

ELEKTRONIKA PRAKTYCZNA

EP.com.pl

● Międzynarodowy magazyn elektroników konstruktorów ● wrzesień ● 9/2023 ●

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

- Płytki rozwojowe do kursu FPGA Lattice • toneCtrl – regulator barwy dźwięku • Sansuix – lampowy wzmacniacz mocy 2×20 W • Solarna ładowarka akumulatorów z MPPT • Kinetyczne wzory na piasku, sterowane przez mikrokontroler ESP32 • Filtr zasilania do Raspberry Pi • Moduł wejść cyfrowych z optoizolacją i interfejsem USB-C • Regulowany zasilacz napięcia ujemnego • Standby killer

podzespoły, sprzęt, aplikacje

- Arduino UNO R4 Minima – wydajna i nowoczesna, ale... • Pozycjonowanie GNSS centymetrowej precyzji, dostępne również do standardowych zastosowań
- Narzędzia do szybkiego prototypowania z 32-bitowymi mikrokontrolerami Microchip • Nie tylko GPS. Przegląd alternatywnych systemów nawigacji satelitarnej

tutoriale

- Precyzja i dokładność w Systemach Globalnej Nawigacji Satelitarnej • Pomiar odległości i prędkości
- Poprawa jakości sygnału uwalnia prawdziwy potencjał transceiverów CAN-FD • Wzmacniaczu operacyjny, nie wzbudza się!

kursy

- Kurs FPGA Lattice. Statyczna analiza czasowa i maksymalna częstotliwość zegara

GPS I INNE SYSTEMY NAWIGACJI

TEMAT NUMERU



POMIAR ODLEGŁOŚCI I PRĘDKOŚCI



**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 826d83c0a0

Krótki zarys historii sklepu AVT

Parę dni temu, 28 sierpnia miało miejsce wydarzenie ważne dla Czytelników EP i wszystkich konstruktorów elektronik – nowa odsłona sklepu internetowego AVT, na nowej platformie, spełniającej najwyższe standardy technologii e-commerce dnia dzisiejszego. Nowa szata graficzna, duża szybkość działania, przejrzystość oferty i łatwość wyszukiwania produktu to nie wszystkie zalety, które zapewne docenią klienci nowego sklepu. To wydarzenie jest dobrą okazją do nostalgicznej retrospekcji ponad 30-letniej historii naszego sklepu. Zaczniemy od prehistorii, z której wykluła się firma AVT i jej sklep.

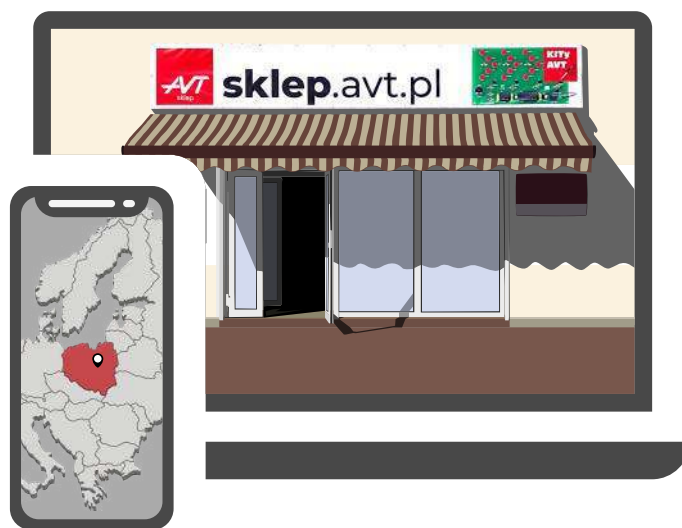
W roku 1984 znalazłem się w zespole 6 osób tworzących redakcję nowego tytułu „Audio – Video”, miesięcznika istniejącego do dziś. Chyba warto uświadomić młodym Czytelnikom, jak inny był tamten świat. Dla naszego zespołu jedyną motywacją do działania była potrzeba aktywności użytecznej społecznie, bez żadnych gratyfikacji. Nie istniały prywatne wydawnictwa. Państwo miało totalną kontrolę nad wszelkim słowem drukowanym. Wszystkie czasopisma techniczne należały do jednego państwowego wydawcy o nazwie Sigma. Cała prasa była drukowana w jednej państwowej drukarni Dom Słowa Polskiego i podlegała prewencyjnej kontroli przez cenzorów urzędu przy ul. Mysiej, a przydział papieru, którego w PRL zawsze brakowało, trzeba było „wychodzić” w gabinetach władz partyjnych. Było mocno pod górkę, a jednak energiczny lider naszego zespołu doktor Jerzy Auerbach dał radę i pojawił się nowy tytuł „Audio – Video”, który przy nakładzie 150 000 egzemplarzy sprzedawał się w 100% (w kioskach zniknął po 2–3 dniach sprzedaży). Gdy w podziale redakcyjnych obowiązków przypadła mi rubryka Hobby, w której publikowaliśmy bardzo ambitne projekty, zrozumiałem, jak bardzo hobbystom brakuje gotowych płytek drukowanych i wielu trudno dostępnych na rynku komponentów. Redakcja nie mogła pomóc Czytelnikom, ale ich potrzeby mogła spełniać inna firma.

Gdy tylko zaistniały ku temu możliwości ustrojowe, **w roku 1990 utworzyliśmy z żoną (Lela Marciniak – mgr inż. elektronik) firmę rodzinną oferującą kity** (zestawy płytek drukowanych i komponentów) do projektów publikowanych w „Audio – Video”. Pierwsze 4 kity, mierniki panelowe i multimetry 3,5 cyfry oparte na układach ICL 8006, ICL 8007, zrobiły furorę. Były to pierwsze produkty sklepu wysyłkowego AVT (wysyłkowego, ale nie internetowego, bo internet wtedy nie istniał). Tak to się zaczęło.

Handel wymagał skali, większych obrotów i jeszcze większych obrotów. Żeby kompletować kity, trzeba mieć magazyn z zapasem komponentów, które leżąc na półkach, zamrażają środki. Zatem zaczęliśmy oferować nie tylko kity, ale i komponenty kupione na zapas do kompletacji tych kitów. A dlaczego tylko te komponenty? Przecież można sprzedawać szerszą gamę komponentów. A dlaczego tylko komponenty? Przecież hobbysci potrzebują również akcesoriów, narzędzi, mierników itp. Musiał więc powstać **sklep stacjonarny**, który w miarę wzrostu zmieniał lokalizację (starsi Czytelnicy zapewne pamiętają nasz sklep przy ul. Granicznej).

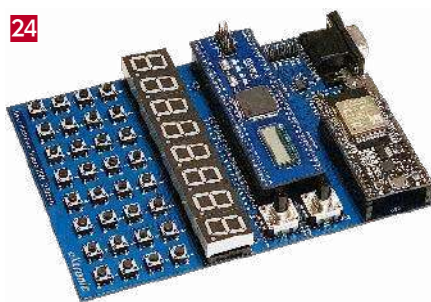
Chcieliśmy też zwiększyć tempo opracowywania i wdrażania kolejnych kitów. Kilkustronicowa rubryka w „Audio – Video” już nie wystarczała. **W styczniu 1993 roku wystartował miesięcznik „Elektronika Praktyczna”**, prowokując proces lawinowego wzrostu firmy AVT. Co miesiąc wdrażaliśmy 5 do 10 nowych kitów. Miesięcznik EP **zadziałał jak magnes** przyciągający do współpracy całą plejadę niezwykle uzdolnionych konstruktorów. Wymienię choćby kilku najwybitniejszych, w których szybko dostrzegłem potencjał na samodzielnych, charyzmatycznych Redaktorów Naczelnych: **Piotr Zbysiński** – wieloletni RN „Elektroniki Praktycznej”, **Piotr Górecki** – uwielbiany przez Czytelników RN „Elektroniki dla Wszystkich”, **Andrzej Kisiel** – RN kultowego magazynu „Audio”, **Andrzej Janeczek** – RN „Świata Radio”, guru polskich krótkofalowców, **Tomasz Wróblewski** – elektronik i muzyk, którego nieograniczona skala talentów wyniosła na pozycję RN magazynu „Estrada i Studio”. W ten sposób, w ciągu paru lat, AVT spontanicznie wyrosło na liczącego się wydawcę wielu tytułów.

Proces lawinowy potoczył się też w części handlowej. Sukces EP i rosnąca siła marki AVT przyciągały przedsiębiorcze osoby, chętne do uruchomienia sklepów filialnych AVT. Jak grzyby po deszczu zaczęły powstawać sklepy AVT w Olsztynie, Krakowie, drugi sklep w Warszawie (przejście podziemne przy gmachu GUS). Jednak zawróciliśmy z tej drogi. Ja odpowiadałem za rozwój firmy AVT i organizowanie sieci sklepów to nie moja bajka. W moim „pierwszym życiu” byłem naukowcem i nauczycielem akademickim, a to słabe kwalifikacje na organizatora rozległej sieci handlowej. Zresztą internet całkowicie zmienił sytuację. Dlatego skupiliśmy się na rozwoju sklepu internetowego z solidnym zapleczem magazynowym. Unikalność naszego asortymentu kitów i wszystkiego, czego potrzebuje konstruktor elektronik, nadaje sklepowi AVT szczególnie, niepowtarzalny charakter. Sklep to ludzie (załoga) plus technologia sprzedaży. Mamy wspaniałą, kompetentną załogę, ale w tej krótkiej historii wspomnę tylko o jednej osobie – przez 30 lat Ostoją naszego sklepu był Pan **Stanisław Dziwota** – tak, tak, ten skromny starszy Pan, zawsze życzliwy i pomocny. Mało kto się domyślał, że w „pierwszym życiu” był pułkownikiem w Wojskowej Akademii Technicznej, cenionym konstruktorem systemów laserowych. Pan Stanisław już nie pracuje, ale z pasją równą jego przykładowi prowadzi sklep świetny, stabilny zespół, który na pewno doloży wszelkich starań, żeby odnowiony sklep internetowy dobrze służył konstruktorom elektronikom.



Stanisław Dziwota

24



Nie przeocz

Konkurs	5
Nowe podzespoły	6
Dodaj do obserwowanych	12
Koktajl niusów	104
Arduino UNO R4 Minima – wydajna i nowoczesna, ale	18

Projekty

Płytki rozwojowe do kursu FPGA Lattice	24
toneCtrl (1) – regulator barwy dźwięku	30
Sansuix – lampowy wzmacniacz mocy 2×20 W (2)	35

Miniprojekty

Filtr zasilania do Raspberry Pi	40
Moduł wejść cyfrowych z optoizolacją i interfejsem USB-C	41
Regulowany zasilacz napięcia ujemnego	45
Standby killer	47

Temat numeru: GPS i inne systemy nawigacji

Precyzja i dokładność w Systemach Globalnej Nawigacji Satelitarnej	50
Nie tylko GPS. Przegląd alternatywnych systemów nawigacji satelitarnej	56

Prezentacje

Pozycjonowanie GNSS centymetrowej precyzji, dostępne również do standardowych zastosowań	54
Narzędzia do szybkiego prototypowania z 32-bitowymi mikrokontrolerami Microchip	60

Elektronika w praktyce

Pomiary odległości i prędkości	63
--------------------------------------	----

Notatnik konstruktora

Poprawa jakości sygnału uwalnia prawdziwy potencjał transceiverów CAN-FD	78
Wzmacniacz operacyjny, nie wzbudź się!	82

Projekty SOFT

Solarna ładowarka akumulatorów z MPPT (2)	84
Kinetyczne wzory na piasku, sterowane przez mikrokontroler ESP32	90

Kursy

Kurs FPGA Lattice (11).	
Statyczna analiza czasowa i maksymalna częstotliwość zegara	95

Prenumerata	2
Od wydawcy	3
Hity następnego numeru	107

30



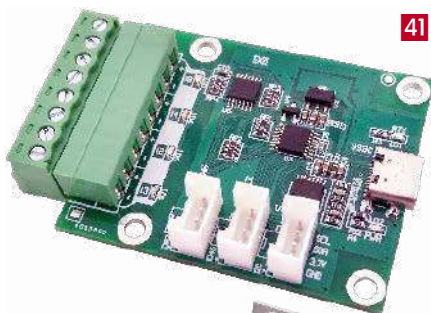
35



40



41



45





Wygraj programator/debugger Microchip ICD 5

MPLAB ICD 5 oferuje zaawansowane funkcje programowania i debugowania w docelowym obwodzie (In-Circuit), dla twórców projektów bazujących na mikrokontrolerach PIC, AVR i SAM oraz cyfrowych kontrolerach sygnału dsPIC (Digital Signal Controller). Jest to narzędzie nowej generacji, wyposażone w komunikację Fast Ethernet i zasilanie Power over Ethernet Plus (PoE+), dzięki czemu oferuje elastyczność i wygodę zdalnego programowania, jednocześnie izolując aplikację od warunków środowiskowych.

MPLAB ICD 5 (DV164055) oferuje różnorodne możliwości i funkcje, które przyspieszą rozwój każdego projektu i skrócą czas debugowania. Urządzenie jest obsługiwane za pomocą wydajnego i łatwego w użyciu graficznego interfejsu użytkownika zintegrowanego ze środowiskiem programistycznym (IDE) MPLAB X. Niezależnie od tego, czy jesteś doświadczonym programistą, czy dopiero zaczynasz, programator/debugger Microchip ICD 5 przyspieszy proces programowania i pomoże Ci przenieść projekty na wyższy poziom.

Kluczowe parametry i cechy MPLAB ICD 5:

- podłączenie do komputera PC za pomocą złącza USB typu C,
- komunikacja poprzez interfejs High Speed USB 2.0,

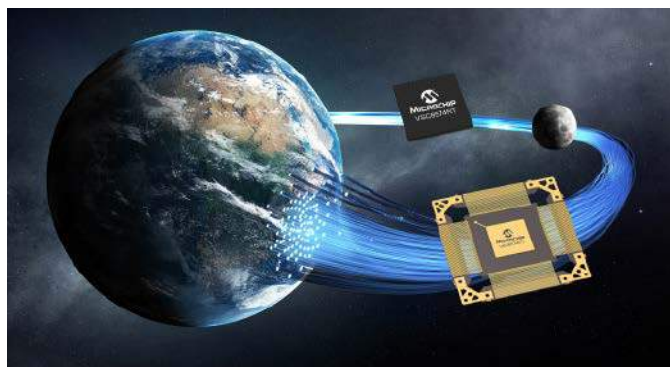
- komunikacja poprzez Ethernet do programowania i debugowania umożliwia zdalną obsługę,
- zasilanie Fast Ethernet z PoE+ zapewnia większą elastyczność i wygodę,
- umożliwia szybkie tworzenie oprogramowania w systemie CI/CD (continuous integration/continuous delivery),
- przechwytyje dane dotyczące zasilania, takie jak wartości prądu oraz napięcia,
- współpracuje z MPLAB Data Visualizer, który graficznie analizuje dane dotyczące zasilania,
- monitorowanie mocy pozwala zoptymalizować projekt poboru energii,
- obsługuje kilka różnych interfejsów debugowania i programowania,
- oferuje możliwość programowania JTAG, SWD, ICSP,
- pracuje z szerokim zakresem napięcia układu docelowego: od 1,2 V do 5,5 V.

Aby mieć szansę na wygraną programatora/debuggera Microchip ICD 5 lub aby +otrzymać kupon rabatowy 15% i bezpłatną wysyłkę, należy wypełnić formularz zgłoszeniowy na stronie: <https://page.microchip.com/E-Prak-ICD5.html>.

Szczegółowe informacje na temat Microchip ICD 5 można znaleźć na: <https://www.microchip.com/en-us/development-tool/dv164055#>.

nowe podzespóły

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



Transceiver Gigabit Ethernet PHY o zwiększonej odporności na promieniowanie jonizujące

Aby usprawnić wdrażanie sieci ethernetowych Klientom z branży lotniczej i obronnej, Microchip rozszerza ofertę kontrolerów Ethernet PHY o nowy układ o symbolu VSC8574RT, zapewniający obsługę standardów SGMII (Serial Gigabit Media Independent Interface) i QSGMII (Quad Serial Gigabit Media-Independent Interface), kompatybilny z kablami miedzianymi i optycznymi. Jest on wyposażony w poczwórny port do obsługi połączeń Ethernet 10, 100 i 1000BASE-T, zapewniający optymalną prędkość transmisji i zasięg w zależności od wymagań urządzenia. Obsługuje SyncE oraz protokół precyzyjnej synchronizacji czasu IEEE 1588v2 Precision Time Protocol (PTP).

VSC8574RT może znaleźć zastosowanie zarówno w satelitach nisko-orbitalnych, jak i urządzeniach pracujących w głębokiej przestrzeni kosmicznej. Zapewnia odporność na oddziaływanie pojedynczych cząstek jonizujących o energii do co najmniej 78 MeV.cm²/mg oraz całkowitą dawkę napromieniowania do 100 krad.

Microchip oferuje też płytkę ewaluacyjną VSC8574-EV, pozwalającą na testowanie kontrolera w różnych konfiguracjach pracy.

www.microchip.com



Monochromatyczne wyświetlacze OLED od firmy Winstar

W ofercie Unisystemu dostępne są różne warianty graficznych wyświetlaczy OLED, których producentem jest firma Winstar. Wyświetlacze OLED, dzięki wysokiemu kontrastowi i szerokim kątom obserwacji, zapewniają doskonałą czytelność prezentowanych na nich treści zarówno w świetle, jak i w mroku, z niemal każdej płaszczyzny. Co więcej, to matryce, które nie wymagają dodatkowego podświetlenia (każdy pojedynczy piksel zbudowany jest z diody OLED, która jest niezależnym źródłem światła), co czyni je rozwiązaniami energooszczędnymi (i chętnie stosowanymi w urządzeniach zasilanych bateryjnie).



Wśród nowości od Winstara jest dziesięć modeli monochromatycznych wyświetlaczy OLED: WEO006448B, WEA009616B, WEO128128G, WEA128128G, WEO128128H (w dwóch wariantach kolorystycznych), WEO012832P (w dwóch wariantach kolorystycznych) i WEO012832N (w dwóch wariantach kolorystycznych) o przekątnych od 0,66 do 2,23 cali. Są to komponenty wykonane w technologii COG oraz COG+PCB – z diodami OLED umieszczonymi na szkle ekranu. W zależności od modelu transfer danych odbywa się za pomocą co najmniej jednego z popularnych interfejsów stosowanych do obsługi wyświetlaczy OLED, tj.: 6800, 8080, SPI, I²C. Niemal wszystkie są przystosowane do pracy w zakresie temperatur od -40 do 80°C, co pozwala na ich implementację w aplikacjach indoorowych i outdoorowych, także w ekstremalnych warunkach, przy występowaniu skrajnie niskich i wysokich temperatur. W tabeli zaprezentowano kluczowe parametry omawianych modeli wyświetlaczy OLED:

Model	Przekątna	Rozdzielczość [px]	Struktura	Obszar aktywny [mm]	Wymiary zewnętrzne [mm]	Interfejs	Kolor treści	Zakres temperatur pracy [°C]
WEO006448B	0,66"	64×48	COG	13,42×10,06	18,46×18,10×1,26	6800, 8080, SPI, I ² C	biały	-30...70
WEA009616B	0,69"	96×16	COG + PCB	17,26×3,18	30,0×10,0×4,71	I ² C	biały	-30...70
WEO128128G	1,12"	128×128	COG	20,076×20,076	25,9×30,1×1,26	6800, 8080, SPI, I ² C	biały	-40...80
WEA128128G	1,12"	128×128	COG + PCB	20,076×20,076	26,4×38,6×5,21	SPI	biały	-40...80
WEO128128H	1,5"	128×128	COG	26,86×26,86	33,80×36,50×2,01	8080, SPI, I ² C	biały/żółty	-40...80
WEO012832P	1,71"	128×32	COG	42,22×10,54	50,50×15,75×2,01	SPI, I ² C	biały/żółty	-40...80
WEO012832N	2,23"	128×32	COG	55,016×13,096	62,0×24,0×2,17	6800, 8080, SPI, I ² C	biały/żółty	-40...80

unisystem.pl

IQD rozszerza ofertę oscylatorów przystosowanych do temperatury od -40 do +125°C

Komponenty do pracy w ekstremalnej temperaturze otoczenia są coraz bardziej poszukiwane i to nie tylko przez inżynierów,



projektujących sprzęt do ekstremalnych zastosowań, na przykład w odwiertach. Trend ten jest częściowo napędzany przez producentów układów scalonych, którzy często projektują jeden typ układu, nadającego się zarówno do zastosowań przemysłowych, jak i motoryzacyjnych, a także przez coraz szersze zastosowanie urządzeń elektronicznych pracujących w ekstremalnych środowiskach. Ponadto samonagrzewanie się elementów elektronicznych, montowanych w pobliżu oscylatorów na płycie drukowanej, może być motywacją do wyboru wariantu, który gwarantuje bezpieczną pracę i zachowanie parametrów katalogowych w temperaturze do +125°C.

Projektanci, używający oscylatorów o tej samej częstotliwości pracy w wielu projektach o różnych wymogach temperaturowych, mogą korzystać z jednego modelu o podwyższonej temperaturze pracy, bez konieczności magazynowania wielu typów komponentów.

Z wymienionych powodów firma IQD wprowadza na rynek oscylatory zamykane w popularnych obudowach ceramicznych SMD rozmiaru 3,2×2,5 mm (CFPX-180) i 2,0×1,6 mm (IQXC-42), mogące obecnie pracować w rozszerzonym zakresie temperatury roboczej od -40 do +125°C. CFPX-180 jest dostępny na zakres częstotliwości od 12,0 MHz, a IQXC-42 od 16,0 MHz w wariantach przystosowanych do pracy z różnymi pojemnościami obciążenia. Oba charakteryzują się stabilnością od ±30 ppm dla zakresu -40...+125°C, a ich rezystancja ESR jest nieco mniejsza niż w przypadku wariantów standardowych, dzięki czemu możliwa jest współpraca z najnowszymi mikrokontrolerami, np. stosowanymi w aplikacjach IoT.

www.iqdfrequencyproducts.com

Energooszczędne akcelerometry MEMS z zaawansowanymi funkcjami przetwarzania danych

STMicroelectronics wprowadza na rynek trzy nowe akcelerometry z zaawansowanymi układami przetwarzania danych, umożliwiające szybsze reagowanie na zdarzenia przy jednoczesnym obniżeniu poboru mocy. LIS2DUX12 i LIS2DUXS12 zostały zrealizowane w technologii



MEMS. Zawierają jednostkę uczenia maszynowego (MLC), automat skończony (FSM – *Finite State Machine*) i funkcję autokonfiguracji adaptacyjnej (ASC). Do tego model LIS2DUXS12 zawiera kanał pomiarowy Qvar. Trzeci z nowych układów, LIS2DU12, jest akcelerometrem o podstawowej funkcjonalności do mniej wymagających zastosowań.

Wszystkie trzy akcelerometry komunikują się przez interfejs I²C. Zawierają funkcje wykrywania zdarzeń oraz filtr antyaliasingowy, stosowany przy małych częstotliwościach próbkowania, pozwalając zwiększyć dokładność wykrywania gestów przy bardzo małym poborze mocy.

Zintegrowana w modelach LIS2DUX12 i LIS2DUXS12 jednostka uczenia maszynowego zwiększa niezawodność w zakresie wykrywania aktywności, a FSM poprawia rozpoznawanie ruchu. Razem zapewniają one autonomiczne przetwarzanie danych w czujniku, zmniejszając obciążenie współpracującego mikrokontrolera i zapewniając szybszą reakcję systemu. Ponadto, dzięki funkcji adaptacyjnej autokonfiguracji (ASC), niezależnie dostosowują parametry, takie jak zakres pomiarowy i częstotliwość, w celu optymalizacji poboru mocy.

Model LIS2DUXS12 zawiera również kanał pomiarowy Qvar do identyfikacji zmian w otaczającym środowisku elektrostatycznym, umożliwiając wykrywanie obecności i odległości. Pozwala on projektantom dodawać nowe funkcje, np. sterowania interfejsem użytkownika, wykrywania cieczy i wykrywania biometrycznego (przydatne np. w czujnikach tętna). W aplikacjach z interfejsem użytkownika, Qvar w połączeniu z sygnałem przyspieszenia może wyeliminować błędne wykrywanie zdarzeń dwu- i wielodotykowych.

Inteligentne akcelerometry zapewniają wykrywanie kontekstu w najnowocześniejszych urządzeniach przenośnych, głośnikach, słuchawkach TWS, smartfonach, aparatach słuchowych, kontrolerach do gier, inteligentnych zegarkach, robotach i aplikacjach IoT. Trzy nowe modele produkcji ST zawierają energooszczędną architekturę i filtr antyaliasingowy, poprawiający parametry systemu dzięki usunięciu niepożądanych szumów z sygnału użytecznego. Gotowe do użycia algorytmy MLC i FSM, dostępne w serwisie ST MEMS GitHub, ułatwiają m.in. rozpoznawanie złożonych gestów i śledzenie zasobów.

REKLAMA



1552 - Ręczne obudowy plastikowe

Dowiedz się więcej:
hammfg.com/1552

eusales@hammfg.com • + 44 1256 812812



LIS2DUX12 i LIS2DUXS12 są zamykane w 12-wyprowadzeniowych obudowach LGA o wymiarach 2×2×0,74 mm. Ich ceny hurtowe zaczynają się odpowiednio od 1,38 USD i 1,43 USD przy zamówieniach 1000 sztuk. Cena hurtowa LIS2DU12, zamykanego w tym samym typie obudowy, wynosi 1,20 USD.

www.st.com

Kontroler 50-watowego nadajnika ładowania bezprzewodowego zgodny z Qi v1.3.x EPP i BPP

WLC1150 to kontroler 50-watowego nadajnika do systemów ładowania bezprzewodowego, zgodny ze specyfikacjami WPC Qi v1.3.x EPP (extended power profile), BPP (basic power profile) i PPDE (proprietary power delivery extension).

Może współpracować z portem USB PD lub z zasilaczem sieciowym, bez konieczności stosowania zewnętrznego konwertera DC-DC. Pracą układu steruje 32-bitowy mikrokontroler ARM Cortex-M0 z wbudowaną pamięcią (128 kB Flash, 16 kB RAM + 32 kB ROM), przetwornikiem A/C, kontrolerem PWM i zestawem timerów.

WLC1150 charakteryzuje się dużą sprawnością energetyczną i małym poziomem generowanych zaburzeń EMI. Zawiera sterowniki bramek tranzystorów wyjściowych, kontroler DC-DC buck do sterowania zewnętrznym wentylatorem, funkcję wykrywania obiektów obcych (FOD) oraz zestaw zabezpieczeń (nadm napięciowe, nadprądowe i termiczne). Jest układem o dużym stopniu integracji, wymagającym niewielu komponentów współpracujących. Ze względu na dużą moc wyjściową może znaleźć wiele zastosowań, w tym również przemysłowych. Przykładem mogą być elektronarzędzia, stacje dokujące, ładowarki smartfonów z obsługą trybu Qi EPP, a także robotyka i drony.

Oprogramowanie układowe oferuje wiele konfigurowalnych opcji, umożliwiających konfigurowanie parametrów pracy, w tym funkcji wykrywania obiektów obcych (FOD), monitorowania, zabezpieczeń itp. Ułatwia to graficzny interfejs użytkownika, eliminujący potrzebę debugowania kodu i usprawniający proces konfiguracji.

Pozostałe cechy:

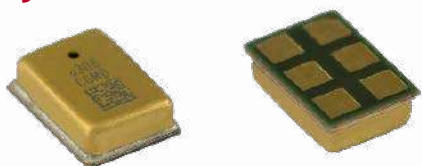
- zakres napięcia wejściowego od 4,5 do 24 V,
- porty komunikacyjne I²C i UART,
- zakres temperatury otoczenia od -40 do +105°C,
- obsługa protokołów QC 2.0/3.0, AFC i Apple Charging,
- kontrola za pomocą częstotliwości, napięcia lub współczynnika wypełnienia,
- obudowa VQFN-68 (8,0×8,0 mm).

www.infineon.com

Cyfrowy czujnik wibracji do poprawy jakości głosu w zaszumionym otoczeniu

Szerokopasmowy cyfrowy czujnik wibracji V2S200D firmy Knowles umożliwia poprawę jakości głosu m.in. w słuchawkach dousznych i innych urządzeniach audio, pracujących w zaszumionym otoczeniu, np. w zatłoczonych przestrzeniach publicznych oraz na zewnątrz pomieszczeń podczas podmuchów wiatru. Jego selektywna charakterystyka czułości pozwala uwypuklić zakres częstotliwości odpowiadający mowie, umożliwiając stłumienie szumu otoczenia nawet do 50 dB.

V2S200D jest zamykany w obudowie SMD o wymiarach 3,3×2,3×0,9 mm. Pracuje z napięciem zasilania od 1,65 do 3,3 V, pobierając w zależności od częstotliwości zegara od 480 do 850 µA



prądu. Charakteryzuje się dobrymi właściwościami szumowymi (SNR 63,5...64,5 dB, PSRR 85 dB, gęstość szumu 9 µg/√Hz), małą zawartością harmonicznych (THD=0,25% @ 1 g, 1 kHz) i punktem przesterowania akustycznego (AOP) powyżej 10 g. Jest łatwy w integracji dzięki zastosowanemu interfejsowi PDM.

www.knowles.com



650-woltowe tranzystory GaN HEMT do zastosowań w systemach zasilania

ROHM rozpoczyna masową produkcję dwóch 650-woltowych tranzystorów GaN HEMT o oznaczeniach GNP1070TC-Z i GNP1150TCA-Z, przeznaczonych do zastosowań w systemach zasilania. Zostały one zaprojektowane we współpracy z firmą Ancora Semiconductors (oddział Delta Electronics), specjalizującą się w produkcji półprzewodników na podłożach z azotku galu. Wyróżniają się bardzo małym iloczynem RDS(ON) × Ciss i RDS(ON) × Coss, co zapewnia dużą sprawność i małe straty mocy w systemach zasilania. Zawierają zabezpieczenie przed wyladowaniami ESD do 3,5 kV. Krótkie czasy przełączania tranzystorów GaN HEMT również przyczyniają się do zmniejszenia ich wymiarów i uproszczenia projektu układu zasilania.

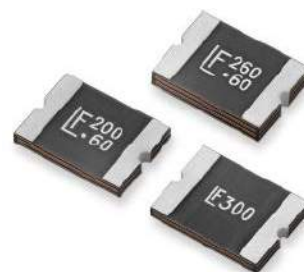
GNP1070TC-Z i GNP1150TCA-Z to tranzystory o napięciu przebicia 650 V, pracujące z dopuszczalnym ciągłym prądem drenu odpowiednio 20 A i 11 A. Różnią się też rezystancją RDS(on), wynoszącą odpowiednio 70 i 150 mΩ, ładunkiem bramki (5,2 nC i 2,7 nC) oraz pojemnością wejściową CISS (200 i 112 pF) i wyjściową COSS (50 i 19 pF). Oba są zamykane w obudowach DFN8080K o wymiarach 8,0×8,0×0,9 mm.

	V _{DS}	I _D @ T _c =25°C	R _{DS(on)}	Q _g	C _{iss}	C _{oss}
GNP1070TC-Z	650 V	20 A	70 mΩ	5,2 nC	200 pF	50 pF
GNP1150TCA-Z	650 V	11 A	150 mΩ	2,7 nC	112 pF	19 pF

www.rohm.com

Resetowalne bezpieczniki SMD do aplikacji wysokonapięciowych

Littelfuse dodaje do oferty resetowalnych bezpieczników PPTC (Polymeric Positive Temperature Coefficient) nowe wersje do układów wysokonapięciowych, zamykane w obudowach SMD rozmiaru 3425 (8,7×6,3 mm). Bezpieczniki serii 3425L, stanowiące rozszerzenie rodziny PolySwitch, realizują zabezpieczenie nadprądowe w układach o napięciu roboczym do 60 V. Mogą być stosowane w elektronarzędziach, motoryzacji, sprzęcie komputerowym, robotyce oraz w centrach danych i systemach telekomunikacyjnych, gdzie realizują m.in. zabezpieczenie portów IEEE



1394 i Ethernet IEEE 802,3 af. Nadają się do montażu przy użyciu standardowych automatów pick-and-place. Są przystosowane do pracy w temperaturze otoczenia od -40 do $+85^{\circ}\text{C}$. Zapewniają odporność na udary termiczne i rozpuszczalniki, zgodnie z wymogami normy MIL-STD-202 oraz na wibracje, zgodnie z MIL-STD-883C.

	I_{hold}	I_{TRIP}	Napięcie robocze	I_{maks}^*	R_{min}	R_{maks}	Maks. czas zadziałania (dla 8 A)	P_d
3425L200/60	2 A	4 A	60 V	20 A	0,04 Ω	0,2 Ω	10 s	2,5 W
3425L260/60	2,6 A	5,2 A	60 V	20 A	0,02 Ω	0,12 Ω	10 s	2,5 W
3425L300/36	3 A	6 A	36 V	20 A	0,01 Ω	0,06 Ω	20 s	2,5 W

www.littelfuse.com



Sterowniki bramek do modułów IGBT/SiC z funkcją odczytu temperatury z czujnika NTC

Firma Power Integrations opracowała nowe sterowniki bramek SCALE-iFlex LT NTC do modułów mocy bazujących na tranzystorach IGBT/SiC. Są one przystosowane do współpracy z popularnymi modułami IGBT rozmiaru 140×100 mm, np. LV100 produkcji Mitsubishi i XHP 2 produkcji Infineon, jak również z wariantami SiC o napięciu blokowania do 3300 V. Umożliwiają odczyt temperatury modułu z wbudowanego termistora NTC za pośrednictwem izolowanego wyjścia PWM. Pozwala to na precyzyjne zarządzanie chłodzeniem konwerterów i jest szczególnie ważne w systemach z wieloma modułami połączonymi równolegle, zapewniając właściwe współdzielenie prądu i wydłużając czas bezawaryjnej pracy.

Sterowniki SCALE-iFlex LT, bazujące na technologii SCALE-2 opracowanej przez firmę Power Integrations, poprawiają dokładność podziału prądu, a tym samym zwiększają obciążalność prądową modułów połączonych równolegle o 20%. Zawierają aktywne zabezpieczenie przepięciowe AAC (Advanced Active Clamping). Są w pełni zgodne ze specyfikacjami IEC 61000-4-x (EMI), IEC-60068-2-x (narażenia środowiskowe) i IEC-60068-2-x (narażenia mechaniczne). Ponadto przeszły testy niskonapięciowe, wysokonapięciowe i termiczne, co pozwala na skrócenie czasu projektowania. Opcjonalnie mogą być dostarczane z pokryciem chroniącym przed kondensacją wilgoci.

www.power.com



Sterowniki linii USB 3,2 do transmisji sygnałów na odległość do 15 m

EQCO510 i EQCO5X31 to 2-kanalowe sterowniki linii (Reclocker/Redriver), umożliwiające dwukierunkową transmisję sygnałów na odległość do 15 m w oparciu o protokół USB 3,2 Gen 1 SuperSpeed

w aplikacjach odpowiednio przemysłowych i motoryzacyjnych. Oba układy mogą pracować z maksymalną szybkością transmisji 5 Gbps. Funkcja reclocking zapewnia regenerację sygnału zegarowego (CDR – Clock-Data Recovery), eliminując akumulowanie się błędów jitteru, a funkcja redriving przywraca poziom i kształt sygnału kierowanego do następnego segmentu, np. kabla lub ścieżki na płycie drukowanej, kompensując degradację sygnału spowodowaną tłumieniem kabli. Kompensacja działa w zakresie od 0 do 24 dB z krokiem 1 dB. Dodatkowo oba układy oferują funkcję MarginLink, umożliwiającą ocenę integralności całej ścieżki sygnałowej w czasie pracy.

EQCO510 i EQCO5X31 obsługują ekranowaną skrętkę dwużyłową i kable koncentryczne. Zawierają układ korekcji CDR, niewymagający współpracy z zewnętrznym rezonatorem, co zmniejsza liczbę komponentów współpracujących i wymaganą powierzchnię płytki drukowanej. Są zamykane w 20-wyprowadzeniowych obudowach QFN. Wersja motoryzacyjna EQCO510 uzyskała kwalifikację AEC-Q100 Grade 2 i może pracować w zakresie temperatury otoczenia od -40 do $+105^{\circ}\text{C}$.

Ceny hurtowe EQCO510 i EQCO5X31 wynoszą odpowiednio 4,82 USD i 4,38 USD przy zamówieniach 1000 sztuk. W ofercie Microchipa są też zestawy ewaluacyjne ekstendera USB-C (EVB-EQCO5X31) i repeatera USB-C (EVB-EQCO5X31).

www.microchip.com



Najmniejsza na rynku pamięć Flash SPI NOR 128 Mb w obudowie o wymiarach $3 \times 3 \times 0,4$ mm

GigaDevice prezentuje najmniejszą na rynku pamięć SPI NOR Flash o pojemności 128 Mb, zamykaną w obudowie USON8 o wymiarach $3 \times 3 \times 0,4$ mm. Dzięki bardzo małej grubości pamięć GD25LE128EXH nadaje się idealnie do przechowywania kodu w urządzeniach IoT, przenośnych akcesoriach medycznych i fitness oraz wszelkiego typu aplikacjach o małych gabarytach, wymagających dużej funkcjonalności i małego poboru mocy.

GD25LE128EXH pracuje z maksymalną częstotliwością taktowania 133 MHz i zapewnia przepustowość 532 Mbps. Pobiera zaledwie

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



6 mA prądu w trybie odczytu, co zmniejsza pobór mocy o 45% w porównaniu z wcześniejszymi wersjami. Obudowa FO-USON8 zajmuje mniejszą o 70% powierzchnię i charakteryzuje się dwukrotnie mniejszą grubością od konwencjonalnych obudów WSON8 o wymiarach 6×5×0,8 mm. Nowa pamięć jest kompatybilna pod względem rozkładu wyprowadzeń z wcześniejszymi odpowiednikami 128-megabitowymi, zamykanymi w obudowach USON8 o wymiarach 4×3×0,6 mm.

www.gigadevice.com



Resetowalne bezpieczniki termiczne TCO z kwalifikacją AEC-Q200

Firma Bourns wprowadza na rynek dwie nowe serie resetowalnych bezpieczników termicznych TCO (*Thermal Cutoff*), mogących znaleźć zastosowanie w układach zabezpieczających grzejników, silników, pakietów akumulatorowych oraz akumulatorów litowo-jonowych w urządzeniach przenośnych. Są to pierwsze tego typu elementy z kwalifikacją AEC-Q200, spełniające wymogi niezawodnościowe instalacji samochodowych. Zawierają bimetaliczny mechanizm o zwiększonej odporności na korozję w porównaniu ze standardowymi bezpiecznikami TCO.

Bezpieczniki serii AD i SD zawierają wyprowadzenia do odpowiednio zgrzewania i lutowania. Ich temperatura progowa jest programowana fabrycznie w zakresie od +55 do +150°C z krokiem co 5°C. W zależności od wersji maksymalne napięcie pracy wynosi 14 V lub 28 V, a maksymalny prąd rozłączany to 8 A lub 35 A. Rezystancja przewodzenia nie przekracza 4 mΩ.

www.bourns.com



Miniaturowe głośniki wodoszczelne serii CMS o poziomie ciśnienia akustycznego do 108 dB

Oddział produktów audio firmy CUI Devices ogłasza wprowadzenie do sprzedaży 24 nowych głośników o stopniu ochrony IPX5, IPX7 i IPX8, zapewniających dużą odporność na działanie cieczy i wilgoci w trudnych warunkach pracy. Wodoodporne głośniki CMS charakteryzują się poziomem ciśnienia akustycznego od 85 do 108 dB. Są produkowane w okrągłych i prostokątnych obudowach o powierzchni od 15×8 mm do 57×57 mm i grubości od 2,5 mm. Mogą zawierać kontakty lutowane lub sprężynowe. Ich częstotliwość rezonansowa wynosi od 180 do 1200 Hz, moc znamionowa 0,8...3 W, a impedancja 4 lub 8 Ω.

Ceny głośników serii CMS zaczynają się od 1,16 USD przy zamówieniach 100 sztuk.

www.cuidevices.com

Moduły Bluetooth dużego zasięgu, bazujące na dwurdzeniowych mikrokontrolerach ARM Cortex M33

Laird Connectivity prezentuje nową serię modułów komunikacyjnych Bluetooth BL5340PA (*Power Amplified*), będących obecnie najbardziej zaawansowanymi i bezpiecznymi modułami tej klasy, wyposażonymi w dwurdzeniowe mikrokontrolery ARM Cortex M33. Są to moduły oparte na układach nRF5340 (SoC) i nRF21540 (front end) produkcji Nordic



Semiconductor, przeznaczone do aplikacji wymagających dużego zasięgu i dużej mocy obliczeniowej. Dodanie nRF21540 poprawia budżet łącza i niezawodność transmisji oraz znacznie zwiększa zasięg sieci bezprzewodowej w porównaniu z samym układem SoC nRF5340. Połączenie nRF5340 i nRF21540 w certyfikowanym module zapewnia szeroki zakres zastosowań, obejmujący m.in. przemysłowe systemy konserwacji predykcyjnej i aplikacje LE Audio dalekiego zasięgu.

Dwurdzeniowy mikrokontroler ARM Cortex M33 umożliwia programistom zastosowanie rdzenia o małym poborze mocy do obsługi łączności bezprzewodowej oraz rdzenia o dużej mocy obliczeniowej do aplikacji końcowej. To dodatkowo rozszerza zakres zastosowań na aplikacje Bluetooth Low Energy (LE), 802.15.4 (Thread/ZigBee) i NFC. Dodatkowo, interfejsy API CryptoCell-312 oferują funkcje m.in. root-of-trust i bezpiecznego przechowywania kluczy.

BL5340PA udostępnia wszystkie znaczące funkcje i możliwości sprzętowe nRF5340 i nRF21540, w tym dostęp przez USB, dużą moc nadajnika i szeroki zakres temperatury pracy. Uzyskane certyfikaty FCC, ISM, RCM i Bluetooth SIG eliminują konieczność przeprowadzania kolejnych testów, skracając czas wprowadzania produktów na rynek.

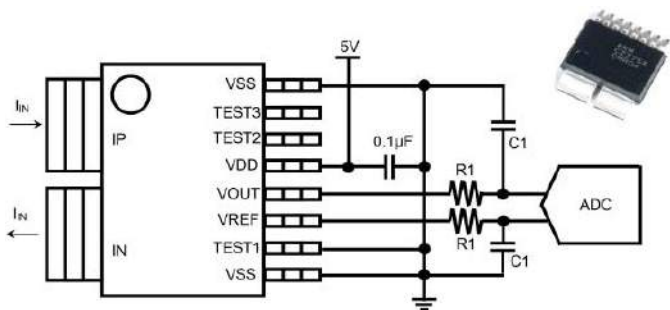
Ważniejsze cechy:

- chipset: nRF5340 + nRF21540,
- moc nadajnika: +18,5 dBm,
- antena: wbudowana + złącze MHF4 do anteny zewnętrznej,
- mikrokontroler: dwurdzeniowy ARM Cortex M33,
- obsługiwane standardy: Bluetooth 5.2 LE, 802.15.4 (Thread/ZigBee), NFC,
- obsługiwane funkcje: Bluetooth LE (Peripheral/Central), LE Coded (long range), AoA/AoD, LE Audio/Isochronous Channels, Mesh,
- bezpieczeństwo: ARM TrustZone, Root of Trust, ARM CryptoCell-312 & KMU, Access Control Lists, System Protection Unit, Encrypted QSPI,
- aktualizacja oprogramowania: Firmware Over the Air (FOTA) poprzez MCUboot i Zephyr,
- zakres temperatury pracy: od -40 do +105°C,
- wymiary: 21×10×2,5 mm.

www.lairdconnect.com

Miniaturowe czujniki prądu SMD o zakresie pomiarowym 100 A rms

Firma Asahi Kasei Microdevices opracowała serię miniaturowych czujników natężenia prądu CZ375, zamykanymi w obudowach SMD, oferujących zakres pomiarowy do 100 A rms (±225 A w szczycie). Są to czujniki bezrdzeniowe o małej emisji ciepła, mogące znaleźć zastosowanie m.in. w stacjach szybkiego ładowania, klimatyzatorach i zasilaczach UPS, spełniające wymogi normy UL61800-5-1.



Charakteryzują się małą rezystancją wewnętrzną, wynoszącą jedynie 0,27 mΩ oraz identycznymi wymiarami (12,7×10,9×2,25 mm), jak czujniki poprzedniej serii CZ370 Series o prądzie szczytowym ±180 A. Ich dokładność pomiaru wynosi ±0,4% FS w zakresie temperatury otoczenia od 0 do +90°C.

Pozostałe parametry:

- napięcie robocze: 550 V rms,
- czas odpowiedzi: 0,5 μs,
- droga upływu/odstęp izolacyjny: min. 8 mm,
- izolacja: 4,3 kV rms (50 Hz, 60 s),
- napięcie zasilania: 5 V,
- zakres temperatury pracy: -40...+105°C.

www.akm.com

Konwertery DC-DC quarter brick o mocy ciągłej do 1600 W i szczytowej do 2320 W

BMR351 to nieizolowany konwerter DC-DC formatu quarter brick (58,4×36,8×12,5 mm) o napięciu wejściowym od 40 do 60 V, charakteryzujący się maksymalną wyjściową mocą ciągłą 1600 W i szczytową do 2320 W. Jego maksymalny prąd wyjściowy



wynosi 136 A przy pracy ciągłej i może wzrosnąć do 200 A przez maksymalnie 500 ms.

Konwerter osiąga sprawność szczytową nawet do ponad 98%. Typowa wartość to 97,8% przy napięciu wejściowym 54 V i połowie obciążenia. Istnieje możliwość łączenia kilku jednostek w aplikacjach większej mocy. Wbudowany interfejs PMBus służy do konfigurowania, monitorowania i kontroli parametrów pracy, w tym umożliwia dostrajanie napięcia wyjściowego w zakresie od 8 do 13,2 V. Do wyposażenia konwertera należą też linie Remote Sense, Power Good i Remote Control. Innowacyjny, wewnętrzny rejestrator danych ułatwia diagnostykę w przypadku wystąpienia błędów.

BMR351 charakteryzuje się górną dopuszczalną temperaturą pracy wynoszącą +125°C. Jest chłodzony przez przewodzenie. Zawiera zabezpieczenie termiczne i wyjście sygnalizujące jego zadziałanie, a także zabezpieczenie podnapięciowe, nadnapięciowe, nadprądowe i zwarciove.

www.flexpowermodules.com

REKLAMA

3DEXPERIENCE vs SOLIDWORKS

COMPUTER CONTROLS

Masz starszą wersję SOLIDWORKS,
a chcesz pracować wydajniej?

Przejdź na zawsze **aktualną**
platformę 3DEXPERIENCE
z rabatem do **50%**.

Autoryzowany dystrybutor 3DEXPERIENCE

SKONTAKTUJ SIĘ NAMI TERAZ!



Altium

arm KEIL

SOLIDWORKS

KEYSIGHT TECHNOLOGIES

SILICON LABS

e-Peas

miromico

SILERGY

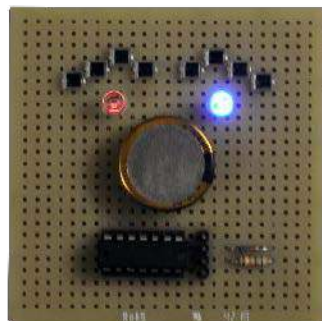
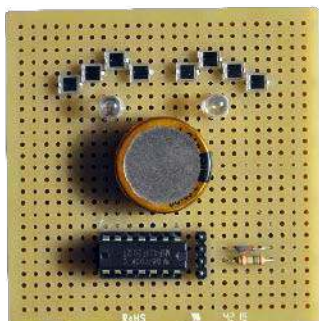
DISPLAY LOGIC

Computer Controls Sp. z o.o. Bielsko-Biała, ul. Budowlanych 1 +48 (33) 485 94 90

info@ccontrols.pl
www.ccontrols.pl

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.



Greenhouse Blinky – wiecznie migająca dioda LED

Opisywany projekt, zawierający osiem fotodiod BPW34 w połączeniu z mikrokontrolerem MSP430F2012, prezentuje innowacyjne podejście do obsługi klasycznego migacza działającego przez cały dzień i noc. Zastosowanie fotodiod BPW34 pozwala na ładowanie kondensatora odpowiedzialnego za zasilanie mikrokontrolera nawet przy słabym oświetleniu. Dzięki połączeniu szeregowemu fotodiod uzyskuje się odpowiednie napięcie do ładowania kondensatora. Sam kondensator został wybrany pod kątem jego niskiego prądu samorozładowania, co przyczynia się do maksymalizacji sprawności systemu. W projekcie zastosowano również diodę Schottky'ego (BAT85), w celu optymalizacji ładowania przy słabym oświetleniu i odciążenia diod zabezpieczających układ przy oświetleniu silnym.

Oprogramowanie układowe dla MSP430 napisane zostało w języku Forth. Jest kompilowane krzyżowo przy użyciu najnowszej wersji kompilatora Mccrisp-Across. W projekcie zaimplementowano również zaawansowane funkcje, takie jak pomiar jasności oświetlenia za pomocą czerwonej diody LED. Wewnętrzny przetwornik analogowo-cyfrowy w mikrokontrolerze umożliwia również pomiar napięcia kondensatora oraz napięcia łańcucha fotodiod w celu określenia, czy kondensator jest aktualnie ładowany.

Projekt „Greenhouse Blinky” udowadnia, że innowacyjne podejście do obsługi migacza bazującego na zaawansowanych algorytmach może przynieść interesujące rezultaty. Zastosowanie fotodiod BPW34 i mikrokontrolera MSP430F2012 pozwala na wydajne ładowanie kondensatora nawet w warunkach słabego oświetlenia.

