intensity: " + String(FWD_intensity));
  if (angle > 5.0){
    analogWrite(mtr, FWD_intensity);
  }else{
    analogWrite(mtr, 0);
  }
  delay(500);
}
```



Fotografia 12. Gotowa konstrukcja urządzenia – wszystkie komponenty zamontowane na kasku

zmienną *STRAIGHT_RESISTANCE* na dowolny opór nieugiętego wąsa, zgodnie z tym, co widzimy na ekranie z monitorem portu szeregowego. Następnie zmienną *BEND_RESISTANCE* ustawiamy na wartość równą *STRAIGHT_RESISTANCE* + 30000,0. Zmienna ta ma odzwierciedlać rezystancję wyjściową czujnika ugięcia przy zgięciu go o 90 stopni. Ponieważ nasze wąsy nie zbliżają się do takiego zgięcia (przynajmniej w typowych sytuacjach), dodanie po prostu 30000,0 Ω do rezystancji bazowej działa dostatecznie dobrze. Można jednak ustawić inną oporność podczas zgięcia, w zależności od potrzeby. Można oczywiście wstawić tutaj także zmierzoną rezystancję sensora przy arbitralnie dobranym kącie. Kąt ten, przy którym zmierzono rezystancję, wpisać należy w zmiennej *BEND_ANGLE*.

Po wprowadzeniu danych kalibracyjnych do kodu, można go ponownie skompilować i wgrać do układu. Jeśli wszystko ustawiono poprawnie, gdy wąż jest rozprostowany, przez port szeregowy wydrukowany zostanie kąt zgięcia 0 stopni (mniej więcej). Analogicznie, wyświetlany powinien być każdy inny kąt, o jaki zgięty zostanie sensor.

Działanie oprogramowania jest bardzo proste. Po zmierzeniu napięcia i wyliczeniu rezystancji przeprowadzane są dwie różne funkcje mapowania. Pierwsza odwzorowuje zmierzoną rezystancję na kąt zgięcia, a druga odwzorowuje kąt zgięcia na natężenie vibracji z silnika (sterowanego przez PWM). Po przetestowaniu urządzenia konieczne może być zmodyfikowanie drugiej funkcji mapowania, aby uzyskać odpowiedni wynik. Parametry zgięcia są początkowo ustawione na 5,0 i 80,0, a parametry wyjściowe silnika na 30,0 i 130,0. Te wartości działały w przypadku autora, ale to bardzo osobnicze preferencje. Wszystko poniżej 5,0 stopni na dolnym końcu powoduje problemy, ponieważ czujnik waha się w przybliżeniu $\pm 3,0$ stopni wokół zera cały czas. Wszystko powyżej 80,0 stopni jest z kolei niepotrzebne, ponieważ czujnik rzadko wygina się poza ten kąt.

Silnik vibracyjny, jaki użył autor, ma próg zadziałania na poziomie około 40,0 jednostek PWM (tj. nie będzie wyczuwalnie wibrował, dopóki wyjście PWM nie osiągnie około 40,0). Rozpoczęcie zakresu

od 30,0 pozwala nieco buforować początkowy poziom vibracji. Górny zakres ustawiono na 130,0, ponieważ ten poziom jest więcej niż wystarczający dla pełnej wyczuwalności.

Obecna wersja oprogramowania, jak pisze autor, nie jest ostateczna. Kolejna wersja miałaby być rozbudowana o ciągłą samokalibrację czujnika ugięcia. Obecnie, gdy używane wąsy są w użyciu, opór może znacznie zmieniać się poza początkowo ustawioną wartością. Może to prowadzić do zmniejszenia czułości wążów. Dodanie ciągłej samokalibracji może pozwolić na wyeliminowanie tego szkodliwego zjawiska. Na razie jednak, od czasu do czasu może być konieczna ponowna kalibracja wążów.

Podsumowanie