<https://hackaday.io/project/190111-greenhouse-blinky>

Płytką dla koprocatora retro NS320xx w formie karty ISA

Opisywana konstrukcja to akcelerator korzystający z procesorów z rodziny NS32000 – układów 32032 lub 32016. Układy te zostały opracowane przez firmę National Semiconductors i były pierwszą linią procesorów, które używały w pełni 32-bitowej architektury. Wprowadzono je na rynek w 1982 roku, z kolejnymi wersjami przez całe lata 80. XX wieku.

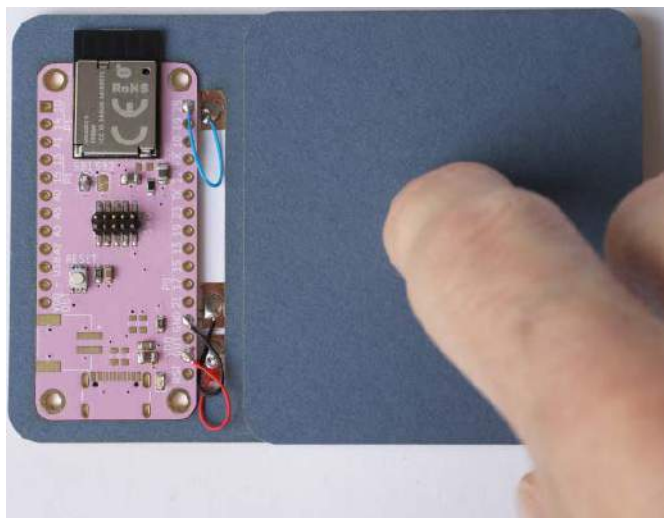
W sierpniu 1985 roku w czasopiśmie „Byte Magazine” opisano konstrukcję modułu koprocatora Definicon DSI-32. Moduł ten



zaprojektowany został jako akcelerator numeryczny dla komputerów PC XT, dlatego też zawierał formę karty ISA, którą można było umieścić w komputerze. Omawiany akcelerator obsługuje układy 32032 i 32016. Wyposażony jest w 256 000 bajtów pamięci DRAM. System ten ma w pełni 32-bitową architekturę. 32-bitowa magistrala danych i dwuportowa pamięć DRAM sprawiły, że obwód DSI-32 był skomplikowany, a płytka drukowana droga (miała cztery warstwy).

Autor projektu zeskanował artykuły z „Byte Magazine” i udostępnił je w artykule. Dalszym celem projektu jest inżynieria wsteczna tego ciekawego koprocatora.

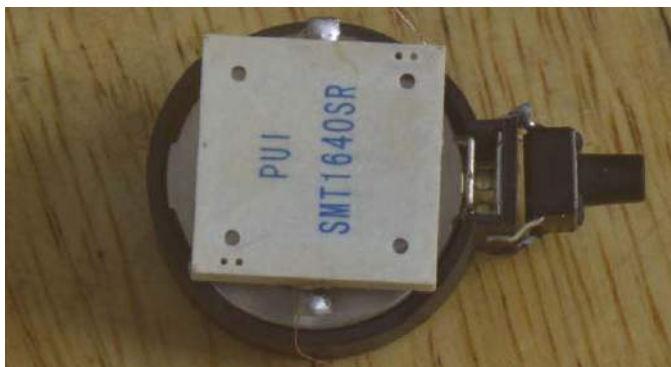
<https://tiny.pl/ccrv4>



Przycisk do kontroli sieci Thread przez MQTT

Zaprezentowany system to minimalistyczny projekt przycisku dla systemów Internetu Rzeczy (IoT), który ma pełnić funkcję prostego cyfrowego wejścia. Przycisk ten to lekka i energooszczędna konstrukcja, dzięki czemu można go zasilić z baterii i umieścić w wygodnym miejscu. Po naciśnięciu przycisku system łączy się z siecią Thread i opublikuje zmianę statusu na serwerze MQTT. Dalsze akcje naszego systemu automatyki domowej zależą już od konkretnej konfiguracji. Typowo wykorzystać można ten przycisk we współpracy np. z aplikacją Home Assistant, gdzie informacja o naciśnięciu przycisku może być użyta do uruchomienia wybranej automatyzacji.

<https://hackaday.io/project/190131-mqtt-thread-button-rev2>



Pozytywka na ATtiny10

Na pewno większość z nas kojarzy kartki z pozytywką. W środku takiej kartki znajduje się zalany w żywicy chip (na ogół w postaci gołej struktury krzemowej – stąd żywica, która go osłania). Układ ten jest przeznaczony do jednej tylko rzeczy – odtwarzania prostej melodyjki. Autor zaprezentowanej konstrukcji zapragnął zbudować kopię tego rozwiązania, ale na bazie taniego mikrokontrolera. „Prawdopodobnie nikt nie lubi tych melodyjek, ponieważ są irytujące, ale odłóżmy ten problem na bok i zadajmy inne pytanie – co by było, gdybyśmy mogli użyć nowoczesnych mikrokontrolerów i ich trybów uśpienia, aby pozytywki te działały naprawdę energooszczędnie i mogły grać przez imponująco długi czas?”, pisze autor we wstępie do projektu. Być może sprawi to tylko, że pozytywki te będą irytujące dłużej, jednak sam problem technologiczny jest bardzo ciekawy – jak zastosować tani, kosztujący zaledwie 36 centów mikrokontroler (ATtiny10) i bardzo niewiele zewnętrznych komponentów do konstrukcji prostego, ekonomicznego urządzenia o niskim poborze prądu.

Do projektu został dołączony plik z kodem assemblera (wraz z komentarzami). Z uwagi na dużą liczbę studentów zmagających się z tym językiem autor dodał obszerne komentarze, które mają ułatwić zrozumienie, zwłaszcza dla osób dopiero zaczynających przygodę z tym językiem programowania. Generalnie techniki, które zastosował autor, powinny działać z większością mikrokontrolerów AVR. Jednak szczegóły, takie jak adresy rejestrów, mogą się różnić w zależności od konkretnego modelu, więc warto sprawdzić je w karcie katalogowej odpowiedniego układu.

Projekt płytki był bardzo prosty, dlatego do jego narysowania zastosowano Inkscape, zamiast odpowiedniego tonera, a następnie wyprodukowano płytki za pomocą znanej metody wykorzystującej drukarkę laserową i żelazko do transferu tonera na PCB.

Pobór energii był jednym z głównych aspektów projektu. Podczas głębokiego uśpienia, kiedy obwód oczekuje na naciśnięcie przycisku przez użytkownika, zużycie wynosiło około 100 nanoamperów, co jest znikomą wartością w porównaniu do samorozładowania baterii CR2032. Przy takim poborze energii, nasz układ mógłby działać nawet przez około 240 lat w trybie oczekiwania na naciśnięcie przycisku.

W czasie odtwarzania utworu mikrokontroler AVR zużywał około 350 mikroamperów przy napięciu 3 V i częstotliwości taktowania 1 MHz. Należy jednak zaznaczyć, że kod wymagał pełnej aktywności mikrokontrolera tylko przez krótki okres. Przeważnie układ pozostawał w trybie uśpienia w bezczynności, gdzie włączony był jedynie zegar (watchdog), co powinno zużywać około 50 mikroamperów (ok. 35 μ A dla trybu bezczynności, ok. 5 μ A dla timera watchdog i 10 μ A dla timera/licznika generujących dźwięk). Dodatkowo generowanie przebiegów wyjściowych odbywało się w sprzętowych urządzeniach peryferyjnych, co pozwoliło na wyłączenie głównego zegara mikrokontrolera. Mimo że bezpośrednio nie zmierzono dokładnie poboru prądu, to wydaje się, że zaprezentowane dane są poprawne – wynikają z wartości zawartych w karcie katalogowej i wcześniejszych projektów z tym samym układem mikrokontrolera, jakie zrealizował autor.

W ramach oszczędzania energii wyłączany jest również cały osprzęt przetwornika ADC w układzie oraz wewnętrzne rezystory

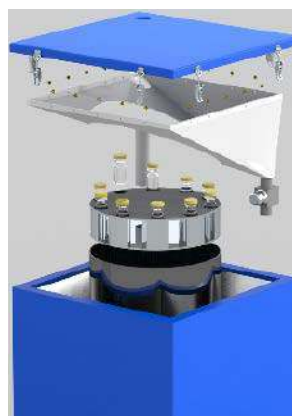
podciągające. Pozostały pobór mocy całkowicie zdominowany jest przez element piezoelektryczny. Zastosowany element okazał się bardzo cichy, dlatego autor projektu zamierza wymienić go na brzęczyk piezoelektryczny z komputera PC. Standardowe piezoelektryczne elementy zużywają maksymalnie do 10 mA, a piny ATtiny10 są w stanie podać tę wartość bezpośrednio, co pozwala zaoszczędzić na konieczności stosowania dodatkowego elementu, np. tranzystora MOSFET.

Dzięki tym optymalizacjom udało się zrównać zużycie energii układu ze zużyciem elementu piezoelektrycznego. Ostatecznie, przy użyciu baterii CR2032, można oczekiwać co najmniej 21 godzin odtwarzania i w zasadzie dowolnie długi czas czuwania, co stanowi znaczące osiągnięcie projektu.

<https://hackaday.io/project/190134-attiny10-chiptunes>

VacSafe – termos z funkcją reakcji endotermicznej do chłodzenia leków

Vacsafe to zaawansowany pojemnik do transportu szczepionek z aktywnym chłodzeniem, stworzony z myślą o dostarczaniu szczepionek na duże odległości, nawet do odległych obszarów. W celu zapewnienia odpowiedniej temperatury i bezpiecznego przechowywania zastosowano innowacyjny mechanizm chłodzenia bazujący na reakcji endotermicznej z użyciem azotanu amonu. Konstrukcja wewnętrzna pojemnika bazuje na twardym plastiku o grubości 1 cm, który został pokryty cienką warstwą miedzi. Wewnątrz tej warstwy znajduje się cienka warstwa aluminium, która spełnia dwie ważne funkcje: zapobiega rozpraszaniu ciepła przez promieniowanie oraz kontroluje emitowanie gazów z twardego plastiku pod wpływem ciepła lub próżni.



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.piekarz.pl

części elektroniczne

sprzedaz@piekarz.pl tel. 22 599 49 70

www.gamma.pl

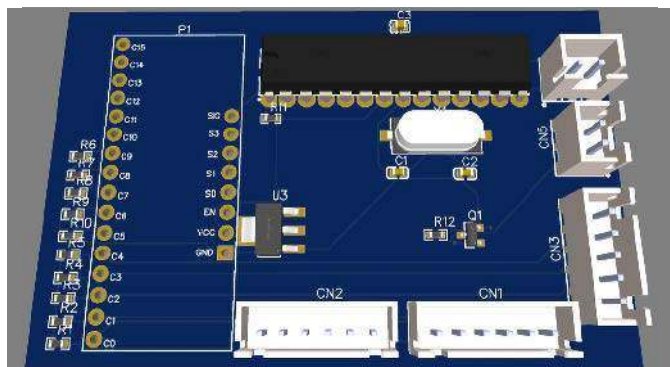
PODZESPOŁY ELEKTRONICZNE

info@gamma.pl

RACK i Eurocarta 19" Wyposażenie szaf 19"

www.obudowa.pl

Producent obudów dla elektroniki tel. 032-230-2301



W środku pojemnika znajduje się plastikowa wanienka, która pełni funkcję izolatora termicznego, dzięki czemu minimalizuje przenikanie ciepła i utrzymuje stałą temperaturę wody. Plastikowa wanienka ma także za zadanie zapobieganie wyciekaniu wody, co jest niezbędne do utrzymania odpowiedniej temperatury. Woda pełni kluczową funkcję w procesie chłodzenia. Układ zawiera prostą reakcję rozpuszczania azotanu amonu w wodzie, co powoduje absorpcję ciepła ze szczepionek. Dla uzyskania optymalnych wyników użyto jednego litra wody i 848 g azotanu amonu.

Reakcja endotermiczna zachodzi w momencie, gdy azotan amonu wchodzi w kontakt z wodą. Proces ten wymaga energii, która jest pobierana z otoczenia i powoduje ochłodzenie roztworu. Chociaż reakcja jest egzotermiczna w pewnym stopniu, to ciepło generowane przez jony azotanu amonu jest znacznie mniejsze od ilości energii potrzebnej do rozpuszczenia substancji. Dlatego cały proces jest reakcją endotermiczną, czyli wymagającą dostarczenia energii z otoczenia.

Kontrola chłodzenia w pojemniku Vacsafe jest możliwa dzięki zastosowaniu zaworu elektromagnetycznego. Elektrozawór, który jest aktywowany przez mikrokontroler ATmega8, otwiera się na określony czas, uwalniając określoną ilość azotanu amonu do wody. Cały proces jest monitorowany przez mikrokontroler, który odczytuje temperatury za pomocą termistorów umieszczonych na fiolkach ze szczepionkami. Mikrokontroler przesyła również informacje o temperaturze przez Bluetooth Low Energy (BLE) do smartfona transportera szczepionek, dzięki czemu personel może na bieżąco monitorować temperaturę.

Ważnym aspektem projektu jest oszczędzanie energii. Proces chłodzenia odbywa się co 15 minut, a przez resztę czasu mikrokontroler znajduje się w trybie wyłączenia, pobierając jedynie 0,5 μ A przy napięciu 3 V i częstotliwości zegara 4 MHz.

Projekt Vacsafe pozwala na utrzymanie odpowiedniej temperatury szczepionek podczas długich podróży, zapewniając bezpieczeństwo i skuteczność dostarczanych dawek. Jego innowacyjne podejście do chłodzenia na bazie reakcji endotermicznej sprawia, że jest to zaawansowane i obiecujące rozwiązanie dla przyszłości transportu szczepionek.

<https://hackaday.io/project/190141-vacsafe>

Przenośny, kieszonkowy beacon dla sieci Helium

Projekt Helium to zdecentralizowana sieć IoT, korzystająca z LoRaWAN. Powiązana jest z kryptowalutą Helium Network Token (HNT). Punkty tej sieci są typowo hostowane przez osoby prywatne w domach czy biurach. HNT jest zapłatą za tworzenie i utrzymywanie tych punktów. Celem sieci jest zapewnienie urządzeniom IoT sposobu na łatwą komunikację, który będzie charakteryzował się dużym zasięgiem.

Zaprezentowany projekt ma na celu zbudowanie uniwersalnej



platformy do korzystania z Helium Network. Pozwala na użycie jej we własnych projektach itp. Dzięki temu, że źródła i projekt hardware w tym module są otwarte, można je dowolnie modyfikować.

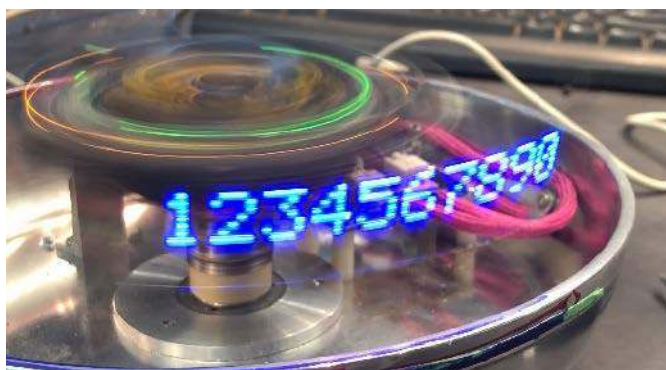
Podstawowe funkcje modułu, jakie wymienia autor, to:

- wysyłanie prywatnych wiadomości między punktami sieci a klientem aplikacji internetowej za pomocą sieci Helium IoT do komunikacji bezprzewodowej,
- uzyskiwanie przybliżonej geolokalizacji (dokładność 500 m). Funkcja ta powinna działać również w pomieszczeniach.

Moduł wyposażony jest w klawiaturę, co pozwala na łatwe pisanie wiadomości itp. Do wyświetlania zastosowano ekran typu e-ink, dzięki czemu urządzenie jest bardzo energooszczędne. Pozwala to na długi czas pracy na baterii (typowo ponad 3 dni). Moduł jest ładowany zwykłym kablem USB.

<https://hackaday.io/project/190142-beacon>

<https://gitlab.com/nlighten-systems/beacon>



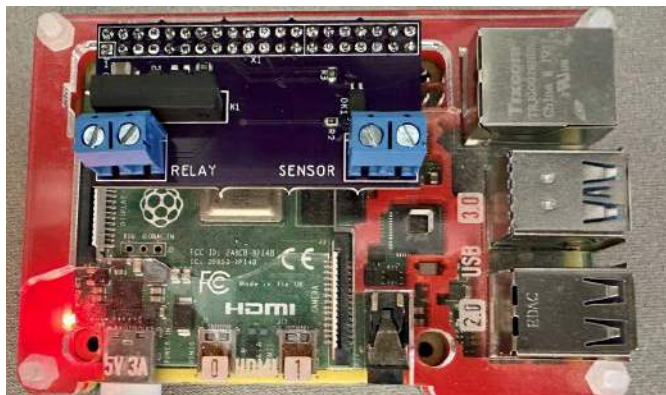
Dwukanałowy wyświetlacz smugowy sterowany przez Arduino

Zaprezentowany wyświetlacz smugowy ma dwa kanały wyświetlania – jeden z diodami białymi, a drugi z niebieskimi. Diod jest osiem, co przekłada się na rozdzielczość pionową. Dane do diod przekazywane są, wraz z zasilaniem 5 V, poprzez pierścienie ślizgowe. Dane przekazywane są szeregowo przez dwie linie – jedną dla niebieskich, a drugą dla białych diod.

Każdy wyświetlacz (niebieski i biały) ma własny mikrokontroler ATmega do obsługi danych szeregowych i synchronizacji oraz sterowania diodami LED. Wszystkie schematy, oprogramowanie napisane w Arduino i więcej szczegółów na temat konstrukcji można znaleźć na stronie z projektem.

<https://hackaday.io/project/190143-dual-channel-pov-display>

<https://shorturl.at/cfqRW>



Interfejs do drzwi garażowych z Raspberry Pi

Projekt ma na celu zapewnienie prostego i efektywnego rozwiązania do automatyzacji otwierania drzwi garażowych przy użyciu Raspberry Pi. W tym celu zastosowano platformę HomeBridge, która umożliwia łatwe konfigurowanie przełączników i czujników.

Do realizacji samego zadania otwierania i zamykania bramy użyto dwóch pinów GPIO w Raspberry Pi – jeden służy jako wyzwalacz przekaźnika do otwierania drzwi, a drugi pełni funkcję wejścia – czujnika, który monitoruje pozycję drzwi.

Ważnym aspektem projektu jest odizolowanie Raspberry Pi od mechanizmu otwierania drzwi garażowych, co ma zapewnić bezpieczne działanie. Zastosowano kontaktron sterowany przez N-MOSFET jako odpowiednik przycisku drzwi garażowych. Kontaktron umożliwia bezpieczne działanie elektronicznego przekaźnika.

Do wykrywania pozycji drzwi użyto typowego czujnika magnetycznego, który działa jako przełącznik zamykający się pod wpływem magnesu, umieszczonego w drzwiach. W celu zapewnienia izolacji pomiędzy czujnikiem a Raspberry Pi zastosowano optoizolator, który przetwarza sygnał z czujnika na odpowiedni poziom logiczny dla wejścia Raspberry Pi.

W zakresie oprogramowania zastosowano platformę HomeBridge, która umożliwia komunikację z usługą za pomocą Chamberlain. Dzięki temu Raspberry Pi może kontrolować mechanizm otwierania drzwi garażowych za pomocą tego popularnego narzędzia. Urządzenie jest w stanie sprawnie i niezawodnie obsługiwać drzwi z użyciem Raspberry Pi, co pozwala na wygodne zdalne sterowanie. Warto zaznaczyć, że układ ten może być łatwo dostosowany do różnych typów drzwi garażowych i różnych producentów mechanizmów otwierania. Użycie Raspberry Pi i HomeBridge pozwala na elastyczne i skalowalne rozwiązanie, które można dostosować do indywidualnych potrzeb użytkownika. W porównaniu z zakupem odpowiedniego modułu do drzwi garażowych z połączeniem do sieci, projekt ten oferuje korzyści w postaci oszczędności oraz większej kontroli nad automatyzacją drzwi garażowych.

<https://shorturl.at/fmvLO>



Sensible Watch – zegarek z interfejsem rozpoznającym gesty

Zaprezentowany projekt to inteligentny zegarek z paskiem składającym się z pojemnościowego i rezystancyjnego czujnika dotykowego, który pozwala na innowacyjną interakcję z interfejsem zegarka. Typowe zegarki mają przyciski z boku, co może być trudne w obsłudze np. dla osób z amputowaną ręką. Takie przyciski wymagają umiejętności motorycznych lub siły uchwytu, co może nie być możliwe w przypadku kończyny protetycznej lub po prostu jej braku.

Nowoczesne zegarki oferują różne funkcje poza określaniem czasu, minutnikami, alarmami czy funkcjami śledzenia kondycji. Poruszanie się po tych funkcjach, dostosowywanie ustawień lub aktywowanie określonych trybów może być trudne przy użyciu niesprawnej ręki.



Osoby po amputacji ręki mogą nosić Sensible Watch na sprawnym ramieniu lub ramieniu po amputacji dłoni itp. Czuły na dotyk pasek może wykrywać gesty i ruchy, umożliwiając użytkownikom interakcję z zegarkiem za pomocą amputowanego ramienia, sprawnej dłoni lub różnych innych części ciała (np. podbródek, łokieć itp.). Pasek współpracuje z protezami.

Działanie systemu jest bardzo proste – dotykając lub przesuwając pasek po częściach ciała, można poruszać się po różnych funkcjach zegarka i menu. Gdy zegarek jest noszony, użytkownicy mogą uzyskać dostęp do różnych funkcji, takich jak czas, stoper, alarm itp., przesuwając palcem. Ale nie mogą zmienić czasu, ustawić alarmów itp. Dopiero po zdjęciu zegarka z nadgarstka użytkownicy mogą zmienić godzinę, alarm, odliczanie itp., wykonując gesty dotykowe na pasku.

W planach jest zastosowanie hybrydowego – rezystancyjnego i pojemnościowego systemu dotykowego. Obie technologie zostaną połączone, aby zapewnić lepszą wydajność i funkcjonalność. Pojemnościowy sensor dotyku będzie używany głównie do wykrywania dotyku. Zapewnia krótszy czas reakcji i może precyzyjnie wykrywać wiele punktów dotyku jednocześnie, umożliwiając wykonywanie zaawansowanych gestów, takich jak przybliżanie przez rozsuwanie lub interakcje za pomocą wielu palców.

Rezystancyjne sensory dotyku są stosowane do pomiaru siły nacisku. Sensor taki może mierzyć poziom siły przyłożonej do powierzchni dotykowej. Łącząc pomiar nacisku z pojemnościowym wykrywaniem dotyku, system może rozróżniać lekkie dotknięcia i mocniejsze naciśnięcia, umożliwiając bardziej szczegółowe interakcje.

Łącząc mocne strony obu technologii, hybrydowe rezystancyjne i pojemnościowe systemy dotykowe mają na celu zapewnienie dokładniejszego, responsywnego i wszechstronnego doświadczenia dotykowego w protetyce. Aby zapobiec przypadkowym naciśnięciom, można zastosować różne strategie, podobne do podwójnego dotknięcia lub blokady ekranu, które są obecnie stosowane w smartfonach itp.

Obecnie sensor znajduje się na pasku, ponieważ funkcje, takie jak zmiana czasu, ustawienie alarmu itp. wymagają dokładnego odbioru wizualnej informacji zwrotnej. Inteligentne zegarki z ekranem dotykowym mają przyciski lub używają podłączonego smartfona do zmiany czasu, ustawiania alarmu itp. właśnie z tych przyczyn. Na stronie projektu znaleźć można wiele informacji na temat technicznych aspektów realizacji poszczególnych sensorów.

<https://hackaday.io/project/191320-sensible-watch>

System zarządzania telewizorami i kamerami

Pracując w domu opieki nad osobami starszymi, autorowi tego projektu dość często zdarza się, że mieszkańcy nie mogą (lub nie chcą) wychodzić ze swoich pokoi na organizowane imprezy czy przedstawienia. Poglębia ten problem jeszcze fakt, że wielu pensjonariuszy





jest przykutych do łóżka lub mieszkają w oddzielnym bezpiecznym skrzydle dla osób z demencją.

Największym i najbardziej popularnym wydarzeniem w tym ośrodku jest cotygodniowy program „Happy Hour” w piątkowe popołudnia, podczas którego muzyk gra na żywo, przez około 1 godzinę, dla mieszkańców domu opieki. Serwowane jest także jedzenie i napoje. Nie wszyscy mogą jednak brać udział w tych spotkaniach.

Dom opieki ma własną instalację telewizyjną, w której zawarto dwa dodatkowe kanały telewizyjne, oba podłączone do odtwarzaczy DVD, początkowo używanych do ręcznego odtwarzania płyt DVD o określonych porach każdego dnia. Ta sytuacja zainspirowała opiekunów do ręcznej transmisji na żywo wspomnianego cotygodniowego wydarzenia z muzyką na żywo w sieci telewizyjnej. Od tego początkowo udanego eksperymentu projekt rozwijał się przez kilka lat, rozbudowywany o kolejne pomysły i eksperymenty oraz kolejny sprzęt i kanały telewizyjne w wewnętrznej sieci. Finalnie, gdy projekt jest już dojrzały i sprawdzony, autor zdecydował się opublikować go jako open source!

Projekt przeszedł wiele zmian i rozszerzeń, a obecnie oferuje:

- integrację z istniejącą siecią MATV dostarczającą wideo 1080p do 170 telewizorów w pokojach,
- możliwość planowania transmisji na żywo,
- kamerę HD w głównym salonie, umieszczoną na sterowanej głowicy PTZ. Głowica może być sterowana za pomocą algorytmów widzenia maszynowego (napisanych z użyciem OpenCV) do śledzenia obiektów,
- możliwość prowadzenia transmisji na żywo na żądanie z tabletu iPad,
- 7 dodatkowych kanałów telewizyjnych za pośrednictwem modulatorów RF o wysokiej rozdzielczości:

- kanał „Wydarzenia na żywo” zawierający zaplanowane i transmitowane na żądanie transmisje na żywo z kamer z wnętrza obiektu, a także inne zaplanowane treści z biblioteki multimedialnych plików,
- kanał przeznaczony do ręcznego odtwarzania płyt DVD zgodnie z pierwotnymi założeniami,
- kanał „TV Shows”, który zawiera playlistę z ponad tygodniem programów telewizyjnych, przedstawieniami i filmami dokumentalnymi nadawanymi 24 godziny na dobę, 7 dni w tygodniu,
- kanał „Relaks” z uspokajającymi treściami dotyczącymi zwierząt/natury z muzyką relaksacyjną, również nadawany 24/7,
- trzy razy więcej kanałów używanych obecnie do odtwarzania wybranych kanałów telewizji kablowej;
- internetowego menedżera list odtwarzania do konfiguracji list odtwarzania dla kanałów „TV Shows” i „Relaks”,
- internetową zdalną kontrolę kanału z wydarzeniami na żywo w celu ręcznej zmiany standardowego harmonogramu,
- trzy komputery PC z systemem Linux do odtwarzania wideo HDMI dla wydarzeń na żywo, programów telewizyjnych i kanałów relaksacyjnych,
- nagrywanie strumieniowe kanałów z wydarzeniami na żywo – są one zapisywane, ponownie kodowane i wysyłane do NASa w celu późniejszego odtworzenia.

Dalsze plany obejmują ulepszenie systemu śledzenia w OpenCV, ponieważ wiele w tym zakresie zmieniło się w technologiach od czasu, gdy obecna wersja została napisana i zaimplementowana. Dodatkowo autor chciałby również ulepszyć aplikację internetową menedżera list odtwarzania tak, aby można było drukować z niej harmonogram, aby przekazać go mieszkańcom, by wiedzieli o tym, co jest grane i o której godzinie.

<https://shorturl.at/ERW26>



Cyfrowy miernik ciśnienia krwi

To ciekawy projekt zawierający mikroczujnik ciśnienia do wykrywania tzw. tonów Korotkowa. Stworzony został przez autora, który ma już bogate doświadczenie w dziedzinie sprzętu medycznego. Dzięki zastosowaniu tej technologii możliwe jest obliczenie skurczowego (maksymalnego) i rozkurczowego (minimalnego) ciśnienia krwi z pomiarów akustycznych. Projekt ten może stanowić interesującą alternatywę dla metody oscylometrycznej, która często wymaga opracowania własnych algorytmów matematycznych w celu dokładnego pomiaru ciśnienia skurczowego i rozkurczowego.

Projekt został udostępniony jako open source, co otwiera możliwość jego wykorzystania nawet w odległych społecznościach, gdzie higiena i zdrowie, zwłaszcza osób niepełnosprawnych, są ważnym wyzwaniem. Warto podkreślić, że autorem projektu jest osoba, która ma pełną wiedzę na temat działania ciśnieniomierzy, a sama implementacja odbyła się przy użyciu narzędzia GNAT Programming Studio. Dodatkowo kod źródłowy, schemat ideowy oraz schemat blokowy zostały dostarczone przez autora, co pozwala na łatwe zrozumienie i modyfikację projektu.

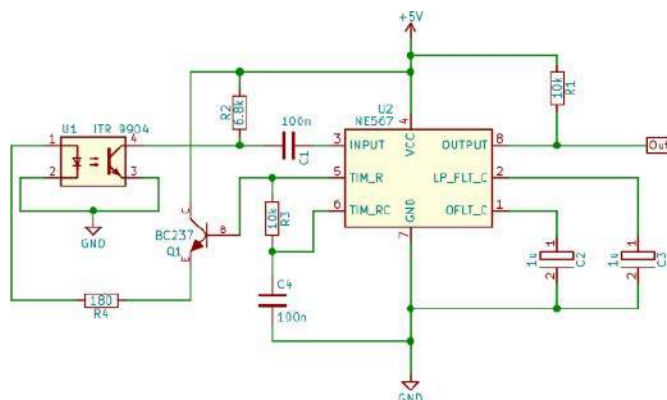


Wśród kluczowych cech tego projektu jest porównanie dokładności pomiarów ciśnienia wykonanych za pomocą omawianego urządzenia z tradycyjnym manometrem aneroidowym. W celu uzyskania optymalnych wyników autor wielokrotnie kalibrował zawór odpowietrzający. Dzięki swojej specyficznej funkcjonalności ten model może okazać się niezwykle przydatny dla lekarzy i pielęgniarek, którzy mogą cierpieć na słaby słuch, co może utrudniać usłyszenie tonów Korotkowa samodzielnie, przy użyciu tradycyjnych ciśnieniomierzy i stetoskopu.

Wyniki testów projektu, które obejmowały porównanie pomiarów z urządzeniem OMRON, dostarczyły autorowi wiele danych i potwierdziły jego skuteczność. Projekt został podzielony na trzy kluczowe sekcje: pompowanie mankietu, opróżnianie mankietu oraz analiza tonów Korotkowa, co pozwala na łatwą identyfikację kolejnych etapów pracy urządzenia.

Podsumowując, ten fascynujący projekt stanowi wartościowy wkład w dziedzinę elektroniki medycznej i może znaleźć szerokie zastosowanie w praktyce medycznej, szczególnie w odległych i trudno dostępnych regionach. Zastosowanie nowoczesnej technologii połączonej z zaawansowanymi algorytmami pozwala na bardziej precyzyjne i wygodne pomiary ciśnienia krwi, a jednocześnie może poprawić opiekę nad pacjentami, zwłaszcza w przypadku ograniczeń związanych z niedosłuchem personelu medycznego.

<https://hackaday.io/project/191313-digital-blood-pressure-monitor>



Kurtyna świetlna do detekcji przejścia przez wiązkę

Aby wykryć przerwanie wiązki światła lub powierzchni odbijającej światło, dioda LED wysyła światło wykrywane przez fototranzystor. Jeśli używany jest tylko pojedynczy statyczny poziom detekcji, rozproszone światło może fałszywie wskazać powierzchnię odbijającą lub nie wykryć przerwy. Omówiony w projekcie obwód zawiera modułowane światło, które jest wykrywane przez pętlę sprzężoną fazowo, więc każde światło o innej częstotliwości czy fazie jest ignorowane przez detektor.

Tranzystor w układzie steruje diodą LED z właściwą fazą z prądem kontrolowanym opornikiem. Prąd fototranzystora można regulować i zmniejszyć, jeśli następuje nasycenie fototranzystora na skutek np. silnego rozproszenia światła.

Wyjście z układu realizowane jest jako wyjście cyfrowe – stan wysoki prezentowany jest, gdy na detektorze obecne jest światło o prawidłowej fazie i częstotliwości. Wyjście może zapewniać do 100 mA (w stanie włączonym) i wytrzymuje przyłożenie do 15 V (w stanie wyłączonym).

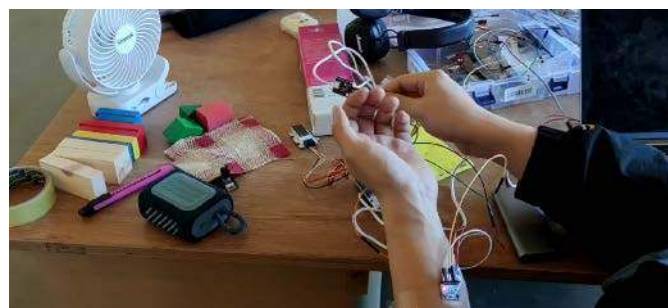
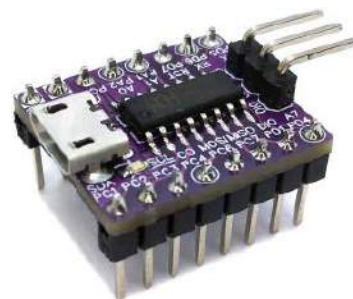
<https://shorturl.at/nuDHI>

Płytki deweloperska do ekonomicznego mikrokontrolera RISC-V CH32V003 A4M6

Seria mikrokontrolerów CH32V003 to rodzina zaawansowanych układów klasy przemysłowej, zawierających nowoczesną architekturę rdzenia QingKe RISC-V2 A z zestawem instrukcji RV32EC. Oferują one różnorodne funkcje, w tym wysoką częstotliwość systemową równą

48 MHz, wsparcie dla szerokiego zakresu napięcia, jednoprzewodowy szeregowy interfejs do debugowania, niskie zużycie energii oraz kompaktowe rozmiary obudowy. Dodatkowo seria układów CH32V003 jest wyposażona we wbudowane komponenty, takie jak kontroler DMA, 10-bitowy przetwornik ADC, komparatory wzmacniaczy operacyjnych, liczne timery oraz standardowe interfejsy komunikacyjne, m.in. USART, I2C i SPI. Dzięki tym zaawansowanym funkcjom mikrokontrolery z serii CH32V003 stanowią doskonałe rozwiązanie do wielu zastosowań przemysłowych.

<https://shorturl.at/mqMZ4>

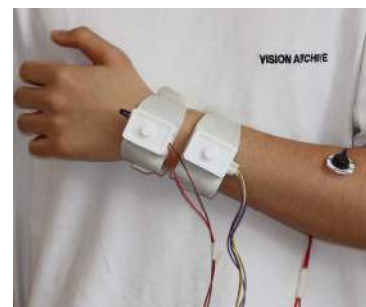


Kontroler reagujący na gesty i pomiary mioelektryczne

GMC (*Gesture and Myoelectric Control*) to niedrogie i bezpieczne urządzenie wspomagające, które można kontrolować za pomocą rąk lub nóg poprzez postawę i siłę mięśni. Dzięki temu osoby z ograniczoną mobilnością lub pewnymi ograniczeniami ruchu kończyn odniosą korzyści z tego urządzenia. GMC składa się z trzech części do noszenia: sygnalizatora korzystającego z podczerwieni, sześcioksięgowego akceleratora i czujnika elektromiograficznego (EMG). Można je nosić na dowolnej części ciała, którą można kontrolować, aby dostosować się do osób o różnych funkcjach fizycznych. Wszystkie te elementy są bardzo lekkie i łatwe w noszeniu.

Głównym celem tego projektu jest umożliwienie sterowania urządzeniami domowymi za pomocą prostej manipulacji kończynami oraz zwiększenie samodzielności osób niepełnosprawnych. Aby osoby niepełnosprawne mogły żyć samodzielnie, muszą mieć możliwość samodzielnego kontrolowania swojego środowiska domowego. Wprowadzenie narzędzi takich jak Alexa i zdalna kontrola aplikacji przez nowoczesną technologię inteligentnego domu uwzględniło już osoby niepełnosprawne. Jednak po rozmowach z rzeczywistymi użytkownikami odkryto, że nadal istnieje pewna grupa demograficzna, którą technologia inteligentnego domu nadal wyklucza. Na przykład spotkano kobietę, która miała udar i nie była w stanie mówić wyraźnie ani w pełni wyprostować rąk i palców. Zdrowie fizyczne wpłynęło na jej interakcje z artykułami gospodarstwa domowego. Mimo to zawsze zachowuje pozytywne nastawienie i stara się jak najlepiej żyć samodzielnie. To zainspirowało twórców do opracowania bardziej dostępnej technologii inteligentnego domu, aby osoby takie jak ta mogły lepiej kontrolować swoje urządzenia gospodarstwa domowego i wzmocnić swoją niezależność i pewność siebie.

<https://shorturl.at/aEHI1>



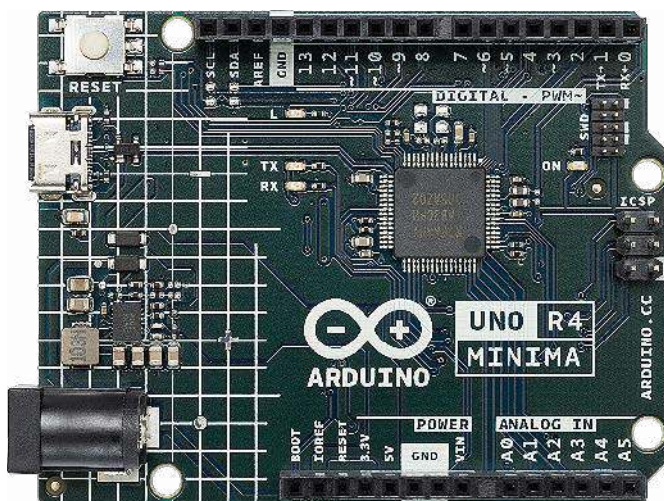
Arduino UNO R4 Minima – wydajna i nowoczesna, ale...

Na krótko przed wakacjami zespół Arduino wprowadził na rynek nową płytę, a w zasadzie dwie wersje płytki zawierające ten sam procesor, są to Arduino UNO R4 Minima i Arduino UNO R4 WiFi. Obie bazują na układzie produkowanym przez Renesasa – 32-bitowym procesorze z rdzeniem ARM Cortex-M4 z rodziny RA4M1.

W założeniach płytka ma zastąpić Arduino UNO R3. Jak zwykle materiały marketingowe obiecują polepszenie wszystkich parametrów, co więcej, R4 to wręcz wprowadzenie „nowego wymiaru prototypowania”, a producent deklaruje zgodność (przynajmniej teoretyczną) zarówno sprzętową, jak i programową z większością aplikacji R3. Dla porównania w **tabeli 1** zestawiono podstawowe parametry UNO R4 i UNO R3. Celowo porównuję wersję R4 Minima, a nie wersję R4 WiFi, gdyż po pierwsze UNO R3 także nie ma żadnych opcji komunikacji bezprzewodowej, a po drugie zakup jej wydawał się rozsądniejszy (była tańsza...).

Mikrokontroler

Wersja UNO R4 WiFi wydaje mi się dosyć nietypowym zestawieniem i przed zakupem wymaga dłuższego zastanowienia. W wersji WiFi do procesora RA4M1 podłączono ESP32, aby zapewnić komunikację bezprzewodową. Takie zestawienie to temat na odrębne



opracowanie, ale już można znaleźć w sieci wiele niezbyt pochlebnych opinii o tym rozwiązaniu. Sytuacja jest tym bardziej zaskakująca, że krótko po premierze R4 WiFi udostępniona została płytka Nano ESP32 (S3).

Arduino UNO R4 Minima zachowuje zgodność mechaniczną z UNO R3. Zastosowany procesor R7FA4M1AB3CFM, co nie dziwi w dzisiejszych czasach, jest wlutowany bezpośrednio w płytke, gdyż

Tabela1. Porównanie parametrów i funkcjonalności płytek UNO R4 i UNO R3

Płytki		Arduino UNO R4 Minima	Arduino UNO R3
Nr katalogowy		ABX00080	A000066
Mikrokontroler		Renesas RA4M1 Arm Cortex-M4 R7FA4M1AB3CFM	ATmega328P
USB		USB-C/HID	USB-B
GPIO	DIO	14	14
	ADC	6 (14-bitowy)	6 (10-bitowy)
	DAC	1 (12-bitowy)	0
	PWM	6	6
	Maksymalny prąd GPIO	8 mA	20 mA
Porty komunikacyjne	UART	1	1
	I ² C	1	1
	SPI	1	1
	CAN	1 (wymaga zewnętrznego transceivera CAN)	0
Zasilanie	VCC	5 V	5 V
	VIN	6...24 V	7...12 V
Taktowanie	CPU CLK	48 MHz	16 MHz
Pamięć	Flash	256 kB	32 kB
	RAM	32 kB	2 kB
	EEPROM	data flash 8 kB	1 kB
Cena		18 €	24 €
Pozostałe		wzmacniacz operacyjny, RTC, złącze SWD	

oczywiście nie jest oferowany w obudowach DIP, a tylko w obudowach LGA, LQFP i QFN o różnej, w zależności od podtypu, liczbie wyprowadzeń (40 do 100). Nie jest to dobra informacja dla konstruktorów, którym czasem udaje się uszkodzić procesor, tym razem nie obejdzie się bez wylutowania lub zakupu następnej sztuki.

UNO R4 zawiera R7FA4M1AB3CFM w obudowie LQFP64, budowę wewnętrzną procesora pokazano na **rysunku 1**. Z racji przemysłowej i motoryzacyjnej specjalizacji Renesas oferuje procesor zakwalifikowany do przemysłowego lub motoryzacyjnego zakresu temperatur, co zawsze gwarantuje większą stabilność rozwiązania, o ile inne komponenty nie zawiodą przy -40°C lub $+85^{\circ}\text{C}$.

Po szybkim sprawdzeniu dostępności procesora u największych dystrybutorów dowiadujemy się, że u niektórych procesory są dostępne, ale niepokojący jest termin 80 tygodni na dostawę większych ilości. Aktualnie w sklepie Arduino wersja Minima jest dostępna, a Wi-Fi niestety nie i nie jest znany termin ponownego pojawienia się w dystrybucji.

Złącza

Oprócz mikrokontrolera najbardziej widoczną zmianą jest zastosowanie wspólnego złącza USB-C, służącego zarówno do zasilania, jak i komunikacji. Zewnętrzny zasilacz w dalszym ciągu połączony jest poprzez złącze DCIN z typowym wtykiem 5,5/2,1 mm. Wysokie złącze zasilania wraz z niestandardowym rozstawem złączy GPIO utrudnia zarówno prototypowanie na płytkach stykowych, jak i kanapkowe łączenie nakładek – opierają się o złącze zasilania. Tutaj bez litości można przytoczyć powiedzenie o tym, że „mylić się jest rzeczą ludzką, ale tkwić w błędzie to już...”. Przynajmniej problem złącza DCIN można było rozwiązać inaczej, tym bardziej że UNO R4 i tak będzie wymagało nowej obudowy, ze względu na złącze USB-C.

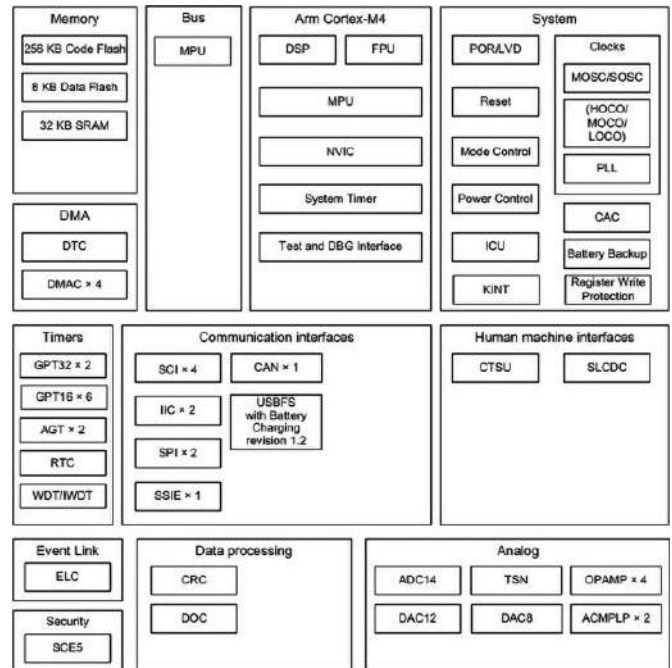
W temacie obudowy pozostaje też kwestia podstawki z tworzywa, która jest dołączona do zestawu (**fotografia 1**). Być może w niektórych zastosowaniach będzie przydatna ale w większości przypadków wyląduje w koszu, ponieważ płytka Arduino trafi do obudowy docelowego urządzenia.

Zasilanie

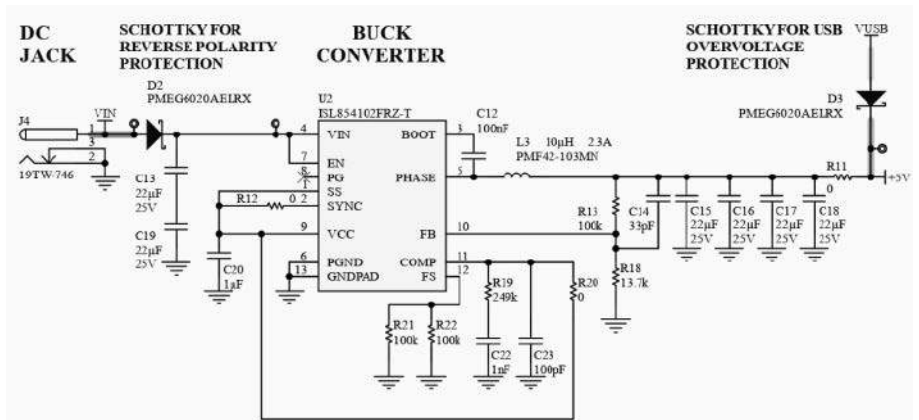
W przypadku układu zasilania poprawiono zakres napięć, dla których poprawnie pracuje wbudowana przetwornica i w UNO R4 sięga on 24 V, choć sama przetwornica może pracować w jeszcze szerszym



Fotografia 1. Arduino UNO R4 Minima w podstawce z tworzywa



Rysunek 1. Budowa procesora rodziny R4M1 (za notą Renesasa)



Rysunek 2. Schemat obwodu zasilania płytki UNO R4

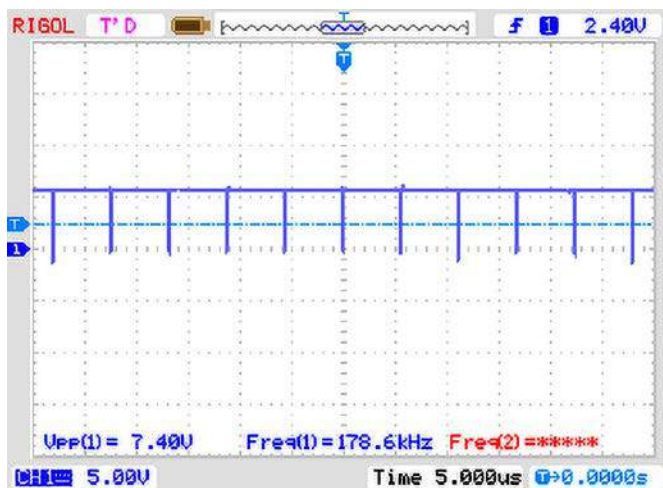
REKLAMA



OBWODY DRUKOWANE

Produkcja, Projektowanie, Montaż

Certyfikat Underwriters Laboratories  94V-0 E480148 TYPE 1	Płytki jednostronne Płytki dwustronne Płytki na podłożu aluminium Płyty czołowe FR4	Serie dowolne Prototypy Maksymalny wymiar płytek 1w 630 mm	
	Zakład produkcyjny: 05-660 Warka ul. M. Ropielewskiej 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	Dokumentacja technologiczna Dokumentacja konstrukcyjna Trawione szablony SMD	Montaż elektroniki Krótkie terminy Wykonania super expresowe
	Aktywny kalkulator prototypów na stronie internetowej	Pokrycie Sn lub SnPb inne na życzenie Maski, opisy montażowe w różnych kolorach	
			



Rysunek 3. Przebiegi w przetwornicy zasilanej z USB-C

zakresie. Problemem może być jednak sposób kluczowania zasilania z USB-C i wbudowanej przetwornicy poprzez dwie diody PMEG6020.

W przypadku poprawnego napięcia z USB-C, wynoszącego 5 V, do zasilania płytki doprowadzone jest napięcie obniżone o spadek napięcia na D3, co pokazano na **rysunku 2**. Nawet gdy UNO R4 zasilane jest z przyzwoitego zasilacza dla Raspberry Pi, z podniesionym napięciem wyjściowym do 5,1 V, napięcie zasilania R4 wynosi ok. 4,8 V i jest na granicy tolerancji niektórych układów i to bez obciążenia GPIO.

Nie jest to jedyny problem obwodu zasilania UNO R4. W przypadku obecności zasilania z USB-C i braku zasilania z gniazda DC, przetwornica U2 jest zasilana tylko od strony wyjścia. W zasilaczu zastosowano układ ISL854102, także produkcji Renesas, który zgodnie z kartą katalogową jest w stanie pochłaniać prąd ze strony wyjściowej, ale po co ma to robić i dodatkowo obciążać zasilanie USB-C, tym bardziej że niepotrzebnie steruje wtedy cewką L3, co pokazuje przebieg z **rysunku 3** zaobserwowany na wyprowadzeniu PHASE. Można to uznać za klasyczny przypadek back-feedingu, czyli niepożądanego wstecznego przepływu zasilania, który powoduje zakłócenia w działaniu układu, a w skrajnych przypadkach może doprowadzić do jego uszkodzenia. Pobór prądu nie jest duży, ale zupełnie niepotrzebny i generuje dodatkowe zakłócenia na szynach zasilania, które naprawdę są zbędne w przypadku zasilania procesora z rozbudowanym blokiem analogowym, tym bardziej że napięcie 5 V jest używane także jako napięcie odniesienia dla ADC.

Problem można było rozwiązać kluczami MOSFET, zmniejszając spadek napięcia zasilania oraz eliminując niepotrzebne oscylacje w przetwornicy lub zmieniając miejsce kluczowania napięć z „za” na „przed” przetwornicą i stosując przetwornicę podwyższająco-obniżającą, zapewniającą stabilne zasilanie. Dodatkowo umożliwiłoby to zasilanie układu z akumulatora litowego lub zestawu baterii 3...4xLR4, jak to ma miejsce w przypadku innych płytek. W skrajnym wypadku można zastosować rozwiązanie zupełnie budżetowe polegające na wstawieniu diody PMEG6020 w miejsce R11, przy równoczesnym podniesieniu napięcia wyjściowego przetwornicy, w celu kompensacji spadku na diodzie.

Sama płytka UNO R4 bez problemu mogłaby pracować też z zasilaniem 3,3 V, bo procesory R4M1 mają bardzo szeroki zakres zasilania 1,6...5,5 V, szkoda że nie przewidziano takiej możliwości i za wszelką cenę zapewniono zgodność tylko z logiką 5 V.

Ostatnim elementem obwodu zasilania jest źródło napięcia 3,3 V, które oczywiście zostało potraktowane oszczędnościowo i pochodzi z wbudowanego w procesor LDO, maksymalna obciążalność zgodnie z kartą katalogową to 100 mA. Czy źródło jest odporne na przeciążenia i zwarcia, karta nie wspomina, więc zalecam daleko idącą ostrożność.

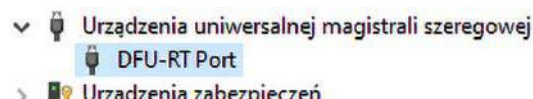
Debugowanie... a raczej jego brak

Kolejnym widocznym elementem jest złącze SWD, służące do debugowania, którego płytka UNO R4 jest... pozbawiona. Początkujący konstruktor raczej nie ma dostępu do programatora/debuggera J-Link, a miganie diodą LED sprawdzi się tylko w najprostszych projektach. Dla bardziej złożonych zadań, dla których polecane jest UNO R4, to chyba słabe rozwiązanie i marna propozycja nabywania poprawnych nawyków podczas prototypowania. W pewnym sensie rozumiem, że programowanie Arduino polega przede wszystkim na znalezieniu właściwej biblioteki, ale za każdym razem, gdy udostępniana jest nowa płytka, ignorowana jest możliwość zmiany takiego podejścia. Tym bardziej, że da się opracować prosty debbuger, co udowodniła fundacja Raspberry, projektując zewnętrzną płytkę Raspberry Pi Debug Probe dla przeciętnie ekstremalnie budżetowego Pi Pico.

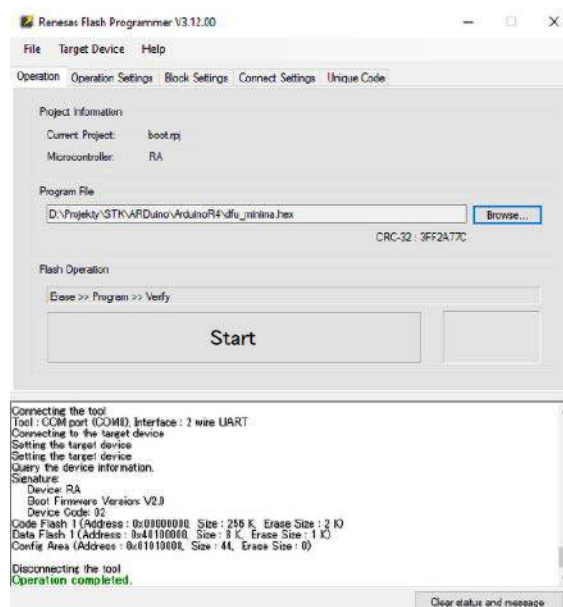
Jest też coś dla wyznawców minimalizmu, czyli niewielka płytka MIKROE RA4M1 Clicker, także z procesorem R7FA4M1AB3CFM i wbudowanym debuggerem oraz złączem rozszerzeń MikroBUS dla ClickBoardów. I oczywiście sam Renesas dla fabrycznej płytki EK-RA4M1, która nie dość, że ma wbudowany programator/debugger, to jeszcze wyposażona jest w procesor z pełną 100-pinową obudową, z wyprowadzonymi wszystkimi pinami GPIO i jest obsługiwana w darmowym E2Studio. Jedyne, do czego można mieć zastrzeżenia, to brak nowego złącza – USB-C. Jej cena jest tylko dwukrotnie wyższa od R4 Minima.

Kompatybilność

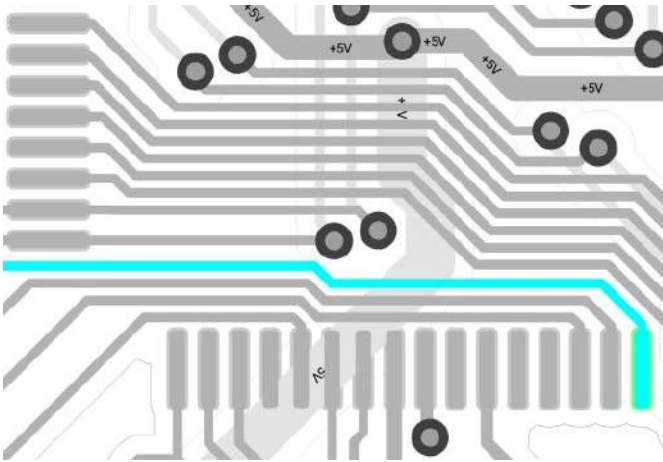
Arduino deklaruje zgodność z oferowanymi dla R3 nakładkami i to nie tylko polegającą na zgodności poziomów napięciowych 5 V dla GPIO, ale na pełnej funkcjonalności programowej. Sprawdzone zostało działanie nakładek 4Relays, Motor3, Ethernet, Edu i 9Axis Motion Shield. Zgodność z nakładkami innych producentów można sprawdzić na stronie <https://docs.arduino.cc/tutorials/uno-r4-minima/shield-compatibility>, niestety te ciekawsze nie są zgodne, m.in. Adafruit Neopixel, NFC, TFT, część niezgodności dotyczy sprzętu, część oprogramowania, zapewne w czasie sytuacji się poprawi. Zatem każdorazowo we własnym zakresie należy sprawdzić zgodność posiadanych rozszerzeń z płytką UNO R4.



Rysunek 4. Poprawna instalacja DFU



Rysunek 5. Oprogramowanie Renesas RFP



Rysunek 6. Ścieżka wejścia analogowego A1

USB

Procesor R4M1 ma wbudowaną obsługę portu USB, która jest używana do komunikacji i programowania pamięci programu poprzez środowisko Arduino. Wspierany jest także tryb USB-HID, który umożliwia emulację klawiatury i myszy, co jest dużą zaletą. W przypadku problemów z wgrywaniem aplikacji należy sprawdzić poprawność instalacji urządzenia DFU-RT Port i odpowiadającego mu urządzenia szeregowego USB(COMxx) – **rysunek 4**. W przypadku problemów w moim przypadku pomogła ponowna instalacja w trybie administratora.

Do wgrania Bootloadera służy tryb BOOT, który jest aktywowany poprzez zwarcie wyprowadzenia BOOT z masą na złączu POWER i ponowne podłączenie UNO R4 do komputera. Bootloader można wgrać, za pomocą Arduino lub oprogramowania Renesas Flash Programmer (RFP) – **rysunek 5**. Pozostaje mieć jednak nadzieję, że w większości przypadków nie będzie to konieczne.

GPIO

Chwilę uwagi należy poświęcić wyprowadzeniom GPIO, a dokładnie ich obciążalności prądowej, która jest ponad 2 razy mniejsza niż w UNO R3. Każdy z pinów może zostać obciążony maksymalnym prądem 8 mA, a prąd sumaryczny wszystkich portów nie może przekraczać 60 mA. To znaczące ograniczenie, szczególnie gdy bezpośrednio sterujemy diodami LED lub multipleksowanymi wyświetlaczami LED, które bez dodatkowego drivera już nie zadziałają prawidłowo.

W przypadku GPIO nasuwa się także pytanie, dlaczego procesor ma 64 wyprowadzenia, a na złączach dostępne jest tylko 20, nie licząc interfejsu SWD, ICSP i wbudowanych LED? Aż 16 wyprowadzeń GPIO pozostało niepodłączonych. A przecież R4M1 dostępne są też w tańszych obudowach o 40 i 48 wyprowadzeniach, które zapewne bez większych konsekwencji mogły być użyte w projekcie.

Analizując budowę wewnętrzną procesora, można zauważyć znaczącą rozbudowę sekcji analogowej. Poprawione zostały osiągi przetwornika ADC, jego maksymalna rozdzielczość to 14 bitów, co razem z 32-bitowym rdzeniem ułatwia tworzenie złożonych aplikacji pomiarowych. Obsługa zwiększonej rozdzielczości odbywa

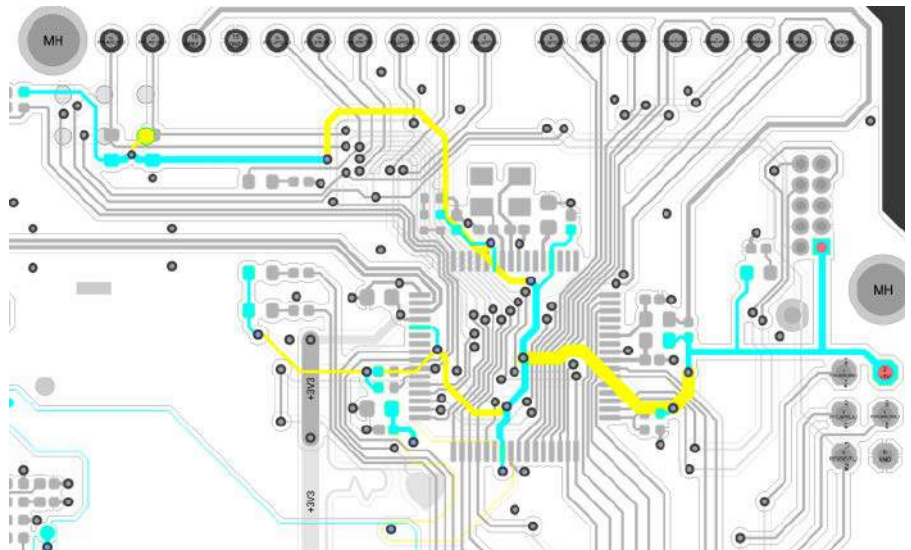
się poprzez nowe funkcje `analogReadResolution(10)`, `analogReadResolution(12)`, `analogReadResolution(14)`, a domyślną rozdzielczością pozostaje 10 bitów. Ciekawe, jak będzie w praktyce wyglądała możliwość używania 14-bitowej rozdzielczości. Filtracja zasilania w UNO R4 jest bardzo oszczędna, a sam obwód drukowany jest zaprojektowany dosyć chaotycznie, np. wejście A1 przetwornika ADC zaznaczone kolorem na sporej długości sąsiaduje ze ścieżkami I²C i SPI/CAN, co pokazano na **rysunku 6**. Ścieżki części analogowej mogły być poprowadzone przy zmianie lokalizacji elementów, tak aby zachować minimalną długość, przecież przypisanie pinów można zmienić programowo.

Równie chaotycznie jest zaprojektowane zasilanie 5 V, pokazane na **rysunku 7**. Pomimo sporej wylewki 5 V przy przetwornicy całe zasilanie przeprowadzone jest jedną przelotką i dosyć chaotyczną siecią ścieżek o przypadkowych szerokościach. Z kolei ścieżka 3 V, o maksymalnej wydajności źródła 100 mA, jest relatywnie gruba, przynajmniej na części trasy, a przecież te 3 V bierze się z 5 V.

Nowości

Nowością w R4 jest wbudowany wzmacniacz operacyjny, dostępny na wyprowadzeniach analogowych A1, A2, A3. Nie jest to nic szczególnego, ponieważ takie rozwiązania są dostępne nawet w niektórych procesorach 8-bitowych. Dodatkowo Arduino słabo dokumentuje jego użytkowanie i trzeba studiować dokumentację procesora.

Dużym plusem jest wbudowany 12-bitowy DAC, którego wyjście dostępne jest na wyprowadzeniu A0. Będzie alternatywą dla wyjścia



Rysunek 7. Ścieżka zasilania 5 V

Listing 1. Modyfikacje kodu dla użycia diod LED RX, TX

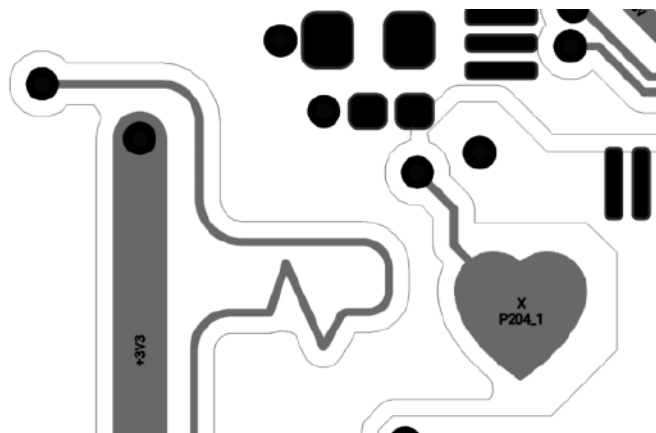
```
#define PORTBASE 0x40040000 /* Port Base */
#define PFS_P012PFS_BY ((volatile unsigned char *) (PORTBASE + 0x0803 + (12 * 4))) // TxLED - VREFH
#define PFS_P013PFS_BY ((volatile unsigned char *) (PORTBASE + 0x0803 + (13 * 4))) // RxLED - VREFL

void code_function_fragment(void) {
    *PFS_P012PFS_BY = 0x04; // OnBoard TxLED on
    *PFS_P012PFS_BY = 0x05; // OnBoard TxLED off

    *PFS_P013PFS_BY = 0x04; // OnBoard RxLED on
    *PFS_P013PFS_BY = 0x05; // OnBoard RxLED off
}
```

Listing 2. Sterowanie diodami LED RX, TX

```
void loop() {
    *PFS_P012PFS_BY = 0x04; // OnBoard TxLED on
    delay(100);
    *PFS_P013PFS_BY = 0x04; // OnBoard RxLED on
    delay(100);
    *PFS_P012PFS_BY = 0x05; // OnBoard TxLED off
    *PFS_P013PFS_BY = 0x05; // OnBoard RxLED off
    delay(1000);
}
```



Rysunek 8. Pole dotykowe w kształcie serca

PWM w wielu zastosowaniach takich, jak generacja sygnału analogowego. Niestety cyfrowy interfejs audio SSIE dostępny jest tylko dla obudowy 100-wyprowadzeniowej.

Absurdem części analogowej w przypadku UNO R4 jest współdzielenie wyprowadzeń I²C z dwoma kanałami ADC. Gdy chcemy użyć magistrali I²C, wyjścia DAC i wzmacniacza operacyjnego, nie pozostaje nam już wolne wejście ADC, a 4 wyprowadzenia ADC procesora pozostają niepodłączone – naprawdę można rozwiązać to lepiej, tym bardziej że drugi kanał I²C dostępny na wyprowadzeniach 1 i 2 procesora pozostaje niepodłączony.

Procesor RA4M1 wyposażony jest w wewnętrzny czujnik temperatury TSN, ale Arduino nie wspomina niestety o możliwości jego użycia.

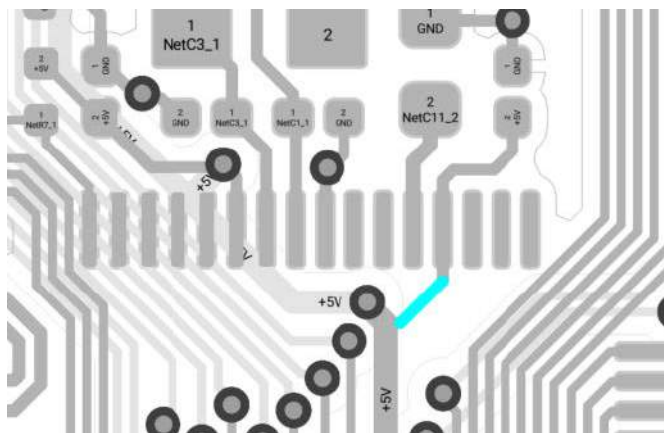
Uno R4 wyposażono w dwa interfejsy komunikacyjne UART, pierwszy realizowany przez USB, drugi sprzętowy dostępny na wyprowadzeniach RX, TX. Kuriozum stanowią diody sygnalizacyjne RX, TX – które niczego nie sygnalizują. Zarówno podczas transmisji USB (Serial), jak i poprzez piny RX, TX (Serial1) pozostają zgaszone. Po szybkiej analizie schematu okazuje się, że są podłączone do wyprowadzeń P012, P013, które nie mają nic wspólnego z transmisją szeregową. Aby było zabawniej, nie są wspierane przez Arduino. Aby je zaświecić, należy zastosować definicję podaną przez użytkownika *susan-parker* z forum arduino: <https://forum.arduino.cc/t/do-rx-tx-leds-glow/1146115/4> zgodnie z **listingi** 1 i wysterować diody we własnym zakresie jak w **listingu 2**.

Diody oznaczone są na płytce RX, TX tylko po to, żeby zachować zgodność opisów z R3, chyba lepiej było je opisać P012, P013. Dzięki własnej definicji przynajmniej można ich do czegoś użyć.

Magistrala I²C jest zgodna z 5 V, płytka nie posiada rezystorów podciągających, sygnały SDA, SCL wyprowadzone są na złączu Analog A4, A5 i powielone na złączu DIGITAL. Drugi interfejs I²C pozostaje nieobecny, chociaż w UNO R4 WiFi podłączony jest do złącza QWIIC, więc po drobnych modyfikacjach mógłby zostać wyprowadzony. Osobiście wolalbym drugi interfejs I²C w miejsce niezbyt praktycznego rozwiązania (żartu projektanta), jakim jest pin dotykowy w kształcie serca i ścieżka z jego pulsem do wyprowadzenia BOOT, pokazane na **rysunku 8**. Tym bardziej że pin dotykowy blokuje wyprowadzenie, na które można było przełączyć SCL0. Można było użyć do obsługi dotyku jednego z kilkunastu wolnych wyprowadzeń procesora. Im dłużej obcuje z płytką R4, tym mniej mam ochoty na żarty, może więc projektant nie poświęcił jej zbyt dużo czasu...

Płytką obsługującą także interfejs CAN dostępny na wyprowadzeniach D5/CANRX0 (RX), D4/CANTX0 (TX), wymaga tylko dodania drivera magistrali np. SN65HVD230 i dołączenia biblioteki *Arduino_CAN.h*. Niestety nie mam jak sprawdzić działania, więc pozostaje mieć nadzieję, że obsługa CAN obędzie się bez niespodzianek.

Ostatnim interfejsem szeregowym jest SPI, które powinno działać, skoro Arduino deklaruje zgodność z Ethernet Shield2 korzystającym



Rysunek 9. Modyfikacja zasilania RTC

z SPI. W sieci jednak można doszukać się informacji o słabej przepustowości SPI, ale wymaga to pewnie dopracowania bibliotek.

Procesor RA4M1 ma wbudowany zegar RTC, wspierany przez biblioteki Arduino. Niestety zasilanie podtrzymujące działanie RTC przy wyłączonym zasilaniu UNO R4 nie zostało zaimplementowane. Może to sposób na zmuszenie użytkowników do zakupu droższej wersji UNO R4 WiFi, gdzie znalazło się miejsce na pin BATT do podłączenia zewnętrznej baterii? Co prawda przeróbka płytki jest banalna, ale wiąże się z utratą gwarancji i koniecznością wylutowania procesora, aby przeciąć zaznaczoną na **rysunku 9** ścieżkę. Zasilanie z baterii CR2032 umieszczonej w pojemniku, np. 1086 KEYSTONE, należy podłączyć bezpośrednio do wyprowadzeń kondensatora C10, zachowując polaryzację. Można użyć oczywiście superkondensatora podładowanego z napięcia 5 V przez diodę i rezystor ograniczający prąd. Przy okazji do sąsiadujących wyprowadzeń 1 i 2 można dolutować złącze z wyprowadzoną drugą magistralą I²C.

Podsumowanie

Podsumowując – płytka UNO R4 Minima w aktualnej wersji ze względu na założenia projektowe nie pozwala korzystać z pełnych możliwości nieco egzotycznego w naszych warunkach procesora RA4M1. Myślę, że jednak należy ją uznać za kolejny falstart w promowaniu procesorów Renesasa w formacie Arduino, po płytce GR-Sakura z procesorem RX63N. Jeżeli najważniejszym kryterium przy projektowaniu była cena, to można było zastosować tańszy procesor w obudowie QFN40 lub zrezygnować z formatu UNO na rzecz wygodniejszych i tańszych w produkcji płytek w formacie Nano lub nawet MKR oraz nie trzymać się kurczowo standardu 5 V, który powoli wymiera, a działający konwerter 3/5 V nie jest ani drogi, ani trudno osiągalny. Byłoby to sensowne nawet, gdyby konwerter wymagał zastosowania dodatkowej płytki bazowej dla zapewnienia mechanicznej zgodności z formatem UNO – takie przejściówki mechaniczne są przecież bez problemu dostępne także w opracowaniach EP.

Pozostawienie na płytce niewylutowanych, ale przewidzianych złącz, a nawet gołych padów dla baterii RTC, drugiej magistrali I²C i kilku dodatkowych analogowych GPIO także nie zwiększyłoby kosztu UNO R4, a poprawiłoby jego funkcjonalność. I chociaż braki uzupełnione są w wersji UNO R4 WiFi, to jest ona zbyt rozbudowana, aby zastępować nią R4 Minima. Jeżeli jednak chcemy faktycznie zapoznać się z procesorami RA4M, a nie tylko migać diodą LED, to rozsądniejszym wyborem jest fabryczny zestaw uruchomieniowy lub płytka MIKROE RA4M1 Clicker. Można też cierpliwie poczekać na następną wersję R4, chociaż bardziej realne wydaje się powolne zniknięcie płytki z rynku, tak jak to stało się z GR-Sakura, Arduino 101, M0 i podobnymi dużo lepiej zaprojektowanymi.

Adam Tatuś, EP

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
dysytrubutora?
Odblokuj bibliotekę
multimediów
samodzielnie.**

Szczegóły na UlubionyKiosk.pl/media



W ofercie AVT*

AVT6006

Podstawowe parametry:

- moduł Mini z układem typu MachXO2-256 zmontowany na płytce w formacie obudowy DIL32 (600 milów),
- moduł Mega z układem typu MachXO2-1200 zmontowany na płytce w formacie szerokiej obudowy DIL64 (900 milów),
- zintegrowane wszystkie elementy niezbędne do działania zastosowanych układów FPGA,
- moduł User Interface ułatwia podłączenie różnych peryferiów – przycisków, wyświetlaczy, enkoderów, modułu z ESP32 do płytki Mega.

* **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]** – 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>

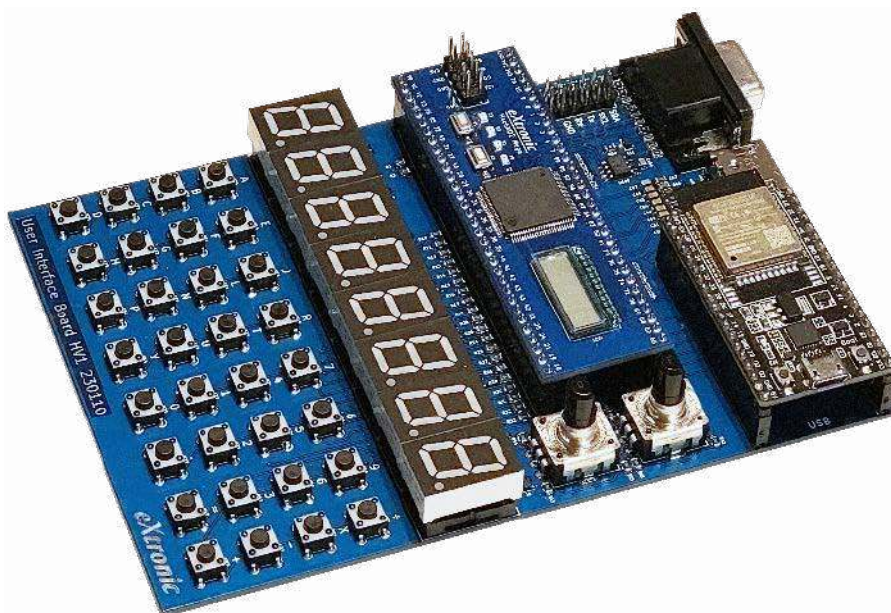
W przypadku braku dostępności na stronie sklepu osoby zainteresowane zakupem płytek drukowanych (PCB) prosimy o kontakt via e-mail: kity@avt.pl.

Płytki rozwojowe do kursu FPGA Lattice

Płytki rozwojowe prezentowane w artykule będą świetną pomocą dydaktyczną do kursu FPGA Lattice mojego autorstwa, który publikowany jest w „Elektronice Praktycznej” od listopada 2021. Ułatwi wyko-
nanie ćwiczeń opisanych w kursie i będą świetną platformą do samodzielnych projektów z zastosowaniem układów FPGA.

MachXO2 Mini

Jest to nieskomplikowany moduł zawierający układ typu MachXO2-256 – najmniejszy, najprostszy i najtańszy układ FPGA z rodziny MachXO2. Choć jego zasoby są nie-
duże, może znaleźć zastosowanie w wielu prostych projektach. Dodatkowym plusem jest to, że występuje w łatwej do przylutowania obudowie QFN32. Ponadto układ zawiera wbudowaną pamięć Flash oraz generator sygnału zegarowego, a żeby rozpocząć pracę, potrzebuje jedynie zasilania i niczego więcej.



Istnieje możliwość, by w module zastosować układ MachXO2-1200, który jest dostępny również w tej niewielkiej obudowie QFN32 i jest kompatybilny pod względem

wyprowadzeń. Dzięki temu zyskamy dużo więcej zasobów logicznych.

Zmontowany moduł Mini (wraz z omawianym w dalszej części artykułu modulem

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

Moduł Mini

Rezystory:

R1...R4: 0 Ω (SMD0603)
R5, R6: 1 kΩ (SMD0603)
R7: 0 Ω (SMD0603) – nie montować

Kondensatory:

C1: 1 μF (SMD0603)
C2...C9: 100 nF (SMD0603)

Półprzewodniki:

U1: LCMXO2-256HC-xSG32x (QFN32)
U2: generator 25 MHz (SMD 3,2x2,5 mm)

Pozostałe:

L1: koralik ferrytowy 600 Ω @ 100 MHz (SMD0603)
J1, J2: listwa goldpin 1x16, raster 2,54 mm
K1: listwa goldpin SMD 2x5, raster 1,27 mm

Moduł Mega

Rezystory:

R1...R4: 0 Ω (SMD0603)
R5, R6: 1 kΩ (SMD0603)
R7: 0 Ω (SMD0603) – nie montować
R21...R32: 4,7 kΩ (SMD0603)
R40...R43: 470 Ω (SMD0603)

Kondensatory:

C1, C2: 1 μF (SMD0603)
C3...C11: 100 nF (SMD0603)
C21...C32: 10 nF (SMD0603)

Półprzewodniki:

D0...D3: dioda LED (SMD0805)
U1: LCMXO2-1200HC-xTG100x (QFP100)
U2: generator 25 MHz (SMD 3,2x2,5 mm)

Pozostałe:

L1: koralik ferrytowy 600 Ω @ 100 MHz (SMD0603)
LCD1: wyświetlacz LCD typu S401M16KR
J1, J2: listwa goldpin 1x32, raster 2,54 mm
K0, K1: mikroszwytk SMD 3x6 mm
K2: listwa goldpin SMD 2x5, raster 2,54 mm

User Interface Board

Rezystory:

R1, R2, R20...R27, R90, R111, R121, R131, R141, R151, R161, R171: 1 kΩ (SMD0805)
R10...R17: 100 Ω (SMD0805)
R30...R33, R40...R45, R60...R65: 10 kΩ (SMD0805)
R50, R51, R52: 270 Ω (SMD0805)
R53, R56, R59, R500: 0 Ω (SMD0805)
R54, R55: 4,7 kΩ (SMD0805)
R57, R58: 75 Ω (SMD0805)

R100, R101, R110, R120, R130, R140, R150, R160, R170: 2 kΩ (SMD0805)

Kondensatory:

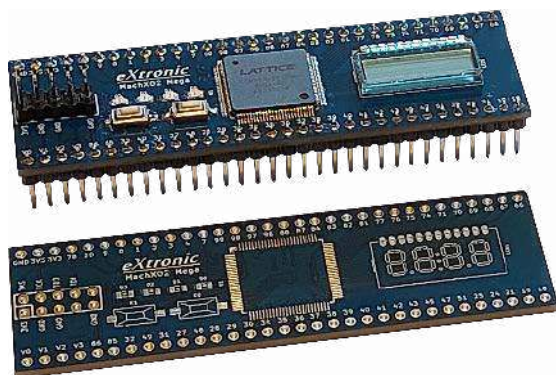
C40...C42, C60...C62, C81, C83, C90, C100: 100 nF (SMD0805)
C80, C82: 10 μF (SMD0805)

Półprzewodniki:

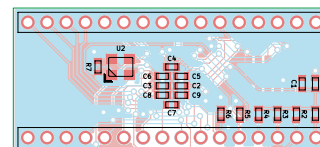
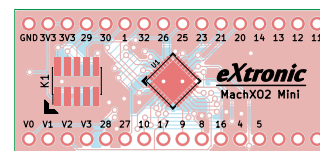
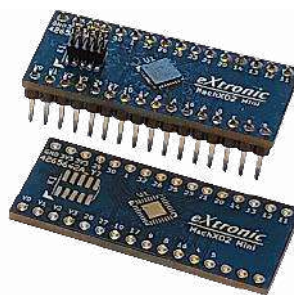
DISP0...DISP7: wyświetlacz 7-seg., 600 milów ze wspólną katodą
T0...T7, T90: BC847 (SOT23)
T50...T51: BSS138 (SOT23)
U80: stabilizator 1117-3,3 (SOT223)
U100: MCP6022 (SO8)

Pozostałe:

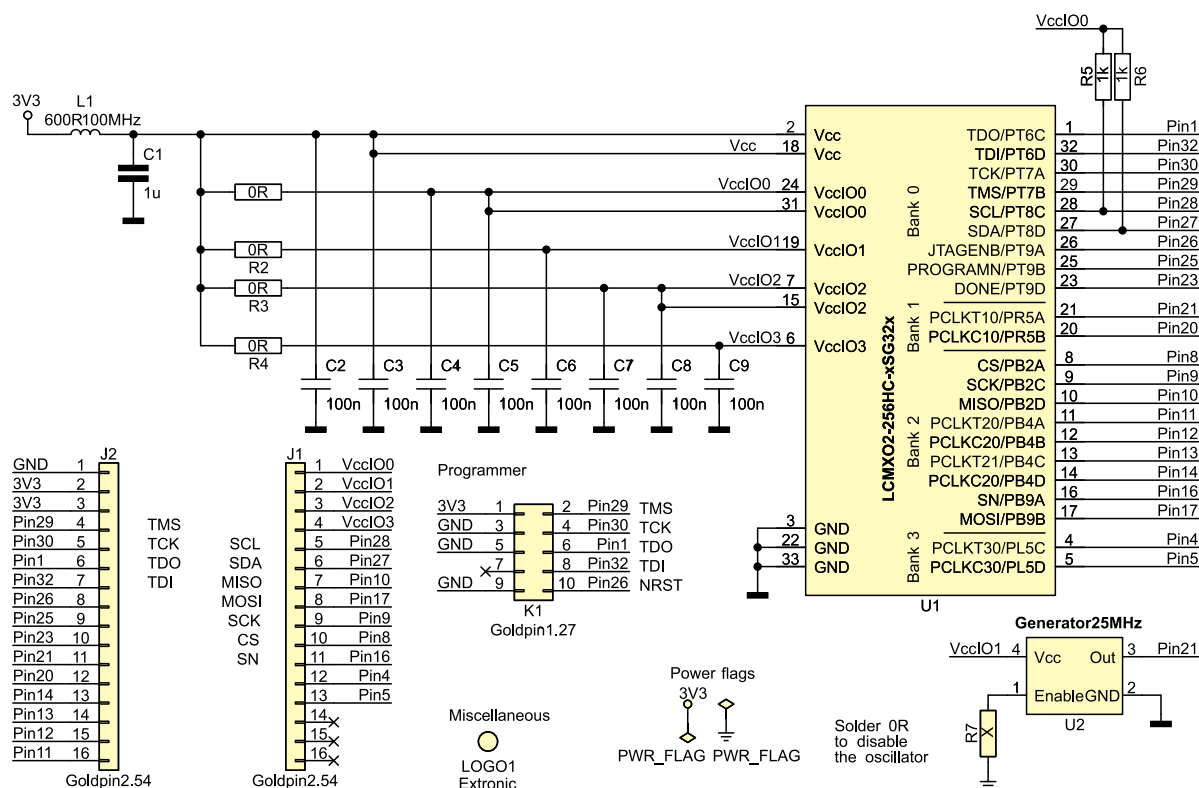
K10...K13, K20...K23, K30...K33, K40...K43, K50...K53, K60...K63, K70...K73, K00...K03 mikroszwytk 6x6 mm
E41, E42: enkoder obrotowy z przyciskiem
K80: gniazdo micro USB
J1, J2: gniazdo goldpin 1x19, raster 2,54 mm
J3, J4: gniazdo goldpin 1x32, raster 2,54 mm
J5: goldpin 2x5, raster 2,54 mm
J50: złącze VGA kątowe
BU90: głośniczek piezo



Fotografia 1. Wygląd modułów Mega i Mini



Rysunek 2. Schemat obwodu PCB modułu MachXO2 Mini



Rysunek 1. Schemat modułu MachXO2 Mini

Mega) pokazano na **fotografii 1**. Kształt i wymiary modułu odpowiadają obudowie DIL32. Dzięki temu możemy moduł łatwo połączyć z innymi elementami, umieszczając go na płytce stykowej lub na innej płytce z użyciem niedrogiej i łatwo dostępnej podstawki DIL32 o szerokości 600 milów. W ten sposób moduł MachXO2 Mini da się łatwo podłączyć oraz wyciągnąć z układu bez konieczności lutowania/rozlutowywania – zupełnie tak, jak układ scalony w obudowie DIL32.

Budowa i działanie

Schemat modułu Mini został pokazany na **rysunku 1**, natomiast schemat obwodu PCB pokazano na **rysunku 2**. Jest on niezwykle prosty. Zasilanie może być doprowadzone z programatora złączem K1 lub poprzez wyprowadzenia przeznaczone do zasilania dostępne na złączach krawędziowych J1 oraz J2.

Układy rodziny MachXO2 można zasilать na różne sposoby. Wyprowadzenia Vcc dostarczają zasilanie do rdzenia układu FPGA. W układach wersji HC (taka została zastosowana w module Mini) napięcie Vcc powinno wynosić 3,3 V lub 2,5 V. Wyprowadzenia VccIO_n (gdzie n jest liczbą od 0 do 3) dostarczają napięcie do czterech banków wyprowadzeń IO. Każdy

z tych banków może być zasilany napięciem o wartości od 1,14 V do 3,6 V. Dzięki takiemu rozwiązaniu układ FPGA może pracować z peryferiami działającymi z różnymi napięciami zasilania i może funkcjonować jako translator napięć.

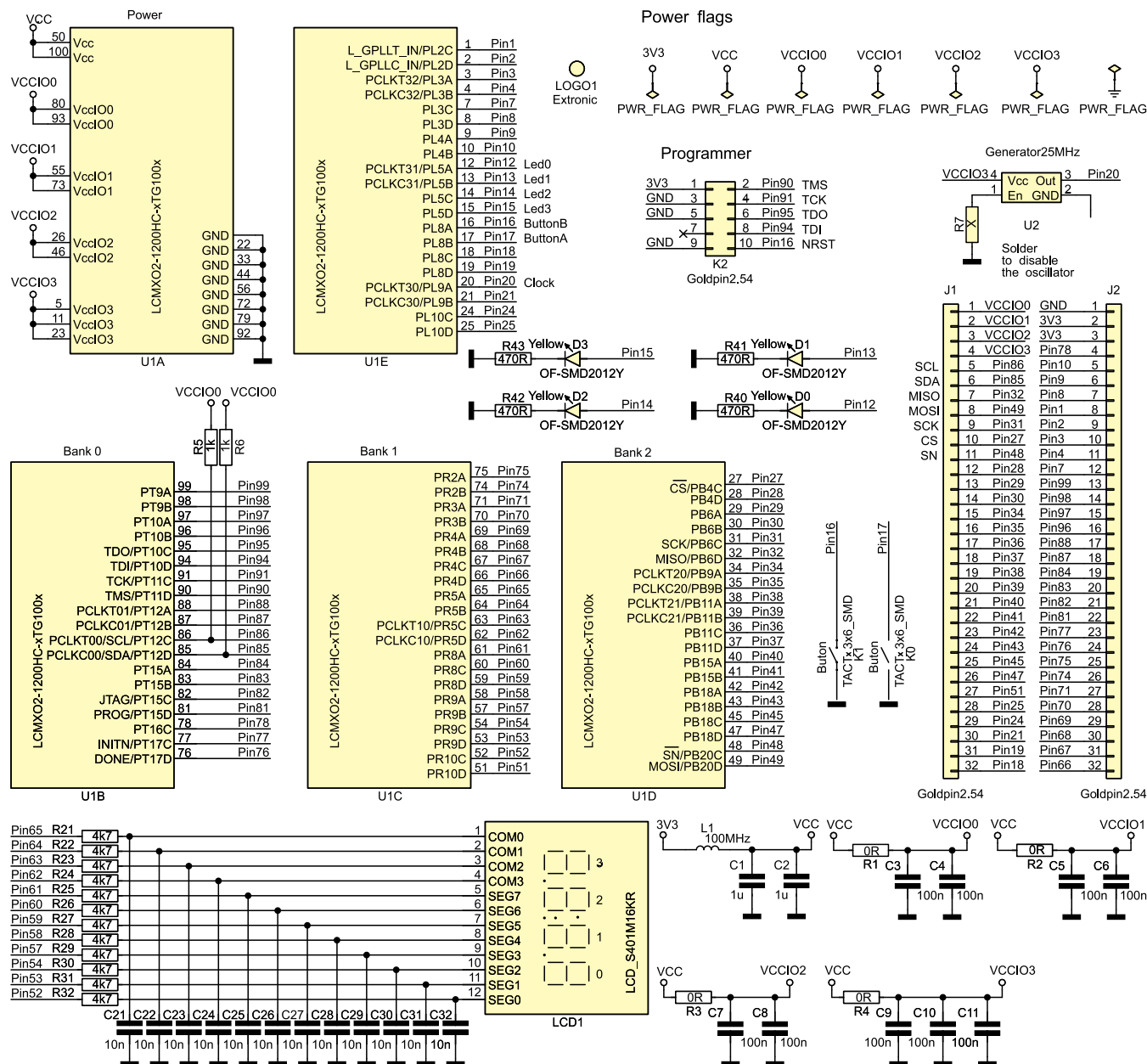
W module Mini domyślnie wszystkie wyprowadzenia Vcc oraz VccIO_n połączone są ze sobą, czyli zasilane są tym samym

Opis programatora JTAG do układów FPGA:

<https://ep.com.pl/projekty/projekty-ep/15365-niedrogi-programator-jtag-do-ukladow-fpga>

Kurs FPGA:

<https://ep.com.pl/kursy/15423-kurs-fpga-lattice-1-wstep>
<https://ep.com.pl/kursy/15444-kurs-fpga-lattice-2-pierwszy-projekt>
<https://ep.com.pl/kursy/15468-kurs-fpga-lattice-3-podstawy-jezyka-verilog>
<https://ep.com.pl/kursy/15515-kurs-fpga-lattice-4-generator-dzielnik-i-licznik>
<https://ep.com.pl/kursy/15555-kurs-fpga-lattice-5-ipexpress-i-inne-gotowce>
<https://ep.com.pl/kursy/15557-kurs-fpga-lattice-6-parametry-i-cwiczenia>
<https://ep.com.pl/kursy/15613-kurs-fpga-lattice-7-analizator-logiczny-reveal>
<https://ep.com.pl/kursy/15685-kurs-fpga-lattice-8-symulacja-w-eda-playground>
<https://ep.com.pl/kursy/15717-kurs-fpga-lattice-9-wyswietlacz-multiplexowany>
<https://ep.com.pl/kursy/15759-kurs-fpga-lattice-10-klawiatura-matrycowa-i-maszyna-stanow>



Rysunek 3. Schemat modułu MachXO2 Mega

napęciem. Jeżeli chcemy skorzystać z możliwości zasilania różnymi napięciami, należy odlutować rezystory o zerowej rezystancji R1...R4, aby odłączyć wybrane wyprowadzenie VccIO od Vcc, a następnie żądane napięcie należy dostarczyć poprzez wyprowadzenia 1...4 na złączu J1.

Koralik ferrytowy L1 oraz kondensatory C1...C9 pełnią funkcję filtrów zasilania. Rezystory R5 i R6 to pull-upy dla interfejsu I²C, który jest dostępny na pinach 27 i 28. Mowa tu o interfejsie I²C, który jest wbudowany w strukturę FPGA i pozwala oszczędzać uniwersalne zasoby logiczne. Oczywiście możemy napisać swój własny interfejs I²C w Verilogu lub VHDL i wtedy można korzystać z dowolnych wyprowadzeń.

Generator sygnału zegarowego wybudowany w strukturę FPGA ma niewielką dokładność. Można go stosować z powodzeniem do zadań, które nie wymagają precyzyjnego sygnału

zegarowego, takich jak na przykład multipleksacja wyświetlaczy LED czy LCD. W przypadku bardziej wymagających aplikacji należy zastosować generator kwarcowy. Na płytce dostępny jest generator sygnału zegarowego o częstotliwości 25 MHz. Jeżeli ta częstotliwość nie odpowiada wymaganiom, można ją pomnożyć lub podzielić za pomocą bloku PLL wbudowanego w FPGA. Generator jest cały czas włączony i dostarcza sygnał zegarowy do pinu 21. Jeżeli chcemy to wyprowadzenie użyć do innego celu, należy w miejscu R7 przylutować rezystor o zerowej rezystancji (lub drukik).

Mach XO2 Mega

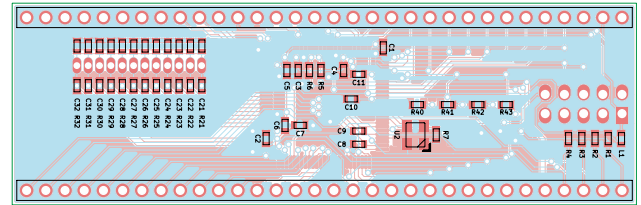
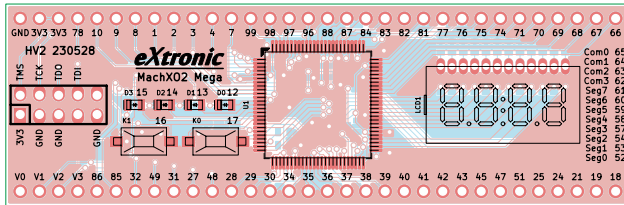
Moduł Mega znajdzie zastosowanie w nieco bardziej zaawansowanych aplikacjach, gdzie płytka Mini nie jest wystarczająca. Schemat modułu Mega został pokazany na **rysunku 3**, a na **rysunku 4** znajduje się

schemat obwodu PCB. Kształt płytki odpowiada obudowie DIL64 o szerokości 900 milów.

Budowa i działanie

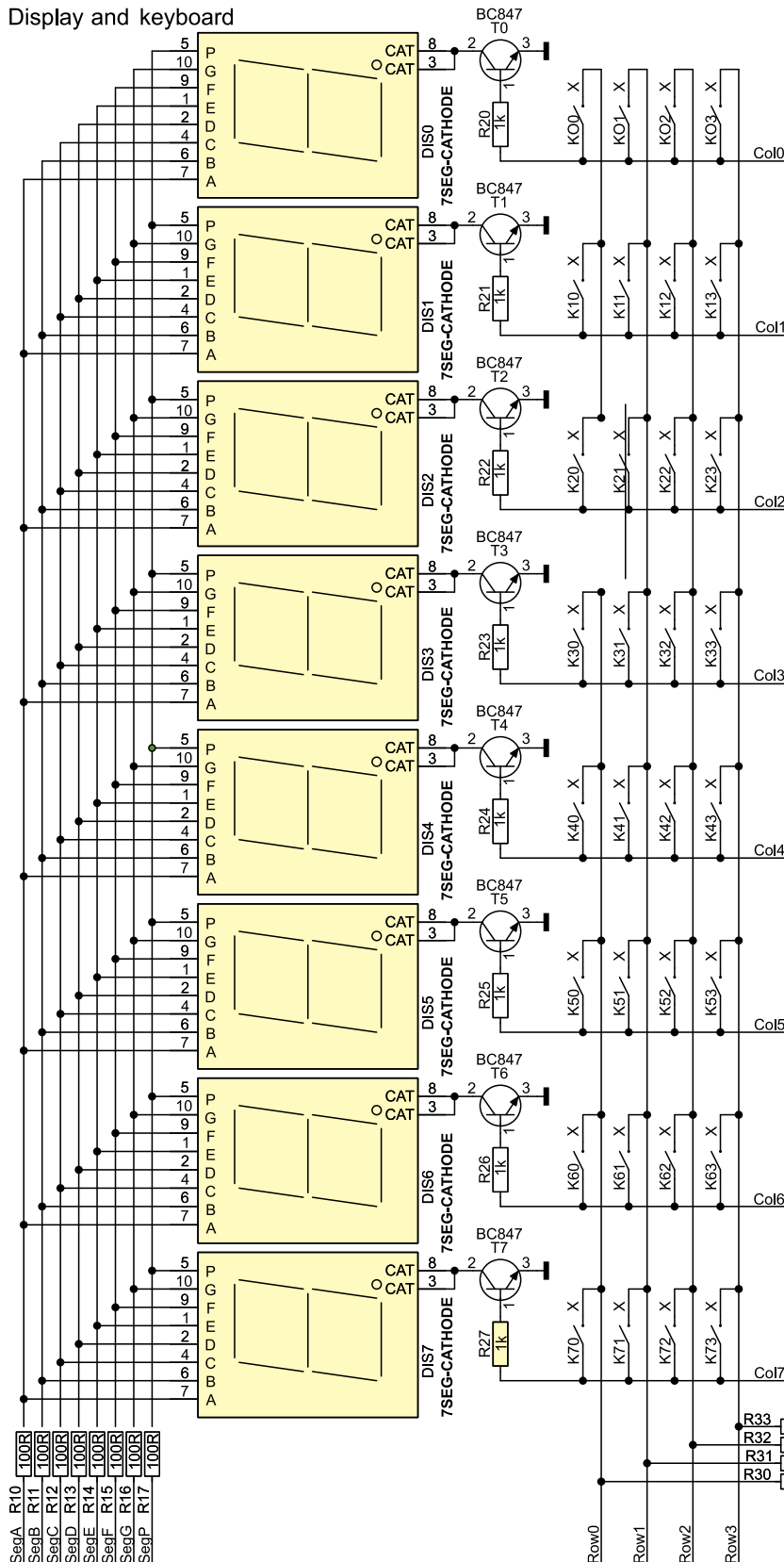
Schemat jest dość podobny do schematu płytki Mini, więc omówimy tylko różnice między nimi. Zastosowany układ FPGA to MachXO2-1200 w obudowie TQFP100. Jest on wyposażony w dużo więcej zasobów logicznych niż maluszek z poprzedniej płytki, a także ma zdecydowanie więcej wyprowadzeń.