Jako, że ludzie nie są wyposażeni w tego rodzaju zmysł, jakiegokolwiek realne zastosowanie tych sensorów wymaga nauki i treningu. Jednak z uwagi na ogromną plastykę mózgu, nauczanie się, jak korzystać z takich wążów nie jest niemożliwe. Taki sensor może być elementem ciekawego przebrania karnawałowego, jak i pomocą dla osób np. niewidomych lub niedowidzących. Spektrum zastosowań tego rodzaju rozwiązania jest szerokie i nawet, jeżeli jego realne użycie jest problematyczne, to sama idea działania, tj. wykorzystanie czuciowych elementów mechanicznych i sprzężenia haptycznego, jest bardzo ciekawa, szczególnie w kontekście osób pozbawionych możliwości korzystania ze zmysłu wzroku.

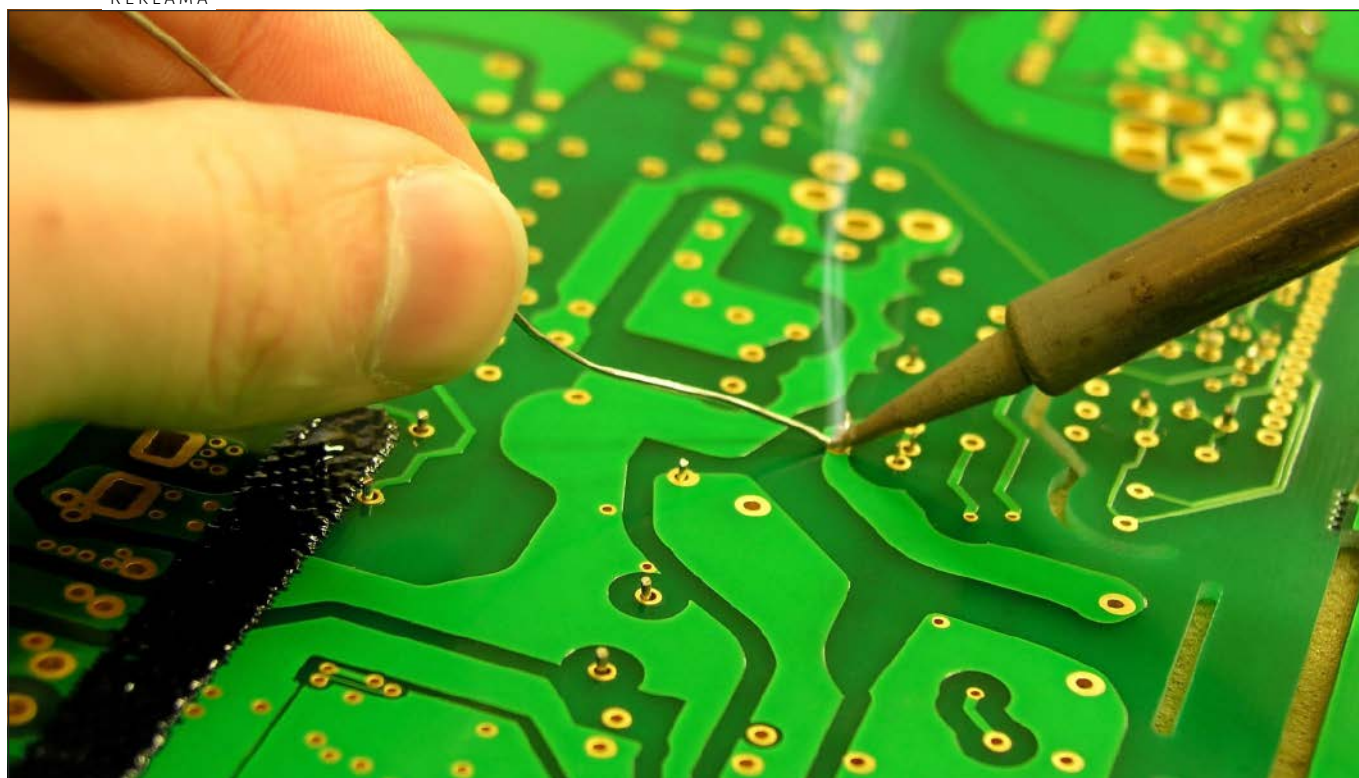
Zaprezentowany projekt nie tylko pozwala na skonstruowanie ciekawego urządzenia, ale także poznanie ciekawych technologii – elastycznych sensorów oraz efektorów haptycznych. Każdy z tych elementów z powodzeniem może być wykorzystany w innych projektach, poszerzając ich funkcjonalność czy wręcz umożliwiając zupełnie nowe zastosowania.

Nikodem Czechowski, EP

Źródło

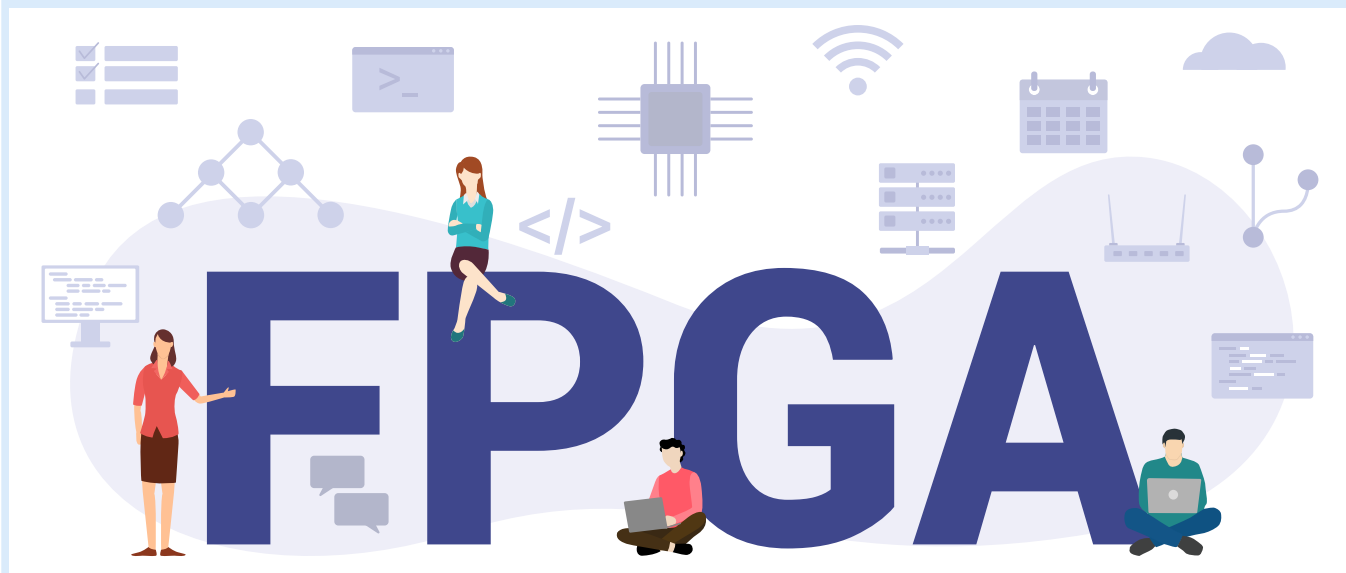
<https://www.instructables.com/Whisker-Sensory-Extension-Wearable/>

REKLAMA



KITy AVT

@KITyAVT <http://bit.ly/2BjVMN7>



Kurs FPGA Lattice (3)

Podstawy języka Verilog

W tym odcinku kursu nauczymy się podstawowych instrukcji języka Verilog i opisemy proste układy kombinacyjne oraz sekwencyjne. Poprzednie części kursu można znaleźć na ep.com.pl.

Moduły i instancje modułów

Kod języka Verilog dzielimy na moduły, podobnie jak kod w językach programowania dzieli się na funkcje czy klasy. Każdy moduł komunikuje się z pozostałymi modułami przy pomocy portów – mogą one być wejściem *input*, wyjściem *output* lub sygnałem dwukierunkowym *inout*. Moduł nadrzędny, od którego zaczyna się cała aplikacja zwyczajowo nazywany jest *top*. Jest odpowiednikiem funkcji *main()* w C++. Każdy inny moduł należy opisać, a następnie powołać do życia tworząc co najmniej jedną instancję tego modułu. Moduły wygodnie jest zapisywać w oddzielnych plikach, których nazwa jest taka sama jak nazwa modułu.

Prześledźmy przykład zaprezentowany na **listingu 1**. W liniach oznaczonych jako **#4** i **#5** definiujemy moduły *BramkaAND* i *BramkaOR*. Każda z nich ma dwa wejścia *x* i *y* oraz wyjście *result*. Samo zdefiniowanie tych modułów jeszcze nie spowoduje wygenerowania ich w układzie FPGA. W module *top* w liniach oznaczonych **#1** i **#2** zostały utworzone dwie instancje modułu *BramkaAND*. W nawiasach podane są nazwy portów, jakie ma moduł *top*. W ten sposób wejścia i wyjścia modułu *top*, które są fizycznymi pinami układu FPGA, zostają połączone z wirtualnymi pinami modułów.

Należy zwrócić uwagę, że utworzyliśmy dwie instancje *AND1* i *AND2*, które działają dokładnie tak samo, jednak połączone są z innymi sygnałami. Istnieją dwa sposoby tworzenia instancji i łączenia ich z sygnałami z innych modułów. Sposób pokazany w liniach **#1** i **#2** jest podobny do kodu C++, w którym tworzony jest obiekt jakiejś klasy, a jego konstruktor przyjmuje kilka argumentów. W tym przypadku kolejność argumentów ma znaczenie. Taki sposób jest przydatny dla bardzo prostych modułów, które mają niewiele portów. Drugi sposób został pokazany w linii **#3**. Każdy z portów modułu



Poprzednie odcinki znajdują się pod adresem:
<https://ulubionykiosk.pl/media>

wymieniony jest po znaku kropki. Następnie w nawiasach podana jest nazwa sygnału, z którym dany port ma być połączony. Taki opis jest bardziej rozwlekły i czasochłonny, jednak ma kilka istotnych zalet.

Listing 1. Przykład podziału kodu na moduły

```
// Plik top.v
module top(
    input ButtonA,
    input ButtonB,
    input ButtonC,
    output LED,
    output Buzzer,
    output Relay
);

// Instancje modułów typu BramkaAND o nazwie AND1 i AND2
BramkaAND AND1(ButtonA, ButtonB, LED); // #1
BramkaAND AND2(ButtonA, ButtonC, Buzzer); // #2

// Instancja modułu typu BramkaOR o nazwie OR1 // #3
BramkaOR OR1(
    .x(ButtonA),
    .y(ButtonB),
    .result(Relay)
);

endmodule

// Plik bramka_and.v // #4
module BramkaAND(
    input x, y,
    output result
);

    assign result = x & y;

endmodule

// Plik bramka_or.v // #5
module BramkaOR(
    input x, y,
    output result
);

    assign result = x | y;

endmodule
```


Zapamiętaj

Układ kombinacyjny to taki, którego stan wyjścia zależy tylko od obecnego stanu wejść.

Układ sekwencyjny to taki, którego stan wyjścia zależy od obecnego stanu wejść oraz od stanów, jakie były w przeszłości. Układ sekwencyjny ma jakiś element pamięciowy, a często sterowany jest sygnałem zegarowym.

Po pierwsze kod staje się bardziej czytelny, ponieważ widzimy nazwę portu wewnątrz modułu i jednocześnie sygnał, z którym jest połączony na zewnątrz. Dodatkowo każdy z nich możemy opatrzyć komentarzem. Dzięki temu, że każdy port jest zapisany w osobnej linii, w razie błędu syntezy powie nam, w której linii jest błąd – w przypadku pierwszego sposobu, gdzie wszystko opisujemy w jednej linijce, komunikatory syntezy mogą być trudniejsze do zrozumienia. Ponadto, kolejność przypisań jest dowolna.

Zmienne

W języku Verilog stosuje się głównie dwa typy zmiennych (jest ich więcej, ale dwa są szczególnie ważne dla syntezy). Są to **wire** oraz **reg**. Oba dostarczają sygnały logiczne, ale reprezentują inne zasoby sprzętowe.

- **Reg** jest elementem, który przechowuje stan logiczny tak długo, aż nie zostanie do niego przypisany inny stan. Najczęściej jest syntezowany jako przerzutnik D, który przepisuje stan swojego wejścia na wyjście w momencie wystąpienia określonego zbocza sygnału zegarowego. Zmienna typu **reg** może tworzyć także przerzutnik RS lub inny, a grupa takich zmiennych może tworzyć pamięć RAM. Jednak należy zaznaczyć, że zmienna typu **reg** wcale nie musi być syntezowana jako przerzutnik – przykład takiej sytuacji jest opisany w dalszej części artykułu. **Reg** może służyć do tworzenia układów zarówno kombinacyjnych jak i sekwencyjnych.
- **Wire** to najprościej mówiąc przewód, który łączy jakieś dwa obiekty. Zmienną typu **wire** musimy przypisać do jakiegoś sygnału źródłowego. W pewnym sensie **wire** przypomina wskaźnik do zmiennej, jednak nie jest to żaden adres zmiennej z pamięci, lecz fizyczne połączenie do jakiegoś wejścia, wyjścia, przerzutnika, itp. **Wire** stosuje się do modelowania układów kombinacyjnych. Porty wejściowe **input**, dwukierunkowe **inout** oraz wyjściowe **output** są sygnałami typu **wire**. Istnieje możliwość, by porty wyjściowe miały możliwość zapamiętywania stanu. Wtedy trzeba je zadeklarować jako **output reg**.

Oprócz **wire** i **reg**, istnieją także typy **time**, **integer**, **real** i wiele innych. Znajdują one zastosowanie głównie w symulacji układów. Wróćmy do nich w następnych odcinkach kursu, kiedy będziemy omawiać kwestie związane z symulacją układów cyfrowych.

Skalary i wektory

Zmienne typu **reg** i **wire** to tylko jeden bit, który może przyjmować wartość 1 lub 0 (w pewnych przypadkach także stan wysokiej impedancji Z lub stan nieokreślony X). Takie zmienne nazywane są skalarami. Kiedy zachodzi potrzeba wykonywania operacji na większej liczbie bitów jednocześnie, możemy utworzyć grupę takich zmiennych, nazywaną wektorem lub magistralą.

Zapoznajmy się z praktycznymi przykładami z listingu 2. W liście portów znajdziemy wejścia i wyjścia wielobitowe. Pierwsza cyfra w nawiasach kwadratowych oznacza numer najstarszego bitu, a druga cyfra to numer najmłodszego bitu. W linii oznaczonej #1 widzimy przykład, jak przypisać jeden wybrany bit sygnału wielobitowego do sygnału 1-bitowego. W linii #2 przypisano sygnał 1-bitowy do zerowego bitu zmiennej 4-bitowej. Istnieje możliwość, by w jednym poleceniu połączyć kilka wybranych bitów. Taki przykład

widzimy w linii #3. W linii #4 jest nowość – nawiasy klamrowe. Jest to operator konkatenacji czyli sklejania ze sobą różnych zmiennych w celu uzyskania jednej zmiennej o większej liczbie bitów. W przykładzie z linii #4 sklejamy fragment zmiennej C z całością zmiennej B i w rezultacie dostajemy sygnał 8-bitowy, który przypisujemy do Z. W linii #4 pominięto nawiasy kwadratowe przy zmiennej Z, pomimo że wpisujemy dane do zmiennej wielobitowej. Kiedy wpisujemy dane do całej zmiennej nie ma potrzeby, by wskazywać na wszystkie jej bity.

Programiści przyzwyczajeni do języka Python powinni zwrócić uwagę, że w Verilogu zapis *zmienna[x:y]* oznacza wybór wszystkich bitów pomiędzy x i y, włączając też y.

Notacja liczbową

W języku Verilog stosuje się dość nietypową notację do zapisu liczb. Najpierw podajemy długość w bitach, następnie jest apostrof, potem litera określająca format liczby i na końcu właściwe dane. Wszystko stanie się jasne, kiedy przeanalizujemy przykłady pokazane na listingu 3. Mamy do wyboru jeden z czterech formatów liczbowych: **b** binarny, **o** ósemkowy, **d** dziesiętny i **h** szesnastkowy.

- W przypadku zmiennych o długości 1 bitu (przykład A i B) format nie ma praktycznie znaczenia i może być dowolny.
- Przykłady C, D i E to różne sposoby przypisania danych 4-bitowych, różniących się formatem, choć zapisywana jest ta sama liczba.
- Przykład F pokazuje, że kiedy stosujemy zapis binarny pomocne może być dodanie separatora **_** aby zwiększyć czytelność długiego szeregu zer i jedynek.
- Przykład G wymaga szerszego komentarza. Na pierwszy rzut oka wydaje się on być taką samą operacją, jaka jest w przykładzie B czyli przypisanie jedynki do zmiennej. Zmienna liczbową podana normalnie bez liczby bitów i formatu traktowana jest jako typ **integer**, który ma 32 bity i jest liczbą ze znakiem. Operacja w przykładzie G jest akceptowalna i zostanie wykonana prawidłowo, ale syntezy zgłosi **warning**, że musi wykonać rzutowanie.
- Przykład H to przypisanie stanu wysokiej impedancji. Ma zastosowanie dla sygnałów **inout** wyprowadzonych na piny układu FPGA.
- Przykład I będzie działał tylko w symulatorze i jest niesyntezowalny. Stan nieokreślony zobaczymy w symulatorze, kiedy nie zainicjujemy jakiejś zmiennej. Przydaje się także w instrukcjach sterujących **if** oraz **case**, aby wskazać bity, których stan jest dla nas nieistotny.

Listing 2. Przykłady operacji na skalarach i wektorach

```
module top(
    input A,           // Wejście 1-bitowe
    input [3:0] B,      // Wejście 4-bitowe
    input [7:0] C,      // Wejście 8-bitowe
    output X,           // Wyjście 1-bitowe
    output [3:0] Y,      // Wyjście 4-bitowe
    output [7:0] Z       // Wyjście 8-bitowe
);

    assign X = B[3];      // #1
    assign Y[0] = A;      // #2
    assign Y[3:1] = B[2:0]; // #3
    assign Z = {C[7:4], B[3:0]}; // #4