W module znalazły się cztery diody LED oraz dwa przyciski. Jeden z nich podłączony jest do linii resetującej NRST połączonej do gniazda programatora. Ponadto płytka została wyposażona w mały wyświetlacz LCD, umożliwiający wyświetlanie czterech cyfr. Obsługa wyświetlacza LCD znacząco różni się od multipleksowanego wyświetlacza LED. Ten

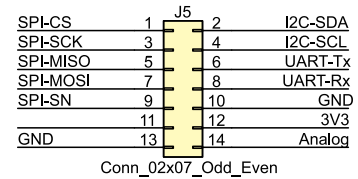
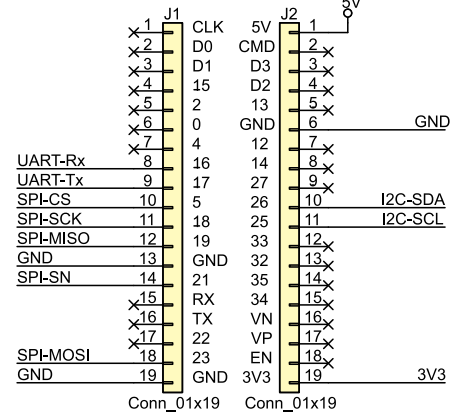


Rysunek 4. Schemat obwodu modułu MachXO2 Mega

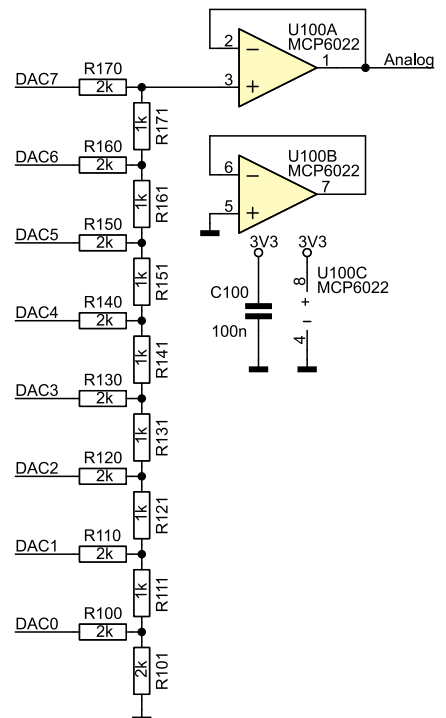
Display and keyboard



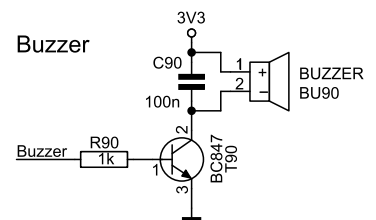
ESP32 DEVKIT C V4



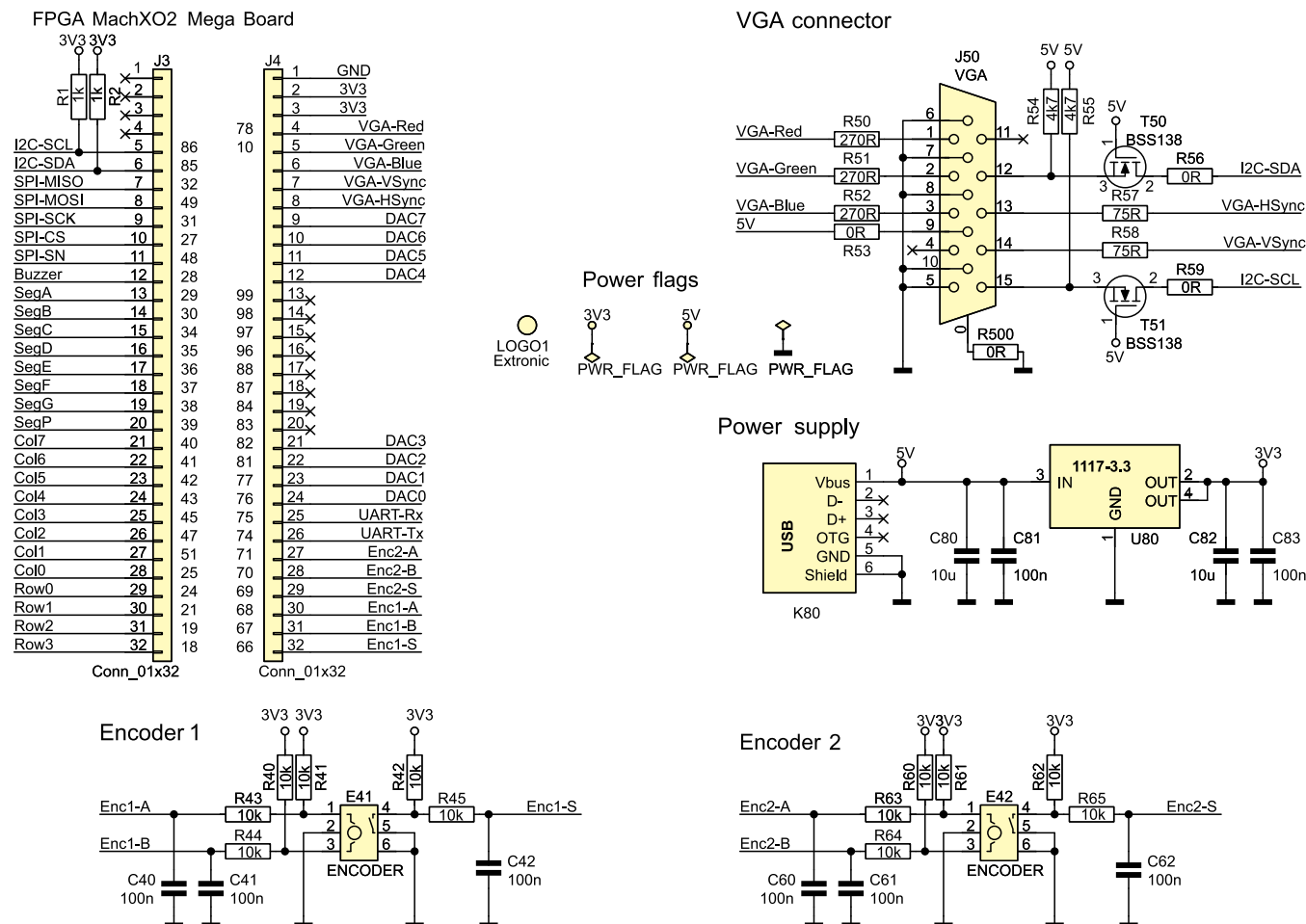
Digital to analog converter



Buzzer



Rysunek 5. Schemat płytki User Interface Board



Rysunek 5. Schemat płytki User Interface Board – cd.

temat będzie omówiony dokładnie w 13 odcinku kursu.

LED multiplexowany. Składa się z ośmiu wyświetlaczy 7-segmentowych ze wspólną katodą o szerokości 600 milsów – czyli tych najbardziej

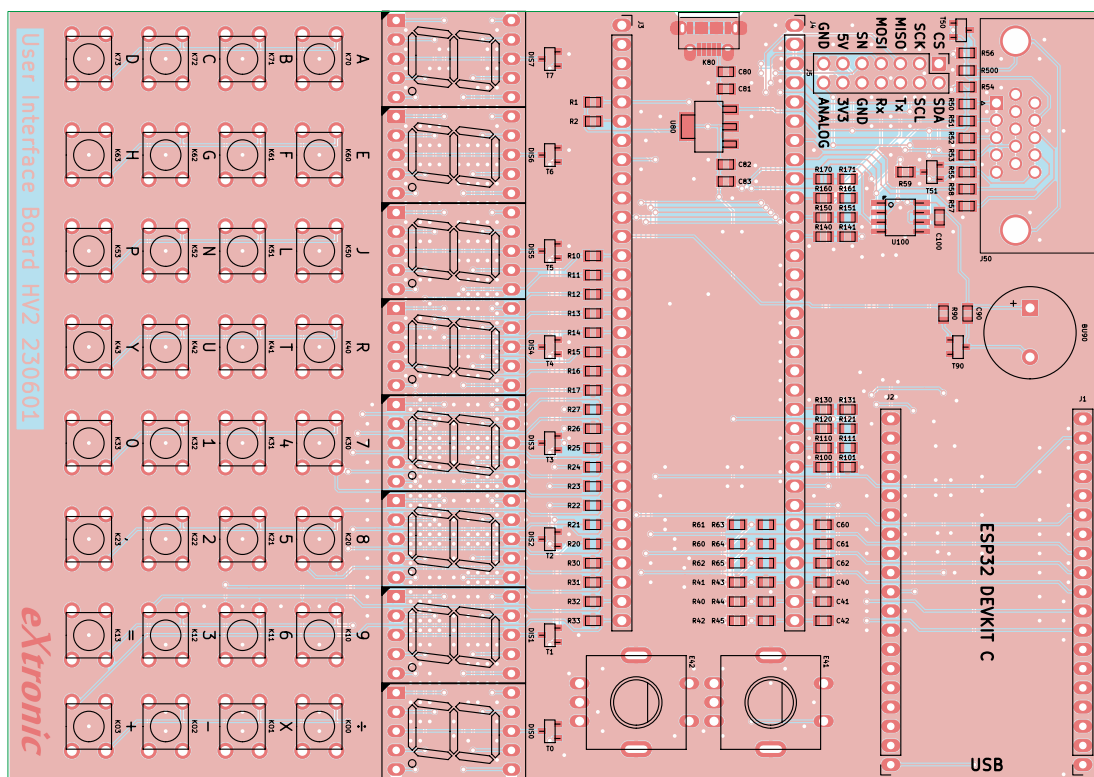
popularnych, które są dostępne we wszystkich kolorach. Obsługa wyświetlacza tego typu została omówiona w 9 odcinku kursu.

User Interface Board

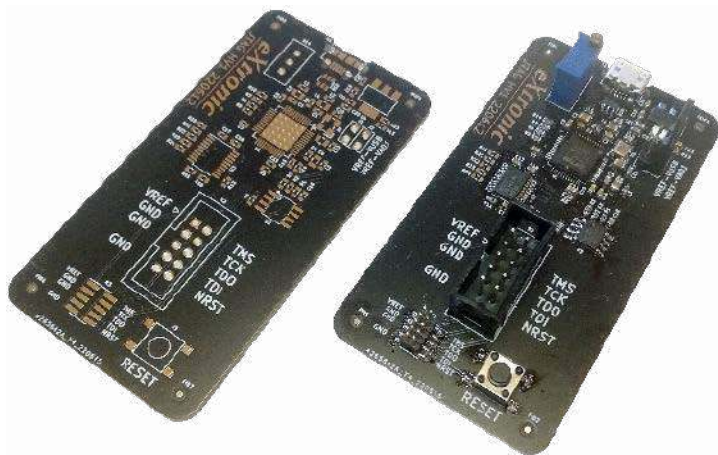
Moduł User Interface powstał, aby ułatwić podłączenie różnych peryferiów do płytki Mega – wystarczy ją umieścić w gnieździe J3 i J4 w taki sposób, jak pokazano na fotografii tytułowej. Korzystając z tej płytki, można łatwo wykonać ćwiczenia opisane w kursie od odcinka 9. Płytkę zawiera szereg różnych peryferiów, które będą omawiane w wielu odcinkach kursu.

Budowa i działanie

Schemat płytki User Interface pokazano na rysunku 5, a schemat obwodu PCB na rysunku 6. W centralnej części płytki umieszczono duży i czytelny wyświetlacz



Rysunek 6. Schemat obwodu PCB płytki User Interface Board



Fotografia 2. Programator JTAG współpracujący z opisanymi płytkami

Nieco poniżej widzimy klawiaturę matrycową, składającą się z 32 przycisków. Temat klawiatury został omówiony w 10 odcinku kursu. Połączenie wyświetlacza multiplexowanego z klawiaturą matrycową sprawia, że korzystając z zaledwie 20 wyprowadzeń FPGA, możemy obsługiwać 8-cyfrowy wyświetlacz i 32 przyciski.

W lewym dolnym rogu schematu widzimy głośniczek BU90, sterowany poprzez tranzystor T90. Sposób generowania sygnałów dźwiękowych zostanie omówiony w 14 odcinku kursu.

Płytkę wyposażono w dwa enkodery obrotowe E41 oraz E42, które omówimy dokładnie w 15 odcinku kursu. Oba enkodery zostały połączone w taki sam sposób – daje to okazję, by utworzyć dwie instancje jednego modułu enkodera opisanego w języku Verilog. Enkoder ma wyjścia A i B, które dostarczają informacji o obracaniu pokrętki. Pokrętło pełni jednocześnie funkcję przycisku – wciśnięcie go powoduje zwarcie wyprowadzenia S do masy.

Rezystory R40, R41, R42 i odpowiednio R60, R61, R62 służą jako pull-upy, które mają zapewnić stan wysoki na wejściach FPGA w momencie, kiedy wewnętrzny mechanizm enkodera nie łączy wyjść A lub B z masą. To samo dotyczy wyjścia S. Pary rezystorów i kondensatorów R43, C40; R44, C41 oraz R45-C42 tworzą filtry dolnoprzepustowe, których zadaniem jest eliminacja drgań styków. Oczywiście wszystkie te rezystory i kondensatory można by pominąć, a zamiast nich zastosować wewnętrzne pull-upy wbudowane w FPGA, a drgania styków wyeliminować w taki sposób, jak to omówiono w 6 odcinku kursu – jednak wymagałoby to poświęcenia większej ilości zasobów logicznych układu FPGA.

W prawym dolnym rogu schematu znajduje się 8-bitowy przetwornik cyfrowo-analogowy w postaci drabinki rezystorowej R-2R. Umożliwia wygenerowanie dowolnego napięcia od zera do napięcia zasilania 3,3 V z rozdzielczością 1/256, czyli 0,013 V. Wzmacniacz operacyjny U100A pracuje jako wtórnik napięcia. Został zastosowany z tego powodu, że drabinka R-2R ma dużą rezystancję i w rezultacie obciążalność wyjścia drabinki jest bardzo

mała. Wtórnik napięcia ma wzmocnienie równe 1, czyli napięcie na jego wejściu i wyjściu jest takie samo, ale umożliwia pobranie dużo większego prądu niż ten, jaki można by pobrać z drabinki. Zastosowano układ MCP6022 – jest to tani wzmacniacz rail-to-rail ogólnego przeznaczenia dostępny w obudowach SO8 oraz DIL8. Wyjście Analog dostępne jest poprzez gniazdo J5 typu goldpin. Można je podłączyć na przykład do oscyloskopu.

W gnieździe J5 znajdziemy jeszcze kilka sygnałów związanych z interfejsami UART, SPI oraz I²C. Są one połączone zarówno do FPGA, jak i do devboarda ESP32. Celem tego rozwiązania jest ułatwienie ćwiczeń z tymi interfejsami. Użytkownik może łatwo zaprogramować ESP32 językiem C lub Python i komunikować się z FPGA. W gniazdo J5 można podłączyć analizator logiczny lub oscyloskop, aby podsłuchiwać i weryfikować, czy komunikacja zachodzi prawidłowo. Alternatywnie, można nie umieszczać devboarda ESP32 i zamiast niego podłączyć kabelkami inną płytkę rozwojową. Ja gorąco polecam moduł ESP32, ponieważ jest tani i ma wielkie możliwości, a za pomocą

MicroPythona można bardzo łatwo i szybko napisać całkiem rozbudowany program.

Ostatnia rzecz wymagająca omówienia to gniazdo VGA. Dzięki niemu będziemy mogli podłączyć monitor. Wydawać się może, że standard VGA obecnie już jest martwy i odszedł do lamusa. Jest w tym dużo prawdy, ale jest to standard bardzo prosty i świetnie nadaje się na początek zabaw z generowaniem obrazu w FPGA. Poza tym dobrze jest opanować VGA przed próbą okiełznania HDMI, które jest bardziej skomplikowane.

Ponadto mamy możliwość użycia pinów 12 i 15 złącza VGA, które zwykle są pomijane przez hobbystów. Pod tymi pinami kryje się interfejs I²C, pozwalający na identyfikację podłączonego monitora. Interfejs I²C w przewodzie VGA działa z napięciem 5 V, a MachXO2 działa przy napięciu 3,3 V. Z tego powodu zastosowano tranzystory MOSFET-N typu BSS138, które funkcjonują jako translatory napięcia. Jeżeli z jakiegoś powodu konieczne będzie odcięcie linii I²C od FPGA, należy wtedy odlutować rezystory zerowe R56 i R59.

To nie wszystko

Płytki Mega oraz Mini mogą być programowane za pomocą programatora JTAG, który opisałem w EP 2022/09 – **fotografia 2**. Zawiera on produkowany przez firmę FTDI układ scalony FT232RL, który jest przejściówką USB/JTAG. Współpracuje z Lattice Diamond, Lattice Radiant, a także wieloma innymi rozwiązaniami open source, jak na przykład OpenOCD.

W dalszych planach jest opracowanie dodatkowych płytek, które umożliwiłyby użycie złącza HDMI, wyświetlacza LCD 14-segmentowego, a także płytek z mocniejszymi układami FPGA zgodnymi z User Interface Board.

Dominik Bieczynski
leonow32@gmail.com

Repozytorium modułów wykorzystywanych w kursie:
<https://github.com/leonow32/verilog-fpga>

REKLAMA

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

WWW.AKSOTRONIK.COM.PL

Aksotronik
ELEMENTY ELEKTRONICZNE

- Magnesy neodymowe oraz ferrytowe
Ceny od 0.10zł
- Przełączniki klawiszowe wodoszczelne, pyłoszczelne
Ceny od 2.40zł
- Druty oporowe od 0.16 do 0.51mm
Ceny od 5.70zł
- Prowadniki do przewodów
Ceny od 11.00zł
- Kostki elektryczne zaciskowe
Ceny od 0.22zł
- Szczetki węglowe do elektronarzędzi
Ceny od 2.60zł/kpl
- Przełączniki do elektronarzędzi zwykłe i elektromagnetyczne
Ceny od 7.00zł
- Złącza hermetyczne Superseal
Ceny od 1.10zł/kpl
- Podkładki/organizery
Ceny od 0.95zł
- Zestawy śrubek M2, M3 z nakrętkami i podkładkami
Ceny od 2.50zł

Uwaga!!! Powyższe ceny dotyczą zakupów minimalnych ilości hurtowych, poprzez nasz sklep internetowy.
W swojej ofercie posiadamy m.in.: półprzewodniki (diody, układy scalone, tranzystory, triaki, elementy optoelektryczne), elementy dystansowe, złącza, przełączniki, elementy akustyczne, rezystory, kondensatory, kwarce, podstawki, moduły Arduino

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

**Podstawowe parametry:**

- zakres regulacji tonów niskich: ± 14 dB; wysokich: ± 14 dB,
- zakres regulacji głośności: $-47...0$ dB; balansu: ± 79 dB,
- zakres regulacji wzmocnienia multipleksera wejściowego: $0...28$ dB,
- separacja kanałów stereo: 90 dB,
- odstęp sygnału od szumu S/N: 106 dB,
- zniekształcenia harmoniczne THD: 0,1%,
- impedancja wejściowa: 100 k Ω ; wyjściowa: 30 k Ω ,
- minimalna impedancja obciążenia: 2 k Ω ,
- maksymalne napięcie wejściowego sygnału audio: 2 VRMS,
- napięcia zasilania: 7...10 V; prąd obciążenia: 120 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:

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

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

- AVT5975 Regulator barwy dźwięku, głośności i balansu (EP 3/2023)
- AVT5873 Stereofoniczny aktywny regulator głośności (EP 8/2021)
- AVT5851 7-pasmowy korektor graficzny (EP 4/2021)
- AVT5816 Regulator balansu tonów (EP 10/2020)
- AVT5637 Wielokanałowy regulator głośności VCA (EP 8/2018)
- AVT5629 Cyfrowy regulator głośności z układem PT2257 (EP 6/2018)
- AVT3222 Sterowany dowolnym pilotem potencjometr audio z przekaźnikiem (EdW 5/2018)
- AVT1979 Korektor barwy dźwięku (EP 11/2017)
- AVT1971 Stereofoniczny regulator barwy tonu zasilany z baterii (EP 9/2017)
- AVT1972 Potencjometr „Panorama” audio (EP 9/2017)

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*

AVT6005

toneCtrl (1)

– regulator barwy dźwięku

Prezentujemy nowoczesny regulator barwy dźwięku wyposażony w dodatkowe funkcjonalności, który może znaleźć zastosowanie zarówno przy implementacji nowoczesnych rozwiązań we współczesnych urządzeniach audio, jak i umożliwić modernizację i poprawę parametrów elektrycznych systemów spod znaku vintage.

Mimo że świat współczesny zdominowany jest przez wszechogarniającą technikę cyfrową, gdzie praca z sygnałem sprowadza się do stosowania zaawansowanych algorytmów z dziedziny DSP, nadal rozwijanych jest wiele nowych urządzeń audio, dla których kluczowa jest implementacja dobrej jakości rozwiązań sprzętowych. Co więcej, istnieje także spora grupa użytkowników, którzy z powodzeniem modernizują lub też utrzymują w doskonałym stanie technicznym urządzenia audio z minionej epoki, podtrzymując wydatnie bieżącą modę na systemy spod znaku vintage. Nie powinno to specjalnie dziwić, gdyż dzisiejsza wirtuozeria w domenie cyfrowej znajduje także swoje odzwierciedlenie w minionej epoce spod znaku analogu, tyle że tym razem w materii sprzętowej.

W gruncie rzeczy czasami zastanawiam się, co było większym osiągnięciem inżynierii i ludzkiej inteligencji, wszak zaawansowane rozwiązania układowe i programistyczne charakterystyczne dla obu tych epok niejednokrotnie zaskakują pomysłowością

i poziomem zaawansowania technologicznego. Sam, projektując różnorakie systemy elektroniczne, nie tak rzadko spotykałem się z koniecznością rozwiązywania nietypowych problemów dla obu wspomnianych domen, cyfrowej i analogowej. Niemniej jednak zdecydowana większość moich projektów związana jest nieodłącznie ze światem mikrokontrolerów, przez co za każdym razem, gdy sięgam do rozwiązań z pogranicza „cyfry i analogu”, odczuwam podwójną satysfakcję. Jako że z rozrównieniem wspominać początki swojej przygody z elektroniką, kiedy to, zapewne jak każdy w tym czasie, konstruowałem proste układy analogowe, tym razem postanowiłem wykonać ukłon w kierunku tego rodzaju systemów, no może z pogranicza obu domen.

Postanowiłem zbudować nowoczesny regulator barwy dźwięku wyposażony w dodatkową funkcjonalność, który w swoich założeniach ma znaleźć zastosowanie zarówno przy implementacji nowoczesnych rozwiązań we współczesnych urządzeniach audio, jak i umożliwić modernizację



i poprawę parametrów elektrycznych systemów spod znaku vintage. Jako że moim celem nie była implementacja systemu złożonego wyłącznie z elementów dyskretnych, proces projektowania rozpocząłem od poszukiwania nowoczesnego układu, przy udziale którego zrealizuję założone cele.

Mając już pewne doświadczenie w realizacji tego rodzaju urządzeń, jak i odpowiednie rozeznanie w kwestii stosowanych peryferiów, dość szybko dokonałem stosownego wyboru, którym stał się zaawansowany, scalony regulator barwy dźwięku pod postacią układu TDA7440 produkowanego przez firmę STMicroelectronics. Prawdę mówiąc, element ten od dawna zagrzewał miejsce w mojej szufladzie pełnej elektronicznych

Wykaz elementów, kupuj na stronie sklep.avt.pl (Warszawa, ul. Leszczyńska 11, tel. +48222578451, e-mail: handlowy@avt.pl)**Rezystory:** (SMD0805)

R1, R2: 10 k Ω
R3, R4: 4,7 k Ω
R5, R6: 5,6 k Ω

Kondensatory:

C1...C5, C7, C8, C13, C14, C17, C18, C22: 100 nF ceramiczny X7R (SMD0805)
C6: 22 μ F/16 V tantalowy (B/3528-21 W)

C9, C10, C20, C21: 1 μ F ceramiczny X7R (SMD0805)C11, C15: 2,2 μ F ceramiczny X7R (SMD0805)

C12, C16: 5,6 nF ceramiczny X7R (SMD0805)

C19: 10 μ F ceramiczny X7R (SMD0805)**Półprzewodniki:**

U1: TDA7440 (SO28)

U2: ATTiny85 (SOIC8)

U3: 78M05 (DPAK)

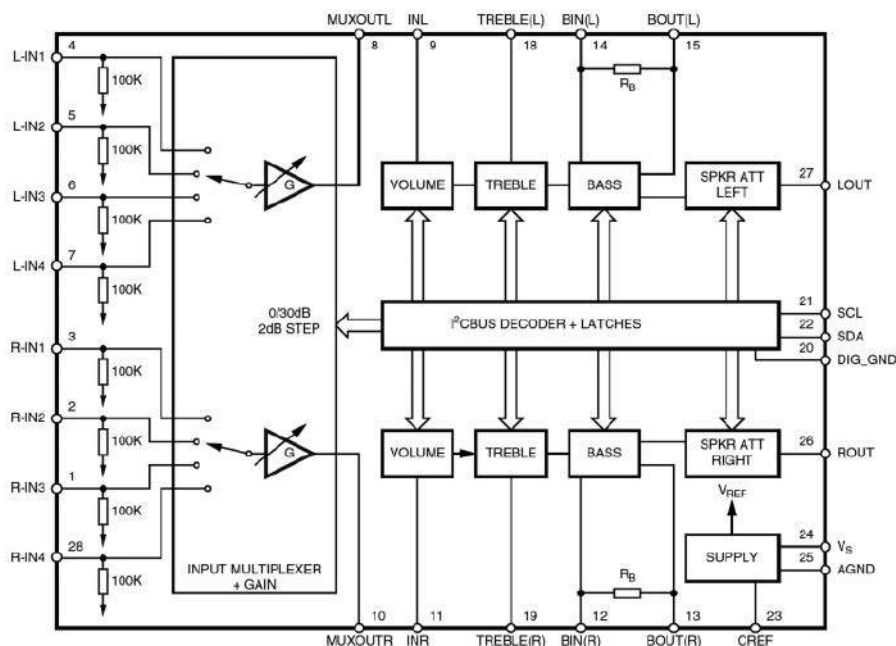
LED1...LED15: dioda adresowalna LED RGB typu OSTW-2020C1E (SMD2020)

Pozostałe:

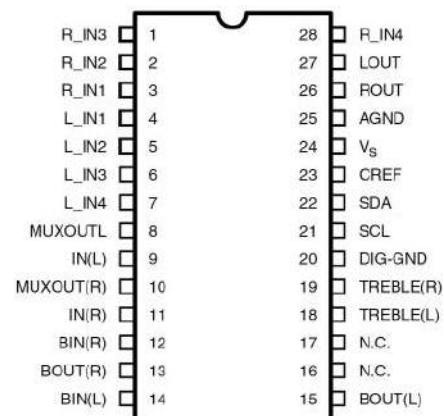
IN, OUT: gniazdo męskie proste 3 piny typu NSL25-3 W

PWR: gniazdo męskie proste 2 piny typu NSL25-2 W

MFK: enkoder SMD z przyciskiem typu ALPS EC11J1524413



Rysunek 1. Uproszczony schemat funkcjonalny układu TDA7440 (za dokumentacją STMicroelectronics)



Rysunek 2. Wygląd obudowy układu TDA7440 wraz z rozmieszczeniem wszystkich wyprowadzeń (za dokumentacją STMicroelectronics)

różności. Układ ten, sterowany dzięki wbudowanemu interfejsowi standardu I²C, umożliwia regulację kilku parametrów audio i dodatkowo zapewnia utrzymanie doskonałych parametrów elektrycznych. Wspomniane peryferium charakteryzuje się następującymi, wybranymi cechami użytkowymi:

- 4-kanałowy, stereofoniczny selektor wejściowy z możliwością regulacji wzmacnienia w zakresie 0 dB...+30 dB (z krokiem 2 dB),
- regulacja tonów niskich w zakresie -14...+14 dB (z krokiem 2 dB),
- regulacja tonów wysokich w zakresie -14...+14 dB (z krokiem 2 dB),
- regulacja głośności w zakresie -47...0 dB (z krokiem 1 dB),
- regulacja balansu każdego kanału w zakresie -79...0 dB (z krokiem 1 dB),
- funkcja wyciszenia (MUTE),
- wysoki, maksymalny poziom sygnału wejściowego 2 VRMS,
- doskonały odstęp sygnału od szumu S/N równy 106 dB,
- niskie zniekształcenia harmoniczne THD rzędu 0,1%,
- niewielka liczba niezbędnych dyskretnych elementów zewnętrznych (w większości konieczne do ustalenia parametrów wbudowanych filtrów).

Budowa i działanie

Jak widać, układ TDA7440 cechuje się doskonałymi właściwościami użytkowymi oraz parametrami elektrycznymi i świetnie wpisuje się w założenia naszego projektu. Aby zrozumieć zasadę działania oraz paletę możliwości regulacyjnych układu TDA7440, najlepiej jest spojrzeć na jego uproszczony schemat funkcjonalny pokazany na **rysunku 1**. Dalej, na **rysunku 2**, pokazano wygląd obudowy układu TDA7440 wraz z rozmieszczeniem wszystkich wyprowadzeń, zaś w **tabeli 1** pokazano z kolei rozkład i opis jego wyprowadzeń.

Co oczywiste, parametry elektryczne filtrów odpowiadających za regulację

Tabela 1. Rozkład i opis wyprowadzeń układu TDA7440		
Nr	Nazwa	Opis
1	R_IN3	Wejście multipleksera wejściowego nr 3, kanał prawy
2	R_IN2	Wejście multipleksera wejściowego nr 2, kanał prawy
3	R_IN1	Wejście multipleksera wejściowego nr 1, kanał prawy
4	L_IN1	Wejście multipleksera wejściowego nr 1, kanał lewy
5	L_IN2	Wejście multipleksera wejściowego nr 2, kanał lewy
6	L_IN3	Wejście multipleksera wejściowego nr 3, kanał lewy
7	L_IN4	Wejście multipleksera wejściowego nr 4, kanał lewy
8	MUXOUT(L)	Wyjście multipleksera wejściowego, kanał lewy
9	IN(L)	Wejście regulatora głośności, kanał lewy
10	MUXOUT(R)	Wyjście multipleksera wejściowego, kanał prawy
11	IN(R)	Wejście regulatora głośności, kanał prawy
12	BIN(R)	Wejście bloku filtra tonów niskich (do podłączenia elementów zewnętrznych), kanał prawy
13	BOUT(R)	Wyjście bloku filtra tonów niskich (do podłączenia elementów zewnętrznych), kanał prawy
14	BIN(L)	Wejście bloku filtra tonów niskich (do podłączenia elementów zewnętrznych), kanał lewy
15	BOUT(L)	Wyjście bloku filtra tonów niskich (do podłączenia elementów zewnętrznych), kanał lewy
16	N.C.	Niepodłączone
17	N.C.	
18	TREBLE(L)	Wyprowadzenie bloku filtra tonów wysokich (do podłączenia elementów zewnętrznych), kanał lewy
19	TREBLE(R)	Wyprowadzenie bloku filtra tonów wysokich (do podłączenia elementów zewnętrznych), kanał prawy
20	DIG_GND	Masa części cyfrowej układu
21	SCL	Sygnał zegarowy magistrali danych I ² C
22	SDA	Sygnał danych magistrali danych I ² C
23	CREF	Wejście do podłączenia kondensatora filtrującego bloku zasilającego
24	VS	Napięcie zasilania (9 V)
25	A_GND	Masa części analogowej układu
26	ROUT	Wyjście układu, kanał prawy
27	LOUT	Wyjście układu, kanał lewy
28	RIN_4	Wejście multipleksera wejściowego nr 4, kanał prawy

tonów niskich i wysokich (w tym topologię filtra, jeśli chodzi o filtr tonów niskich) możemy modyfikować, dobierając odpowiednie wartości (i konfigurację) elementów zewnętrznych, jednak w moim urządzeniu zastosuję parametry domyślne sugerowane przez aplikację producenta układu. Osoby chcące poeksperymentować z tego rodzaju ustawieniami odsyłam do noty aplikacyjnej, gdzie w sposób szczegółowy opisany został sposób doboru (i stosowne wzory) zewnętrznych elementów dyskretnych. Tyle na ten moment, jeśli chodzi o nasz element „regulacyjny”, który, co oczywiste, wymaga współistnienia jakiegogo systemu odpowiedzialnego za przeprowadzanie i wyświetlanie wyników regulacji.

Do sterowania, jak zapewne się domyślicie, użyję jakiegogo niewielkiego mikrokontrolera, ale co z wyświetlaniem wprowadzonych nastaw? Najprościej byłoby zastosować jakiś niewielki wyświetlacz, ale że nasze urządzenie ma znaleźć zastosowanie również przy modyfikacjach istniejących systemów audio, implementacja jakiegokolwiek wyświetlacza nie zawsze będzie możliwa, a już na pewno nie zawsze pożądana, jeśli chcemy zachować wygląd w duchu retro.

Postanowiłem ostatecznie, że jako element interfejsu użytkownika prezentujący wprowadzone nastawy użyję szeregu diod LED zgrupowanych w postaci wycinka okręgu wokół elementu regulacyjnego (enkodera) co powinno wyglądać bardzo efektownie a jednocześnie nie nazbyt nowocześnie. Ale jak na linijce diod LED pokazać szereg wartości regulacyjnych? Najprostszym sposobem jest zastosowanie kilku odrębnych kolorów, każdy dla innego parametru, by już po samym kolorze użytkownik urządzenia wiedział, jaki parametr podlega regulacji. Prawda, że proste? Co oczywiste, można również zastosować kilka pokręteł, każde dla innego parametru i każde wyposażone w szereg diod LED, ale przecież chcemy skonstruować urządzenie kompaktowe, które łatwo zaadaptujemy do istniejących rozwiązań, nieprawdaż? W takim razie najprościej zastosować diody LED RGB, których kolor możemy dowolnie modyfikować przy użyciu mikrokontrolera.

W tym miejscu stanąłem przed wyzwaniem wyboru odpowiednich elementów wykonawczych, a więc sterownika diod LED, jak i samych diod. Jak wiemy, aby płynnie sterować kolorem diody LED typu RGB, należy zastosować 3-kanalowy sterownik PWM. Wynika z tego, że skoro przewiduję zastosowanie 15 elementów tego typu, to liczba niezbędnych kanałów wzrasta nam do 45. Trudno wyobrazić sobie mikrokontroler, który sprostałby tym wymaganiom, a jeszcze trudniej wyobrazić sobie projekt płytki drukowanej o niewielkiej wielkości, na której upakowalibyśmy tyle odrębnych ścieżek. Co oczywiste, można byłoby zastosować jakiegogo rodzaju sterowanie matrycowe, by ograniczyć liczbę koniecznych

połączeń, ale biorąc pod uwagę liczbę wymaganych kanałów PWM i oczekiwaną rozdzielczość takiego sygnału (8 bitów), trudno mi sobie wyobrazić efektywne sterowanie bez użycia dość dużych prądów, by uzyskać wynikową jasność na akceptowalnym poziomie. Zresztą nawet w takim przypadku nie rozwiązuje problemu ze skomplikowaniem rysunku obwodu drukowanego. Pat? Otóż nie.

Dość szybko zdałem sobie sprawę, iż jedynym sensownym rozwiązaniem tego rodzaju problemu konstrukcyjnego będzie zastosowanie adresowalnych diod LED RGB, których konstrukcja pozwala na uniknięcie wszystkich wspomnianych problemów. Diody takie, oprócz wyprowadzeń zasilających, wyposażone są w jakiś szeregowy interfejs komunikacyjny, przy użyciu którego dokonujemy ustawień koloru jej świecenia. Interfejs, o którym mowa, zaimplementowany jest w taki sposób (zarówno w kwestii sprzętowej, jak i logicznej), że diody takie, połączone w łańcuchy, mogą być indywidualnie adresowane, a co za tym idzie, każda z nich ma niezależne sterowanie. Sam przebieg PWM niezbędny do regulacji koloru jej świecenia generowany jest sprzętowo dzięki sterownikowi zabudowanemu w takie element.

Przejdźmy zatem do konkretów. Pierwszą myślą, jaka przyszła mi w tym czasie do głowy, było zastosowanie bardzo popularnych i tanich elementów tego typu pod postacią diod z rodziny WS2811/WS2812, jednak elementy te produkowane są w dość dużych, jak na potrzeby naszej aplikacji, obudowach o rozmiarze 5×5 mm. Kierując się potrzebą znalezienia elementów możliwie najmniejszych, dość szybko natknąłem się na diody adresowalne RGB typu OSTW2020C1E firmy OptoSupply produkowane w niewielkich 4-wyprowadzeniowych obudowach SMD typu 2020 o wymiarach 2,2×2 mm, które idealnie wpisują się w założenia naszego projektu. Dioda tego rodzaju, podobnie jak popularne diody WS2811/WS2812, wyposażona jest w 4 wyprowadzenia:

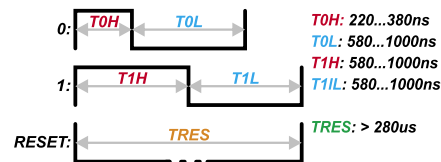
- wyprowadzenia zasilające: GND i VCC,
- wejście asynchronicznego interfejsu komunikacyjnego DIN,
- wyjście asynchronicznego interfejsu komunikacyjnego DOUT.

Wygląd obudowy diody typu OSTW2020C1E z zaznaczeniem nazw wyprowadzeń pokazano na **rysunku 3**.

Jak zapewne się domyślicie, podobnie jak ma to miejsce w przypadku diod WS2811, diody typu OSTW2020C1E łączy się



Rysunek 3. Wygląd obudowy diody typu OSTW2020C1E z zaznaczeniem nazw wyprowadzeń

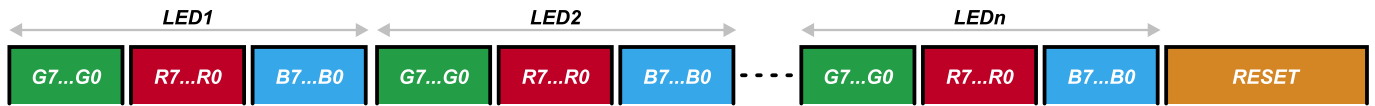


Rysunek 4. Przebiegi sygnałów interfejsu komunikacyjnego diody OSTW2020C1E w trakcie transmisji bitu logicznej „1”, logicznego „0” i sygnału RESET

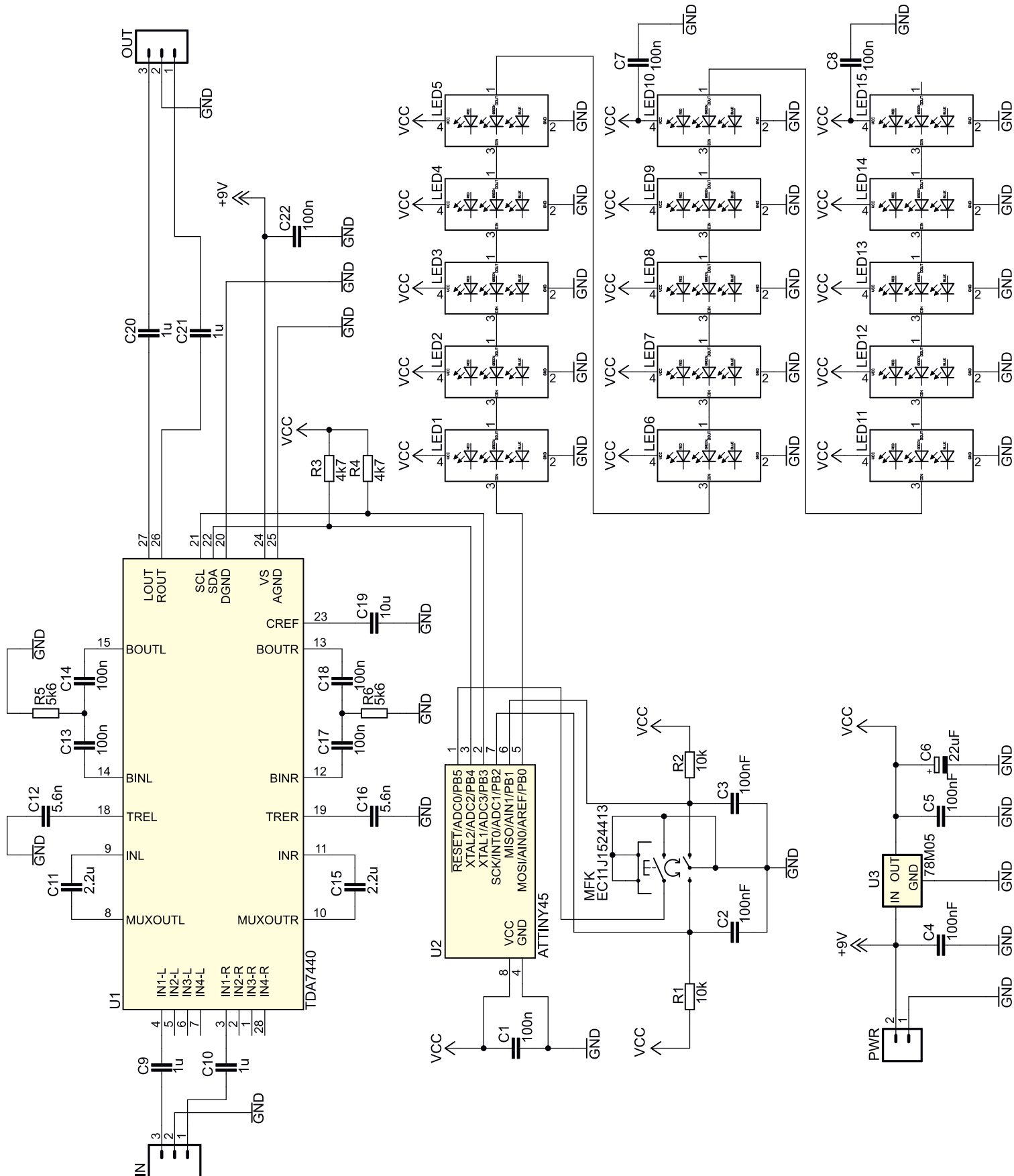
w łańcuchy, łącząc wyjścia interfejsu komunikacyjnego diody bieżącej z wejściami interfejsu komunikacyjnego diody kolejnej i tak dalej. Wejście interfejsu komunikacyjnego diody pierwszej, co oczywiste, łączy się z wyjściem tegoż interfejsu w mikrokontrolerze, zaś sama konstrukcja ramek danych i sposób działania sterownika zabudowanego w strukturze diody zapewnia odpowiednią synchronizację transmisji i niezbędne adresowanie. Zaczniemy więc od podstaw. Jak już wspomniałem, mamy do czynienia z interfejsem asynchronicznym, gdzie nie ma wyprowadzenia sygnału zegarowego, w związku z czym dane przesyłane przy jego użyciu muszą być w pewien sposób zakodowane, by możliwe stało się ich proste zdekodowanie i by były one odporne na zakłócenia i artefakty.

W rozwiązaniu firmy OptoSupply zastosowano mechanizmy dobrze znane z interfejsów bezprzewodowej transmisji danych stosowanych w torach podczerwieni, gdzie stany logiczne „1” i „0” zakodowane zostały długością impulsu. Dodatkowo wprowadzono tak zwany sygnał RESET (również zakodowany długością impulsu), który powoduje zresetowanie interfejsów komunikacyjnych sterowników diod LED i ich oczekiwanie na nowe dane. Na **rysunku 4** pokazano przebiegi sygnałów interfejsu komunikacyjnego diody OSTW2020C1E w trakcie transmisji bitu logicznej „1”, logicznego „0” i sygnału RESET. Co ważne, pojedyncze bity danych zgromadzone w bajty danych przesyłane są w kolejności od bitu najstarszego (MSB) do najmłodszego (LSB) a każda dioda LED w łańcuchu oczekuje na 3 bajty danych odpowiedzialnych za składowe jej koloru przesyłane w kolejności G, R, B.

W tym miejscu zapewne zadacie sobie pytanie, skąd każda z diod w łańcuchu „wie”, które z przesyłanych danych użytecznych przeznaczone są właśnie dla niej, a nie dla innej? Zrealizowano to w bardzo prosty, acz skuteczny sposób. Każda z diod LED w łańcuchu po włączeniu zasilania (jak również po odebraniu sygnału RESET) oczekuje na 3 bajty danych przeznaczonych wyłącznie dla niej. Do tego czasu jej wyjściowy interfejs komunikacyjny (wyjście DOUT) jest „nieprzezroczysty” dla nadchodzących danych (a dokładnie rzecz biorąc, przenosi bez zmian wyłącznie sygnał RESET). Po odebraniu wspomnianych 3 bajtów danych dioda ta staje się „przezroczysta” dla kolejnych



Rysunek 5. Kompletna ramka transmisji łańcucha diod OSTW2020C1E



Rysunek 6. Schemat ideowy urządzenia toneCtrl

nadchodzących danych, co znaczy ni mniej, ni więcej, że retransmituje je do kolejnych diod w łańcuchu (z małym opóźnieniem rzędu 300 ns). Biorąc pod uwagę, iż dokładnie tak samo zachowuje się każda dioda w łańcuchu, dość szybko zdamy sobie sprawę, że kolejne dane użyteczne przesyłane przez tak skonstruowany interfejs komunikacyjny trafiają kolejno do następujących po sobie (w sensie elektrycznym) diod w łańcuchu. Skuteczne i zarazem genialne w swojej prostocie, nieprawdaż?

Niemniej jednak, i co widać na **rysunku 5**, prezentującym kompletną ramkę transmisji łańcucha diod tego typu, przyjęty sposób transmisji stanowiący podstawę adresowania poszczególnych diod LED w łańcuchu powoduje, że nie da się zaadresować (czyli przesłać do niej danych) na przykład diody czwartej

w łańcuchu bez przesłania wcześniejszych (i zdawałoby się niepotrzebnych w tym momencie) danych dla diody trzeciej, drugiej i pierwszej. Mimo tego ograniczenia jest to rozwiązanie bardzo wygodne i chętnie stosowane przez producentów wszelkiego rodzaju peryferiów o takim przeznaczeniu.

Tyle, jeśli chodzi o szczegóły dotyczące naszych podstawowych elementów wykonawczych zaangażowanych w projekt urządzenia toneCtrl.

Pora na omówienie schematu ideowego naszego urządzenia, który to pokazano na **rysunku 6**. Jest to bardzo prosty system mikroprocesorowy, którego sercem jest niewielki mikrokontroler firmy Microchip (dawniej Atmel) typu ATTiny85 taktowany wewnętrznym, wysokostabilnym generatorem RC o częstotliwości 8 MHz. Jest on odpowiedzialny

za programową implementację interfejsu I²C, przy użyciu którego realizuje obsługę układu TDA7440. Ponadto mikrokontroler nasz realizuje obsługę szeregu adresowalnych diod LED RGB stanowiących element graficznego interfejsu użytkownika oraz obsługę switcha MFK używając w celu eliminacji drgań styków wbudowany układ czasowo-licznikowy Timer0 i stosowne przerwanie systemowe (przy okazji obsługiwane jest krótkie i długie naciśnięcie tegoż przycisku). Dodatkowo, dzięki skorzystaniu z przerwanienia zewnętrznego INT0, możliwa stała się efektywna obsługa kierunku obrotów enkodera, co było niezbędne podczas implementacji mechanizmów obsługi interfejsu użytkownika. Tyle w kwestii schematu.

Robert Wołgajew, EP

REKLAMA

Sięgnij po archiwalne wydania „ELEKTRONIKA PRAKTYCZNA”



Przesyłka **GRATIS**

Zamów wygodnie na
www.UlubionyKiosk.pl

**Podstawowe parametry:**

- wzmacniacz jest inspirowany jednym z kultowych wzmacniaczy lampowych – Sansui AU70, lecz nie jest jego wierną kopią,
- zrezygnowano z oryginalnych lamp mocy 7189A, zastępując je pentodami 6P14P-EV, zrezygnowano również z lampy sterującej 6AN8, zamiast niej zastosowano radziecką lampę 6F1P,
- wzmacniacz nie zawiera części przedwzmacniacza, nie ma żadnych regulacji barwy, ale zachowano z oryginalnej konstrukcji funkcję filtra presence,
- moc wzmacniacza wynosi ok. 2×20 W.

* **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:

- 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+] – płytkę drukowaną [A] + zaprogramowany układ [UK] i dokumentacja,
 - wersja [UK] – zaprogramowany układ.