endmodule
```

Listing 3. Przykłady zapisu danych liczbowych

```
assign A = 1'b0;        // 1-bitowe binarne 0
assign B = 1'b1;        // 1-bitowe binarne 1
assign C = 4'b1111;     // 4-bitowe binarne 1111
                        // (czyli 15 dziesiętnie)
assign D = 4'hF;        // 4-bitowe szesnastkowe F
                        // (czyli 15 dziesiętnie)
assign E = 4'd15;       // 4-bitowe dziesiętne 15
assign F = 32'b00001111_11110000_11001100_10101010;
                        // Długa liczba
assign G = 1;           // 32-bitowe dziesiętne 1
assign H = 1'bZ;        // Stan wysokiej impedancji
assign I = 1'bx;        // Stan nieokreślony
assign J = 4'b01xZ;
```



Rysunek 1. Symulacja licznika, którego działanie definiuje kod z listingu 4

Układy sekwencyjne

Układy sekwencyjne mają jakiś rodzaj pamięci. Na ogół jest to przerzutnik D taktowany globalnym sygnałem zegarowym. Opcjonalnie może mieć wejście asynchroniczne resetujące i ustawiające. Przeanalizujmy **listing 4**. Jest to prosty licznik 4-bitowy liczący w górę z możliwością asynchronicznego zerowania. Deklaracja portów jest zgodna z tym, co już wcześniej poznaliśmy podczas kursu. Nowością jest instrukcja *always*, rozpoczynająca blok proceduralny, w którym opisany jest sposób działania licznika.

Blok *always* @ wykonuje się, kiedy zostanie spełniony jeden z warunków opisanych w nawiasach. Jest to tzw. lista wrażliwości (*sensitivity list*). Nasz licznik ma reagować na zbocze rosnące sygnału zegarowego (zmianę z 0 na 1), więc etykieta **Clock** poprzedzona jest instrukcją *posedge*. Natomiast sygnał resetujący ma działać w taki sposób, że kiedy ustawiony jest w stanie 1 to licznik ma działać normalnie, a kiedy przejdzie w stan 0, wówczas licznik ma się wyzerować i ma pozostawać wyzerowany tak długo, aż reset znów przyjmie stan 1. Wykrywanie zbocza opadającego sygnału zapewnia instrukcja *negedge*. W naszym przykładzie sygnały w liście wrażliwości rozdzielone są przecinkiem, ale można zamiast niego stosować operator *or*.

Następnie mamy instrukcję warunkową *if-else*, która działa dokładnie tak samo jak w C++. Warunkiem sprawdzanym w instrukcji *if* jest zanegowany sygnał resetujący. Negację zapewnia operator *!* – tak samo jak w C++ – czyli kiedy sygnał **Reset** ma stan 0, wówczas warunek jest prawdziwy i wykonuje się linia #1, a w przeciwnym wypadku linia #2.

W liniach #1 i #2 widzimy jak można przypisać nową wartość do zmiennej typu *reg*. W odróżnieniu od zmiennych *wire*, zmienna *reg* przechowuje jakieś dane. W przypadku *reg* nie stosujemy instrukcji *assign*. Dane do zmiennych *reg* można zapisywać poprzez przypisanie

blokujące *=* i przypisanie nieblokujące *<=*. W układach sekwencyjnych należy stosować *<=*, a w kombinacyjnych *=*. Jest to dobra praktyka.

W linii #1 zerujemy licznik, a w linii #2 zwiększamy jego wartość o 1. Licznik jest 4-bitowy, a więc liczy od 0 do 15. Kiedy przekroczy swoją maksymalną wartość, wtedy wyzeruje się i będzie liczył od początku.

W przeciwieństwie do C++, w języku Verilog nie ma operatora inkrementacji *++* ani dekrementacji *--* (niestety!). Nie ma także operatorów *+=* i *-=*. Kiedy chcemy zwiększyć lub zmniejszyć zmienną, niestety musimy posłużyć się taką konstrukcją jaką zastosowano w linii #2. Przydatne operatory inkrementacji i dekrementacji zostały dodane dopiero w standardzie SystemVerilog.

Jeżeli po instrukcjach takich jak *if*, *else*, *always* występuje więcej niż jedno polecenie, wówczas musimy te polecenia objąć słowami *begin* i *end*, które są odpowiednikami nawiasów klamrowych z C++.

Kiedy przetestujemy nasz licznik w symulatorze, uzyskamy przebiegi takie, jak zaprezentowano na **rysunku 1**.

Układy kombinacyjne

Układ kombinacyjny to taki, którego stan wyjść zależy tylko od obecnego stanu wejść. Oznacza to, że nie ma żadnej pamięci. Jego działanie jest teoretycznie natychmiastowe, ale w rzeczywistości jest niewielkie opóźnienie pomiędzy zmianą stanu wejść, a ustaleniem się stanu wyjść, nazywane czasem propagacji. Przykładem układu kombinacyjnego jest multiplexer. Zaprezentujemy trzy sposoby realizacji multiplexera, który ma cztery wejścia danych, dwa wejścia adresowe i jedno wyjście.

Przeanalizujmy teraz **listing 5**. Pierwszy sposób zawiera instrukcję *assign* oraz operator warunkowy *?:*, który działa dokładnie tak samo jak w C++. Ma on postać *warunek_logiczny ? wartość_jeżeli_prawda : wartość_jeżeli_fałsz*. W ten sposób cała logika sprowadza się do jednej instrukcji *assign*, która przypisuje do wyjścia **Out** jedno z czterech wejść **DataX**. W pierwszej kolejności sprawdzane jest czy **Address** jest równy *2'd0*. Jeśli tak, to zmienna **Out** (która jest typu *wire*) łączona jest z **Data0** i na tym koniec. Jeżeli nie, wówczas sprawdzany jest kolejny warunek i tak dalej. Cała logika de facto ogranicza się do pojedynczej instrukcji, która tylko dla zachowania czytelności kodu została zapisana w czterech liniach.

Na **listingu 6** zostało pokazane prostsze rozwiązanie, zawierające operator selekcji. Wejście **Data** zostało zrealizowane jako wektor 4-bitowy, a nie cztery niezależne sygnały. Dzięki temu możemy z niego wybrać jeden z sygnałów od 0 do 3 umieszczając go w nawiasach kwadratowych. Numer wybieranego sygnału może być podany poprzez zmienną. Działa to podobnie jak odczytanie zmiennej z tablicy w C++.

Ostatni przykład z **listingu 7** pokazuje, w jaki sposób można zastosować blok *always* w układzie kombinacyjnym. Na początku zwróć

Listing 4. Kod prostego licznika 4-bitowego liczącego w górę z możliwością asynchronicznego zerowania

```
module Counter(
    input Clock,           // Wejście zegarowe
    input Reset,           // Wejście zerujące
    output reg [3:0] Out   // 4 przerzutniki D tworzące licznik
);

// Opis behawioralny licznika
always @(posedge Clock, negedge Reset) begin
    if(!Reset)
        Out <= 4'd0;      // #1
    else
        Out <= Out + 1'b1; // #2
    end
endmodule
```

Listing 5. Multiplexer z użyciem operatora warunkowego *?:*

```
module Mux1(
    input Data3, Data2, Data1, Data0,
    input [1:0] Address,
    output Out
);

// assign zmienna = warunek ? wartość_jeżeli_prawda : wartość_jeżeli_fałsz

assign Out = (Address == 2'd0) ? Data0 : // tu dwukropek
              (Address == 2'd1) ? Data1 : // tu dwukropek
              (Address == 2'd2) ? Data2 : // tu dwukropek
              Data3; // tutaj średnik!

endmodule
```

Listing 6. Multiplexer z zastosowaniem selekcji

```
module Mux2(
    input [3:0] Data,
    input [1:0] Address,
    output Out
);

assign Out = Data[Address];

endmodule
```


Listing 7. Multiplexer w postaci bloku always

```

module Mux3(
    input [3:0] Data,
    input [1:0] Address,
    output reg Out           // Zwróć uwagę na reg
);

always @(Address, Data) begin // Można też always @(*)
    case(Address)
        2'b00: Out = Data[0];
        2'b01: Out = Data[1];
        2'b10: Out = Data[2];
        2'b11: Out = Data[3];
        default: Out = 1'b0; // W tym przypadku
                             // niepotrzebne
    endcase
end
endmodule

```

uwagę na to, że **Out** jest zadeklarowane jako **output reg**, jednak mimo to, nie zostanie zsyntezowany żaden przerzutnik! Układ kombinacyjny zależy od stanu wejść, dlatego w liście czułości bloku **always** są jedynie nazwy sygnałów, na jakie blok ma reagować, ale nie ma żadnych instrukcji **posedge** ani **negedge**. Układ ma reagować natychmiast na zmianę sygnału wejściowego **Data** i **Address**. Tak się składa, że w bloku **always** nie ma żadnych innych sygnałów wejściowych. W tej sytuacji w liście czułości można wstawić gwiazdkę ***** która oznacza: reaguj na wszystko. Przykład ten demonstruje użycie

instrukcji **case**, która działa podobnie do **switch-case** z C++. Różnica jest taka, że nie trzeba pisać **break** na końcu każdej możliwości. Podobnie jak w przypadku **if-else**, jeżeli ma zostać wykonana więcej niż jedna instrukcja, należy je objąć słowami **begin-end**. W tym przypadku wpisujemy tylko wartość do jednej zmiennej, więc **begin-end** pominięto.

W instrukcji **case** można łatwo wpaść w pułapkę i stworzyć układ sekwencyjny zamiast kombinacyjnego. Stanie się to wtedy, kiedy instrukcja **case** nie będzie obejmowała wszystkich możliwych kombinacji sygnału **Address**. Wówczas układ będzie musiał pamiętać swój ostatni stan na wypadek wystąpienia przypadku nieokreślonego i z tego powodu zostanie zsyntezowany przerzutnik. Aby uniknąć takiej sytuacji należy wypisać wszystkie możliwe kombinacje lub zastosować instrukcję **default**, która obejmuje wszystkie przypadki, jakie nie zostały opisane.

Podsumowanie

Wystarczy teorii na dziś! W tym odcinku zaprezentowałem absolutnie elementarne informacje o Verilogu. Zachęcam do dalszego zgłębiania wiedzy. Polecam strony chipverify.com i www.fpga4fun.com, a w kolejnym odcinku wykonamy coś praktycznego.

Dominik Bieczyński
leonow32@gmail.com

REKLAMA

Nie przegap styczniowego wydania „Elektroniki dla Wszystkich”, w której przeczytasz m.in.:

PROJEKTY dla elektroników

- ▶ Lampowy moduł przesterowania i zniekształceń do gitary
- ▶ Ściemniacz nowoczesnego oświetlenia zasilanego z sieci, część 2
- ▶ Zdalna stacja nadzoru (monitoringu)

DIY dla wszystkich

- ▶ Szkoła Konstruktorów
- ▶ Poziomy logiczne, część 1
- ▶ Zwrótnica do mini monitora PE, część 1
- ▶ Silniki krokowe w praktyce, część 2: Wybór i identyfikacja silników krokowych
- ▶ Edukacja w EdW dla szkół i uczelni: Wykład 2 „Układy scalone”
- ▶ Pokój Nauczycielski

TUTORIALE

- ▶ Sterowanie natężenia światła diod LED przy pomocy techniki PWM
- ▶ Lampa-sygnalizator przelotu Międzynarodowej Stacji Kosmicznej (ISS)
- ▶ Prosta ładowarka bezprzewodowa do smartfonów
- ▶ Wykrywanie i klasyfikacja obiektów za pomocą Raspberry Pi i uczenia maszynowego wykorzystującego platformę Edge Impulse

przejrzyj i kupisz na www.ulubionykiosk.pl



koktajl niusów



Zestaw ewaluacyjny EXMU-X261 firmy Innodisk

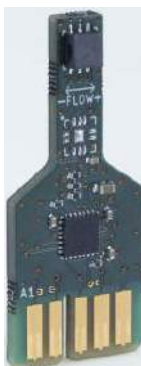
Firma Innodisk wprowadziła do oferty unikalny zestaw ewaluacyjny EXMU-X261. Zastosowano w nim układ SoM o oznaczeniu Kria K26, którego producentem jest firma Xilinx. Zestaw przeznaczony jest np. do aplikacji widzenia maszynowego i może być częścią systemu zautomatyzowanej kontroli produktów w fabrykach. Urządzenie oferuje funkcje zarządzania zdalnego dzięki modułowi InnoAgent, natomiast dzięki pakietowi AI Suite SDK firmy Innodisk możliwa jest szybka budowa aplikacji. Pakiet ten stanowią w szczególności 2 rozwiązania: iCAP (Innodisk Cloud Administration Platform) oraz iVIT (Innodisk Vision Intelligence Toolkit). Środowisko iVIT przeznaczone do sprawnego tworzenia oraz wdrażania aplikacji bez pisania kodu. Urządzenie zawiera porty LAN i 4×USB 3.1 oraz złącza M.2. i jest przeznaczony do pracy w temperaturze otoczenia 0...70°C.

<https://bit.ly/3I1b24B>



Zestaw ewaluacyjny wyposażony w czujnik przepływu powietrza

Dostępny w ofercie firmy Flusso zestaw FLS122 gwarantuje precyzyjne pomiary temperatury oraz prędkości przepływu powietrza. Obowiązujące dla tych pomiarów zakresy wartości wynoszą -40...85°C dla temperatury oraz 0...20 m/s dla przepływu powietrza, oba przy niepewności pomiarowej ±5%. To oryginalne rozwiązanie zawiera miniaturowy czujnik o wymiarach 3,5×3,5 mm, który jest konfigurowany i kalibrowany poprzez interfejs I²C. Powstały dzięki Flusso zestaw można łatwo integrować z wieloma systemami – czujnik można skalibrować tak, aby badał przepływ powietrza przy blisko zerowym spadku ciśnienia, istnieje możliwość



dołączania zestawu ewaluacyjnego do komputerów przenośnych, itd. Chociaż FLS122 miał pierwotnie służyć aplikacjom związanym z monitorowaniem filtrów, to ze względu na niski pobór mocy i łatwość stosowania, zestaw jest odpowiedni do zasilanych bateryjnie aplikacji testowych. Jak dodaje wiceprezes ds. sprzedaży i marketingu w firmie Flusso, Thomas Negre: „Firma Flusso otrzymała tak sporo pozytywnych opinii na temat zestawu FLS122, że nadeszła pora wprowadzić go do masowej produkcji tak szybko jak to tylko możliwe. Zestaw ten będzie produkowany przez nas, w ramach dostępnego łańcucha dostaw, który jest w stanie sprostać rocznej produkcji zestawu na poziomie 100 milionów egzemplarzy”.

<https://bit.ly/3WraTMk>



Scalone wzmacniacze audio klasy D firmy STMicroelectronics do zastosowań w motoryzacji

Wzmacniacze typu FDA803S i FDA903S zapewniają wierne odtwarzanie głosu, muzyki i komunikatów ostrzegawczych przy maksymalnej mocy wyjściowej 10 W. Zawierają 3 interfejsy: I²C, I²S i TDM, a także przetwornik cyfrowo-analogowy o rozdzielczości 24 bitów. Dzięki cyfrowemu przetwarzaniu gwarantowana jest wysoka jakość dźwięku, o której stanowi m.in. dynamika na poziomie prawie 100 dB. Układy wyposażone są w zabezpieczenia termiczne i nadprądowe oraz detekcje zwarcia. Wzmacniacz FDA903S zawiera monitor prądu, który umożliwia przeprowadzanie autodiagnostyki zgodnej z normą ASIL A. Oba wzmacniacze są dostępne w obudowie QFN32 o rozmiarach 5×5 mm i nie wymagają stosowania radiatorów. Kluczową cechą jest też niski pobór mocy w stanie spoczynku. Każdy z nich oferuje częstotliwości próbkowania z zakresu 8...96 kHz.