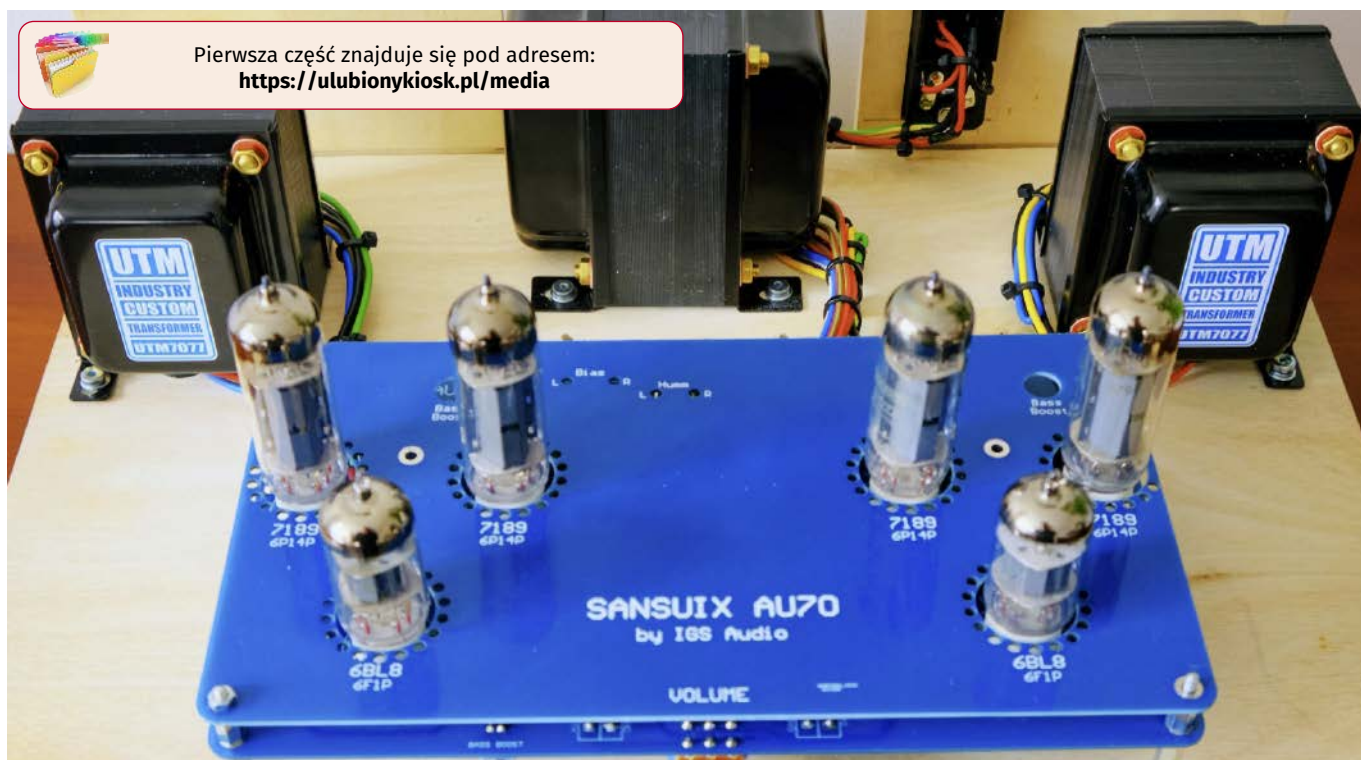
Dodatkowe materiały do pobrania ze strony www.ulubionykiosk.pl/media

- Projekt 252 Hybryda Bis – wzmacniacz hybrydowy z niekonwencjonalnym zasilaniem (EP 6/2021)
- Projekt 252 Hybrydowy wzmacniacz lampowy (EP 1/2021)
- AVT5827 Przedwzmacniacz lampowy do gramofonu (EP 12/2020)
- Automatyczny wyciszacz dźwięku po zaniku zasilania (EP 9/2020)
- Wzmacniacz lampowy z regulacją barwy dźwięku (EP 5/2020)
- Projekt 248 Wzmacniacz lampowy CFB (cathode feedback) (EP 11/2019)
- AVT5727 Hybrydowy wzmacniacz słuchawkowy na lampie Nutube 6P1 (EP 11/2019)
- AVT5719 Przedwzmacniacz na lampie Nutube 6P1 (EP 10/2019)

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*

AVT6001

Sansuix

– lampowy wzmacniacz mocy 2×20 W (2)

Prezentujemy drugą część artykułu o wzmacniaczu lampowym Sansuix – projekcie inspirowanym jednym z kultowych wzmacniaczy lampowych – Sansui AU70. Nie bez znaczenia jest fakt, że w układzie zrezygnowano z egzotycznych, niedostępnych transformatorów czy kosztownych lamp elektronowych. Wszystko można kupić za umiarkowaną cenę jak na konstrukcję lampową. W pierwszej części artykułu zostały opisane rozwiązania układowe i schemat elektryczny. W tej części zostanie omówiony proces montażu i uruchomienia wzmacniacza.

Montaż wzmacniacza

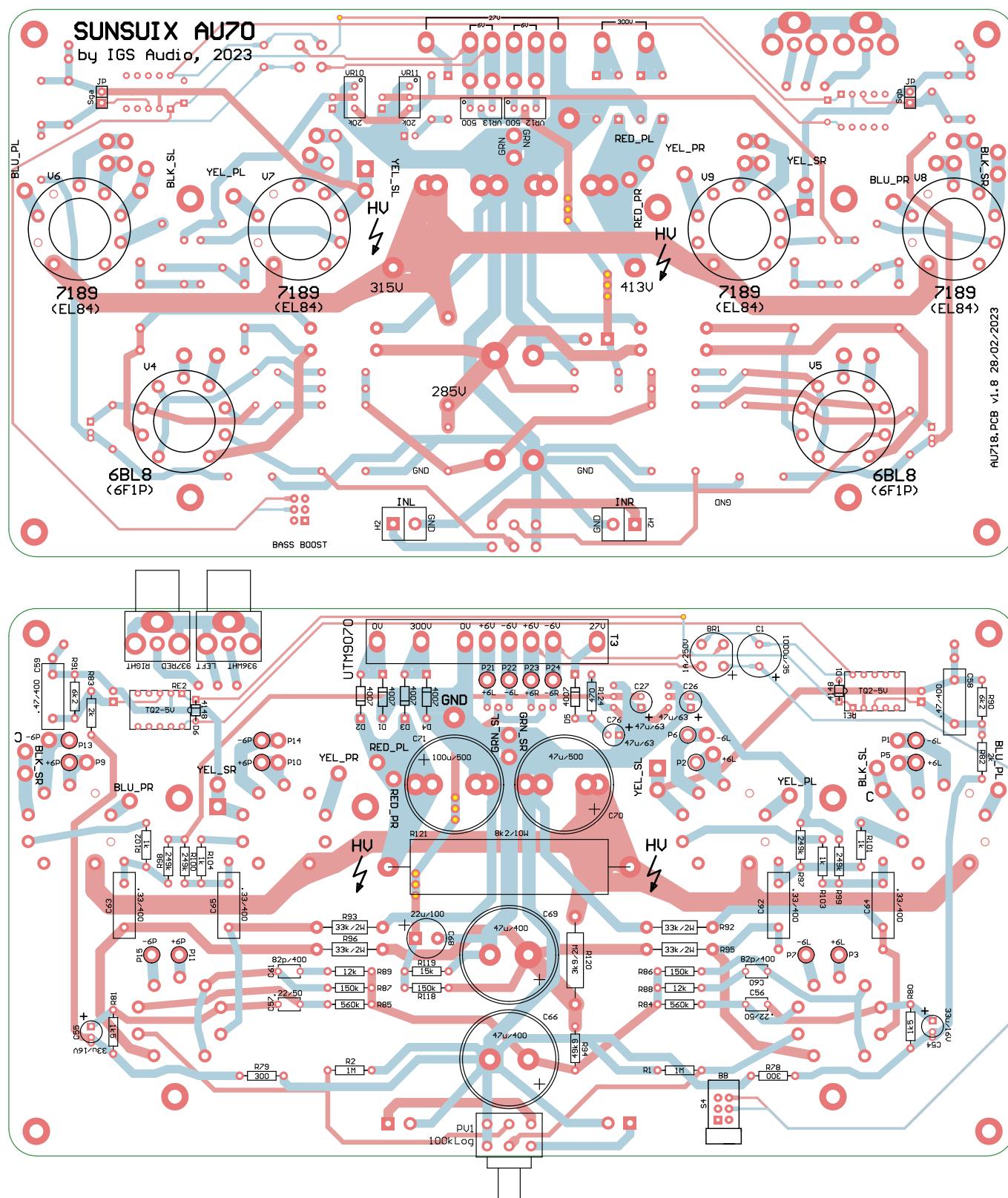
Cały układ elektroniczny wzmacniacza i zasilacz są umieszczone na jednej płytce drukowanej, której schemat został pokazany na rysunku 6. Większość elementów montujemy od strony TOP – umownej strony elementów. Elementy są dokładnie opisane na płytce

– jest podany identyfikator elementu ze schematu, jego wartość oraz dodatkowe informacje typu moc rezystorów czy napięcie kondensatorów – **fotografia 2**.

Montaż rozpoczynamy od wlutowania elementów najniższych, poczynawszy od rezystorów o mocy 0,5 W. Potem wlutowujemy

diody prostownicze, rezystory o większej mocy, kondensatory od najmniejszych do największych i potencjometr. Rekomendujemy, żeby każdy element sprawdzić (zmierzyć) przed wlutowaniem. Może to zapobiec czasochłonnemu szukaniu nieprawidłowości w działaniu wzmacniacza spowodowanej zamontowaniem błędnego lub uszkodzonego elementu.

Po wlutowaniu wszystkich elementów na stronie elementów przechodzimy do drugiej strony płytki. Tam montowane są podstawki pod lampy, potencjometry do regulacji biasu i eliminacji przydźwięku oraz okablowanie układu żarzenia. Należy zwrócić uwagę na równe (proste) wlutowanie podstawek pod lampy.



Rysunek 6. Schemat płytki PCB wzmacniacza

Napięcie żarzenia jest doprowadzane do każdej z lamp za pomocą pary skręconych przewodów. Ma to na celu zminimalizowanie ewentualnego przydźwięku sieciowego. Punkty do prowadzenia przewodów na płytce są oznaczone +6L i -6L dla kanału lewego i +6R oraz -6R dla kanału prawego. Należy zwrócić szczególną uwagę na to, aby połączyć przewodami punkty o takich samych oznaczeniach dla każdego z kanałów.

Wygląd zmontowanej płytki wraz z poprowadzonymi przewodami napięcia żarzenia pokazano na **fotografiach 3 i 4**.

Podłączenie transformatora sieciowego

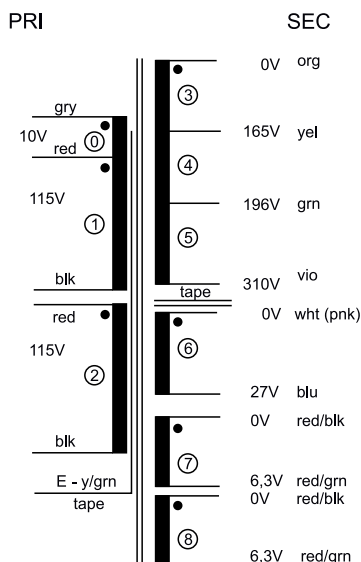
Na **rysunku 7** pokazano schemat transformatora z oznaczeniem kolorów przewodów wyprowadzeń uzwojeń pierwotnych (PRI) i wtórnych (SEC). Uzwojenie pierwotne

składa się z trzech uzwojeń: dwóch na napięcie 115 V i jednego na napięcie 10 V. Po odpowiednim połączeniu uzwojeń można zasilać transformator napięciem sieciowym o napięciach 115 VAC, 230 VAC i 240 VAC. Nas interesuje zasilanie z sieci 230 VAC. W tym celu trzeba połączyć szeregowo dwa uzwojenia na napięcie 115 V, zwracając uwagę na początki i końce uzwojeń. Wyprowadzenie uzwojenia

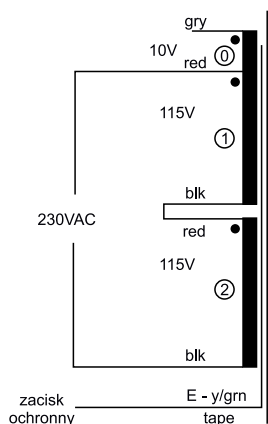
ekranującego podłączamy do zacisku ochronnego wtyczki sieciowej. Połączenie uzwojeń pierwotnych do pracy z napięciem 230 VAC pokazano na **rysunku 8**. Napięcia uzwojeń wtórnych łączymy z odpowiednio opisanymi polami lutowniczymi na płycie drukowanej – **rysunek 9**.

Podłączenie transformatorów głośnikowych

Wyprowadzenia uzwojeń transformatorów głośnikowych są oznaczone kolorami w taki sposób, jak pokazano na **rysunku 10**. Przy podłączaniu wyprowadzeń do płytki należy zwrócić szczególną uwagę na to, że kolorem żółtym są oznaczone wyprowadzenia uzwojeń po stronie pierwotnej i po stronie wtórnej, podobnie jest z kolorem niebieskim. Aby ułatwić poprawną identyfikację punktów lutowniczych, na płycie wzmacniacza przeznaczonych do podłączenia wyprowadzeń transformatora zostały one opisane według klucza XX_YY, gdzie XX oznacza kod koloru, a YY stronę



Rysunek 7. Wyprowadzenia transformatora UTM9070

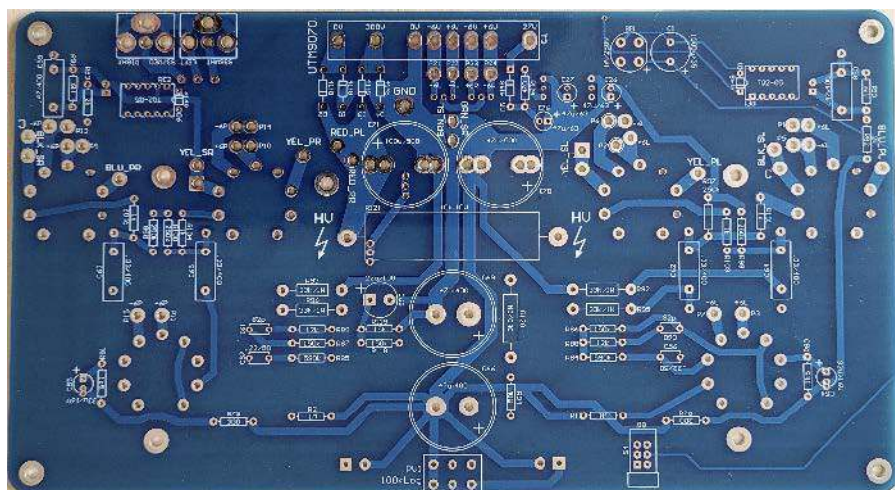


Rysunek 8. Podłączenie uzwojeń pierwotnych transformatora UTM9070 do pracy z siecią 230 VAC

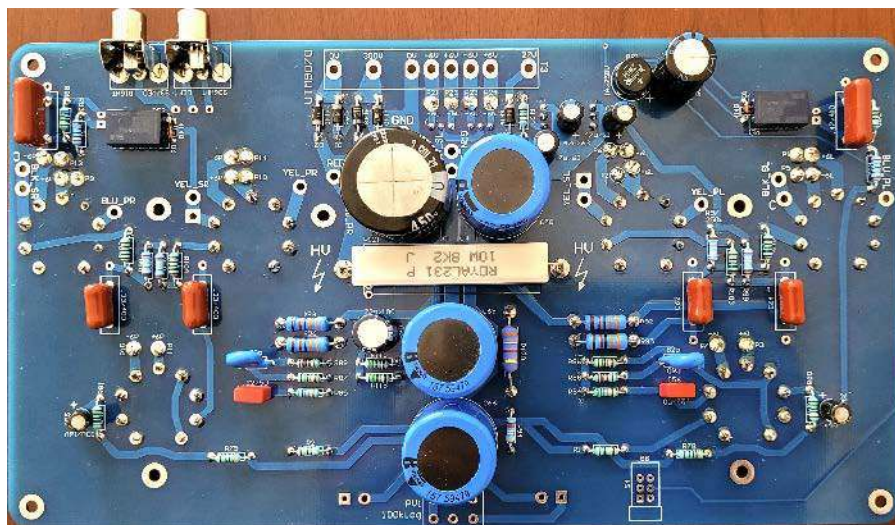
transformatora i kanał stereo. Na przykład BLU_PL oznacza kolor niebieski (*blue*), uzwojenie pierwotne (*primary*) i kanał lewy (*left*), a oznaczenie YEL_SP oznacza kolor żółty, uzwojenie wtórne (*secondary*) i kanał prawy. Na **rysunku 11** pokazano fragment płytki z zaznaczonymi opisami.

Punkty BLK_Sx i YEL_Sx są zdublowane. Należy do nich podłączyć wyprowadzenia przewodów transformatora i drugie kable,

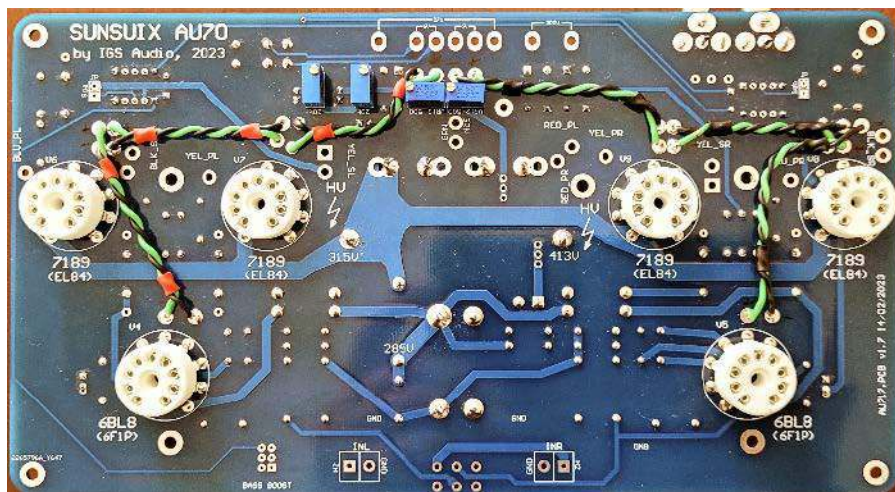
najlepiej o tym samym kolorze przeznaczone do podłączenia zacisków głośnikowych. Schemat połączeń pokazano na **rysunku 12**. Przewód czarny uzwojenia wtórnej jest połączony z katodą lampy V6 i jednocześnie jest przewodem wspólnym wyjścia głośnikowego. Podobnie jest z przewodem żółtym uzwojenia wtórnej. Na **fotografii 5** pokazano okablowanie prototypu w tej części konstrukcji.



Fotografia 2. Umowna strona elementów



Fotografia 3. Wygląd zmontowanej płytki od strony elementów (umownej strony)



Fotografia 4. Wygląd zmontowanej płytki od strony lamp

Uruchomienie wzmacniacza

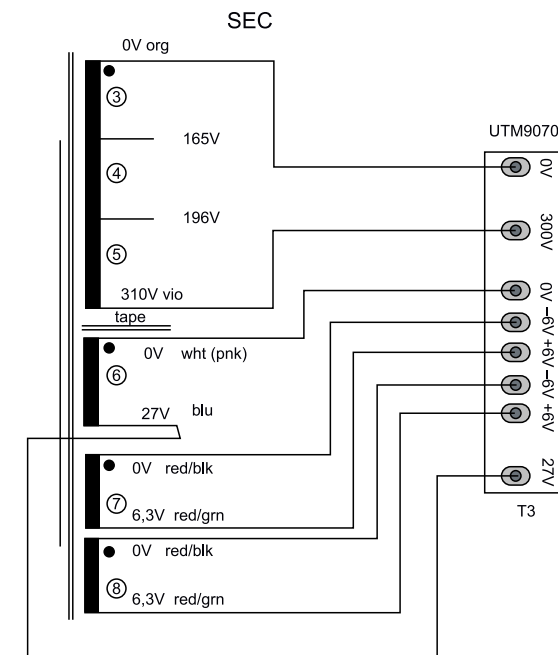
Uwaga! W układzie wzmacniacza występuje szereg napięć, które mogą być źródłem poważnych zagrożeń dla życia i zdrowia. Jest to napięcie przemienne sieci energetycznej 230 VAC oraz napięcia stałe o wartościach od 200 VDC do 400 VDC. W pewnych okolicznościach napięcia stałe mogą się utrzymywać nawet po wyłączeniu wzmacniacza (naładowane kondensatory elektrolityczne). Nieprawidłowy montaż i nieumiejętne prace serwisowe mogą być przyczyną porażenia lub wystąpienia pożaru. Z tego powodu zaleca się bezwzględnie, żeby montaż i uruchomienie wykonywały osoby o odpowiedniej wiedzy i doświadczeniu.

Żelazną zasadą stosowaną przy konstrukcjach tego typu jest praca „jedną ręką” – czyli tak, aby nie zamykać obwodu prądu poprzez obie dłonie i klatkę piersiową. Przy pracy z wysokimi napięciami wymagane jest stosowanie mierników o odpowiedniej klasie odporności na wysokie napięcia (min. 600 VDC). Dotyczy to również sond pomiarowych, ale też izolacji narzędzi, na przykład wkrętek. Nie należy też pozostawiać włączonego wzmacniacza ze zdemontowaną obudową bez nadzoru.

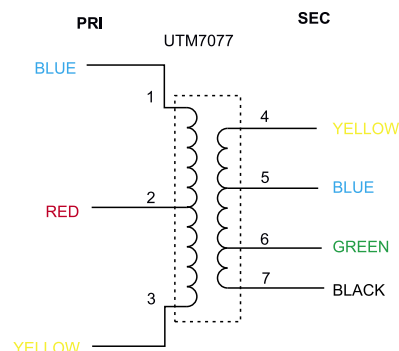
Prezentowany wzmacniacz, jeżeli jest prawidłowo zmontowany z elementów dobrej jakości, nie sprawia żadnych

problemów uruchomieniowych. Pracuje stabilnie od razu po zmontowaniu. Prototyp został mechanicznie zmontowany na kawałku sklejk. Transformatory głośnikowe i sieciowy są przesunięte wobec siebie o kąt 90 stopni tak, by zminimalizować oddziaływanie pola magnetycznego 50 Hz na transformatory głośnikowe – fotografia tytułowa.

Po zmontowaniu płytki i podłączeniu transformatora sieciowego i transformatorów głośnikowych, a przed włączeniem zasilania należy wszystko dokładnie jeszcze raz sprawdzić. Jeżeli wszystko jest w porządku, to wkładamy lampy w podstawki i możemy



Rysunek 9. Podłączenie uzwojeń wtórnych transformatora sieciowego do płytki wzmacniacza

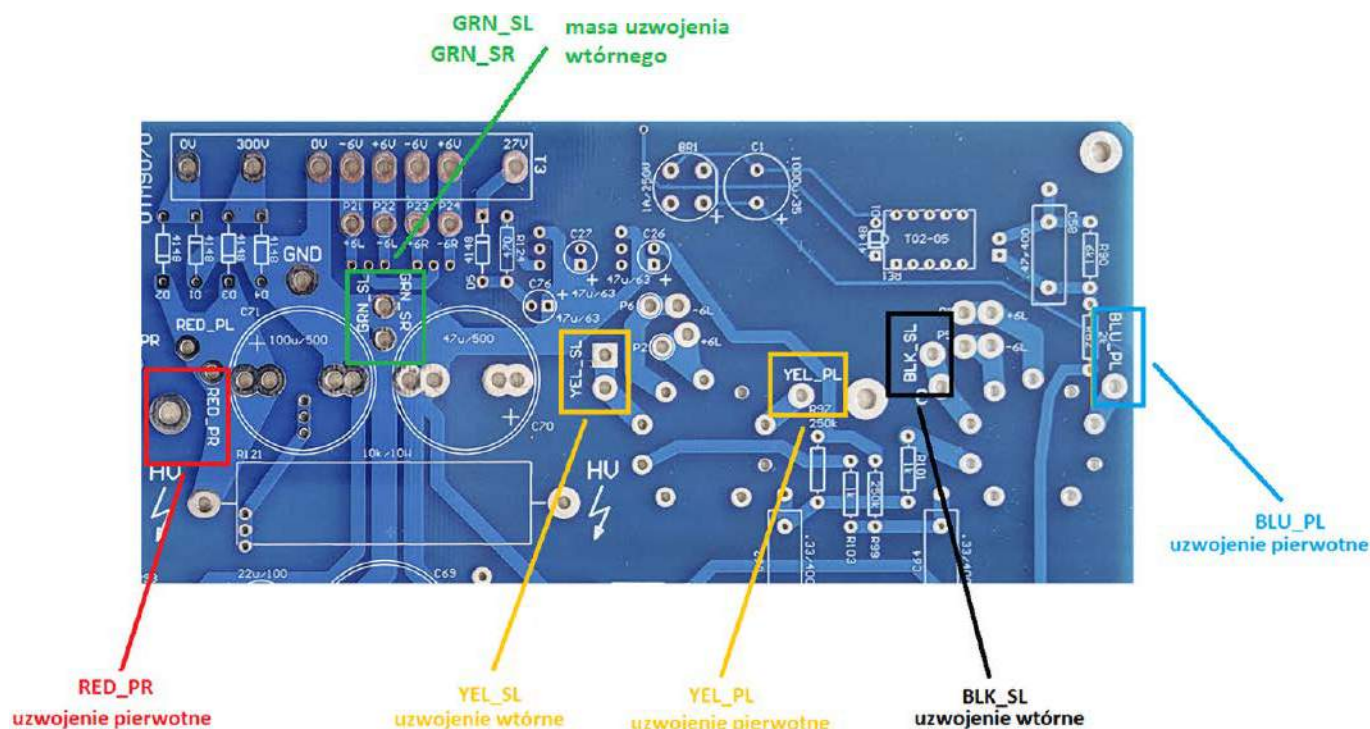


Rysunek 10. Wyprowadzenia transformatora głośnikowego UTM7077

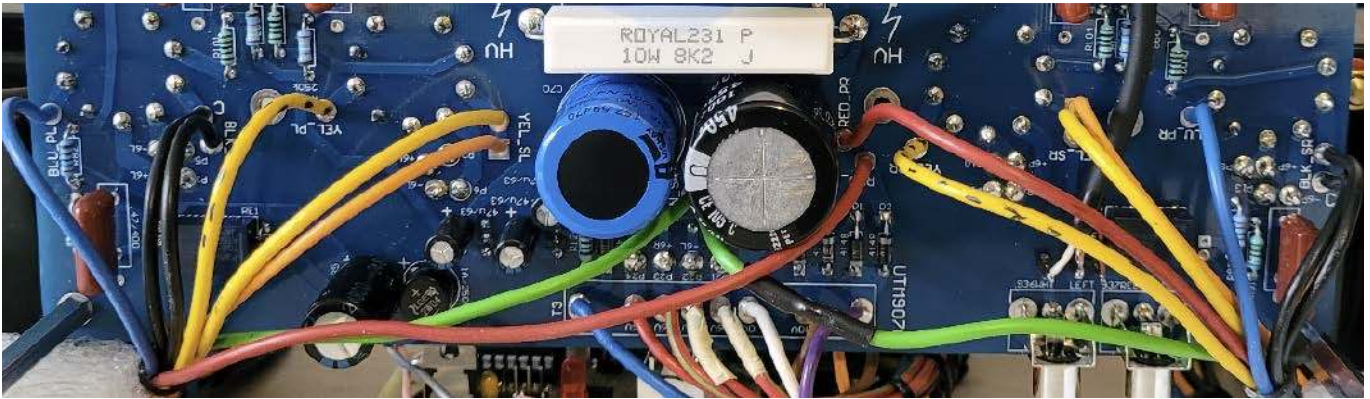
włączyć zasilanie 230 VAC. Należy pamiętać, że lampy mocy w tego typu wzmacniaczach powinny mieć parametry dobierane czwórkami. Lampy różniące się parametrami powodują niesymetrię w pracy stopnia końcowego Push-Pull i wzrost zniekształceń.

Konieczne jest stosowanie lamp 6P14P w wersji militarnej na przykład serii EV. Zwykle lampy 6P14P jak i EL84 nie są przystosowane do pracy z napięciem anodowym +400 V i jeżeli się nie uszkodzą natychmiast, to nie popracują długo w tym układzie. Można próbować obniżyć napięcie anodowe lamp mocy poniżej 300 VDC i godząc się na niższą moc wyjściową, stosować lampy EL84. My takich testów nie robiliśmy.

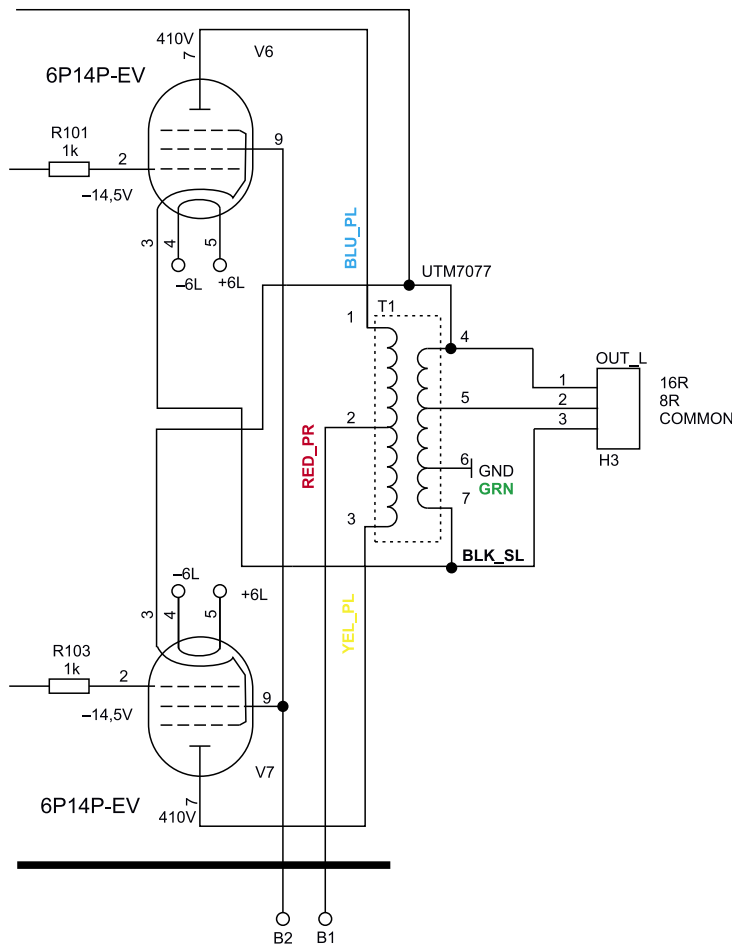
Prawidłowo zmontowany wzmacniacz nie wymaga wielu czynności uruchomieniowych. Przed pierwszym włączeniem do sieci należy obciążyć wyjścia głośnikowe rezystorami 8 Ω/50 W. Włączanie i użytkowanie wzmacniacza bez obciążenia strony wtórnej może spowodować uszkodzenie transformatora głośnikowego. Jeżeli chcemy używać



Rysunek 11. Opisy na płytce z identyfikacją kolorów przewodów transformatora głośnikowego



Fotografia 5. Podłączenie przewodów transformatorów głośnikowych w prototypie wzmacniacza



Rysunek 12. Fragment schematu z zaznaczeniem połączeń z transformatorem głośnikowym

oscylloskopu do diagnostyki wzmacniacza, trzeba pamiętać, że ze względu na topologię stopnia wyjściowego nie można podłączać jednocześnie dwu sond oscylloskopu dwukanałowego do wyprowadzeń głośnikowych. Spowoduje to zwarcie wyprowadzeń common obu kanałów przez masę sond oscylloskopu i błędne działanie wzmacniacza.

Podobnie jest z generatorem sygnałowym. Jeżeli zimny zacisk (masa) generatora jest połączony z kołkiem uziemiającym i masa oscylloskopu jest również połączona z kołkiem uziemiającym, to takie połączenie spowoduje zwarcie wyjścia głośnikowego common do masy i wzmacniacz nie będzie pracował prawidłowo.

Ustawienie punktu pracy (biasu) pentod mocy

Należy podłączyć woltomierz pomiędzy masę i nóżkę 2 jednej z pentod mocy regulowanego kanału. Na wejście nie podajemy żadnego sygnału (potencjometr siły głosu jest w lewym skrajnym położeniu). Potencjometrem VR10 dla kanału lewego i VR11 dla kanału prawego ustawiamy napięcie w zakresie od $-14,5\text{ V}$ do -18 V , ale równo dla obu kanałów. Należy to robić powoli, ponieważ suwaki potencjometrów są zabezpieczone kondensatorami elektrolitycznymi i napięcie na nich musi się ustabilizować.

Regulacja eliminacji przydzwięku

Po upewnieniu się, że wzmacniacz się nie wzbudza z powodu wadliwego montażu i po ustawieniu biasu, możemy podłączyć kolumny głośnikowe. Jeżeli w głośnikach jest słyszalny przydzwięk sieciowy, to można go wyeliminować, regulując potencjometrami VR12 dla kanału lewego i VR13 dla kanału prawego.

Uwagi końcowe

Jak już wspominaliśmy, wzmacniacz jest stabilną konstrukcją i nie stwarza żadnych problemów z uruchomieniem. Jakość brzmienia jest na bardzo dobrym poziomie. Jest to zasługa bardzo dobrej topologii układu pochodzącego ze złotej ery wzmacniaczy lampowych, ale też zastosowania transformatorów głośnikowych o świetnych parametrach technicznych. Połączenie tych dwu elementów daje świetny efekt końcowy. Nie bez znaczenia jest fakt, że nie ma tu egzotycznych niedostępnych transformatorów czy lamp elektronowych. Wszystko można kupić za umiarkowaną cenę jak na konstrukcję lampową.

Igor Sobczyk
Tomasz Jabłoński, EP

REKLAMA

EP.com.pl

Odwiedź stronę z mnóstwem doskonałych projektów

**Podstawowe parametry:**

- możliwość zasilania z zewnętrznego zasilacza 5 V ze złączem USB-C,
- zawiera połączone trzy filtry z elementami LC,
- płytka PCB zgodna z formatem Pi Zero, co nie wyklucza zastosowania układu z innymi wersjami i komputerkami.

*** 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+] – płytka drukowana [A] + zaprogramowany układ [UK] i dokumentacja,
 - wersja [UK] – zaprogramowany układ.

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

- Ekspander GPIO RPi z taśmą FPC (EP 8/2023)
- Sterownik unipolarnego mikrosilnika krokowego dla Pi Pico (EP 7/2023)
- Sterownik dwóch silników krokowych do Raspberry Pi (EP 6/2023)
- Expander wyjść z PWM na bazie układu PCA9624 (EP 6/2023)
- Moduł redundancji zasilania dla Raspberry Pi Zero (EP 5/2023)
- Sterownik dwóch mikrosilników krokowych do Pi Zero (EP 3/2023)
- Sterownik taśm LED RGB CCT 12 V dla RPi Zero (EP 3/2023)
- Eliminatory drgań styków mechanicznych (EP 1/2023)
- Moduł redundancji zasilania do komputerów SBC (EP 1/2023)

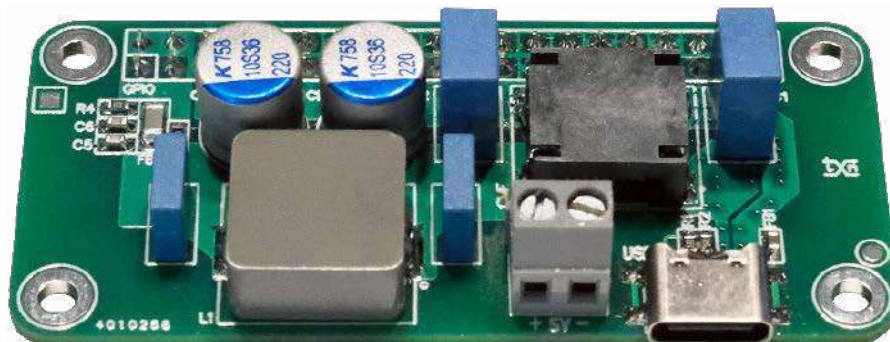
Nie każdy zestaw AVT występuje we wszystkich wersjach! Każda wersja ma załączony ten sam plik PDF! Podczas składania zamówienia upewnij się, którą wersję zamawiasz!
<http://sklep.avt.pl>

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

Filtr zasilania do Raspberry Pi

Zaprezentowana nakładka do Raspberry Pi umożliwia dodatkową filtrację napięcia zasilania. Filtr przydatny jest, gdy komputerka używany w aplikacjach audio lub pomiarowych, gdzie od jakości zasilania silnie zależy uzyskany efekt końcowy. W wielu przypadkach filtr może być tańszą alternatywą zasilacza liniowego.

Zastosowanie gniazda USB-C umożliwia ujednolicenie kabli zasilających i stosowanie nowoczesnych ładowarek bez dodatkowych przejściówek (w przypadku użycia z RPi Zero/RPi 3+). Przy okazji można rozwiązać problem nieprawidłowej współpracy z ładowarkami pierwszych wersji RPi 4. Zastosowane gniazdo USB-C typu 4135 (GTC) przeznaczone jest tylko do aplikacji zasilających. Pozbawione jest wszystkich pinów komunikacyjnych interfejsu USB-C, pozostawiono tylko wyprowadzenia zasilania i sygnały CCx, służące do detekcji współpracującego urządzenia. Ogranicza to liczbę wyprowadzeń

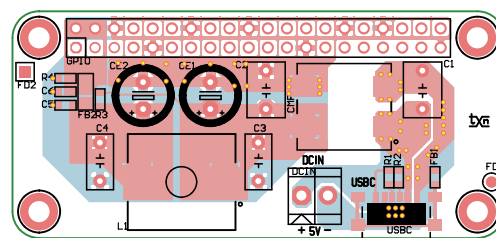


do 6 i znacząco ułatwia montaż, który jest możliwy bez specjalistycznych narzędzi.

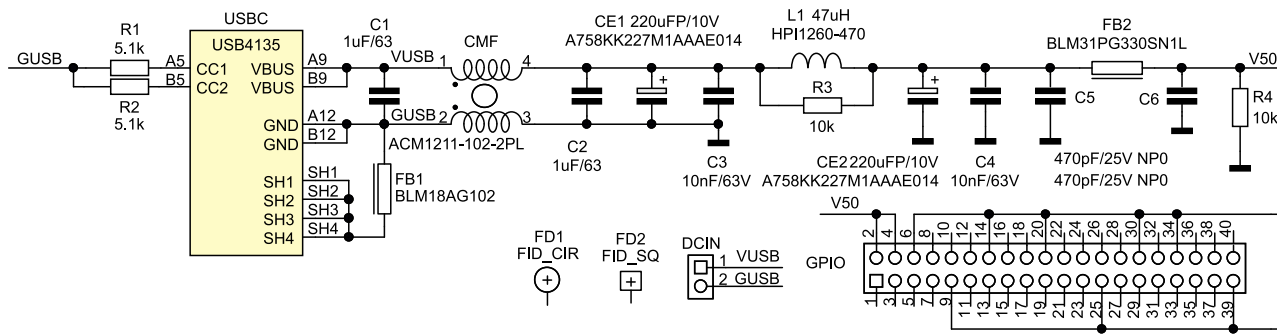
Budowa i działanie

Schemat układu został pokazany na rysunku 1. Zasilanie filtra z zewnętrznego zasilacza 5 V doprowadzone jest do złącza śrubowego DCIN lub gniazda USB-C. Układ zawiera połączone trzy filtry z elementami LC. Pierwszy filtr C1, C2, CMF tłumi zakłócenia wspólne, drugi – C3, C4, L1 zakłócenia różnicowe o niskich częstotliwościach, trzeci – C5, C6, FB2 zakłócenia

różnicowe o częstotliwościach wysokich. Dwa polimerowe kondensatory CE1, CE2 o niskim ESR odsprężają zasilanie. Napięcie V50 po filtracji doprowadzone jest do złącza GPIO.



Rysunek 2. Schemat płytki PCB filtra



Rysunek 1. Schemat ideowy układu

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

Rezystory: (SMD0603, 1%)

R1, R2: 5,1 kΩ
R3, R4: 10 kΩ

Kondensatory:

C1, C2: 1 μF/63 foliowy (C7.2X5.0P5.0)

CE1, CE2: 220 μF/10 V polimerowy (CE8.0P3.5)

C3, C4: 10 nF/63 V foliowy (C7.2X3.5P5.0)

C5, C6: 470 pF/25 V NPO (SMD0603)

Pozostałe:

CMF: dławik CMF ACM1211-102-2PL

DCIN: złącze DG 3,5 mm 2 piny (DG381-3.5-2)

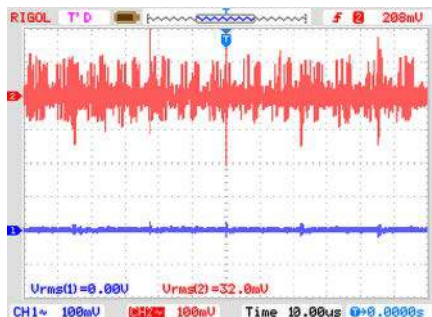
FB1: koralek ferrytowy BLM18AG102 (SMD0603)

FB2: koralek ferrytowy BLM31PG330SN1L (SMD1206)

GPIO: złącze żeńskie IDC40 do druku

L1: dławik mocy 47 μH (HP11260)

USB-C: gniazdo USB-C Power GTC (USB4135)



Rysunek 3. Przykładowy efekt działania filtra

Montaż i uruchomienie

Układ zmontowany jest na dwustronnej płytce drukowanej zgodnej z formatem Pi Zero, co oczywiście nie wyklucza stosowania z innymi modelami Raspberry. Schemat płytki został pokazany na **rysunku 2**. Montaż układu nie wymaga dokładnego opisu, a po zmontowaniu nie wymaga uruchamiania. Po doprowadzeniu zasilania 5 V do jednego z gniazd DCIN lub USB-C układ jest gotowy do pracy.

Przykładowe zmierzone efekty działania filtra pokazano na **rysunku 3**. Przebieg 2 (górny) to napięcie zasilania z zasilacza

impulsowego 5 V współpracującego z Raspberry Pi i nakładką audio. Zasilacz obciążony jest prądem ok. 2,5 A, przebieg 1 (dolny) to napięcia na wyprowadzeniach zasilania 2, 6 złącza GPIO. Filtracja jest bardzo skuteczna, pozostaje tylko niewielki ślad częstotliwości kluczkowania zasilacza. Elementy filtra można potraktować jako uniwersalne, poprzez zmianę ich wartości układ można zoptymalizować pod kątem własnej aplikacji – zachęcam do eksperymentów.

Adam Tatuś, EP



Podstawowe parametry:

- 4 wejścia cyfrowe z optoizolacją oraz 4 wejścia bez optoizolacji,
- komunikacja z modułem odbywa się poprzez interfejs USB z układem FT230,
- interfejs GPIO jest zrealizowany poprzez układ SC18IM704 firmy NXP,
- każde z wyprowadzeń GPIO może pracować w trzech trybach: wejścia, wyjścia z otwartym drenem i wyjścia Push-Pull.

* **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:

- 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

- | | |
|---------|--|
| AVT5988 | Konwerter USB-RS485 (EP 7/2023) |
| --- | Konwerter USB-UART w standardzie Grove (EP 5/2023) |
| AVT5717 | Konwerter USB-UART z ekstenderem (EP 10/2019) |
| AVT5648 | Izolowana przejściówka USB/UART (EP 9/2018) |
| AVT1780 | USB_FT230XQ Miniatury konwerter USB/UART (EP 11/2013) |
| AVT1775 | Miniatury konwerter USB/UART z układem FT230XS (EP 9/2013) |
| AVT1595 | Miniatury konwerter USB/UART (EP 10/2010) |

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>

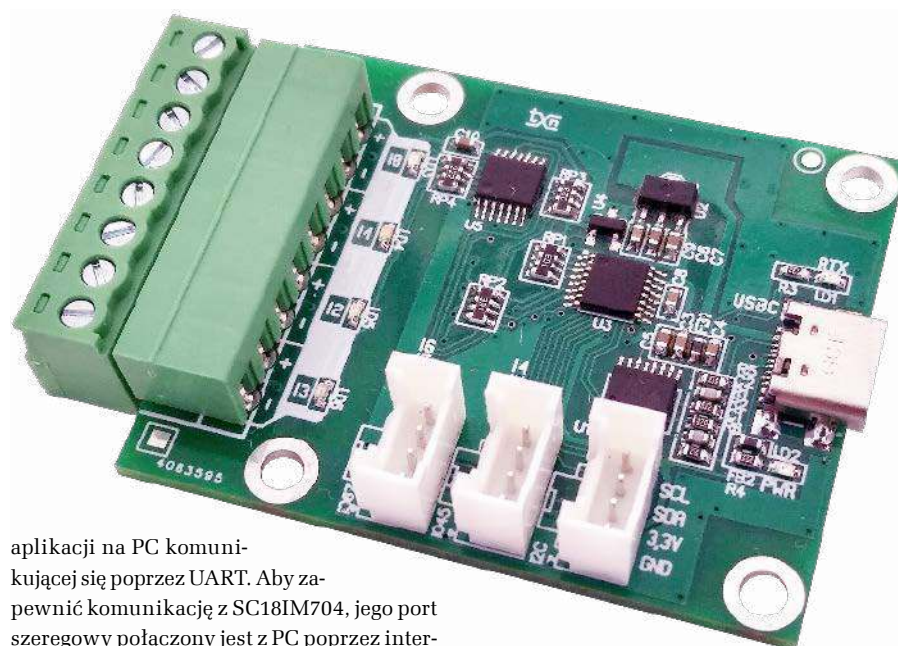
W przypadku braku dostępności na stronie sklepu osoby zainteresowane zakupem płytek drukowanych (PCB) prosimy o kontakt via e-mail: kity@avt.pl.

Moduł wejść cyfrowych z optoizolacją i interfejsem USB-C

Moduł wejść cyfrowych z optoizolacją będzie przydatny w domowej automatyce. Komunikacja z modułem odbywa się poprzez interfejs USB, a zastosowanie fabrycznego mostka UART/GPIO/I²C z gotowym protokołem komunikacyjnym zwalnia nas z potrzeby tworzenia aplikacji dla mikrokontrolera.

Zaprezentowany układ interfejsu zawiera dwa specjalizowane układy, pierwszy to konwerter USB/UART typu FT230, drugi to układ mostka UART/GPIO/I²C typu SC18IM704 firmy NXP, którego strukturę wewnętrzną pokazano na **rysunku 1**. Układ SC18IM704 komunikuje się z komputerem nadrzędnym poprzez standardowy interfejs szeregowy UART, używając komunikacji znakowej ASCII.

Interfejs GPIO dostępny jest poprzez rejestry wewnętrzne układu, każde z wyprowadzeń GPIO może pracować w trzech trybach: wejścia, wyjścia z otwartym drenem i wyjścia Push-Pull. Układ obsługuje także magistralę I²C w trybie odczytu i zapisu. Cała praca programistyczna podczas korzystania z SC18IM704 ogranicza się do opracowania



aplikacji na PC komunikującej się poprzez UART. Aby zapewnić komunikację z SC18IM704, jego port szeregowy połączony jest z PC poprzez interfejs konwertera USB/UART.

Budowa i działanie

Schemat interfejsu został pokazany na **rysunku 2**. Magistrala USB-C doprowadzona jest do gniazda typu USB4110 GCT, które ma ograniczoną do trybu zgodności USB liczbę wyprowadzeń. Dostępny jest jedynie interfejs

USB 2.0, zasilanie oraz wyprowadzenia konfiguracji interfejsu. Zmniejszona liczba wyprowadzeń USB-C ułatwia montaż, który nie wymaga specjalistycznych narzędzi.

Układ U1 typu FT230XS pełni funkcję interfejsu USB/UART, dioda LD1 sygnalizuje aktywną transmisję szeregową (dla obu

Układ U3 typu SC18IM704 pełni funkcję mostka UART/GPIO/I²C, za poprawny restart po włączeniu zasilania odpowiada U4 typu MCP100T-315. Mostek U3 ma 8 wyprowadzeń GPIO, w modelu cztery z nich GPIO3...0 służą do monitorowania stanów optoizolowanych wejść IO3...0. Sygnały wejściowe w standardzie 24 VDC doprowadzone są do złącza INPUT. Poczwośny transoptor IS1 typu TCMT4600 separuje je od części logicznej modułu. Za sygnalizację stanów wejść IN1...4 odpowiadają diody LD3...LD6 buforowane przez bramki U5.

Pozostałe cztery GPIO układu SC18IM704 wraz z zasilaniem 3,3 V doprowadzone są do złączy zgodnych z Grove IO45, IO67, do których można bezpośrednio podłączyć czujniki lub sygnalizatory systemu Grove zgodne z 3,3 V.

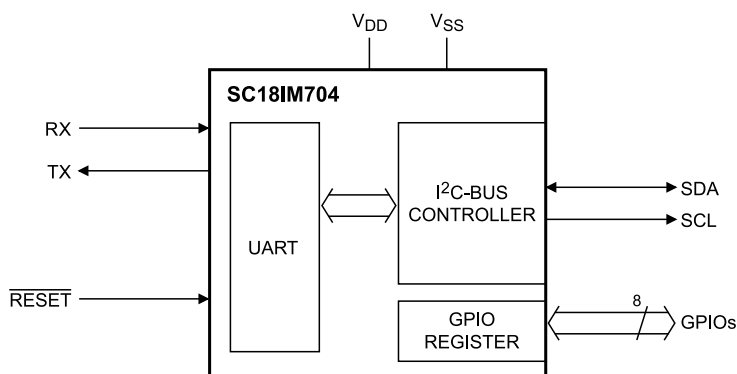
Szczegolnej uwagi wymaga używanie GPIO6, który w momencie resetu SC18IM704 nie może przyjmować stanu niskiego i musi być podciągnięty do potencjału zasilania, co w modelu realizowane jest przez RP1-RB2. Obwody wejściowe U3, mają wbudowane bramki Schmitta, pozwalające na filtrację sygnału. Podanie napięcia 24 V na wejścia złącza INPUT powoduje zaświecenie LED i ustawia podłączone GPIO w stan niski (aktywny stan niski, negacja w transporcie).

Na złącze I²C wyprowadzona jest magistrala I²C wraz z zasilaniem 3,3 V. Umożliwia to rozszerzenie funkcjonalności modułu o obsługę ekspanderów GPIO np. PCF8574 lub czujników zgodnych z I²C.