<https://bit.ly/3FSWS30>

Programowalne źródła sygnału zegarowego z serii VersaClock 7 od Renesas Electronics

Tytułowe komponenty zawierają zintegrowane oscylatory kwarcowe i są przeznaczone m.in. dla komputerów typu high-end, infrastruktury bezprzewodowej i centrów danych. Źródła sygnału zegarowego VersaClock 7 oferują wyjścia różnicowe, dla których średni jitter wynosi mniej niż 150 fs. Każde źródło można zasilać jednym z napięć: 1,8, 2,5 i 3,3 V. Układy obsługują standard PCIe Gen 1...6 oraz Ethernet. Dostępne są też 3 interfejsy: I²C, SPI i SMBUS, a cała struktura jest zamknięta w obudowie QFN o rozmiarach: 5×5 mm bądź 6×6 mm.

Produkowane przez Renesas Electronics źródła sygnału zegarowego zawierają pamięci OTP do przechowywania 27 konfiguracji



ustawień. Mowa tu m.in. o częstotliwości wspomnianego sygnału oraz poziomach napięć wyprowadzeń. Do konfigurowania i zarządzania źródłami serii VersaClock 7 powstała aplikacja Renesas IC ToolBox (RICBox) dla systemów operacyjnych Windows. Umożliwia ona też korzystanie z bibliotek języka Python i jednocześnie nie wymaga specjalistycznej wiedzy na ten temat. Utworzone konfiguracje mogą zostać pobrane, bądź można je przysłać do usługi Lab on the Cloud autorstwa Renesas Electronics. Jak mówi wiceprezes działu Timing Products Division w firmie Renesas Electronics, Zaher Baidas: „Wszystkie potrzeby związane z synchronizacją czasową mogą być zróżnicowane. Źródła sygnału zegarowego z serii VersaClock 7 oferują dużą elastyczność konfiguracji kluczowych parametrów, oferując najlepszą wydajność dostępną dla wielu wymagań”.

<https://bit.ly/3Gk7IRa>



Układ NXH3675 kompatybilny ze standardem Bluetooth 5.3 LE

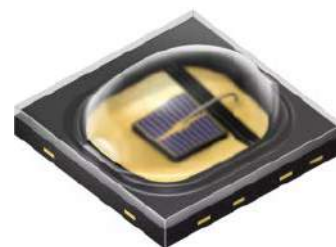
Do oferty NXP Semiconductors dodano układ NXH3675 – energooszczędny transceiver do aplikacji audio. Oferuje niskie opóźnienia przy transmisji, dzięki specjalnemu protokołowi przesyłania danych, opracowanemu przez NXP Semiconductors. Zastosowania układu obejmują m.in. słuchawki do gier, mikrofony oraz soundbary. W ich przypadku standard Bluetooth 5.3 LE zapewnia współdzielenie dźwięku, co stwarza nowe zastosowania, niewyobrażalne dla starszych aplikacji Bluetooth. Układ jest dostępny w obudowie WL-CSP i pobiera zaledwie 20 mW mocy.

Wierną transmisję dźwięku umożliwia procesor DSP o nazwie CoolFlux, który wspiera 2 kodeki audio: LC3 i LC3+. Układ NXH3675 obsługuje do 7 strumieni audio broadcast i multicast, również przy opcji Auracast. Udostępnia takie funkcje jak: korekcja (EQ), kompresja, kontrola wzmocnienia, czy miksowanie. Jak mówi menadżer ds. produktów w firmie NXP Semiconductors, Karolien De Baere: „Standard Bluetooth 5.3 LE odменяia dotychczasowy sposób słuchania dźwięku. W tym celu korzysta z transceiverów o zaawansowanych możliwościach przetwarzania audio, które są w pełni energooszczędne i odznaczają się niskimi opóźnieniami. Jednym z nich jest, rzecz jasna, układ NXH3675 – funkcjonalne rozwiązanie, które obsługuje wiele strumieni audio naraz”.

<https://bit.ly/3GhvJrK>

Nowe diody LED firmy ams OSRAM z rodziny OSLO Black

Diody OSLO Black typu SFH 47267AS A01 i SFH 47278AS A0 emitują promieniowanie podczerwone o długości fali 940 nm. Przeznaczone są do różnych aplikacji w branży automotive, w tym: detekcji rozpraszania uwagi kierowcy, wykrywania



osób wewnątrz, a także określania zapiecia pasów bezpieczeństwa. Nowe diody spełniają normę AEC-Q102 i mogą współpracować z kamerami w pojazdach – mowa tu w szczególności o unikalnych rozwiązaniach w samochodach autonomicznych. Cechują się szczytowym prądem przewodzenia o wartości 5 A i prostopadłościenną wiązką promieniowania, dzięki czemu zapewniają równomierny rozkład mocy optycznej w każdym punkcie przestrzeni. W przeciwieństwie do klasycznych systemów monitorowania wnętrza pojazdów, diody: SFH 47267AS A01 i SFH 47278AS A01 nie wymagają użycia dodatkowych soczewek. Chodzi tu szczególnie o dopasowanie wiązek promieniowania podczerwonego do pól widzenia kamer pojazdowych, które uwieczniają obraz w proporcji 4:3. Dzięki temu dopasowaniu, systemy monitorowania wnętrza pojazdów mogą stać się prostsze i efektywniejsze. Przekłada się to na lepszą wydajność obrazowania i niezawodne działanie aplikacji w pojazdach. Jak dobitnie wyjaśnia starszy menadżer marketingowy w firmie ams OSRAM, Firat Sarialtun: „Diody LED o oznaczeniach: SFH 47267AS A01 i SFH 47278AS A01 są w pełni sprzyjające motoryzacyjnym zastosowaniom. Po raz kolejny wiedza z zakresu inżynierii optycznej czy know-how firmy ams OSRAM umożliwiły opracowanie użytecznych podzespołów”.

<https://bit.ly/3WtisSH>



Niebieskie diody laserowe firmy ams OSRAM dla rozwiązań Convergent Photonics

Convergent Photonics – włoski producent modułów laserowych stosuje niebieskie diody laserowe o długości fali 445 nm produkowane przez firmę ams OSRAM. Są to podzespoły o konstrukcji CoS (Chip-on-Submount) o małych rozmiarach diod, np. 1×1,3×0,2 mm dla diody PLPCOS 450D, czy 4×3×0,3 mm dla diody PLPCOS 450E. Każda ze wspomnianych diod odznacza się mocą wyjściową blisko 5 W i może działać w trudnych warunkach. Używane przez Convergent Photonics diody znalazły zastosowanie m.in. w wieloemiterowych modułach z mikrosoczewkami. Dzięki tym soczewkom, matryce diod laserowych mogą zostać dołączone do wspólnego światłowodu. Umożliwia to zasadniczą kontrolę nad wiązką lasera w modułach Convergent Photonics i pozwala na skalowanie mocy. Jest to rozwiązanie zwłaszcza dla spawarek przemysłowych lub aplikacji medycznych. W ostatnim przypadku diody laserowe z ams OSRAM stanowią kluczowe źródła wiązek tnących dla chirurgów. Obudowa tych diod jest w całości zoptymalizowana pod kątem termicznym. Jak w skrócie mówi dyrektor ds. diod laserowych w firmie Convergent Photonics, Roberto Paoletti: „Stosowane dotychczas obudowy diod laserowych

poważnie ograniczały zastosowania tych diod. Dzięki opracowaniu obudowy CoS, problem ten nie istnieje. Mamy teraz podzespoły, które mogą być stosowane w rozwiązaniach wysokich mocy. Odpowiada to oczekiwanym przez nas rozmiarom, wadze oraz charakterystyce mocy diod, w widzialnym i podczerwonym zakresie widma promieniowania elektromagnetycznego”.

<https://bit.ly/3I1apZ5>



Wielorzędowe złącza BTB od firmy Interplex

Aktualna oferta firmy Interplex uwzględnia wielorzędowe złącza BTB (Board-to-Board) o wyprowadzeniach miniPLX. Rozstaw tych wyprowadzeń to zaledwie 0,4 mm, a ich konstrukcja powoduje, że nie wymagają lutowania i mogą pracować w temperaturze od -40 do 150°C. Wszystkie złącza firmy Interplex są trwałe i funkcjonalne. Wyprowadzenia zostały wykonane z miedzi i charakteryzują się rezystancją na poziomie 1 mΩ. Każde z wyprowadzeń przewodzi prąd poniżej 3 A. Dzięki zastosowaniu technologii IndiCoat złącza są odporne na powstawanie węgla cynowych. Dostępne są warianty o wysokościach: 7...30 mm, przy 1...6 rzędach, po maksymalnie 30 wyprowadzeń. Te wytrzymałe produkty znoszą wysokie poziomy wilgotności (do 10% RH), wstrząsów (do 35 g) oraz wibracji (co najmniej 8 godzin pracy). Jak wyjaśnia menadżer ds. linii produktów w firmie Interplex, Ralph Semmeling: „Uwzględnienie wystarczającej liczby wyprowadzeń jest bardzo ważne z punktu widzenia klientów. Chcą oni mieć dostęp do gotowych produktów w dogodnych cenach. Nasze złącza BTB zapewniają to co najlepsze, gwarantując budowę rozwiązań przy niskich kosztach”.

<https://bit.ly/3Wq9t4A>



Nowe urządzenie zabezpieczające z serii Reyrolle firmy Siemens

Zadaniem urządzenia 7SR46 jest zapewnianie ochrony nadprądowej oraz przeciwzwarciowej dla stacji transformatorowych średniego napięcia. Wykrywa ono uszkodzenia kabli i błyskawicznie rozłącza zasilanie. Urządzenie 7SR46 zawiera solidną i wytrzymałą obudowę o wymiarach 10,4×18,5×7,9 cm. Z przodu znajduje się niewielki wyświetlacz LCD wraz z przyciskami, które służą do konfiguracji, a także przeglądania zapisów usterek urządzenia. Przewidziane zostały 4 diody LED wyrażające stany pracy oraz port USB, który pozwala na pobieranie danych i zasilanie. W tylnej części znajduje się port RS485, obok innych cyfrowych wyprowadzeń. Uwzględniono w nich wyjście impulsowe przeznaczone dla wyłączników w stacjach

transformatorowych. Urządzenie współpracuje z oprogramowaniem Reydisp z firmy Siemens. Dostępny jest klarowny schemat, który bardzo ułatwia obsługę. Urządzenie 7SR46 można z łatwością dołączać do źródła zasilania, ale zawiera także wewnętrzną baterię.

<https://bit.ly/3WoeNFM>

Dobra weryfikacja osłon i zderzaków w autach dzięki zrobotyzowanemu rozwiązaniu Rohde & Schwarz oraz Löhnert Elektronik

Konstrukcja pokazana na rysunku korzysta z urządzenia R&S QAR50-R i pozwala dogłębnie sprawdzać duże fragmenty osłon lub zderzaków aut. Precyzyjne pozycjonowanie anten pomiarowych



w R&S QAR50-R gwarantuje dokładne i stabilne testowanie, niezależnie od okoliczności i warunków. Do niedawna rewizja sporych i złożonych w kształcie podzespołów aut było skomplikowanym zadaniem. Aby odmienić ten stan rzeczy, firmy Rohde & Schwarz i Löhnert Elektronik zrobiły w krótkim czasie zrobotyzowane rozwiązanie, które wyznacza tłumienia, fazy i odbicia wiązek z pasma RF. Na tej podstawie stwierdzane są różne defekty osłon i zderzaków w autach. Dzięki elektronicznemu ogniskowaniu urządzenie spełnia najwyższe wymagania produkcji. Istnieje opcjonalny moduł, który pozwala stwierdzać zgodność otrzymywanych wyników z krajowymi i międzynarodowymi normami. Wszystko to przy zachowaniu kompaktowych wymiarów i małej wagi urządzenia R&S QAR50-R. Ze względu na swoją wyjątkową konstrukcję R&S QAR50-R sprawdza się tam, gdzie inne systemy testowe zawodzą.

<https://bit.ly/3Q78og1>

Niskonapięciowe złącza magnetyczne marki MULTICOMP PRO w ofercie firmy Farnell

Firma Farnell rozszerzyła ofertę produktów o magnetyczne złącza marki MULTICOMP PRO. Każde z tych złączy odpowiada za wytrzymałe oraz samoza-ciskowe połączenia. W razie potrzeby złącza można łatwo demontować. Dostępne na rynku niskonapięciowe złącza magnetyczne stają się cenionym rozwiązaniem. Nie inaczej jest ze złączami marki MULTICOMP



PRO, które zapobiegają uszkodzeniom kabli. Spełniają wysokie standardy jakości i mogą być stosowane m.in. w przemyśle, komputerach, a także w medycznych urządzeniach. Jednym ze wspomnianych złączy jest model MP002495, który umożliwia obrót złącza o 360° bez utraty połączenia. Model spełnia wymogi klasy IP67 i jest przeznaczony np. do ładowania baterii przy prądzie znamionowym 5 A oraz napięciu 30 V. Jak wyjaśnia starszy menadżer w firmie Farnell, Gareth James: „Jesteśmy wyłącznym dystrybutorem marki MULTICOMP PRO. Bardzo się cieszymy z faktu, że możemy zaoferować złącza magnetyczne, które są innowacyjne oraz zmieniają branżę złączy. Proste w użyciu złącza są świetnym rozwiązaniem szczególnie gdy potrzebne jest rozwiązanie, które może obracać się o 360°”.

<https://bit.ly/3PVXvNY>

Jakub Tyburski
jakub.tyburski@wat.edu.pl

powerMonitor

W praktyce każdego elektronika amatora czy profesjonalisty dość często zachodzi potrzeba monitorowania parametrów elektrycznych odbiornika mocy DC, czyli parametrów napięcia, prądu czy mocy pobieranej przez badane urządzenie. Właśnie do takich zastosowań powstał projekt o nazwie powerMonitor. Kluczowym problemem implementacyjnym był wybór odpowiedniego przetwornika, przy pomocy którego udałooby się zrealizować precyzyjny pomiar napięcia, prądu i ładunku. Przy realizacji jednego z wcześniejszych projektów poznałem dość dobrze specjalizowany przetwornik pomiarowy pod postacią układu INA226 i dlatego zdecydowałem się na jego użycie również w tym urządzeniu.

Termostat MIN/MAX