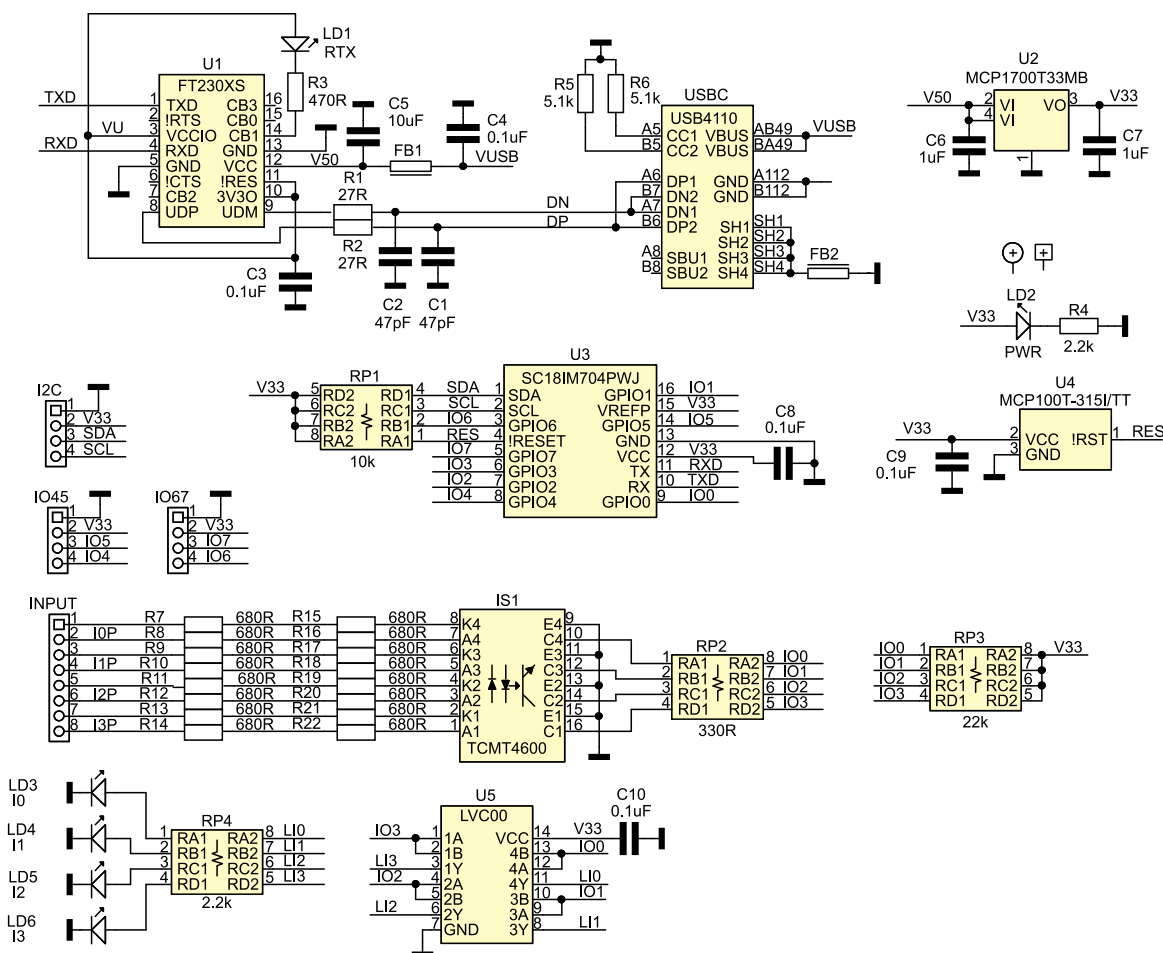
Montaż i uruchomienie

Moduł zmontowany jest na dwustronnej płytce drukowanej. Schemat płytki został pokazany na **rysunku 3**. Montaż jest typowy i nie wymaga opisu. Poprawnie zmontowany układ nie wymaga uruchamiania. Po podłączeniu do portu USB, konwerter FT230 powinien zostać wykryty i automatycznie zostaną zainstalowane sterowniki. Przy użyciu oprogramowania narzędziowego FT Prog należy skonfigurować prąd pobierany z magistrali zgodnie z **rysunkiem 4**. Następnie należy zmienić konfigurację wyprowadzenia C1 na sygnalizację sumy sygnałów RXD/TXD zgodnie z **rysunkiem 5**. Jeżeli konfiguracja została ustalona i zapisana, należy układ zrestartować poprzez odłączenie i ponowne podłączenie do magistrali USB.

Dla sprawdzenia modułu należy zainstalować terminal portu szeregowego z możliwością wysyłania sekwencji znaków np. Realterm oraz źródło sygnału 24 V i kilka czujników Grove 3,3 V. Komunikacja z układem SC18IM704 odbywa się poprzez wysłanie polecenia w postaci komendy ASCII, zestaw



Rysunek 1. Schemat wewnętrzny układu SC18IM704 (za nota NXP)



Rysunek 2. Schemat ideowy modułu

rozpoznawanych komend został pokazany w tabeli 1. Komendy nierozpoznane są ignorowane. Układ ma wbudowany licznik time-out, który kasuje zawartość buforów, jeżeli odbiór kolejnych dwóch znaków zajmuje więcej niż 655 ms.

Po restarcie układ jest skonfigurowany wstępnie – standard transmisji ustawiony jest na 9600,8,N,1. Konfiguracja może zostać zmieniona w rejestrach układu, których zestawienie pokazuje rysunek 6. Szczegółowy opis dostępny jest w dokumentacji SC18IM704.pdf dołączonej do materiałów dodatkowych. Po uruchomieniu terminala i resecie układu zwrócone zostaną znaki OK, potwierdzające połączenie z układem. Dla sprawdzenia poprawnej komunikacji z układem SC18IM704 odczytujemy jego sygnaturę, wysyłając:

56 50

Układ powinien zwrócić ciąg:

53 43 31 38 49 4D 37 30 34 20 31 2E0 30 2E 32 00

Czyli typ i wersję układu:

SC18IM704 1.0.2

Do zmiany konfiguracji służą sekwencje poleceń zapisu rejestru wewnętrznego:

W <rejestr 0> <dana 0>.....<rejestr N> <dana N> P

lub odczytu:

R <rejestr 0>.....<rejestr N> P

po którym SC18IM704 zwróci ciąg wartości: <dana 0>..... <dana N>

Odczyt rejestru 00 to sekwencja:

52 00 50

która po resecie zwraca wartość domyślną 0xF0.

Przed użyciem wyprowadzeń GPIO w rejestrach PortConf1/2 należy ustawić wymaganą konfigurację wyprowadzenia, zgodnie z tabelą 2. Do monitorowania stanów sygnałów doprowadzonych do złącza INPUT konieczna jest konfiguracja GPIO3...0 jako wejść (0x). Wymaga to zapisu rejestrów wewnętrznych PortConf1 0x02 (GPIO3...0), które muszą zawsze być ustawione jako wejścia! Konfiguracja GPIO7...4 odbywa się poprzez rejestr PortConf2 0x03 (GPIO7...4) i zależy od aplikacji (z uwzględnieniem uwagi o GPIO6):

57 02 55 50 – ustawia w rejestrze PortConf1 (0x02) piny GPIO3...0 jako wejścia,

57 03 55 50 – ustawia rejestrze PortConf2 (0x03) piny GPIO7...4 jako wejścia.

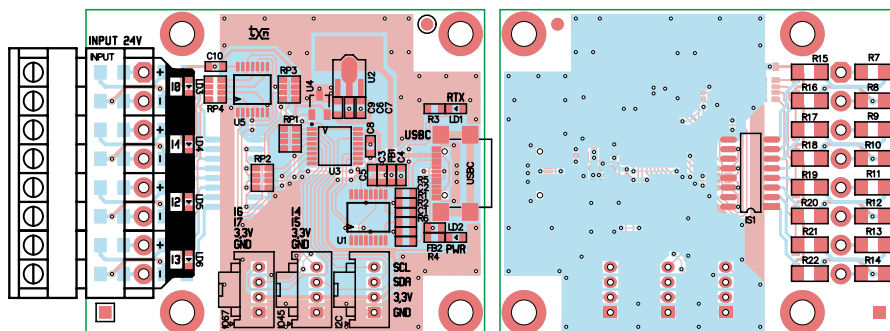
Domyślnie GPIO układu SC18IM704 po włączeniu zasilania skonfigurowane są jako wejścia. Weryfikację zapisu rejestrów można wykonać poleceniem odczytu rejestru wewnętrznego:

52 02 50 – która powinna zwrócić zapisaną wartość 0x55,

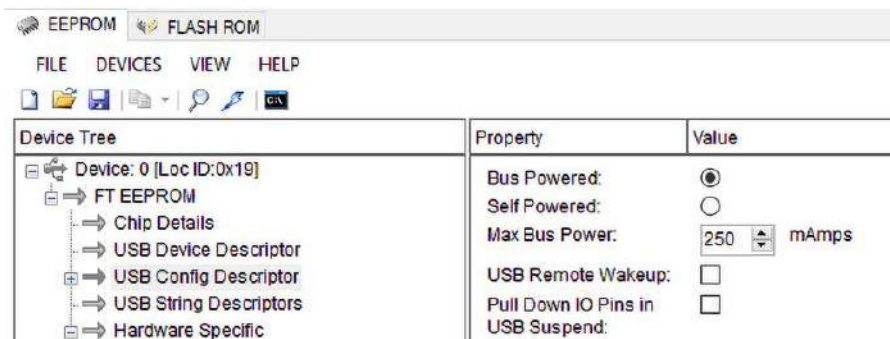
52 03 50 – która powinna zwrócić zapisaną wartość 0x55.

Do zmiany stanu GPIO służą sekwencje poleceń odczytu I, zapisu O:

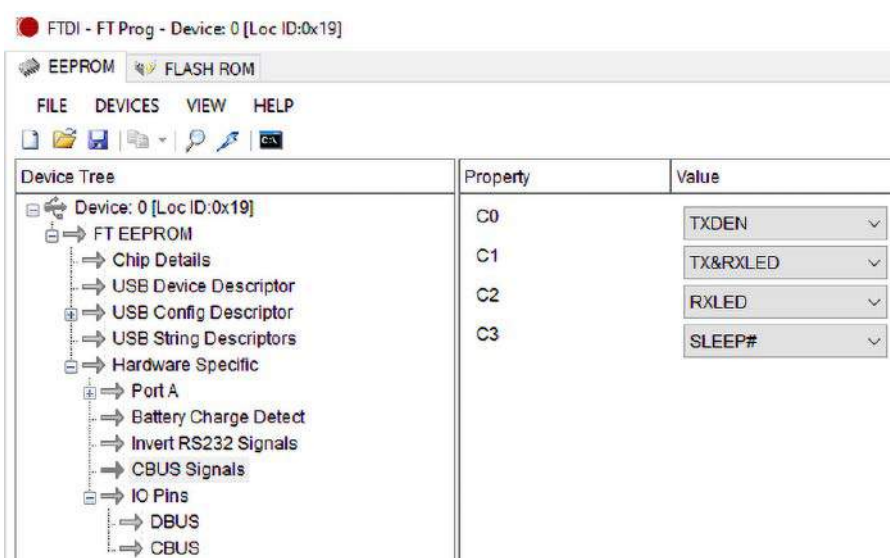
49 50 – po której układ zwróci <stan GPIO>



Rysunek 3. Schemat płytki PCB



Rysunek 4. Konfiguracja prądu USB



Rysunek 5. Konfiguracja wyprowadzenia C0

Tabela 1. Komendy ASCII układu SC18IM704		
Komenda ASCII	Kod hex	Opis
S	0x53	I ² C start
P	0x50	I ² C stop
R	0x52	Odczyt rejestrów wewnętrznych
W	0x57	Zapis rejestrów wewnętrznych
I	0x49	Odczyt portu GPIO
O	0x4F	Zapis portu GPIO
Z	0x5A	Tryb obniżonego poboru mocy

Tabela 2. Konfiguracja wyprowadzeń GPIO		
GPIOx.1	GPIOx.0	Konfiguracja
0	x	Wejście
1	0	Wyjście Push-Pull
1	1	Wyjście OD

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

LD1: dioda LED żółta 2 mA (SMD0603)
 LD2: dioda LED czerwona 2 mA (SMD0603)
 LD3...LD6: dioda LED zielona 2 mA (SMD0603)
 U1: FT230XS (SSOP16)
 U2: MCP1700T-3302MB (SOT-89)
 U3: SC181M704PWJ (TSSOP16_065)
 U4: MCP100T-3151/TT (SOT-23)
 U5: LVC00 (TSSOP14)

Kondensatory: (SMD0603)

C1, C2: 47 pF/25 V

C3, C4, C8, C9, C10: 0,1 µF/10 V
 C5: 10 µF/10 V
 C6, C7: 1 µF/10 V

Rezystory: (SMD0603, 1%)

R1, R2: 27 Ω
 R3: 470 Ω
 R4: 2,2 kΩ
 R5, R6: 5,1 kΩ
 R7, R8, R9, R10, R11, R12, R13, R14, R15, R16, R17, R18, R19, R20, R21, R22: 680 Ω, 1% (SMD1206)

Pozostałe:

FB1, FB2: koralek ferrytowy BLM18EG101T1N1D (SMD0603)
 INPUT: złącze MC 3,81 mm, śrubowe, rozłączane, komplet IO45, 67, I²C: złącze Grove proste (110990030)
 IS1: transoptor TCM4600 (SO16)
 RP1: drabinka rezystorowa 4×10 kΩ (CRA06S08)
 RP2: drabinka rezystorowa 4×330 Ω (CRA06S08)
 RP3: drabinka rezystorowa 4×22 kΩ (CRA06S08)
 RP4: drabinka rezystorowa 4×2,2 kΩ (CRA06S08)
 USB-C: złącze USB-C GCT (USB4110)

Register address	Register	Bit 7	Bit 6	Bit 5	Bit 4	Bit 3	Bit 2	Bit 1	Bit 0	R/W	Default value
General register set											
0x00	BRG0	bit 7	bit 6	bit 5	bit 4	bit 3	bit 2	bit 1	bit 0	R/W	0xF0
0x01	BRG1	bit 7	bit 6	bit 5	bit 4	bit 3	bit 2	bit 1	bit 0	R/W	0x02
0x02	PortConf1	GPIO3.1	GPIO3.0	GPIO2.1	GPIO2.0	GPIO1.1	GPIO1.0	GPIO0.1	GPIO0.0	R/W	0x55
0x03	PortConf2	GPIO7.1	GPIO7.0	GPIO6.1	GPIO6.0	GPIO5.1	GPIO5.0	GPIO4.1	GPIO4.0	R/W	0x55
0x04	IOState	GPIO7	GPIO6	GPIO5	GPIO4	GPIO3	GPIO2	GPIO1	GPIO0	R/W	- [1]
0x05	reserved	bit 7	bit 6	bit 5	bit 4	bit 3	bit 2	bit 1	bit 0	-	0x00
0x06	I2CAdr	bit 7	bit 6	bit 5	bit 4	bit 3	bit 2	bit 1	bit 0	R/W	0x26
0x07	I2CClkL	bit 7	bit 6	bit 5	bit 4	bit 3	bit 2	bit 1	bit 0	R/W	0x13
0x08	I2CClkH	bit 7	bit 6	bit 5	bit 4	bit 3	bit 2	bit 1	bit 0	R/W	0x00
0x09	I2CTO	TO7	TO6	TO5	TO4	TO3	TO2	TO1	TE	R/W	0x66
0x0A	I2CStat	1	1	1	1	I2CStat[3]	I2CStat[2]	I2CStat[1]	I2CStat[0]	R	0xF0

Rysunek 6. Rejestry konfiguracyjne układu SC181M704

Bit 7	Bit 6	Bit 5	Bit 4	Bit 3	Bit 2	Bit 1	Bit 0	I ² C-bus status description
1	1	1	1	0	0	0	0	I2C_OK
1	1	1	1	0	0	0	1	I2C_NACK_ON_ADDRESS
1	1	1	1	0	0	1	0	I2C_NACK_ON_DATA
1	1	1	1	1	0	0	0	I2C_TIME_OUT

Rysunek 7. Rejestr kontroli stanu magistrali I²C

4F <dana> 50 – po której układ zapisze stan GPIO.

Oczywiście zapis pod GPIO3...0 nie jest dozwolony.

Do sprawdzenia magistrali I²C do modułu (złącze I²C) podłączony będzie układ PCF8574, który ma ustawiony adres 0x20, odczyt stanu wyprowadzeń wykonany jest sekwencją:

53 41 01 50 – która odczytuje jeden bajt spod adresu 0x41, czyli stan wyprowadzeń PCF8574 z ustawionym siedmiobitowym adresem 0x20. Ze względu na adresowanie 8 bitowe najmłodszy bajt ustawiony jest na 1 dla operacji odczytu (0100 000 1).

Ustawienie wyjść wykonywane jest sekwencją:

53 40 01 00 50 – która zapisuje jeden bajt o wartości 0x00 pod adres 0x40, czyli stan wyprowadzeń PCF8574 z ustawionym siedmiobitowym adresem 0x20 zostanie zmieniony na 0x00. Ze względu na adresowanie 8-bitowe najmłodszy bajt ustawiony jest na 0 dla operacji zapisu (0100 000 0).

Poprawność wykonania operacji sprawdzamy, odczytując rejestr statusu I²CStat: 0x0A: 52 0a 50, która powinna zwrócić wartość 0xF0, zgodnie z dokumentacją, której fragment pokazuje **rysunek 7**.

Po sprawdzeniu działania wszystkich wyjść i magistrali I²C moduł można zastosować we własnej aplikacji.

Adam Tatuś, EP

REKLAMA

świat radio
Magazyn wszystkich użytkowników eteru
KRÓTKOFALARSTWO CB RADIOELEKTRONIKA

przejrzyj i kup na
www.ulubionykiosk.pl

Nowości sprzętowe, prezentacje ds. Ham Radio
świat radio
 Magazyn wszystkich użytkowników eteru
Yaesu FTM-6000

**Podstawowe parametry:**

- wytwarzanie ujemnego napięcia stałego przy użyciu przetwornicy impulsowej,
- regulowana wartość napięcia wyjściowego w zakresie -26...-1,23 V,
- wejście wyłaczające przetwornicę,
- prąd wyjściowy około 250 mA przy UWE=12 V i UWY=-12 V,
- zasilanie napięciem stałym z przedziału 5...35 V,
- pobór prądu około 10...30 mA, zależnie od napięcia wejściowego i wyjściowego.

* **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 wzlutować 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 wzlutowane 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

- AVT5977 Stabilizator napięcia symetrycznego z regulacją współbieżną (EP 3/2023)
- AVT5963 Zasilacz warsztatowy cz. 1 i 2 (EP 12/2022 01/2023)
- 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)

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*

AVT6002

Regulowany zasilacz napięcia ujemnego

Ujemne napięcie bywa przydatne, zaś jego źródeł mamy niewiele – zwłaszcza kiedy zasilamy urządzenie z akumulatora bądź nie mamy możliwości rozbudowy już istniejącego zasilacza sieciowego. Zaprezentowany układ nie dość, że potrafi „odwrócić” polaryzację napięcia, to na dodatek pozwala wyregulować jego wartość w szerokim zakresie.

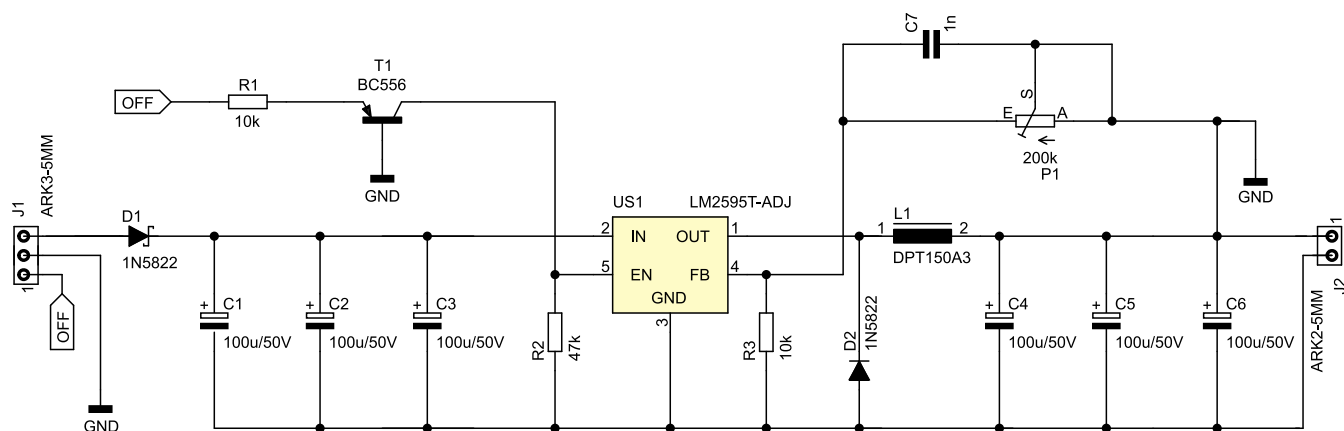
Szukałem gotowego modułu, który dawałby napięcie niższe niż 0 V i o wydajności kilkuset miliamperów, do ładowania zestawu ogniw Li-Ion, zasilających jedną z dwóch gałęzi układu analogowego. Ładowanie wszystkich ogniw (zarówno dla gałęzi ujemnej, jak i dodatniej) odbywało się z pojedynczego zasilacza wtyczkowego, więc wspomnianego napięcia ujemnego próżno szukać. Trzeba je było wytworzyć w układzie.

Poszukiwania nie przyniosły efektów – większość układów wytwarzających napięcie ujemne względem masy bazuje na pompie ładunkowej. To sprytnie i tanie rozwiązanie, lecz cechuje je bardzo niska wydajność prądowa, ponadto ma ono silne ograniczenia dotyczące napięcia wyjściowego. Opracowany przeze mnie moduł

bazuje na sprytnie podłączonej przetwornicy impulsowej, więc jego wydajność prądowa jest znacznie wyższa. Potencjometrem można wygodnie regulować napięcie wyjściowe w bardzo szerokim zakresie, zaś kolejnym atutem jest łatwość uruchomienia i duża dostępność podzespołów.

**Budowa i działanie**

Schemat ideowy omawianego układu znajduje się na **rysunku 1**. Najważniejszym



Rysunek 1. Schemat ideowy regulowanego zasilacza napięcia ujemnego

Wykaz elementów, kupuj na stronie sklepu avt.pl (Warszawa, ul. Leszczyńska 11, tel. +48222578451, e-mail: handlowy@avt.pl)**Rezystory:** (o mocy 0,25 W)

R1, R3: 10 kΩ
R2: 47 kΩ
P1: 200 kΩ potencjometr pionowy 3296W

Kondensatory:

C1...C6: 100 μF 50 V Low ESR (10 mm, raster 5 mm),

np. EGF107M1HG1BRRSHP SAMXON
C7: 1 nF MKT 5 mm

Półprzewodniki:

D1, D2: 1N5822
T1: BC556
US1: LM2595T-ADJ (TO220-5)

Pozostałe:

J1: ARK2/500
J2: ARK3/500
L1: DPT150A3 Ferrocore
Ew. radiator (opis w tekście)

elementem jest scalona przetwornica impulsowa US1 typu LM2595 w obudowie TO220-5 i w wersji z regulowanym napięciem wyjściowym (LM2595T-ADJ). Warto zwrócić uwagę, że zacisk GND tego układu nie jest połączony z masą zasilania, lecz prowadzi do zacisku wyjściowego złącza J2. Z kolei napięcie wyjściowe, generowane przez tę przetwornicę pracującą w klasycznym układzie buck, jest... połączone z masą! Nie, to nie błąd. W ten sposób możemy uzyskiwać napięcie ujemne, gdyż dla układu US1 potencjał masy jest tym wyższym, zaś potencjał jego linii GND – niższym. Zatem z punktu widzenia US1 wszystko jest w porządku. Z kolei napięcie wejściowe, przykładane do zacisków złącza J1, trafia między zaciski IN oraz wyjście tej przetwornicy.

Zatem jak ona jest zasilana? Wystarczy przyjrzeć się pojemnościowemu dzielnikowi napięcia, jaki tworzy się z dwóch pojemności: równolegle połączonych C1...C3 oraz równolegle połączonych C4...C6. Ich pojemności elektryczne są – w przybliżeniu – równe, więc napięcie zasilające dzieliłoby się po połowie między obie te grupy kondensatorów, jednak C4...C6 są zwarte przez P1 i R3, więc nie naładują się prawie wcale – większość napięcia wejściowego odłoży się na C1...C3. W takich warunkach układ US1 może rozpocząć działanie. Nawiasem mówiąc, ta przetwornica pracuje inaczej niż w tradycyjnym układzie buck – wartość międzyszczytowa tętnień prądu jest wyższa niż w układzie buck.

Potencjometr P1 wraz z rezystorem R3 tworzą dzielnik napięcia wyjściowego. Ich zadaniem jest zamknięcie pętli ujemnego sprzężenia zwrotnego, aby napięcie wyjściowe mogło zawsze znajdować się na tym samym poziomie, ustalonym uprzednio przez potencjometr P1. Maksimum napięcia wyjściowego, które da się w tym układzie ustawić, to -1,23 V – tyle wynosi napięcie referencyjne układu LM2595. Z kolei minimum, jakie wynika z dzielnika R3+P1, to około -26 V. Kondensator C7 przyspiesza działanie pętli ujemnego sprzężenia zwrotnego dla wysokich częstotliwości, poprawiając stabilność pracy przetwornicy.

W tym układzie zastosowano po trzy kondensatory wejściowe i trzy wyjściowe. Taka decyzja nie była podyktowana rozrzutnością, lecz chęcią zmniejszenia natężenia składowej zmiennej prądu, jaki płynie przez te elementy. Użyte w prototypie elementy firmy SAMXON mają impedancję 120 mΩ, zaś dopuszczalny prąd tętnień

wynosi 760 mA. Połączenie równoległe trzech kondensatorów po trzykroć poprawia te parametry: zastępcza impedancja wynosi około 40 mΩ, zaś dopuszczalny prąd tętnień wzrasta do ponad 2,2 A.

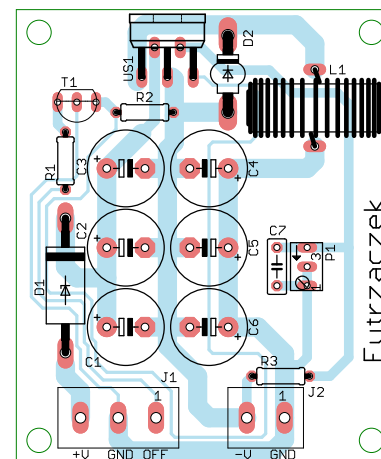
Ostatnim obwodem, który nie był dotychczas omówiony, jest zagadkowy tranzystor T1 wraz z rezystorami R1 i R2. Na dodatek T1 jest wpięty w układzie wspólnej bazy, po co...? Otóż ta przetwornica jest normalnie włączona: rezystor R2 polaryzuje wejście EN układu US1 potencjałem równym potencjałowi nóżki GND tego układu. Aby go wyłączyć, trzeba podać tam potencjał wyższy. Tylko jak to zrobić, jeżeli wszystko rozgrywa się na potencjałach niższych niż 0 V? Właśnie w tym może nam pomóc T1, który służy jako sterowane prądem źródło prądowe. Z jego kolektora płynie w stronę R2 prąd o takim samym (z dokładnością do prądu bazy) natężeniu, jakie wpływa do jego emitera. Ten prąd jest ograniczany przez rezystor R1 do ułamka miliampera, jeżeli podane na J1 napięcie wyłączające wynosi kilka woltów – na przykład 5 V albo 3,3 V z mikrokontrolera. Ponieważ rezystancja R2 jest kilkukrotnie większa niż R1, to odłożone na nim napięcie również będzie kilkukrotnie wyższe niż na R1. US1 na pewno się wyłączy, lecz nie grozi mu uszkodzenie, bo natężenie prądu wypływającego z kolektora T1 jest niewielkie.

Montaż i uruchomienie

Układ został zmontowany na niewielkiej, jednostronnej płytce drukowanej o wymiarach 50 × 60 mm. Jej 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 najniższych obudowach, czyli rezystorów i diody D1. Dioda D2 jest montowana pionowo, zwrócona katódą w stronę układu US1. Pod układem US1 znajduje się fragment ścieżki odsłonięty spod maski lutowniczej, polecam dodatkowo pogrubić go cienkim drutem lub chociaż warstwą spoiwa lutowniczego – prowadzi on prąd wejściowy dla US1, zatem jego impedancja powinna być możliwie niska.

Prawidłowo zmontowany układ nie wymaga żadnych czynności uruchomieniowych i jest od razu gotowy do pracy. Zanim jednak zostanie wpięty w docelowe miejsce, polecam podłączyć go do zasilacza o napięciu kilku woltów i ustawić żądane napięcie wyjściowe przy pomocy potencjometru P1, bez



Rysunek 2. Schemat płytki PCB

żadnego obciążenia. Dlaczego? Dopuszczalne napięcie wejściowe układu LM2595T-ADJ wynosi 40 V. W układzie z rysunku 1 układ US1 „widzi” na swoim wyprowadzeniu IN potencjał napięcia wejściowego (pomniejszony o niewielki spadek na diodzie D1), zaś na zacisku GND potencjał napięcia wyjściowego. Różnica potencjałów daje napięcie, więc w tym układzie musi być spełniony warunek $U_{WE} - U_{WY} < 40$ V. Jeżeli ustawione napięcie wyjściowe byłoby wysokie, to US1 może ulec uszkodzeniu po przyłożeniu wysokiego napięcia wejściowego.

Dużą zaletą tego układu jest możliwość ustawienia dowolnego napięcia wyjściowego: zarówno (co do wartości bezwzględnej) niższego, jak i wyższego względem napięcia wejściowego. Trzeba jedynie pamiętać o powyższym warunku, jak i o tym, że przy niskich napięciach wyjściowych spada sprawność układu oraz dopuszczalny prąd wyjściowy, ponieważ rośnie natężenie prądu płynącego przez klucz wbudowany w US1. Źródło zasilania dla tego układu powinno mieć wysoką wydajność prądową, a jeszcze lepiej, gdyby było możliwe uruchomienie tej przetwornicy z pewnym opóźnieniem względem podania napięcia zasilającego. Ma to na celu szybkie naładowanie kondensatorów w układzie.

Na koniec – jedna uwaga dotycząca ewentualnych modyfikacji. Przestrzegam przed bezwiednym wlutowaniem układu LM2596T-ADJ w tę płytkę – jego układ wyprowadzeń jest inny niż użytego w projekcie LM2595T-ADJ. Mimo, że jest to mocniejszy bliźniak, to dwa jego wyprowadzenia (IN oraz OUT) są zamienione miejscami. Warto mieć to na uwadze, by nie puścić układu z dymem zaraz po pierwszym włączeniu.

Michał Kurzela, EP

**Podstawowe parametry:**

- wyłączenie odbiornika po wykryciu obniżonego poboru prądu,
- załączanie przyciskiem monostabilnym na zadany czas,
- podtrzymywanie załączenia przez cały czas wykrycia podwyższonego poboru prądu,
- całkowity brak poboru mocy przez układ po wyłączeniu,
- maksymalna moc odbiornika: ok. 1000 W,
- podtrzymanie styków przełącznika: około 30 s z możliwością regulacji,
- minimalna moc wymagana do podtrzymania zasilania: regulowana w zakresie 0,7...40 W,
- zasilanie napięciem przemiennym 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 wzlutować 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] – 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

- AVT5937 Automatykny wyłącznik zestawu audio (EP 6/2022)
 ---- Automatykny wyciszacz dźwięku po zaniku zasilania (EP 9/2020)
 AVT5717 Opóźniacz dołączenia głośników zasilany 230 V (EP 9/2019)
 AVT2854 Opóźniacz dołączenia głośników (EdW 2/2008)

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>

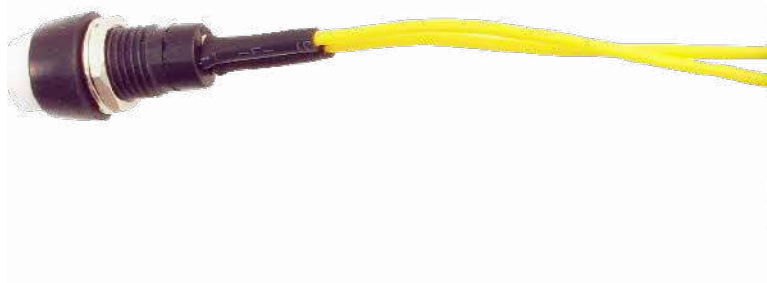
W przypadku braku dostępności na stronie sklepu osoby zainteresowane zakupem płytek drukowanych (PCB) prosimy o kontakt via e-mail: kity@avt.pl.

W ofercie AVT*

AVT6003

Standby killer

Zaprezentowany układ to prawdziwy pogromca! Jednak nie trzeba trzymać go za kratami, gdyż na jego celowniku znajduje się wyłącznie tryb standby popularnych urządzeń RTV. Stojąc na straży niskiego zużycia energii w gospodarstwie domowym, dba też o nasze portfele.



Ten, kto wymyślił tryb obniżonego poboru mocy (popularnie znany jako „standby”), na pewno miał na uwadze ludzką wygodę. Tak uśpione urządzenie można przecież łatwo i szybko załączyć dotykowym przyciskiem lub pilotem zdalnego sterowania. Często też jest zachowywana większość ustawień, więc szybciej można korzystać z naszego sprzętu. Jest jednak pewne „ale”...

Tradycyjne wyłączniki sieciowe dawały możliwość całkowitego odłączenia urządzenia od sieci. Ile ono wtedy pobiera? Nic, zero,

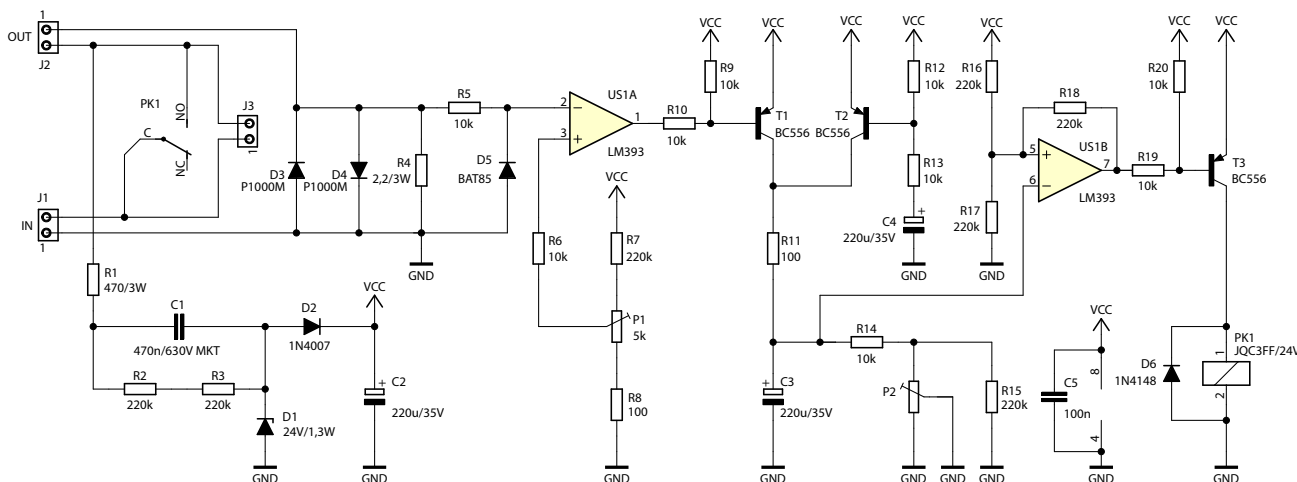
jest zupełnie wyłączone. Zdecydowana większość nowoczesnych urządzeń oferuje nam w ramach wyłączenia jedynie wspomniany tryb standby, fizyczne przerwanie obwodu zasilania jest już rzadko kiedy możliwe. Zatem te urządzenia pobierają energię cały czas.

Ile jej wtedy zużywają? Tego tak naprawdę nie wiemy, producenci sprzętu też nie kwapią się, by podawać ów parametr w sposób jednoznaczny i czytelny. Często okazuje się, że tryb standby sprowadza się do wyłączenia

wyświetlacza i modułów wyjściowych, więc pobór mocy wcale nie staje się taki niski. Jaki jest tego skutek? Ta energia, pochłaniana w trybie standby, jest pobierana całą dobę, 7 dni w tygodniu, co licznik energii elektrycznej skrupulatnie rejestruje. Można to zmienić bez potrzeby wyciągania wtyczki z gniazdka.

Budowa i działanie

Schemat ideowy omawianego układu znajduje się na **rysunku 1**. Napięcie pochodzące



Rysunek 1. Schemat ideowy układu „standby killera”

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

C1: 470 nF 630 V raster 22,5 mm MKT
C2...C4: 220 µF 35 V raster 3,5 mm
C5: 100 nF raster 5 mm MKT

Półprzewodniki:

D1: dioda Zenera 24 V, 1,3 W
D1: 1N4007
D3, D4: P1000M
D5: BAT85

D6: 1N4148
T1...T3: BC556
US1: LM393 (DIP8)

Rezystory: (THT o mocy 0,25 W, jeżeli nie napisano inaczej)

R1: 470 Ω 3 W
R2, R3, R7, R15...R18: 220 kΩ
R4: 2,2 Ω 3 W
R5, R6, R9, R10, R12...R14, R19, R20: 10 kΩ
R8, R11: 100 Ω

P1: 5 kΩ jednoobrotowy montażowy leżący
P2: opis w tekście

Pozostałe:

J1...J3: ARK2/500
PK1: JQC3FF/2412S
Jedna podstawka DIP8
Przycisk monostabilny (opis w tekście)

z sieci elektrycznej (230 V AC) podłącza się do zacisków złącza J1, zaś odłączane urządzenie czerpie energię z zacisków złącza J2. Pomiędzy nimi znajdują się styki przekaznika PK1, który w normalnych warunkach rozłącza obwód. Do uruchomienia tego przekaznika, jak i zasilanego urządzenia, służy przycisk monostabilny, który powinien zostać podłączony do zacisków złącza J3. Pozwala on na chwilę obejść rozwarne styki, co zasilą pozostałą część układu, więc możliwe będzie pobranie energii do zasilenia cewki przekaznika. Po około dwóch sekundach trzymanie palcem przycisku stanie się zbędne, jego rolę przejmie PK1.

Do pomiaru natężenia prądu pobieranego przez obciążenie służy rezystor R4, na którym ów prąd wywołuje spadek napięcia – zgodnie z prawem Ohma. Jednak nie musimy mierzyć prądów o bardzo wysokim natężeniu, zależy nam jedynie na informacji, czy urządzenie jest w trybie normalnej pracy, czy też czuwania (standby). Dlatego równolegle do tego rezystora zostały dołączone antyrównolegle dwie diody krzemowe, więc spadek napięcia na tymże rezystorze nie przekroczy około 0,8 V w którąkolwiek stronę. Nawet gdyby zasilane urządzenie przez chwilę pobierało bardzo wysoki prąd, to R4 nie zostanie przegrzany, bo napięcie odłożone na nim zostanie zredukowane do wartości równej napięciu przewodzenia jednej z tych diod. Trzeba jednak pamiętać o ciepłe wydzielanym w tychże diodach, więc sytuacja przeciążenia nie może trwać zbyt długo.

Napięcie odkładające się na R4 trafia na wejście komparatora US1A. Ponieważ ten komparator jest zasilany asymetrycznie, rozpatrywana będzie jedynie dodatnia połówka tego napięcia. Rezystor R5 ogranicza natężenie prądu płynącego przez wejście komparatora, zaś w szczególności chroni diodę D5 przed przepływem przez nią prądu o wysokim natężeniu w czasie występowania ujemnych połówek napięcia pochodzącego z R4. Dioda ta otwiera się przy napięciu około -0,2 V, przez co nie pozwala na otwarcie złącza baza-kolektor tranzystora wejściowego układu US1A.

Komparator US1A sprawdza, czy natężenie pobieranego prądu jest dostatecznie wysokie do potrzymania działania przekaznika. Napięcie referencyjne pobiera z dzielnika napięcia R7+P1+R8, który dzieli napięcie zasilające układ (około 24 V) w stosunku ustalonym przez użytkownika. Minimalne napięcie wychodzące z tego dzielnika wynosi około

10 mV, zaś maksymalne około 540 mV. To przekłada się na przedział mocy czynnej pobieranej przez odbiornik, wynoszący 0,7...40 W. Jeżeli pobierana moc jest wyższa niż ustalony potencjometrem próg, wyjście komparatora US1A załącza się cyklicznie (w każdym okresie napięcia sieciowego, najdłużej na pół okresu) i podtrzymuje działanie przekaznika. Niedostateczna amplituda impulsów napięcia wejściowego nie powoduje reakcji tego podzespołu, przez co uruchamiane jest odliczanie czasu do wyłączenia. R6 zgrubnie wyrównuje rezystancje „widziane” przez oba wejścia tego komparatora.

Załączające się wyjście komparatora otwiera tranzystor T1, ponieważ potencjał jego bazy ulega obniżeniu. Rezystor R10 ogranicza jej prąd do bardzo bezpiecznej dla niej wartości, zaś R11 utrzymuje ten tranzystor w stanie zatkania po wyłączeniu wyjścia US1A. Prąd wypływający z kolektora T1 doładowuje kondensator C3 za każdym razem, gdy wyjście komparatora załączy się. R11 ogranicza prąd tego ładowania, aby nie występowały silne impulsy prądu pobieranego z zasilania.

Kondensator C3 pełni zatem funkcję elementu pamiętającego – cyklicznie doładowywany będzie podtrzymywał napięcie na swoich zaciskach, zaś w chwili zaprzestania zacznie się rozładowywać. O to dba rezystor R15, którego wartość ustala czas podtrzymania na około 30 sekund. Jeżeli byłoby to dla kogoś za dużo lub zbyt mało, może zamiast R15 włączyć potencjometr P2 o dowolnej wartości i ten czas wyregulować samodzielnie. Dzięki R14 nawet skrajnie mała rezystancja potencjometru P2 nie spowoduje uszkodzenia układu.

Skąd jednak układ ma wiedzieć, że właśnie nastąpiło jego załączenie i trzeba podtrzymać działanie PK1 nawet pomimo niskiego poboru prądu przez odbiornik? Temu służy tranzystor T2, który również może naładować C3 prądem wypływającym ze swojego kolektora. Kondensator C4, który jest wstępnie rozładowany, po załączeniu zasilania przyciskiem monostabilnym zaczyna się ładować prądem pochodzącym w dużej mierze z bazy tranzystora T2. To otwiera ten tranzystor i daje jednorazowy impuls ładujący C3. R13 ogranicza prąd tego ładowania, natomiast R12 utrzymuje T2 w stanie zatkania już po naładowaniu oraz przyczynia się do rozładowania C4 po wyłączeniu zasilania. Bez tego prostego obwodu konieczne byłoby trzymanie

przycisku oraz jednoczesne uruchomienie odłączanego urządzenia. To mało wygodne rozwiązanie, dlatego można najpierw wcisnąć przycisk i przytrzymać go przez chwilę, a potem spokojnie uruchomić urządzenie, na przykład telewizor przy użyciu pilota.

Wysokie napięcie na zaciskach C3 powoduje załączenie wyjścia komparatora US1B. Jego napięciem referencyjnym jest połowa napięcia zasilającego, za co odpowiadają rezystory R16 i R17. Został objęty pętłą głębokiego, dodatniego sprzężenia zwrotnego (poprzez R18) tak, aby jego progi przełączenia różniły się od siebie znacząco. W ten sposób nie dochodzi do cyklicznego załączania i wyłączania przekaznika PK1, który jest sterowany za pośrednictwem tranzystora T3. Jeżeli napięcie na C3 spadnie do odpowiednio niskiej wartości, PK1 rozłącza się, przez co odłącza od napięcia sieciowego zarówno podłączone do złącza J2 urządzenie, jak i sam układ „standby killera”.

Pozostało omówienie źródła niskiego napięcia dla tego układu. Aby zredukować jego cenę do możliwie niskiej wartości, postanowiono na prosty obwód zasilacza beztransformatorowego z elementem reaktancyjnym w postaci kondensatora C1. Rezystor R1 ogranicza prąd jego ładowania, by nie dochodziło do iskrzenia styków przełącznika monostabilnego. Zadaniem R2 i R3 jest rozładowywanie tego elementu po wyłączeniu zasilania. Na diodzie Zenera D1 odkłada się tętniące napięcie o wartości szczytowej około +24 V. Te impulsy doładowują kondensator elektrolityczny C2 za pośrednictwem zwykłej diody półprzewodnikowej D2. Z zacisków C2 pobierane jest zasilanie dla całego układu. Takie rozwiązanie nie daje izolacji galwanicznej od źródła wysokiego napięcia, lecz obwód pomiaru prądu i tak wymusza połączenie masy układu z jednym z przewodów zasilających 230 V.

Montaż i uruchomienie

Układ został zmontowany na jednostronnej płytce drukowanej o wymiarach 80×60 mm. Jej schemat został pokazany na rysunku 2. W odległości 3 mm od krawędzi płytki znalazły się cztery otwory montażowe, każdy o średnicy 3,2 mm.

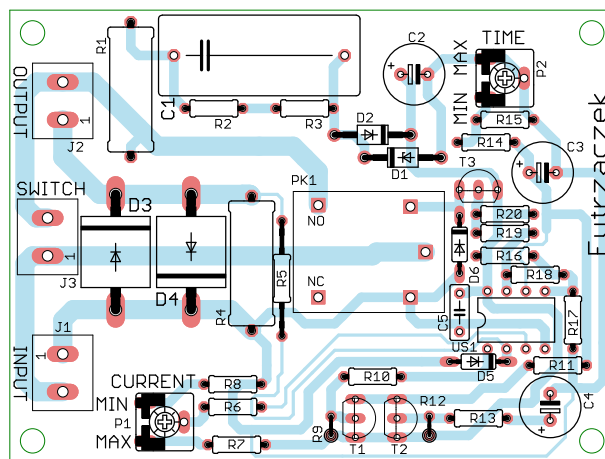
Montaż proponuję rozpocząć od elementów o najmniejszej wysokości obudowy, czyli rezystorów i diod półprzewodnikowych. Pod układ scalony US1 proponuję zastosować podstawkę, aby ułatwić jego wymianę

w razie ewentualnego uszkodzenia. Na fotografii tytułowej można zobaczyć ten układ z dołączonym przyciskiem monostabilnym. Dobierając konkretny element, należy mieć na uwadze, że przez jego styki płynie prąd zasilający odbiornik, który w tym układzie może wynosić nawet ponad 4 A. W prototypie zastosowano PRZEL 35 z oferty sklepu AVT, który może przewodzić prąd o natężeniu do 2 A. Większe natężenie mogą przewodzić inne przełączniki, jak PS507A-BR (do 3 A) lub PRZEL392 (do 5 A).

Poprawnie zmontowany układ jest gotowy do działania po podaniu zasilania. Jediną czynnością regulacyjną, jaką należy wykonać, jest ustalenie progu zadziałania przy użyciu potencjometru P1. Skręcenie go w stronę MIN zwiększa czułość (mniejsza moc pobierana przez odbiornik będzie podtrzymywać przełącznik), zaś w stronę MAX – zmniejsza, czyli urządzenie pobiera w trybie standby większą moc. Najprostszą metodą na ustalenie tego progu jest ustawienie P1 w połowie i załączenie zasilania

przełącznikiem monostabilnym. Jeżeli urządzenie zostanie załączone do normalnej pracy i PK1 nie rozłączy zasilania po około 30 sekundach, zaś kiedy wejdzie w tryb standby i PK1 odłączy zasilanie po upływie tego czasu, to wszystko jest w porządku. Gdyby urządzenie było wyłączone w czasie normalnej pracy, trzeba skrócić P1 w stronę MIN, zaś gdyby nie następowało jego wyłączenie po przejściu w tryb standby, trzeba skrócić P1 w stronę MAX. Kiluminutowa regulacja wystarczy nawet na wiele lat.

Jeżeli podtrzymanie styków przełącznika przez 30 s to za mało lub zbyt dużo, można wymienić R15 na rezystor o innej wartości (mniejsza – krótszy czas, większa – dłuższy czas) lub wymontować go i włutować



Rysunek 2. Schemat płytki PCB

do układu potencjometr P2 o dowolnej wartości, na przykład 1 MΩ – pozwoli to na regulację w zakresie od około sekundy do nawet kilku minut.

Michał Kurzela, EP

REKLAMA

sklep.avt.pl

- Nauka elektroniki
- AVT Kits
- Elektronika
- Sprzęt pomiarowy i zasilanie
- Warsztat
- Dom i ogród





Precyzja i dokładność w Systemach Globalnej Nawigacji Satelitarnej

W specyfikacjach modułów i systemów do nawigacji satelitarnej podawana jest dokładność pozycjonowania, jednak większość z nich nie wyjaśnia, czym tak naprawdę jest ta dokładność. W artykule wyjaśnimy znaczenie parametrów DOP, które często opisują moduły GNSS. Ponadto postaramy się zbadać, od czego zależy dokładność pozycjonowania i w jaki sposób projektować i użytkować systemy z odbiornikami nawigacji satelitarnej, aby były one najdokładniejsze.

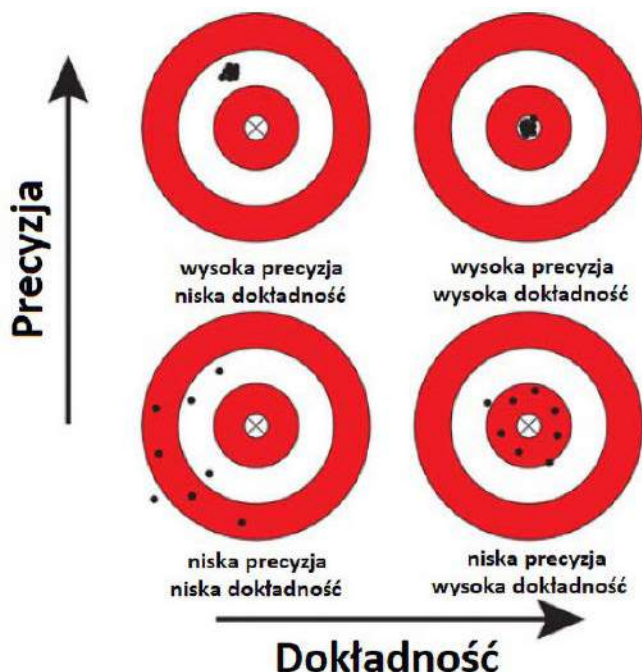
W dzisiejszych czasach systemy globalnej nawigacji satelitarnej (GNSS) stanowią kluczowy element naszego codziennego życia, zapewniając precyzyjne i dokładne informacje o pozycji, prędkości i czasie na całym świecie. Jednak nawet w przypadku zaawansowanych technologicznie systemów GNSS, tematy określania i maksymalizacji precyzji i dokładności pozostają nadal istotne dla inżynierów elektroników oraz specjalistów z dziedziny nawigacji i geodezji.