W znacznej liczbie termostatów można ustawić temperaturę średnią oraz dopuszczalną histerezę. Dla wielu użytkowników nie jest to zbyt intuicyjne rozwiązanie – lepiej ustawiać temperaturę załączenia i wyłączenia sterowanego urządzenia. W zaprezentowanym układzie zadajemy temperaturę załączenia przełącznika oraz temperaturę jego wyłączenia. Na podstawie naszych nastaw układ sam rozszyfruje, czy będzie miał do czynienia z obiektem chłodzonym (czyli takim, którego temperatura będzie samoczynnie wzrastała) czy też z ogrzewanym, który ma naturalną tendencję do stygnięcia. Wszystko odbywa się w bardzo prosty sposób, przy użyciu czterech przycisków.

Prosty generator sygnału PWM

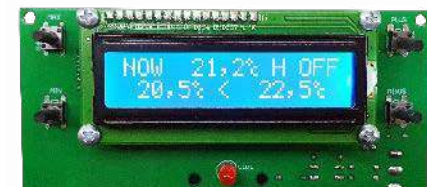
Sygnał PWM jest przydatny w wielu zastosowaniach, chociażby przy wykonywaniu testów układów wykonawczych średniej i dużej mocy. Zamiast wyciągać duży i nieporęczny generator sygnałowy, lepiej użyć niewielkiego urządzenia. Może ono służyć do płynnej regulacji jasności oświetlenia lub mocy dostarczanej do silnika prądu stałego.

Sterownik taśm LED dla RPi Zero

Urządzenie umożliwia rozszerzenie funkcjonalności RPi Zero, o sterowanie taśmami LED 12 V RGB z dodatkowymi diodami CCT (W/WW) w kolorze białym z odcieniem ciepłym i zimnym. Dodatkowe diody umożliwiają uzyskanie nie tylko szerszej palety kolorów, ale też światła białego o zmienianej płynnie temperaturze barwowej – co daje dodatkowe możliwości w aranżacji oświetlenia. Driver sterowany jest poprzez magistralę I²C, pozwala obciążyć wyjścia sumarycznym prądem do 10 A, pozwalając zasilać, w zależności od mocy, do kilkunastu metrów taśmy LED. Każdy kolor ma możliwość indywidualnej regulacji PWM w 8-bitowej skali.

a ponadto tematy wiodące EP 2/2023:

- Automatyczne systemy pomiarowe
- Elektronika w aplikacjach militarnych



Wykaz firm ogłaszających się w tym numerze „Elektroniki Praktycznej”

AKSOTRONIK.....	41
ARMEL	15
BORNICO.....	13
COMPUTER CONTROLS.....	7
ELMAX.....	17
FERYSTER.....	15
GAMMA	15
HAMMOND.....	11
PIEKARZ	15

Miesięcznik „Elektronika Praktyczna” (12 numerów w roku) jest wydawany przez AVT-Korporacja Sp. z o.o. we współpracy z wieloma redakcjami zagranicznymi.



Wydawnictwo:
AVT-Korporacja Sp. z o.o.
03-197 Warszawa, ul. Leszczynowa 11
tel. 22 257 84 99, e-mail: avt@avt.pl

Wydawca:
Wiesław Marciniak

Adres redakcji:
03-197 Warszawa, ul. Leszczynowa 11
e-mail: redakcja@ep.com.pl, www.ep.com.pl

Redaktor Naczelny:
Damian Sosnowski

**Redaktor Programowy,
Przewodniczący Rady Programowej:**
Piotr Zbysiński

Menedżer Magazynu:
Katarzyna Gugala

Szef Pracowni Konstrukcyjnej:
Jakub Sobański

Zespół marketingu i reklamy:

Katarzyna Gugala, tel. 22 257 84 64
Bożena Krzykawska, tel. 22 257 84 42
Grzegorz Krzykowski, tel. 22 257 84 60

Stali współpracownicy:

Lucjan Brynda, Nikodem Czechowski, Jarosław Doliński,
Andrzej Gawryluk, Krzysztof Górski, Tomasz Jabłoński,
Henryk Kowalski, Rafał Kozik, Michał Kurzela, Przemysław
Musz, Szymon Panecki, Sławomir Skrzyński, Ryszard
Szymaniak, Adam Tatuś, Jakub Tyburski, Robert Wołgajew

Uwaga!

Kontakt z wymienionymi osobami jest możliwy via e-mail,
według schematu: imię.nazwisko@ep.com.pl

DTP i okładka:

MAD Sp. z o.o.

Redakcja strony internetowej www.ep.com.pl

MAD Sp. z o.o.

Prenumerata w Wydawnictwie AVT

www.ulubionykiosk.pl lub tel. 22 257 84 22
(godz. 10:00–14:00)
e-mail: prenumerata@avt.pl

Prenumerata w RUCH S.A.

www.prenumerata.ruch.com.pl
lub tel. 801 800 803, 22 717 59 59
e-mail: prenumerata@ruch.com.pl



Wydawnictwo
AVT-Korporacja Sp. z o.o.
należy do Izby Wydawców Prasy

Copyright AVT-Korporacja Sp. z o.o.

03-197 Warszawa, ul. Leszczynowa 11

Projekty publikowane w „Elektronice Praktycznej” mogą być wykorzystywane wyłącznie do własnych potrzeb. Korzystanie z tych projektów do innych celów, zwłaszcza do działalności zarobkowej, wymaga zgody redakcji „Elektroniki Praktycznej”. Przedruk oraz umieszczanie na stronach internetowych całości lub fragmentów publikacji zamieszczanych w „Elektronice Praktycznej” jest dozwolone wyłącznie po uzyskaniu zgody redakcji. Redakcja nie odpowiada za treść reklam i ogłoszeń zamieszczanych w „Elektronice Praktycznej”.

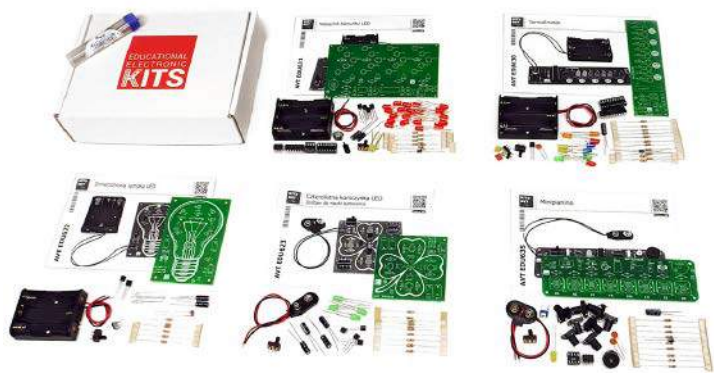


AVT EDU

Innowacyjna seria zestawów do nauki lutowania:

- większe pady
- duże odstępy między punktami lutowniczymi
- atrakcyjna grafika
- praktyczne zastosowanie

Zestawy dostępne pojedynczo i w pakietach.



Choinka LED RGB

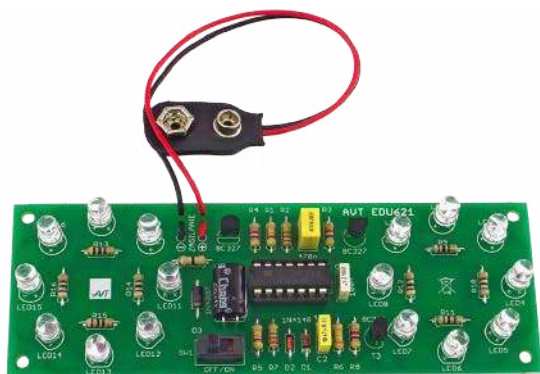
Świąteczny nastrój buduje mnóstwo pozornie drobnych szczegółów – migoczące lampki, klimatyczne melodie, zapach herbaty z pomarańczą i goździkami. Czego jeszcze brakuje? Może choinki! W te Święta spraw sobie prezent, który pozwoli Ci na rozwijanie swoich umiejętności w lutowaniu!

SPECYFIKACJA:

- źródło światła – płynnie zmieniające kolor diody LED RGB,
- bardzo prosty montaż,
- zasilanie: 3 VDC [2×AA] – zestaw nie zawiera baterii,
- wymiary płytki: 68×83 mm

kod handlowy: **AVTEDU640**

cena: **24zł**



Stroboskop policyjny LED

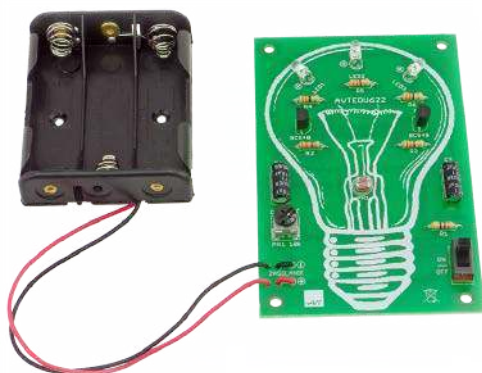
Efektom wizualnym generowanym przez moduł jest imitacja świateł pojazdu uprzywilejowanego.

SPECYFIKACJA:

- 2 pola świetlne z diodami LED (czerwone i niebieskie),
- 8 diod LED w każdym polu,
- wymiary płytki: 125×44 mm,
- napięcie zasilania: 9 VDC [6F22] – zestaw nie zawiera baterii

kod handlowy: **AVTEDU621**

cena: **24zł**



Zmierzchowa lampka LED

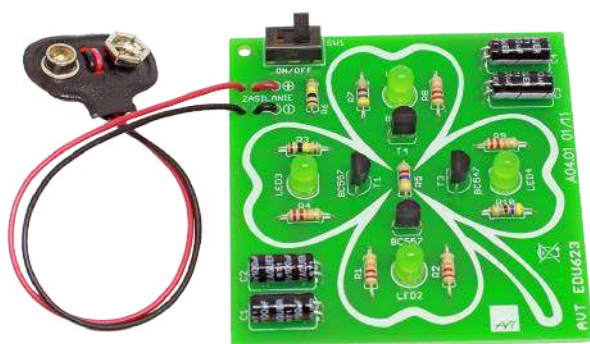
Praktyczna, wyróżniająca się designem lampka nocna z czujnikiem zmierzchu, która po zapadnięciu zmroku, rozbłyśnie jasnym światłem diod LED.

SPECYFIKACJA:

- źródło światła: 3 białe diody LED,
- płynna regulacja czułości zadziałania,
- napięcie zasilania: 5 VDC [3×AA] – zestaw nie zawiera baterii,
- wymiary płytki: 92×60 mm

kod handlowy: **AVTEDU622**

cena: **24zł**



Czterolistna koniczynka LED

Urokliwy i prosty w montażu gadżet będzie cieszył oko dzięki dwóm parom diod LED, które migają w zmiennym rytmie.

SPECYFIKACJA:

- 4 diody LED wysokiej jasności (pure green),
- automatyczna regulacja częstotliwości błysków,
- napięcie zasilania: 9 VDC [6F22] – zestaw nie zawiera baterii,
- mały pobór prądu (ok. 9 mA dla zasilania 9 V),
- wymiary płytki: 65×65 mm

kod handlowy: **AVTEDU623**

cena: **22zł**