W artykule skupimy się na zrozumieniu kluczowych pojęć takich jak dokładność (*Accuracy*) i precyzja (*Precision*) w kontekście

systemów GNSS. Będziemy analizować, jak te dwa parametry odnoszą się do jakości pozycjonowania przy użyciu satelitarnych sygnałów nawigacyjnych, zwracając uwagę na ważne terminy, takie jak VDOP (*Vertical Dilution Of Precision*), HDOP (*Horizontal Dilution Of Precision*), które odzwierciedlają błędy pozycjonowania w trzech wymiarach.

DOP – *Dilution Of Precision*, jest szczególnie istotne dla dokładności pozycjonowania w systemach nawigacyjnych, również tych satelitarnych. W dalszej części artykułu sprawdzimy, jak różne czynniki, takie jak geometria satelitów, wpływają na DOP i jakie konsekwencje może to mieć dla ostatecznej precyzji pozycjonowania. W artykule omówimy nie tylko teoretyczne aspekty precyzji i dokładności w systemach GNSS, ale także zbadamy praktyczne metody redukcji *Dilution Of Precision* w rzeczywistych zastosowaniach. Przyjrzymy się różnym technikom i strategiom, które inżynierowie mogą zastosować, aby zminimalizować wpływ DOP na dokładność pozycjonowania w gotowych systemach nawigacyjnych. Wśród przykładów takich metod znajdują się techniki filtrowania danych, optymalizacja wyboru satelitów oraz zastosowanie danych z wielu źródeł nawigacyjnych.

Zapraszam do lektury dalszych sekcji artykułu, aby zgłębić tę fascynującą tematykę – od matematycznych definicji poszczególnych wartości, poprzez praktyczne strategie poprawy dokładności i precyzji



Rysunek 1. Dokładność a precyzja na przykładzie celowania do tarczy

w systemach GNSS, które mają kluczowe znaczenie dla dzisiejszych zaawansowanych aplikacji nawigacyjnych, lokalizacyjnych i nie tylko.

Dokładność i precyzja

Zanim przejdziemy do dalszej części i przyjrzymy się bardziej złożonym parametrom, warto uściślić, czym w ogóle jest dokładność i precyzja. Często te dwa słowa stosuje się niemalże wymiennie w kontekście omawiania pomiarów z użyciem systemów globalnej nawigacji satelitarnej, jednak w rzeczywistości dotyczą one dwóch różnych aspektów pomiaru (nie tylko pomiaru pozycji z użyciem GNSS – to ogólna definicja), odpowiednio, bliskości wartości prawdziwej oraz powtarzalności pomiaru.

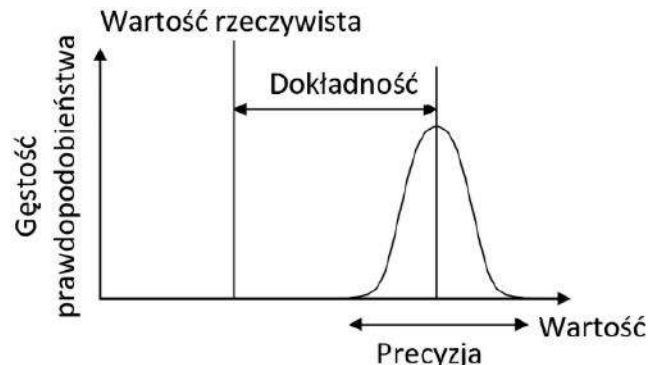
W życiu codziennym słowa dokładność i precyzja są również często używane zamiennie. Jednak, jako terminy w metrologii, mają one dwie odmienne definicje. To, że pomiar jest dokładny, nie oznacza, że jest precyzyjny, i odwrotnie. Zarówno dokładność, jak i precyzja są kluczowymi aspektami dobrych pomiarów, ale czym dokładnie są te wartości? Spójrzmy na ich definicje i różnice, które z nich wynikają.

Jaka jest różnica między dokładnością a precyzją?

Dokładność i precyzja są dwiema formami opisu pomiaru, które określają, jak blisko jest on trafienia w cel. Dokładność ocenia, jak blisko wartość pomiaru jest prawdziwej wartości mierzonej zmiennej, podczas gdy precyzja pokazuje, jak blisko siebie są poszczególne wartości pomiarów.

Przykład z tarczą jest najczęstszym sposobem na pokazanie różnicy między dokładnością a precyzją. Rozważmy strzelanie do tarczy albo rzucanie lotkami. Celem jest oczywiście osiągnięcie zarówno wysokiej dokładności, jak i precyzji. Innymi słowy, trafiać w centrum tarczy jak najczęściej.

Jeśli nasze strzały były dokładne, oznacza to, że trafienia były blisko środka tarczy, ale generalnie lądowały na całej jej powierzchni. Jeśli strzały były z kolei precyzyjne, to oznacza, że punkty trafienia znajdują się blisko siebie, ale niekoniecznie blisko środka tarczy. Zobrazowane jest to na **rysunku 1**, korzystając z analogii tarczy strzeleckiej. Dopiero kiedy strzały są dokładne i precyzyjne, będą znajdować się w większości w środku tarczy, jak widać na **rysunku 1**. Aspekty te mają również istotne znaczenie w przypadku nie tylko strzelania do tarczy, ale innych pomiarów, w tym także pomiaru



Rysunek 2. Matematyczna reprezentacja dokładności i precyzji pomiaru

lokalizacji z użyciem GNSS. Oczywiście, mają one również definicje bardziej ścisłe i ujęcie nie jakościowe, a ilościowe.

Czym jest dokładność?

Dokładność wskazuje, jak blisko rzeczywistej wartości jest uzyskany w pomiarze wynik. Jest miarą poprawności pomiaru. Matematycznie mówiąc, miarą dokładności jest odległość pomiędzy wartością oczekiwaną pomiaru (np. wartością średnią) a wartością rzeczywistą, jak pokazano na **rysunku 2**. Oczywiście w realnym przypadku nie znamy wartości rzeczywistej, jednak oszacowanie jej położenia jest możliwe dzięki danym statystycznym.

W przypadku systemów nawigacyjnych dokładność określana jest za pomocą DOP – Dilution Of Precision (rozmycie dokładności), które wyznaczone jest geometrycznie, jak opisano w dalszej części artykułu.

Czym jest precyzja?

Precyzja mierzy, jak bliskie siebie są pomiary. Pokazano to na **rysunku 2**. Miarą precyzji może być np. odchylenie standardowe pomiarów. Odchylenie standardowe jest pierwiastkiem kwadratowym z wariancji. Oblicza się ją dla wielu pomiarów. Wariancja jest miarą zmienności danych i wynika z błędów losowych, popełnianych przez nasze urządzenia pomiarowe oraz rozkładu mierzonej zmiennej.

Normy

Znaczenia tych terminów zostały ściśle określone wraz z opublikowaniem norm z serii ISO 5725 w 1994 roku, co znalazło również odzwierciedlenie w wydaniu Międzynarodowego Słownika Metrologicznego (VIM) BIPM (Międzynarodowe Biuro Miar i Wagi) z 2008 roku, w punktach 2.13 i 2.14. Zgodnie z ISO 5725-1, ogólny termin „dokładność” jest używany do opisanego zbliżenia pomiaru do rzeczywistej wartości. Kiedy termin jest stosowany do zestawów pomiarów tego samego obiektu mierzalnego, obejmuje składnik błędu losowego oraz składnik błędu systematycznego. W tym przypadku trafność oznacza zbliżenie średniej wyników pomiaru do wartości rzeczywistej (prawdziwej), a precyzja oznacza zbliżenie uzgodnienia między zestawem wyników.

ISO 5725-1 i VIM unikają również użycia terminu obciążenie (*bias*), wcześniej określonego w normie BS 5497-1, ponieważ ma on różne konotacje poza dziedzinami nauki i inżynierii, na przykład w medycynie, prawie czy systemach sztucznej inteligencji.

Dilution Of precision

Termin Dilution Of Precision (DOP) to parametr opisujący wpływ geometrii konstelacji satelitów na wyznaczenie pozycji w systemie nawigacji satelitarnej. Koncepcja DOP wywodzi się od użytkowników systemu nawigacji Loran-C. Idea geometrycznego wyznaczania DOP polega na określeniu, w jaki sposób błędy w pomiarze wpłyną na ostateczne oszacowanie stanu. Można to zdefiniować jako:

$$DOP = \frac{\Delta(\text{zmierzona pozycja})}{\Delta(\text{zmierzone dane})}$$

Loran-C – hiperboliczny system nawigacji radiowej, który umożliwia określenie pozycji poprzez nasłuch sygnałów radiowych o niskiej częstotliwości nadawanych przez nadajniki naziemne o znanej pozycji. Po wprowadzeniu w 1957 roku z uwagi na wysoki koszt odbiorników Loran-C miał zastosowanie głównie militarne.

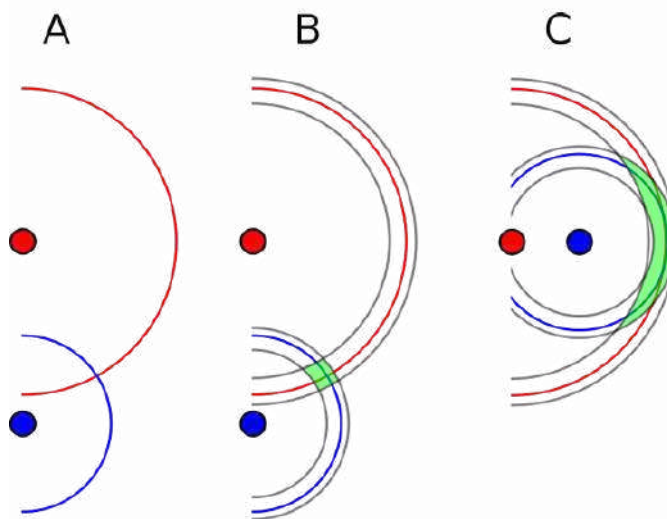
Koncepcyjnie oznacza to, że błędy w pomiarze skutkują zmianą Δ (zmierzonych danych). Dla idealnego systemu o niskim DOP niewielkie zmiany w danych pomiarowych nie powinny prowadzić do dużych zmian w wynikowym położeniu. Przeciwnością tej sytuacji jest rozwiązanie bardzo wrażliwe na błędy pomiarowe. Interpretacja geometryczna (dla przykładu 2D) jest pokazana na **rysunku 3**, gdzie pokazano dwa możliwe scenariusze z akceptowalnym i słabym DOP.

W systemie nawigacji satelitarnej to, co jest de facto mierzone, to odległość od satelitów nawigacyjnych. Na przecięciu się okręgów o promieniu równym zmierzonej odległości znajduje się odbiornik GNSS. Na rysunku 3a pokazano idealny przypadek, tj. gdy pomiary obu odległości nie są obciążone błędem. Wtedy wynikiem pomiaru pozycji jest punkt – miejsce przecięcia okręgów o promieniu równym odległości pomiędzy satelitą a odbiornikiem (w rzeczywistości są to dwa punkty, jednak tylko jeden z nich leży w sensownym miejscu, np. na powierzchni Ziemi, drugi znajduje się np. w przestrzeni kosmicznej, co pozwala na jego łatwe odrzucenie).

Niestety w rzeczywistości nie jest tak łatwo – wszystkie pomiary obciążone są pewnym błędem, także pomiary odległości, jak pokazano na rysunkach 3b i 3c. Na tej wizualizacji widać, że w momencie, gdy pomiary odległości obciążone są błędem, punkt przecięcia (punkt, gdzie znajduje się odbiornik) ulega rozmyciu (obszar zakolorowany na zielono na rysunku 3). Ten sam błąd pomiaru odległości przekładać może się na różny błąd pozycjonowania, zależnie od wzajemnego położenia satelitów względem siebie i względem pozycji odbiornika – o tym właśnie mówi DOP. Na rysunku 3b system wykazuje niski DOP, a na rysunku 3c DOP jest wysoki – mimo takich samych błędów pomiaru odległości od satelitów.

Z uwagi na względną geometrię dowolnego danego satelity względem odbiornika, precyzja pseudoodległości satelity przekłada się na odpowiadającą jej składową w każdym z czterech wymiarów pozycji mierzonych przez odbiornik (tj. x, y, z i t). Precyzja wielu satelitów widocznych z odbiornika łączy się zgodnie z położeniem względnym satelitów, aby określić poziom precyzji w każdej składowej pomiaru odbiornika. Gdy widoczne satelity nawigacyjne są blisko siebie na niebie (kąty pomiędzy nimi są niewielkie), geometria jest słaba, a wartość DOP wysoka. Gdy kąt pomiędzy satelitami jest duży, geometria sprzyja uzyskaniu niskiego DOP.

Przyjrzyjmy się dwóm nachodzącym na siebie okręgom różnych środków, jakie pokazano na rysunku 3. Jeśli nachodzą na siebie pod kątem prostym, największy zakres nachodzenia jest znacznie mniejszy niż w przypadku nachodzenia prawie równoległego. Niska wartość DOP reprezentuje lepszą precyzję pozycji dzięki szerszemu kątowemu rozdzielaniu między satelitami używanymi do obliczania pozycji jednostki. Inne czynniki, które mogą efektywnie zwiększyć



Rysunek 3. Geometryczna reprezentacja DOP dla pomiaru 2D z pomocą dwóch satelitów (czerwone i niebieskie kropki). Wynik pomiaru zaznaczony jest na zielono – to obszar przecięcia się czerwonego i niebieskiego okręgu z uwzględnieniem błędów pomiarowych (szare okręgi). a) Przypadek idealny, bez uwzględnienia błędów; b) Niski DOP; c) Wysoki DOP

DOP, to przeszkody takie jak pobliskie góry lub budynki, które ograniczają widoczność satelitów.

Rozróżniamy następujące rodzaje DOP:

- GDOP – parametr geometryczny opisujący dokładność położenia punktu w 4 wymiarach (3 wymiary przestrzenne + czas),
- HDOP – horyzontalny DOP – dla współrzędnych płaskich,
- VDOP – wertykalny DOP – dla wysokości,
- TDOP – DOP dla pomiaru czasu,
- PDOP – DOP pozycji w trzech wymiarach; współczynnik opisujący stosunek między błędem pozycji użytkownika a błędem pozycji satelity.

Parametry DOP podawane są przez odbiornik. Nie musimy ich wyznaczać, ale należy sprawdzać je, gdy prowadzimy pomiary z pomocą GNSS. W **tabeli 1** zawarto podsumowanie różnych przedziałów DOP wraz z informacjami, jak na danym poziomie DOP należy traktować pomiar.

Poprawa parametru DOP

Istnieje kilka strategii, które można zastosować, w celu minimalizacji DOP w zastosowaniach GNSS.

Optymalny wybór satelitów

Jednym ze sposobów minimalizacji DOP jest optymalny wybór satelitów. Poprzez staranne wybieranie, które satelity mają być używane do pozycjonowania, można zredukować błąd wprowadzany przez słabą geometrię satelitów. Wymaga to dobierania satelitów, które są równomiernie rozmieszczone na niebie, mają dobre kąty elewacji oraz szeroki rozkład kątów pomiędzy nimi. Dzięki temu DOP może zostać zmniejszony, a dokładność pomiaru poprawiona.

Tabela 1. Klasyfikacja wartości DOP i jego znaczenie

Wartość DOP	Ocena	Opis
<1	Idealny	Najwyższy poziom zgodności pomiaru. W takiej sytuacji można stosować wyniki do najbardziej wymagających aplikacji
1...2	Doskonały	Na tym poziomie DOP pomiary uznaje się za dostateczne do większości najbardziej wymagających aplikacji
2...5	Dobry	Odpowiada minimalnemu poziomowi DOP, który gwarantuje pomiary dostatecznie dokładne do podejmowania większości decyzji np. w systemach nawigacji satelitarnej
5...10	Średni	Przy tym DOP pozycjonowanie może być używane do pewnych obliczeń, jednak należy dążyć do poprawy pomiaru
10...20	Kiepski	Bardzo niski poziom DOP, który przekłada się na kiepską jakość pomiarów. Przy tym DOP pozycja powinna być używana jedynie do wskazywania zgrubnego położenia
>20	Zły	Przy tym poziomie DOP pomiar należy odrzucić

Zwiększenie widoczności satelitów

Innym sposobem minimalizacji DOP jest zwiększenie liczby obserwowalnych satelitów. Można to uzyskać poprzez zwiększenie zysku anteny, dzięki czemu będzie ona w stanie odbierać słabsze sygnały lub poprawę jej pozycjonowania czy orientacji, tak aby więcej satelitów było widocznych dla anteny odbiornika. Dodatkowo, wiele odbiorników GNSS może śledzić wiele konstelacji satelitarnych naraz (takich jak GPS, GLONASS, Galileo czy BeiDou). Zapewnienie, że odbiornik śledzi jak najwięcej satelitów, może zredukować DOP.

Maska kąta elewacji

Kąt maski elewacji to minimalny kąt elewacji nad horyzontem, jaki musi mieć satelita, aby być brany pod uwagę przy pozycjonowaniu. Poprzez ustawienie niższego kąta maski elewacji będzie dostępnych więcej satelitów do śledzenia, co może prowadzić do zmniejszenia DOP. Jednak ważne jest, aby mieć świadomość tego, że zbyt niskie ustawienie kąta maski elewacji może wprowadzić więcej zakłóceń sygnału, co może pogorszyć ogólną dokładność pozycjonowania. Dodatkowo, jeżeli nasz odbiornik umiejscowiony jest relatywnie nisko, wiele satelitów o niskiej elewacji może być zasłaniane lub częściowo przysłaniane przez drzewa czy pobliskie budynki, co utrudniać będzie uzyskanie dobrej jakości sygnału z tych satelitów.

Położenie odbiornika i środowisko

Fizyczne umieszczenie odbiornika oraz otoczenie mogą również wpływać na DOP. Unikanie przeszkód, takich jak wysokie budynki, góry czy gęste korony drzew, może poprawić widoczność satelitów i zredukować DOP. Zaleca się umieszczenie odbiornika na otwartej przestrzeni z wolnym widokiem na niebo.

Uśrednianie wielu pomiarów

Uśrednianie wielu pomiarów pozycji w czasie może pomóc w redukcji szumu pomiarowego i poprawić dokładność oszacowań pozycji. Techniki filtracji, takie jak filtr Kalmana, mogą również być zastosowane w celu poprawy dokładności oszacowań poprzez uwzględnienie czasowej korelacji pomiarów.

Podsumowanie

Podstawową miarą dokładności systemów nawigacji satelitarnej są parametry DOP. Są one obliczane i podawane przez moduły odbiorników GNSS, należy jednak uwzględniać je przy implementacji takich parametrów. W powyższym artykule opisano, czym są w metrologii dokładność i precyzja, a także jak definiowany jest parametr DOP w nawigacji.

Konieczne jest utrzymanie DOP na odpowiednim poziomie. Im niższy DOP, tym precyzyjniejsze są pomiary. DOP powyżej 20 oznacza, że uzyskane pomiary są zupełnie bezwartościowe. Dla uzyskania pomiarów zdalnych do wykorzystania w większości aplikacji DOP powinien być mniejszy niż 5, a w przypadku najbardziej wymagających zastosowań poniżej 2.

Jakkolwiek DOP wynika z geometrii znajdujących się w kosmosie konstelacji satelitów, istnieje wiele sposobów na poprawę DOP działaniami na ziemi. Krytyczne jest tutaj zastosowanie czułych anten (o wysokim zysku), które będą miały dobrą widoczność nieba – nie zasłonięte przeszkodami terenowymi, drzewami czy budynkami. Dodatkowo w uzyskaniu lepszego DOP pomocne jest zastosowanie różnych konstelacji satelitów – oprócz GPS dołączyć można satelity np. z układu Galileo czy BeiDou, w zależności od tego, które w danym momencie zapewniają lepszy sygnał.

Nikodem Czechowski, EP

Bibliografia:

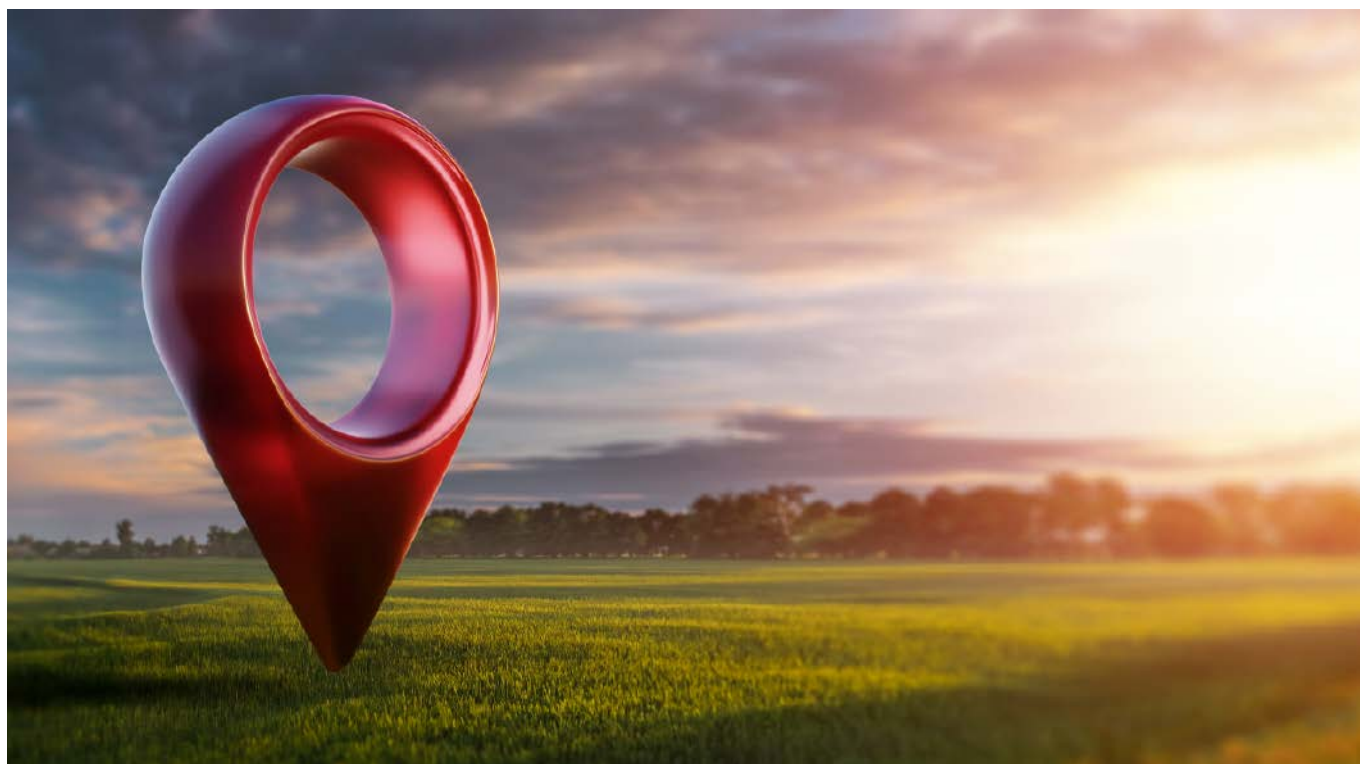
1. <https://en.wikipedia.org/>
2. <https://gisgeography.com/>
3. Richard B. Langley „Dilution of Precision”, GPS World, maj 1999
4. M. Tahsin, S. Sultana, T. Reza and M. Hossam-E-Haider „Analysis of DOP and its preciseness in GNSS position estimation” 2015 International Conference on Electrical Engineering and Information Communication Technology (ICEEICT), Savar, Bangladesh, 2015

REKLAMA

m.technik
Ciekawi świata są zawsze młodzi

w prezencie na każdą okazję
przejrysz i kupisz na
www.ulubionykiosk.pl





Pozycjonowanie GNSS centymetrowej precyzji, dostępne również do standardowych zastosowań

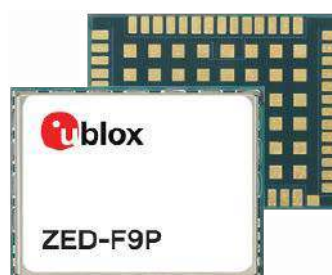
Zapotrzebowanie na systemy lokalizacji o wysokiej precyzji – HPG (High Precision GNSS) – występuje od dawna. Wysoka cena, zużycie energii i duże wymiary odbiorników oraz złożoność implementacji ograniczały stosowanie HPG do aplikacji niszowych.

Nawigacja precyzyjna to złożony system, co do tej pory generowało wysokie koszty implementacji w urządzeniu końcowym. Jednak w ostatnim czasie pojawiły się rozwiązania spełniające wysokie wymagania dostępne dla tak popularnych aplikacji jak np.:

- pojazdy autonomiczne, np. kosiarki operujące na precyzyjnie wyznaczonym terenie, trafiające do stacji bazowej,
- roboty autonomiczne,
- pojazdy rolnicze wymagające precyzji np. nawożenia czy siewu,
- samochody, ciężarówki, minibusy, w tym autonomiczne,
- maszyny budowlane,
- drony.

SPG vs. HPG

Nawigacja standardowej precyzji – SPG, oznacza dokładność rzędu 1,0...2,5 m CEP, przy sygnałach dobrej jakości i markowych odbiornikach. Słabe sygnały z satelitów (100 tysięcy razy słabsze od sygnału GSM) przebywają dystans tysięcy kilometrów i są zniekształcane przez zmienne, niemożliwe do przewidzenia czynniki (np. wilgotność i ciśnienie w troposferze, wyładowania w jonosferze, odbicia itp.).



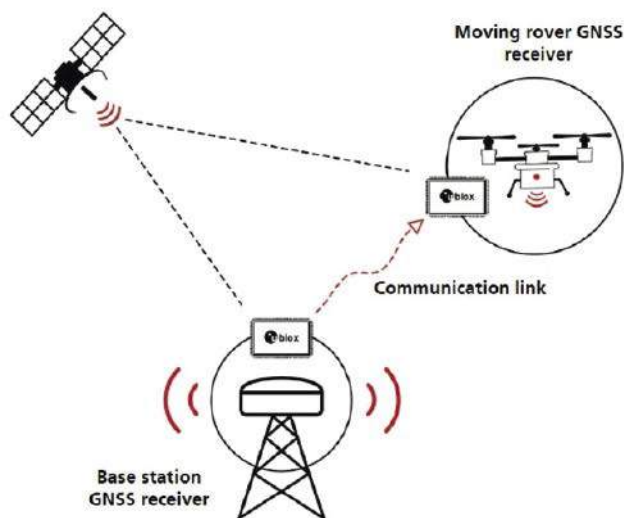
Nawigacja wysokiej precyzji – HPG, czyli precyzji centymetrowej, wymaga systemu korygującego, eliminującego wpływ opisanych czynników. Stacja referencyjna o znanych współrzędnych na bieżąco porównuje pozycję GNSS z rzeczywistą i przesyła korekty do odbiorników. W ten sposób eliminuje się błędy wynikające z przebiegu sygnału i uzyskuje dokładność centymetrową. Dane korygujące muszą być dostarczane w czasie rzeczywistym, mogą być generowane na kilka sposobów:

- Własna stacja referencyjna, przekazująca poprawki RTK do odbiornika w pobliżu (odległość <30 km);
- Regionalna sieć stacji referencyjnych, oferujących dane RTK najczęściej odpłatnie. Stacja powinna znajdować się blisko odbiornika, pracuje w standardzie OSR i wymaga komunikacji dwukierunkowej – odbiornik informuje o położeniu, żeby uzyskać korekty od najbliższej stacji, korekty są pomocne tylko lokalnie, w jej pobliżu;
- Stacja referencyjna PPP-RTK, o pokryciu kontynentalnym lub globalnym, komunikacja jest jednokierunkowa od stacji do odbiornika np. w formacie SPARTN, csRR. Pracuje w najnowszym standardzie SSR, dane są ważne na kontynencie, odbiornik nie musi znajdować się w pobliżu stacji.

Prosty w implementacji system HPG

Niedawno przystępny cenowo i prosty w implementacji system HPG stał się wreszcie dostępny, jego składniki można opisać, posługując się produktami pioniera technologii HPG – szwajcarskiej firmy u-blox. Są to:

1. Sygnały z kilku konstelacji GNSS: GPS, Galileo, Glonass, BeiDou, QZSS;
2. Odbiornik HPG gwarantujący najwyższej jakości parametry, ale także dopracowywane przez wiele generacji i sprawdzone algorytmy. Moduły NEO-F9P oraz ZED-F9P firmy u-blox obsługują równolegle cztery konstelacje GNSS oraz dwa pasma: L1/L2 lub L1/L5. Stosując technologię opartą na doświadczeniu firmy u-blox dostarczają informację o położeniu wysokiej precyzji. Zoptymalizowana cena, wymiary 12×16 mm i niski pobór energii umożliwiają zastosowanie tych odbiorników w aplikacjach dotychczas niedostępnych. Moduły wspierają wszystkie wymienione rodzaje korekt.
3. Poprawki ze stacji referencyjnych. PointPerfect to wygodny system niedawno wprowadzony przez u-blox. Pracuje w standardzie SSR, z formatem danych SPARTN. Ta nowoczesna implementacja pracuje z jednostronną komunikacją do odbiornika. PointPerfect jest natywnie wspierany przez moduły NEO-F9P oraz ZED-F9P. Poprawki mogą być pobierane z serwera poprzez dostęp IP i z satelitów przez dostęp L-Band (np. przy braku dostępu do internetu i sieci komórkowej). 40 godzin miesięcznie danych IP jest



darmowe, a stawka 3,90 USD przy 60 godzinach czy np. 6,80 USD przy 120 godzinach to bardzo przystępne ceny, dane są ważne na całym kontynencie.

4. Wysokiej jakości dwupasmowa antena HPG, na przykład ANN-MB u-blox, dopełnia całości.

Ważnym parametrem HPG jest czas dostrajania się odbiornika do korekt (np. po zaniku sygnału). W przypadku połączenia modułów firmy u-blox oraz PointPerfect są to zaledwie sekundy.

Dostępna jest także funkcja tzw. lokalnej bazy, która otwiera cię-kawe możliwości, np.:

- określanie precyzyjnego (0,4 stopnia) kierunku ruchu pojazdu,
- *moving base* – precyzyjne określenie wektora pomiędzy odbiornikiem a ruchomą stacją, przydatne np. przy śledzeniu pojazdu przez drona (*follow me*) albo lądowaniu na poruszającej się platformie.

Co ciekawe, dostępne są pinowo zgodne moduły ZED-F9R, integrujące HPG z IMU (żyroskop, akcelerometr), określające pozycję nawet bez widoczności satelitów.

Podsumowanie

Określanie położenia z centymetrową dokładnością stało się przystępne. Pomimo skomplikowanej technologii implementacja jest prosta, jednak wybór optymalnej architektury oraz funkcjonalności wymaga doświadczenia. Sugerujemy kontakt z zespołem wyszkolonych przez producenta specjalistów, używając adresu: ubloxFAE@microdis.net.

Robert Panufnik
Microdis Electronics
ubloxFAE@microdis.net

REKLAMA

Centymetrowa dokładność GNSS teraz dostępna dla aplikacji budżetowych, proste i szybkie wdrożenie



Sprawdzony na rynku system GNSS wysokiej precyzji:

- 1) NEO-F9P i ZED-F9P - najnowsze odbiorniki GNSS HPG
 - ⊗ cenioma technologia firmy u-blox: własne algorytmy i chipsety
 - ⊗ wbudowane wsparcie dla poprawek PointPerfect
 - ⊗ Protection Level: 95% pewności położenia
 - ⊗ ZED-F9R: wersja z IMU (żyroskop, akcelerometr)



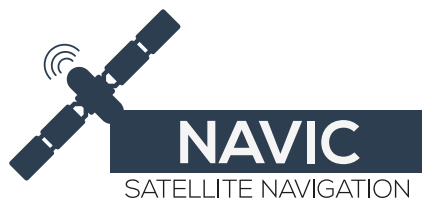
- 2) PointPerfect - system poprawek
 - ⊗ wiarygodny, zoptymalizowany kosztowo i wygodny w użyciu
 - ⊗ pobierane przez internet lub z satelity

Typowe zastosowania:



WWW.MICRODIS.NET

MARKETING@MICRODIS.NET



Nie tylko GPS

Przegląd alternatywnych systemów nawigacji satelitarnej

Jest wiele modułów przeznaczonych do nawigacji satelitarnej. Większość z nich oferuje współpracę z różnymi systemami poza GPS – GLONASS, Galileo itd. Czym te systemy nawigacji satelitarnej (GNSS) różnią się od GPS i jakie mogą być zalety ich stosowania? W artykule przyjrzymy się innym GNSS, w szczególności GLONASS i Galileo.

Przed powstaniem GPS działało kilka innych systemów nawigacji. Wcześniej były to rozwiązania bazujące na naziemnych stacjach radiowych, jednak miały one wiele ograniczeń. Dlatego na potrzeby militarne zdecydowano o opracowaniu systemu satelitarnego. Początkowo system GPS miał jedynie militarne zastosowanie, ale w końcu stał się pierwszym, jaki trafił do użytku cywilnego.

Nad naszymi głowami znajduje się szereg konstelacji satelitów, tworzących różne systemy nawigacji satelitarnej (GNSS). Różnią się one parametrami technicznymi i pewnymi detalami dotyczącymi rozwiązań technologicznych, ale co do zasady oferują porównywalne możliwości. Obecnie najbardziej użyteczne są następujące GNSS:

- GPS, a właściwie Navstar-GPS – to pierwszy system nawigacji satelitarnej. Został utworzony i jest utrzymywany przez Amerykanów od lat 70. XX wieku. Obecnie w jego skład wchodzi 38 satelitów, z czego 32 działające.
- GLONASS to rosyjski odpowiednik GPS, budowany od lat 80. XX wieku, z przerwami w latach 90. Finalnie doprowadzony do operacyjności w 2011 roku. Konstelacja ta składa się z 26 satelitów, z czego 24 działające i 2 w fazie uruchamiania.
- Galileo to europejska sieć nawigacyjna, której koncepcja powstała na przełomie XX i XXI wieku. Konstelacja rozpoczęła funkcjonowanie w 2016 roku. Składa się na nią obecnie 28 satelitów – 23 aktywne i 5, które z różnych przyczyn aktualnie nie funkcjonują (wliczają się w to awarie, wycofanie z użytkowania oraz okresowe planowane wyłączenia). Jest to najdokładniejszy system z dostępnych cywilnie.
- BeiDou to najnowszy system, tworzony przez Chińczyków od początku XXI wieku – początkowo jako system lokalny (BeiDou-1),



Fotografia 1. Moduł u-blox NEO-6M

który rozbudowany został do globalnego (BeiDou-3) w 2020 roku, kiedy to ostatni satelita z planowanych trafił na orbitę. Obecnie konstelacja składa się z 35 funkcjonalnych satelitów.

Oprócz tych systemów istnieją także konstelacje działające lokalnie, takie jak QZSS, działający na obszarze Japonii, EGNOS, który obejmuje zasięgiem Europę i północną Afrykę, indyjski GAGAN czy WAAS, który działa na terenie Ameryki Północnej. W artykule nie będziemy się skupiać na tych lokalnych sieciach, szczególnie, że tylko jedna z nich – EGNOS – obejmuje swoim zasięgiem Polskę.

Jakkolwiek GPS był jednym z pierwszych na świecie systemów nawigacji satelitarnej, to z różnych – głównie politycznych – przyczyn szybko zaczęły pojawiać się inne systemy. Jak wspomniano, początkowo GPS był systemem przeznaczonym do zastosowań militarnych – wiele technologii na świecie rodzi się w ten sposób. Podobnie było z większością wymienionych tutaj systemów. Nawet jeśli motywacja do tworzenia tych konstelacji nie była typowo militarna, to wynikała z chęci uniezależnienia się od GPS, który mógłby być używany jako metoda politycznego nacisku.

Obecnie na rynku większość modułów odbiorników GNSS potrafi komunikować się z większą liczbą sieci, nawet często łącząc dane



Fotografia 2. Moduł u-blox Neo-M8

z różnych konstelacji, aby uzyskać lepsze rezultaty nawigacji. W dalszej części artykułu przyjrzymy się dwóm głównym alternatywom dla GPS i dostępnym na rynku modułom, które umożliwiają ich używanie.

Przegląd modułów

Jednym z chyba najpopularniejszych wśród hobbystów modułem odbiornika GNSS jest NEO-6M firmy u-blox (**fotografia 1**). Z uwagi na tę popularność należy o nim tutaj wspomnieć i jednocześnie zaznaczyć – moduł ten jest ograniczony do systemu GPS.

Jeśli chcemy skorzystać z modułów u-blox i odbierać dane z większej liczby systemów GNSS, sięgnąć musimy po wyższe modele,

korzystające z chipsetu M8, np. NEO-M8P (**fotografia 2**) lub F9, np. NEO-F9, które są w stanie śledzić jednocześnie, odpowiednio, dwa i cztery zestawy różnych satelitów.

Układ u-blox M8 równocześnie obsługuje dwa systemy nawigacji satelitarnej – domyślnie jest to GPS i GLONASS. Równoczesne odbieranie GPS i BeiDou, a nawet równoczesne odbieranie GLONASS i BeiDou można wybrać w ustawieniach. Zoptymalizowane układy odbioru sygnału w połączeniu z algorytmami programowymi oraz zaawansowanymi silnikami śledzenia i wyszukiwania uwzględniają jakość sygnału, a nie tylko liczbę satelitów używanych do dostarczania informacji o pozycji.

W ofercie firmy Telit większość modułów jest zdolna do odbioru sygnałów z wielu konstelacji. Przykładem tego mogą być moduły z serii SE868xx, moduł SE871L lub SE873K5, pokazane, odpowiednio na **fotografiach 3a, 3b** oraz **3c**. W przypadku modułów Telita domyślnie włączony jest odbiór GPS i GLONASS. W przypadku modułów 871x, jeśli włączony zostanie odbiór GPS i BeiDou, nie będzie możliwy odbiór systemów GLONASS i Galileo. Pozostałe moduły są w stanie odbierać sygnał nawet ze wszystkich czterech konstelacji, aby zwiększyć bezpieczeństwo i dokładność pozycjonowania.

Firma Simcom w zeszłym roku przygotowała nową serię modułów, które zdolne są do pracy z wieloma konstelacjami – **rysunek 1**. Dostępna jest szeroka gama modułów o różnych parametrach i aplikacjach. Moduły takie jak SIM65M obsługują usługę pozycjonowania na wielu konstelacjach, w tym GPS, GLONASS, BeiDou, Galileo, ale również QZSS oraz SBAS (*Satellite-based Augmentation System*). Ze względu na większą liczbę dostępnych opcjonalnych sygnałów satelitarnych znacznie zmniejsza się ryzyko blokowania lub niedokładności sygnałów, co pozwala generować bardziej precyzyjne dane dotyczące pozycji.

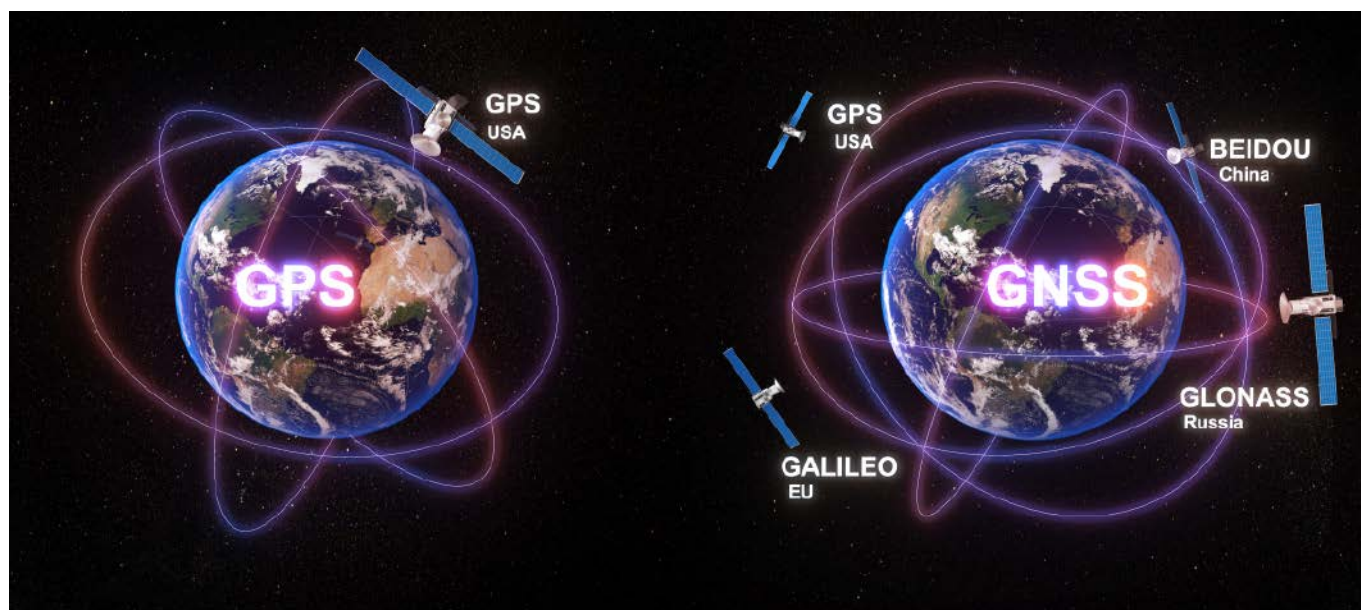


Fotografia 3. Moduły firmy Telit: a) SE868K5-DI; b) SE871L; c) SE873K5

SIMCom GNSS Modules

High Precision Dual Frequency		Standard Precision Dual Frequency		Standard Precision Single Frequency		Standard Precision Embedded Antenna	
SIM66MD	SIM68AT	SIM68MD	SIM68D	SIM65MD	SIM68M	SIM68V	SIM33ELA
- GPS/BDS/GLONASS/Galileo/QZSS/SBAS - L1+L5 - RTK	- GPS/BDS/GLONASS/Galileo/QZSS/SBAS - L1+L5 - RTK - DR	- GPS/BDS/GLONASS/Galileo/QZSS/SBAS - L1+L5	- GPS/BDS/GLONASS/Galileo/QZSS/SBAS - L1+L5	- GPS/BDS/GLONASS/Galileo/QZSS/SBAS - L1+L5	- GPS/GLONASS/Galileo/QZSS/SBAS - L1	- GPS/GLONASS/Galileo/QZSS/SBAS - L1	- GPS/GLONASS/Galileo/QZSS/SBAS - L1

Rysunek 1. Nowe moduły Simcom



Galileo

Galileo, jak wspomniano wcześniej, jest europejskim systemem nawigacji satelitarnej, który jest rozwijany przez Unię Europejską poprzez Europejską Agencję Kosmiczną (ESA). System ten został zaprojektowany w celu dostarczania precyzyjnych i niezawodnych informacji o pozycji, prędkości i czasie na całym świecie. Galileo powstał, aby pozwolić uniezależnić się Europie od innych systemów nawigacji satelitarnej, wspomnianych w tym artykule.

Jedną z kluczowych cech odróżniających Galileo od systemu GPS jest jego wysoka precyzja. Galileo oferuje sygnały o wyjątkowo precyzyjnym czasie i dokładności pozycji, co sprawia, że jest szczególnie przydatny w dziedzinach wymagających bardzo dokładnego pozycjonowania, takich jak nawigacja lotnicza, morska, geodezja czy badania naukowe. Jest to osiągane, między innymi, dzięki zastosowaniu aż pięciu pasm radiowych, w których nadawane są różne sygnały nawigacyjne.

Galileo ma również unikalną funkcję wspomagania działań ratowniczych – *Search and Rescue* (SAR), która umożliwia określenie lokalizacji i usprawnia ratunek w przypadku np. wypadków na morzu, lądzie lub w powietrzu. Ta funkcja jest ważnym elementem, który może przyczynić się do zwiększenia bezpieczeństwa użytkowników Galileo w różnych sytuacjach.

Inną cechą wyróżniającą Galileo jest jego pełne zintegrowanie z systemami EGNOS (*European Geostationary Navigation Overlay Service*). EGNOS to europejski system korekcji sygnałów GNSS, który poprawia precyzję pozycjonowania, korzystając z systemu stacji naziemnych. EGNOS dostarcza trzy główne usługi:

1. Usługa poprawki różnicowej – stacje naziemne monitorują sygnały z satelitów nawigacyjnych i identyfikują błędy w sygnale. Następnie te poprawki są przesyłane do użytkowników EGNOS, co pozwala na poprawę dokładności pozycjonowania.
2. Usługa poprawki oszacowanej – ta usługa dostarcza korekty sygnałów nawigacyjnych, które mogą być użyteczne w sytuacjach, gdzie nie ma dostępu do poprawek różnicowych.
3. Usługa monitorowania jakości sygnałów – ta usługa monitoruje sygnały z satelitów nawigacyjnych i udostępnia informacje o ich integralności. W razie wykrycia problemów lub zakłóceń w sygnale użytkownicy mogą otrzymać ostrzeżenia, co jest szczególnie ważne w zastosowaniach, gdzie kluczowe jest bezpieczeństwo.

W przeciwieństwie do innych systemów GNSS, Galileo działa na bazie cywilnej infrastruktury, co oznacza, że nie będzie podlegał militarnym ograniczeniom dostępu ani zakłóceniom w określonych regionach. To sprawia, że Galileo jest systemem otwartym i dostępnym dla różnych sektorów gospodarki oraz badań naukowych.

W skrócie, Galileo wyróżnia się swoją wysoką precyzją, unikalnymi funkcjami, integracją z innymi systemami nawigacji oraz otwartością na różnorodne zastosowania. Według danych ESA, w typowych warunkach Galileo oferuje dokładność poziomą poniżej 20 cm i pionową poniżej 40 cm w trybie wysokiej dokładności (HAS) i odpowiednio, 4 m i 8 m, w trybie otwartego dostępu (OS). Wszystkie te cechy sprawiają, że jest on ważnym elementem europejskiej infrastruktury kosmicznej i przyczynia się do rozwijania nowych technologii oraz poprawy bezpieczeństwa i dokładności nawigacji na całym świecie.

GLONASS

Jest to rosyjski globalny system nawigacji satelitarnej, który został zaprojektowany w ZSRR i jest obecnie rozwijany przez Rosję. Podobnie jak system Galileo i GPS, GLONASS umożliwia użytkownikom na całym świecie precyzyjne pozycjonowanie i nawigację przy użyciu sygnałów emitowanych przez satelity.

Satelity GLONASS krążą w trzech różnych płaszczyznach (w odróżnieniu od dwóch dla GPS), co może wpływać na ich widoczność i dokładność w różnych miejscach na Ziemi. Jest to szczególnie odczuwalne w rejonach podbiegunowych, gdzie GLONASS oferuje o wiele wyższe pokrycie.

System GPS używa dwóch pasm do komunikacji (L1 i L2), a GLONASS używa trzech częstotliwości (L1, L2 i L3). Wpływa to na dokładność pozycjonowania w różnych warunkach. GLONASS jest systemem wojskowym, więc ma te same zagrożenia, co GPS – może być okresowo zakłócany, szczególnie w pobliżu miejsc konfliktów zbrojnych. Może być też stosowany jako metoda wywierania politycznego nacisku przez Rosję (co jednak jest w dużej mierze już mało problematyczne z uwagi na opracowanie europejskiego, cywilnego systemu Galileo).

GLONASS oferuje precyzję pozycjonowania na poziomie około 5...6 metrów (wysokość) oraz 2...3 metrów (szerokość i długość geograficzna), w zależności od lokalizacji i liczby obserwowanych w danym momencie satelitów.

Zalety korzystania z wielu systemów

Kiedy używanych jest jednocześnie kilka konstelacji GNSS, poprawia się wiele parametrów systemu. Główne korzyści z użycia systemu multikonstelacyjnego to:

- poprawia się dostępność sygnału – dzięki większej liczbie dostępnych satelitów, łatwiej jest uzyskać dobrej jakości sygnał od odpowiedniej ich liczby;
- lepszy dostęp do usługi – ograniczenia widzialności niektórych satelitów lub ich tymczasowe problemy z działaniem nie uniemożliwiają zastosowania systemu, jeśli ten korzysta z wielu konstelacji;

- zwiększona precyzja – zastosowanie wielu konstelacji pozwala na uzyskanie wyższej precyzji, szczególnie jeżeli do tego dodatkowo dołącza się systemy korekcyjne, takie jak np. EGNOS;
- dodatkowe bezpieczeństwo – zastosowanie wielu konstelacji, pracujących na różnych pasmach częstotliwości, zwiększa odporność systemu na awarie satelitów, zakłócenia czy intencjonalnie transmitowane fałszywe dane, mające zakłócać działanie GNSS;
- zwiększona integralność danych – zastosowanie wielu konstelacji zwiększa odporność systemu na zakłócenia atmosferyczne czy wynikające z otoczenia np. w mieście.

Dodatkowo, wspomnieć należy regionalne systemy wspomagania nawigacji bazujące na satelitach (SBAS), które wspomagają systemy globalne. Często wspomniane odbiorniki mają możliwość korzystania z części lub wszystkich z nich:

- System Rozszerzenia Obszaru (WAAS) w Ameryce Północnej i Południowej,
- Europejski Geostacjonarny System Nawigacyjny (EGNOS) w Europie,
- Nawigacja z Użyciem Ulepszeń przez Satelity GPS (GAGAN) w Indiach,
- MTSAT Satellite-Based Augmentation System (MSAS) w Japonii.

Podsumowanie

Korzystanie z wielokonstelacyjnych systemów nawigacji satelitarnej, odbierających sygnały systemów, takich jak GPS, Galileo i GLONASS, niesie ze sobą wiele korzyści oraz pewne wyzwania. Pozwala na zwiększenie dokładności i ciągłości aktualizacji pozycji, co jest szczególnie istotne w dzisiejszych zastosowaniach nawigacyjnych. Dzięki temu użytkownicy mogą cieszyć się lepszą dostępnością sygnału, większym zakresem działania i dokładniejszym pozycjonowaniem.

Jednak należy mieć na uwadze pewne kompromisy, zwłaszcza jeżeli chodzi o zużycie energii. Wybór odpowiednich konstelacji do śledzenia w zależności od lokalizacji może pomóc w zminimalizowaniu zużycia energii, co jest istotne szczególnie w przypadku urządzeń o niewielkiej mocy, takich jak małe urządzenia IoT. Używanie wielu konstelacji naraz przekłada się na wyższe zużycie mocy przez moduł odbiornika GNSS.

Ogólnie rzecz biorąc, wielokonstelacyjne systemy nawigacji satelitarnej stają się coraz bardziej wszechstronne i efektywne, a dodatkowo na rynku pojawia się coraz więcej modułów, które są w stanie korzystać z dwóch lub więcej konstelacji naraz. Dzięki temu użytkownicy mogą osiągnąć wyższą dokładność pozycjonowania i wyższą niezawodność w różnych warunkach środowiskowych. Jednocześnie wciąż istnieją wyzwania związane ze zużyciem energii, które wymagają zrównoważonego podejścia w wyborze konstelacji i zoptymalizowanym zarządzaniu mocą w urządzeniach nawigacyjnych.

Nikodem Czechowski, EP

Bibliografia

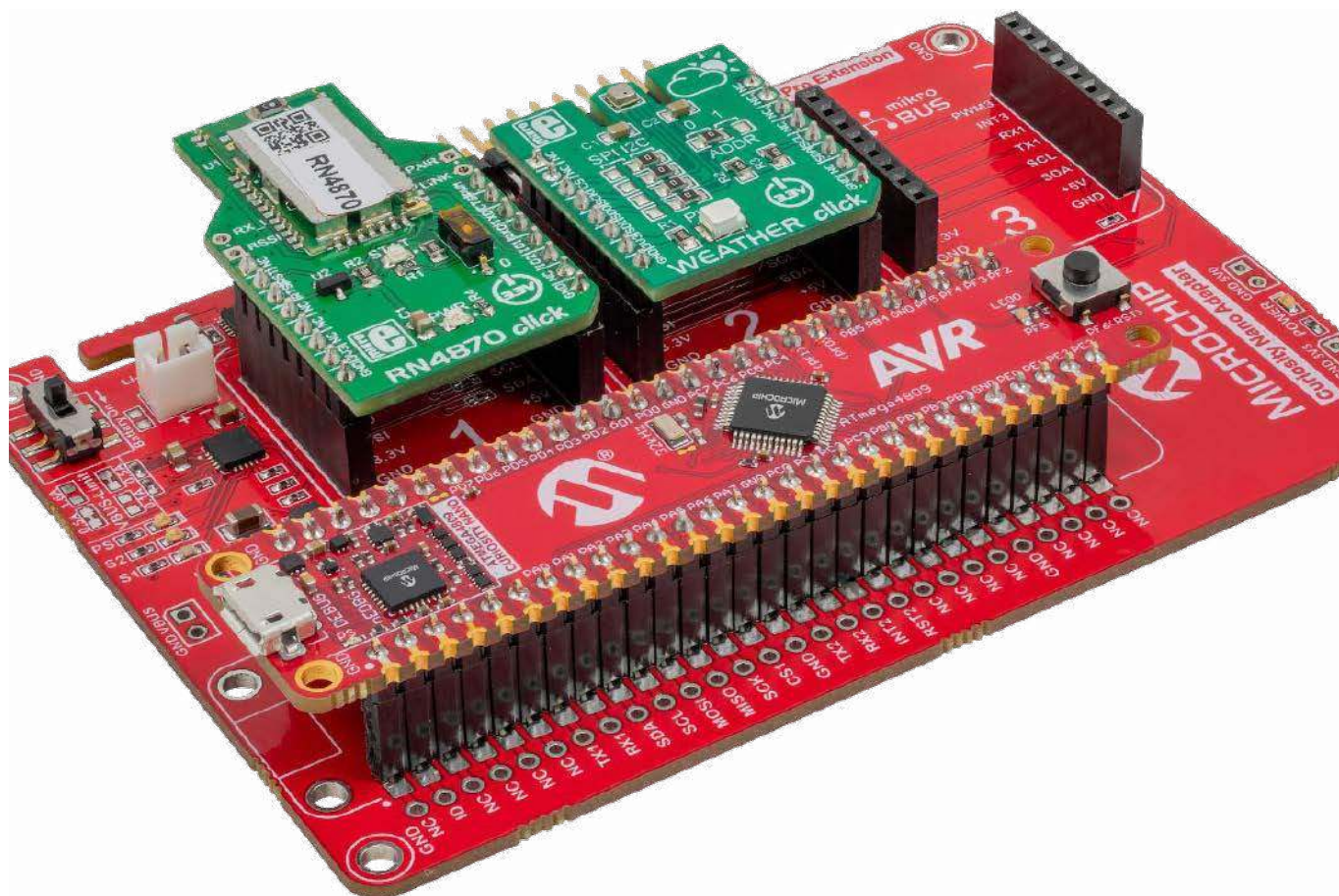
1. <https://www.gps.gov/>,
2. <https://www.glonass-iac.ru/en/>,
3. <http://en.beidou.gov.cn/>,
4. <https://www.gsc-europa.eu/>,
5. <https://egnos-user-support.essp-sas.eu/>,
6. gianmarco Zanda „What Is the Galileo Global Navigation Satellite System?”, Telit Cinterion Blog 14,07,2021,
7. gianmarco Zanda „The Business Advantages of a Multi-GNSS Setup”, Telit Cinterion Blog 06,04,2022,
8. eSA Navipedia – <https://gssc.esa.int/navipedia/index.php>.

REKLAMA

www.ep.com.pl/EPwtoku

Czytaj artykuły
zanim zostaną
wydane
w formie
papierowej





Narzędzia do szybkiego prototypowania z 32-bitowymi mikrokontrolerami Microchipa

Microchip jest znanym producentem mikrokontrolerów, ale oferuje również zaawansowane narzędzia programowe oraz bogaty asortyment gotowych rozwiązań sprzętowych. Każda rodzina mikrokontrolerów ma zapewnione projekty przykładowych programów oraz zestawy do prototypowania przeznaczone do takich segmentów przemysłu i technologii, jak łączność i komunikacja, sterowanie obwodami małej mocy, bezpieczeństwo funkcjonalne i nie tylko. Gotowe komponenty sprzętowe i programowe pozwalają szybko tworzyć prototypy aplikacji bez faktycznego projektowania płytek PCB.

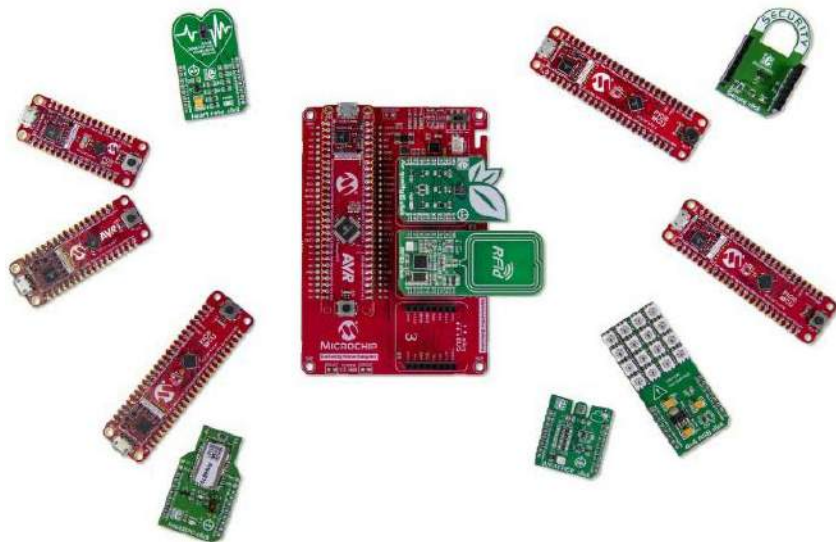
W ramach swojej oferty narzędzi sprzętowych firma Microchip oferuje płytki rozwojowe Curiosity Nano i płytki rozwojowe Curiosity. Zestawy rozwojowe Curiosity Nano to niedroge zestawy rozwojowe o niewielkich rozmiarach z wbudowanymi funkcjami debugowania i programowania. Są dostosowane do adaptera o nazwie Nano Base Board (**fotografia 1**) i korzystają z dodatkowych płytek Click Board oraz innych profesjonalnych płytek rozszerzeń, które zwiększają możliwości



Fotografia 1. Płytki Nano Base Board z zamontowanym modułem serii Nano

płytek deweloperskich poprzez dodanie czujników, interfejsów dotykowych, wyświetlaczy lub innych interfejsów łączności wymaganych do opracowania dowolnej aplikacji (**fotografia 2**).

Zestawy Nano mają spójną konstrukcję w całej rodzinie mikrokontrolerów 8-bitowych, 16-bitowych i 32-bitowych firmy Microchip, dzięki czemu aplikacje wszystkich typów są skalowalne w obrębie tych rodzin mikrokontrolerów. Płytki rozwojowe



Fotografia 2. Różne płytki rozszerzeń

Curiosity są wyposażone w dodatkowe funkcje, takie jak złącza audio, interfejsy graficzne, złącza interfejsów komunikacyjnych i do łączności z Internetem. Oprócz złączy MikroE dostępne są złącza XPRO, natomiast niektóre płytki Curiosity mają również złącza Arduino, umożliwiające zastosowanie nakładek Arduino do realizacji prototypów aplikacji.

Narzędzia programowe uzupełniające sprzęt

Jeśli chodzi o narzędzia programowe, Microchip oferuje cztery główne rozwiązania. Przede wszystkim **MPLAB X IDE** – zintegrowane środowisko programistyczne używane do tworzenia aplikacji osadzonych na cyfrowych kontrolerach sygnału i mikrokontrolerach firmy Microchip.

Kolejnym narzędziem są kompilatory **MPLAB XC**. Kompilator XC umożliwia tłumaczenie kodu języka wysokiego poziomu C lub C++ na język assemblera mikrokontrolera.

Innym ważnym narzędziem jest platforma programistyczna **MPLAB Harmony**, część większego ekosystemu MPLAB X firmy Microchip, która umożliwia tworzenie aplikacji wbudowanych na 32-bitowych mikrokontrolerach firmy Microchip. Harmony ma kilka cech i funkcjonalności. Zawiera kilka przykładów demonstracyjnych ułatwiających rozpoczęcie tworzenia aplikacji i prototypów, a także zapewnia biblioteki warstwowe i modułowe. Ma również biblioteki peryferyjne, które zapewniają dostęp do rejestrów sprzętowych i oferują biblioteki oprogramowania pośredniego oraz abstrakcyjne sterowniki i usługi systemowe.

Istnieje również kilka narzędzi do tworzenia grafiki, umożliwiających tworzenie projektów i generowanie kodu za pomocą graficznego interfejsu użytkownika. Wszystko to można pobrać za pośrednictwem platformy GitHub, a całe środowisko programistyczne MPLAB Harmony v3 jest dostępne w GitHub.

Jednym z programów do tworzenia grafiki jest **MCC** – narzędzie konfiguracji kodu firmy Microchip, które obsługuje teraz 32-bitowe mikrokontrolery, zapewniając w ten sposób klientom wspólne narzędzie konfiguracyjne dla wszystkich rodzin mikrokontrolerów. Za pomocą niego możesz skonfigurować swój

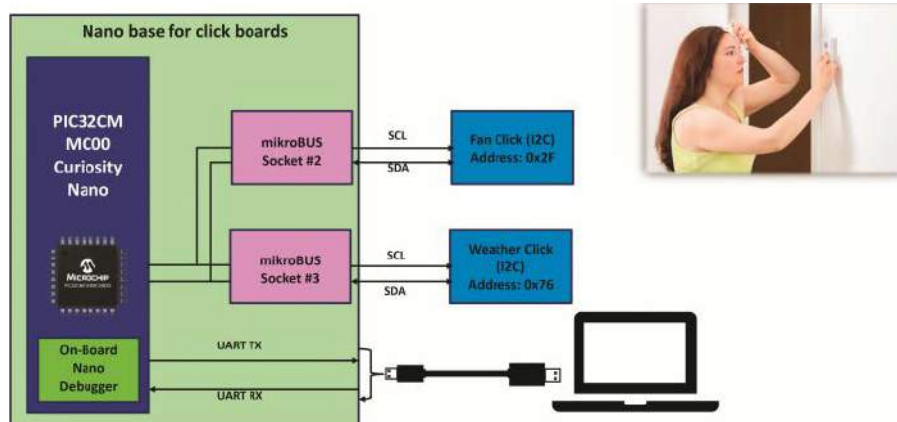
kod, tworzyć projekty i generować gotowe pliki przy użyciu bibliotek oprogramowania.

Przyjrzymy się trzem przykładowym projektom, ale najpierw musimy poznać kilka zasobów, które są istotne dla tych przykładów. Pierwszym z nich jest pakiet aplikacji MPLAB Harmony Reference Apps, dostępny jako repozytorium w serwisie GitHub. To repozytorium zawiera wiele samodzielnych aplikacji demonstrujących funkcje i możliwości 32-bitowych mikrokontrolerów. Przykłady w naszych aplikacjach referencyjnych obejmują aplikacje dla początkujących i aplikacje demonstracyjne, a także znacznie bardziej złożone i bogate w funkcje projekty oraz przykłady korzystające z płytek MikroElektronika Click Board, płytek XPRO itp. W ramach samego pakietu aplikacji referencyjnych dostępne są także przykłady programów dla MikroElektronika Click Board, które pokazują, jak z nich korzystać w środowisku programistycznym Harmony.

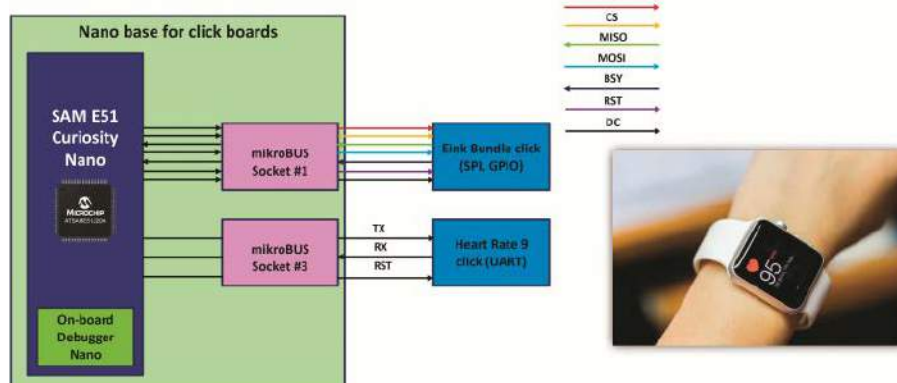
Przykłady szybkiego prototypowania

Naszym pierwszym przykładem jest sterowanie pracą wentylatora na podstawie wartości temperatury w pomieszczeniu, jak pokazano na **rysunku 1**. Wentylator pracuje z niską, średnią lub dużą prędkością, można go włączać i wyłączać w zależności od temperatury w pomieszczeniu. W tym przykładzie używamy zestawu ewaluacyjnego PIC32CM MC00 Curiosity Nano, który zawiera mikrokontroler Cortex M0+ i łączy się z płytką MikroElektronika Click Board za pomocą linii I²C.

Mamy tutaj dwie płytki Click Board – płytkę Click odczytującą parametry pogodowe i płytkę Click do sterowania wentylatorami. Kontrolowane parametry to temperatura, ciśnienie i wilgotność,



Rysunek 1. Schemat blokowy aplikacji 1 – sterowanie wentylatorem z kontrolą temperatury



Rysunek 2. Schemat blokowy aplikacji 2 – pomiar pulsu i wyświetlanie wyniku na wyświetlaczu eINK

a elementem wykonawczym będzie podłączony wentylator. Obie płytki połączone są poprzez interfejs I²C z mikrokontrolerem pod dwoma różnymi adresami. Aplikacja inicjuje te płytki, a następnie odczytuje pomiar temperatury w pomieszczeniu. Wartość jest porównywana z wartością zadaną i odpowiednio sterowany jest wentylator – włącza się, wyłącza i pracuje z różnymi prędkościami.

Przykład drugi to aplikacja do monitorowania tętna człowieka, co jest pomocne w analizie wzorców snu lub wyników podczas różnych zajęć sportowych. W tym przykładzie zastosowano płytke SAM E51 Curiosity Nano, która jest wyposażona w mikrokontroler Cortex M4 i jest połączona z modułami Heart Rate Click i eINK Bundle Click. Płytkę Heart Rate 9 podaje dane dotyczące tętna, które są przesyłane przez linie UART. Do prezentacji wyników mamy wyświetlacz o niskim poborze mocy, znany jako eINK Click Bundle, połączony przez linie SPI (rysunek 2).

Po zainicjowaniu tych modułów i wyświetleniu wartości domyślnej na ekranie eINK Click Bundle użytkownik uruchamia pomiar, naciskając przycisk, a następnie kładąc palec na czujniku tętna. Mikrokontroler odczytuje to, a następnie wyświetla wartość na wyświetlaczu.

Przykład trzeci jest rozszerzeniem przykładu 1, w którym dodajemy możliwość wyświetlania informacji o sterowaniu wentylatorem na podstawie temperatury pokojowej.

Tutaj mamy zakropkowany obszar, który dodano do schematu blokowego z przykładu pierwszego – jest to moduł eINK Click (rysunek 3). Struktura aplikacji pozostaje taka sama. Gdy mamy już temperaturę, wentylator nadal jest ustawiony na niską, średnią lub wysoką prędkość, ale jednocześnie wartość temperatury i prędkość wentylatora są wyświetlane na wyświetlaczu.

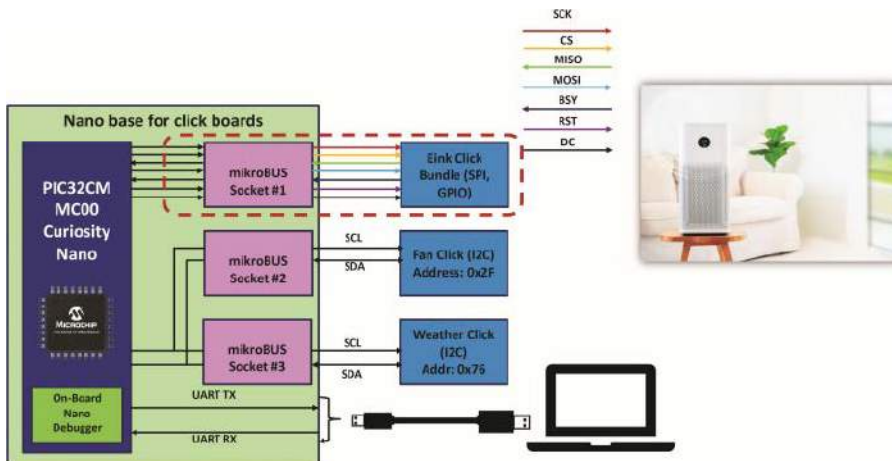
Część 2 tego przykładu jest rozbudowana o funkcje bezprzewodowe, dając możliwość sterowania za pomocą aplikacji na smartfona z systemem Android i BLE. Wysyłając polecenia przez telefon, masz możliwość wyłączenia wentylatora lub uruchomienia w trybie kontroli temperatury. Przykład 3, części 1 i 2 można również zintegrować z aplikacją sterowaną inteligentnym urządzeniem, z pełnymi możliwościami zdalnego sterowania i wyświetlania.

Rozwijanie aplikacji

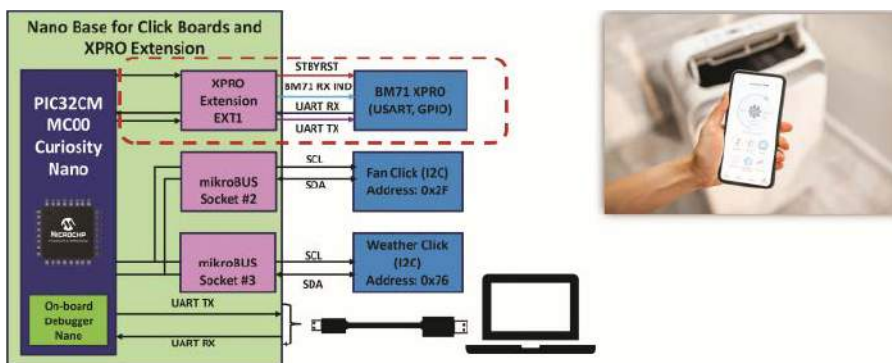
Warto poznać pewne działania, które należy wykonać, aby opracować te aplikacje. Pierwszym krokiem jest utworzenie projektu MPLAB X IDE dla wybranego mikrokontrolera, na przykład SAM E51 przy użyciu MCC. Następnie skonfigurujemy zegar, urządzenia peryferyjne i odpowiednie piny. Jeśli używamy istniejącego przykładu i chcemy go rozszerzyć, to nie tworzymy nowego projektu, ale raczej używamy istniejącego i po prostu otwieramy go w MPLAB X IDE.

Krok 2 to dodanie kodu aplikacji. Wcześniej korzystaliśmy z istniejących przykładów MikroElektronika Click Board. W tych przykładach znajdują się pewne procedury, których można użyć do zaimplementowania funkcji kliknięcia w aplikacjach końcowych – zasadniczo można je dodać do swojej aplikacji, a następnie zaimplementować aplikację zgodnie z wymaganiami.

Trzeci krok polega na uruchomieniu aplikacji i ocenie wyników. Aby to zrobić, uruchom aplikację na Androida, znaną jako MBD, a jeśli poprosi Cię o włączenie lokalizacji Bluetooth, zrób to. Wybierz ikonę BLE UART, która pokazuje typ urządzeń, które można skanować.



Rysunek 3. Schemat blokowy aplikacji 3 – sterowanie wentylatorem z kontrolą temperatury i wyświetlaniem parametrów pracy



Rysunek 4. Schemat blokowy aplikacji 4 – sterowanie wentylatorem z kontrolą zdalną

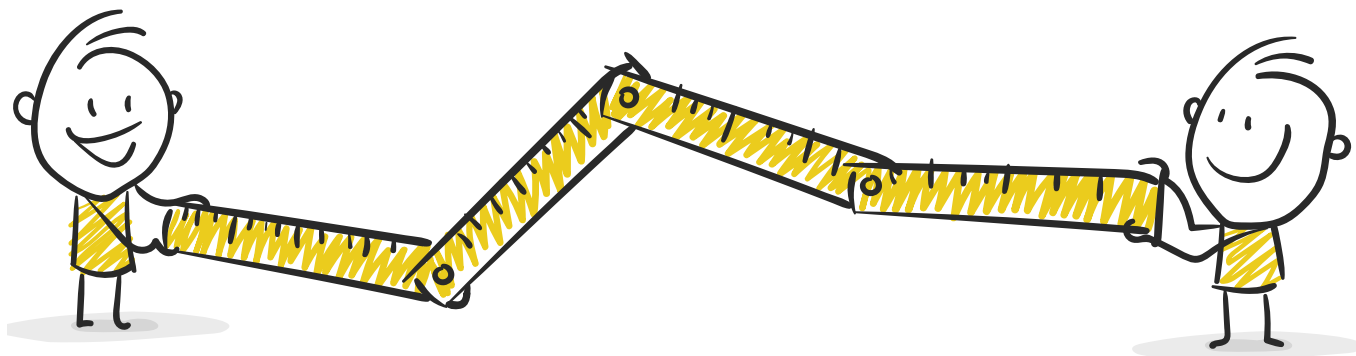
Wybierz skanowanie BM70, a wyświetli się lista wszystkich dostępnych urządzeń z BLE. Zwróć uwagę na urządzenie demonstracyjne UART, a kiedy je zobaczysz, możesz anulować trwające skanowanie i wybrać to urządzenie. Wówczas zostanie nawiązane połączenie z urządzeniem i udostępniona zostanie ścieżka transmisji danych. Wystarczy dotknąć ikony przesyłania danych i włączyć odpowiedni transfer.

Dotarliśmy do ekranu, na którym można wpisać polecenia sterujące urządzeniem – wyłącz wentylator, włącz go lub uruchom z określoną prędkością. Polecenia te są realizowane poprzez wiadomości tekstowe. Jeśli chcesz wyłączyć wentylator, wpisz polecenie *BLE control fan off* i naciśnij przycisk wysyłania; aby uruchomić go w trybie kontroli temperatury, wpisz polecenie *temp control* i naciśnij przycisk, a wentylator uruchomi się w zależności od temperatury pokojowej. Można także zmieniać temperaturę, przykładając palec do czujnika, a wentylator odpowiednio zmieni prędkość, pracując z niską, średnią lub dużą prędkością, w zależności od temperatury w pomieszczeniu.

Platforma programistyczna MPLAB Harmony

Microchip oferuje kilka zasobów pomocnych w tworzeniu aplikacji, w tym stopy graficzne, audio, biblioteki kryptograficzne i kilka innych. Oprócz płytek Curiosity i Curiosity Nano do szybkiego prototypowania, Microchip oferuje również kilka znacznie bardziej wszechstronnych platform, takich jak seria płytek rozwojowych Xplained Pro lub XPLORE, a także seria płytek rozwojowych Curiosity Ultra. Więcej informacji na ten temat: <http://www.microchip.com/harmony>.

Syed Thaseemuddin,
inżynier personelu technicznego,
Microchip Technology Inc.
Shridhar Channagiri,
menedżer ds. marketingu produktów,
Microchip Technology Inc.



Pomiary odległości i prędkości

Pomiary odległości zostały opanowane przez ludzkość jako jedne z pierwszych, a przez kolejne stulecia zmieniały się nie tylko metody pomiarowe, ale także i stosowane w praktyce wzorce, a co za tym idzie – jednostki. Jednak dopiero rozwój elektroniki doprowadził do prawdziwej rewolucji w dziedzinie pomiarów wielkości fizycznych, a postęp nie ominął także głównych bohaterów tego wydania „Elektroniki Praktycznej”. W artykule zajmiemy się metodami oraz czujnikami, stosowanymi do pomiaru odległości jako takiej, ale także jej pochodnej po czasie, czyli prędkości. Wykażemy też, że – pomimo iż obydwie te wielkości są ze sobą ściśle powiązane (właśnie za sprawą rachunku różniczkowego) – rzadko zdarza się, by sposoby ich pomiaru bazowały na takiej samej technologii.

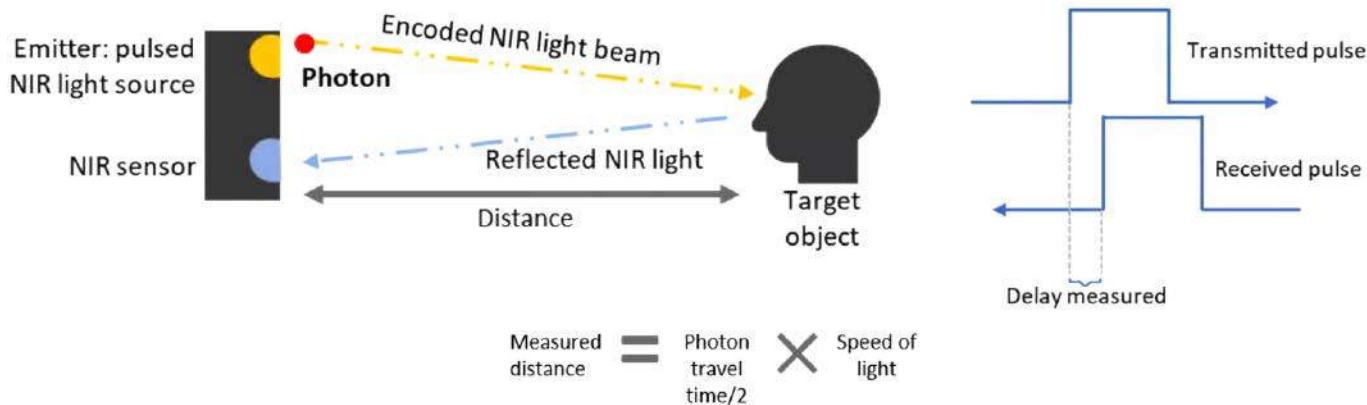
Pomiary odległości oraz prędkości (zarówno liniowej, jak i obrotowej) są szeroko stosowane niemal we wszystkich gałęziach techniki, choć na prowadzenie wysuwa się zdecydowanie kilka z nich – automatyka przemysłowa, robotyka mobilna, motoryzacja, branża wojskowa i morska czy wreszcie nowoczesne, autonomiczne i półautonomiczne drony. Przykłady – czasem wysoce zaawansowane technologicznie

i imponujące, np. pod względem zakresu czy dokładności pomiaru, a innym razem banalnie proste i tanie w realizacji – można jednak znaleźć także w wielu innych dziedzinach: aparaturze lotniczej i kosmicznej, sprzęcie AGD czy nawet... wyposażeniu sanitarnym. W artykule nie będziemy jednak dokonywać przeglądu wszystkich możliwych zastosowań czujników odległości i prędkości, skupimy się natomiast na tym, co najbardziej interesuje inżynierów – czyli na aspektach technologicznych. Pokażemy zarówno wybrane moduły OEM oraz gotowe urządzenia pomiarowe, jak i kompaktowe czujniki scalone, a także niektóre wyspecjalizowane komponenty, przeznaczone do budowy specjalizowanych rozwiązań klasy high-end.

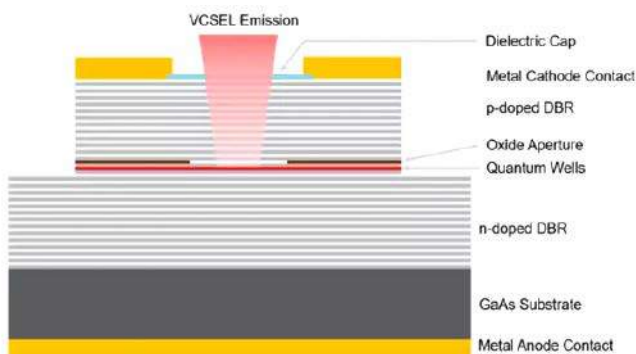
Pomiary odległości metodąToF

Na początek przyjrzyjmy się metodom oraz urządzeniom, służącym do pomiaru odległości obiektu (np. przeszkody znajdującej się na drodze robota mobilnego) od czujnika. Do dyspozycji mamy, jak zawsze, szeroką gamę metod, a większość z nich opiera się w gruncie rzeczy na tym samym zjawisku – pomiarze czasu przelotu wiązki fal określonego typu (ToF – *Time of Flight*). Wiązka jest emitowana przez nadajnik, a po odbiciu od przeszkody jej część powraca w kierunku odbiornika (**rysunek 1**). Różnica czasu Δt – przy znanej prędkości propagacji fali c w danym ośrodku – jest przeliczana na odległość d z doskonale znanego wzoru (1):

$$d = \frac{c \cdot \Delta t}{2} \quad (1).$$



Rysunek 1. Ogólna ilustracja zasady działania systemów ToF (<https://t.ly/RUZE9>)



Rysunek 2. Przekrój poprzeczny przez strukturę lasera typu VCSEL (<https://t.ly/RUZE9>)

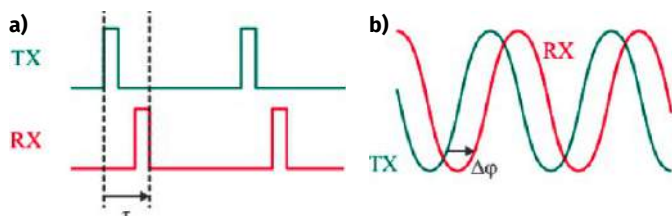
Równanie 1 można zastosować do dowolnej fali, używanej w celu pomiaru odległości – i tak, w praktyce stosuje się następujące zjawiska:

- światło (zwykle lasera półprzewodnikowego lub – rzadziej – diod LED),
- ultradźwięki (wytwarzane i odbierane najczęściej przez przetworniki piezoelektryczne),
- mikrofale (generowane za pomocą anten pojedynczych lub – w przypadku bardziej zaawansowanych systemów – z użyciem macierzy anten pracujących w szyku fazowanym).

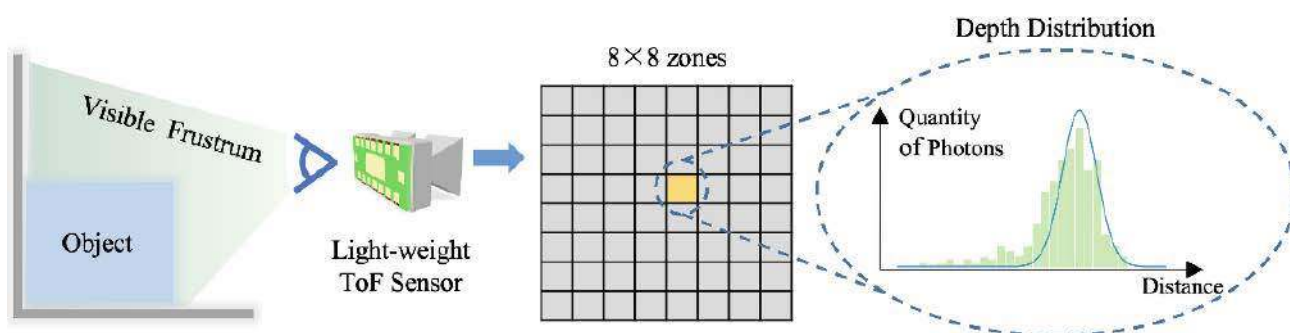
W dalszej części artykułu zwrócimy uwagę na modyfikacje oryginalnej metody ToF, istotne zwłaszcza w przypadku rozwiązań laserowych i radarowych – często bowiem wykorzystuje się nie tyle samo opóźnienie (jako takie) w propagacji fali, ile powiązane z nim zależności fazowe pomiędzy sygnałami nadanymi i powracającymi.

Techniczne realizacje optycznej metody ToF

Metoda pomiaru odległości z użyciem technologii Time-of-Flight zyskała niebywałą popularność w ciągu ostatnich kilku lat, a to głównie za sprawą upowszechnienia niedrogich czujników scalonych, zbudowanych w oparciu o podczerwone lasery z pionową wnęką rezonansową VCSEL (*Vertical Cavity Surface Emitting Laser* – rysunek 2). Zanim jednak przejdziemy do opisu przykładowych rozwiązań dostępnych na rynku, musimy zaprezentować dwie, zasadniczo różne, realizacje ToF: bezpośrednią i pośrednią (rysunek 3).



Rysunek 3. Metoda dToF (a) korzysta z pomiaru opóźnienia pomiędzy momentami nadania i odbioru krótkiego impulsu laserowego; metoda iToF (b) bazuje na pomiarze różnicy fazy pomiędzy nadaną a odebraną wiązką światła zmodulowanego amplitudowo (<https://t.ly/WXIXH>)



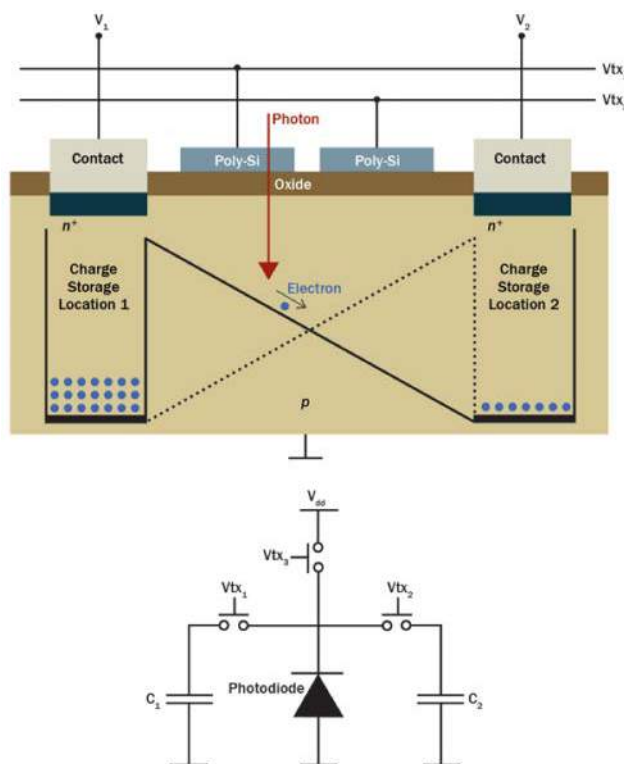
Rysunek 4. Obliczenia statystyczne, zastosowane w celu wyznaczenia odległości na podstawie danych z pojedynczego piksela wielosegmentowego czujnika ToF (<https://t.ly/XMvql>)

Metoda dToF

Ta pierwsza (dToF, *direct Time-of-Flight*) bazuje bezpośrednio na praktycznej implementacji równania 1 – układ wysyła bardzo krótki impuls światła laserowego (zwykle o szerokości rzędu 0,2...5 ns) w kierunku przeszkody i jednocześnie rozpoczyna pomiar czasu, upływającego od rozpoczęcia emisji do momentu odebrania fali odbitej (powrotnej). Z uwagi na znikomą ilość odebranych fotonów, a także ze względu na wymogi dotyczące dynamiki (czasu reakcji), w roli detektorów stosuje się fotodiody lawinowe (APD) lub detektory jednofotonowe typu SPAD (szczegółowy opis obydwu tych odmian fotodetektorów można znaleźć w obszernym opracowaniu pt. *Fotoelementy – serce optoelektroniki*, zamieszczonym w „Elektronice Praktycznej” nr 3/2023). Aby zminimalizować wpływ szumu i zakłóceń zewnętrznych na wynik pomiaru, konieczna jest wielokrotna akwizycja próbek i obróbka ich na zasadzie analizy histogramu. Metoda nadaje się do zastosowania w prostszych czujnikach z pojedynczym polem widzenia (FOV) lub polem podzielonym na niewielką liczbę pikseli (czujniki wielosegmentowe – rysunek 4).

Metoda iToF

Diametralnie inne założenia leżą u podstaw metody pośredniej ToF (iToF, *indirect ToF*). W tym przypadku stosuje się falę ciągłą,



Rysunek 5. Uproszczony przekrój struktury mieszczącej fotonicznego (PMD) – na górze – oraz jego schemat zastępczy – na dole (<https://t.ly/K51XI>)

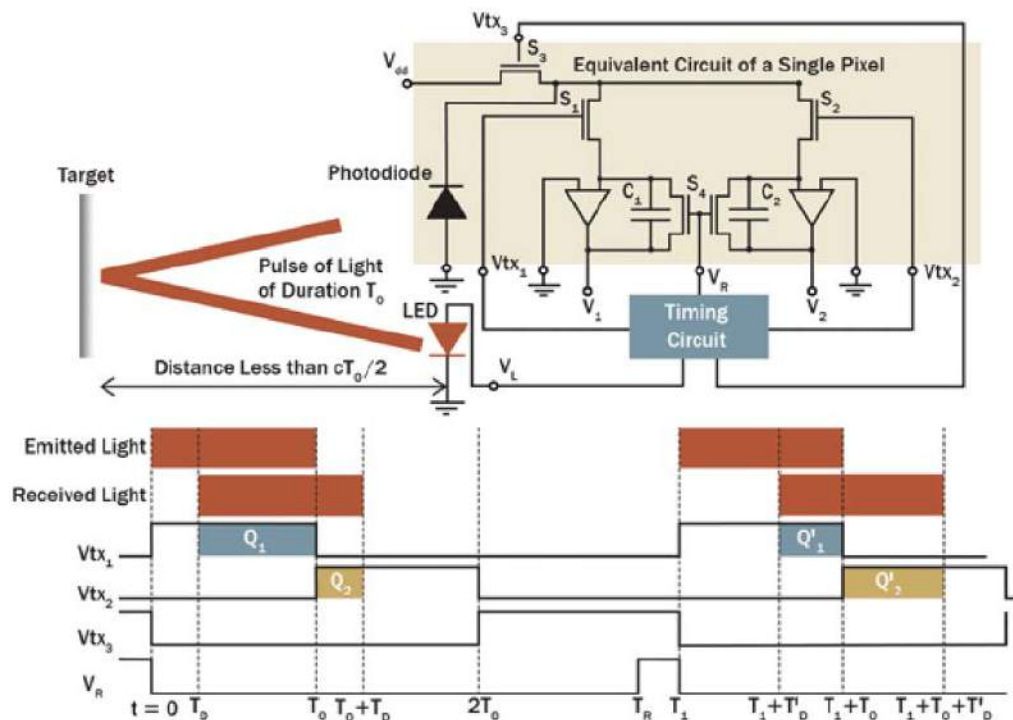
a ściślej rzecz biorąc – wiązkę zmodulowaną amplitudowo sygnałem o relatywnie wysokiej częstotliwości (zwykle 20...100 MHz). Pośrednią miarą odległości – przeliczaną następnie na konkretną jednostkę, np. milimetry – jest tutaj zatem nie tyle opóźnienie czasowe, co... różnica fazy pomiędzy sygnałem nadanym a falą odebraną przez detektor. Metoda ta pozwala znacząco zwiększyć rozdzielczość przestrzenną detektora nawet do poziomu porównywalnego ze standardowymi modułami kamer CMOS. W tym przypadku jednak znajdują zastosowanie specjalne detektory, określane mianem PMD (od *Photonic Mixer Device*, co w dosłownym tłumaczeniu oznacza fotoniczny mieszacz – **rysunek 5**). Ich niezwykłą zaletą jest zdolność do demodulacji fazy odebranej wiązki światła w sposób... całkowicie organiczny, tj. na poziomie poszczególnych pikseli, bez udziału zewnętrznych układów kondycjonujących.

Ilustrację zasady działania PMD pokazano na **rysunku 6** – foton padający na powierzchnię struktury krzemowej detektora powoduje wygenerowanie ładunku elektrycznego, który – w zależności od podanego z zewnątrz sygnału modulującego (zsynchronizowanego z oświetlaczem) – przerzuca go do jednego z dwóch kondensatorów. W zależności od przesunięcia fazowego pomiędzy sygnałem nadanym a odebranym przez dany piksel zmienia się stosunek napięć (V_1 , V_2), będący efektem gromadzenia ładunków w poszczególnych kondensatorach.

Warto odnotować, że z uwagi na fakt, iż najważniejsza część procesu demodulacji jest wykonywana na poziomie struktury detektora, do jego pracy jest konieczny w zasadzie tylko zewnętrzny układ synchronizacji oraz przetworniki ADC, odpowiedzialne za próbkowanie napięć z półpikseli. Za cenę większego stopnia złożoności samej matrycy (w porównaniu do zwykłych czujników obrazu, np. CMOS) otrzymujemy jednak niebywale cenną możliwość mapowania głębi sceny, czyli... obrazowania odległości poszczególnych jej elementów od kamery.

Scalone czujniki laserowe ToF

W ostatnich latach mamy do czynienia z istną ekspansją sensorów dToF i to zarówno w urządzeniach konsumenckich, jak i aplikacjach przemysłowych. Opanowanie technologii produkcji niedrogich, kompaktowych czujników w obudowach SMD sprawiło, że kolejne firmy dołączają do półprzewodnikowego wyścigu, udostępniając klientom coraz bardziej zaawansowane układy o zaskakujących możliwościach pomiarowych. Głównym graczem na rynku jest firma STMicroelectronics, która wdrożyła całą serię sensorów odległości, bazujących na opatentowanej odmianie technologii dToF nazwanej FlightSense. Czujniki zawierają oświetlacze laserowe VCSEL oraz detektory bazujące na strukturach SPAD, zaś dane pozyskane z analogowego front-endu są przetwarzane statystycznie i udostępniane poprzez interfejs I²C. Rodzina czujników, oznakowanych jako VL61... lub VL53..., zawiera zarówno proste, podstawowe sensory odległości, jak i rozbudowane, wielostrefowe czujniki, przystosowane do aplikacji AI związanych m.in. ze zliczaniem obiektów (np. osób w pomieszczeniu). Niektóre modele mają ponadto wbudowany czujnik oświetlenia zewnętrznego (ALS), co ma znaczenie w przypadku



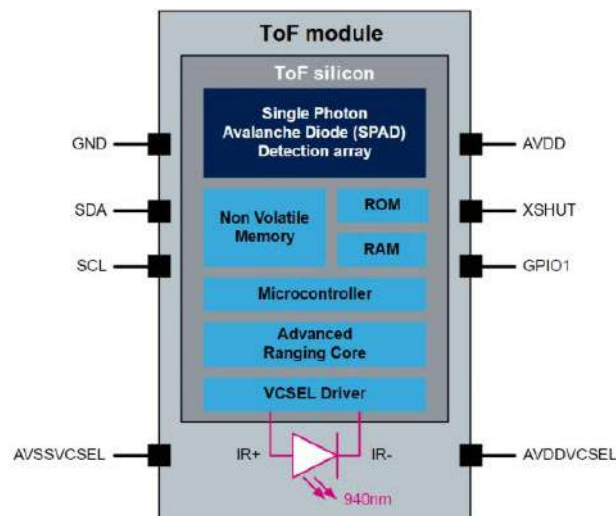
Rysunek 6. Zasada działania mieszacza fotonicznego (<https://t.ly/K51XI>)

zastosowań w urządzeniach mobilnych, głównie smartfonach.

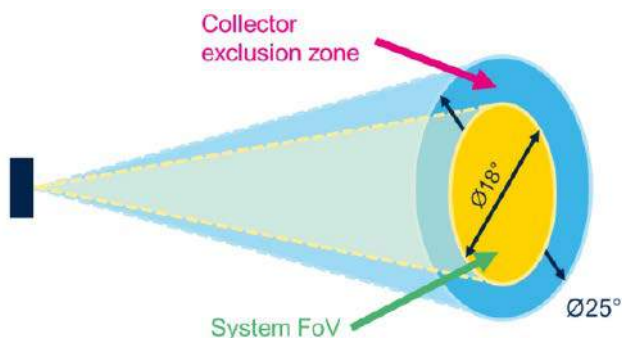
Jako przykład jednego z bardziej zaawansowanych czujników ToF marki STMicroelectronics warto podać układ o oznaczeniu VL53L4CX (**fotografia 1, rysunek 7**). Dzięki wąskiemu polu widzenia (efektywny obszar FoV to zaledwie 18° – **rysunki 8 i 9**) konstruktorom udało się uzyskać spory zasięg, dochodzący aż do 6 metrów (!), choć rzeczywista wartość zależy w dużej mierze od refleksyjności powierzchni – w przypadku białej planszy (reflektancja 88%) zasięg istotnie osiąga 500...600 cm, ale przy jasnoszarych obiektach (54%) parametr ten spada do 420...460 cm, zaś dla refleksyjności 17% wynosi już zaledwie 210...250 cm. Co więcej – wszystkie te pomiary dotyczą tylko pracy wewnątrz pomieszczeń, gdyż przy oświetleniu słonecznym zasięg



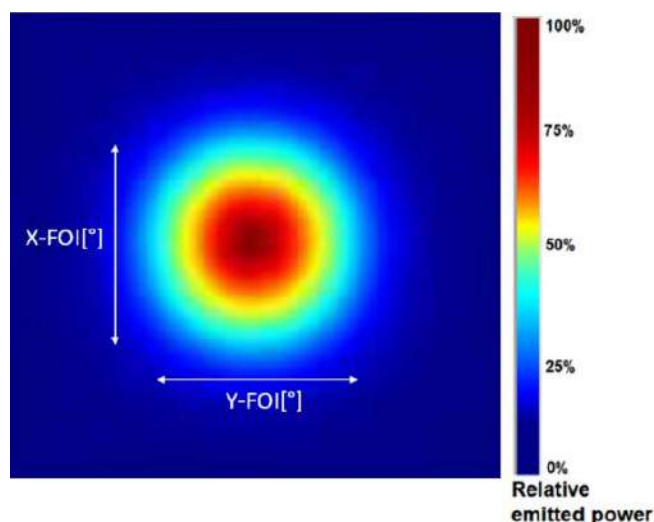
Fotografia 1. Czujnik VL53L4CX marki STMicroelectronics (<https://t.ly/8DY3r>)



Rysunek 7. Uproszczony schemat blokowy czujnika VL53L4CX (<https://t.ly/8DY3r>)

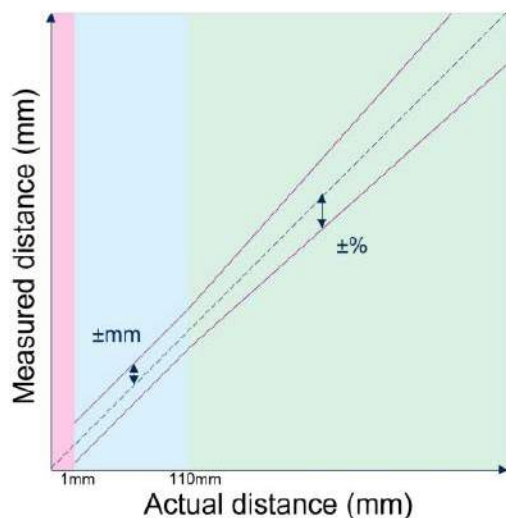


Rysunek 8. Pole widzenia czujnika VL53L4CX (<https://t.ly/8DY3r>)

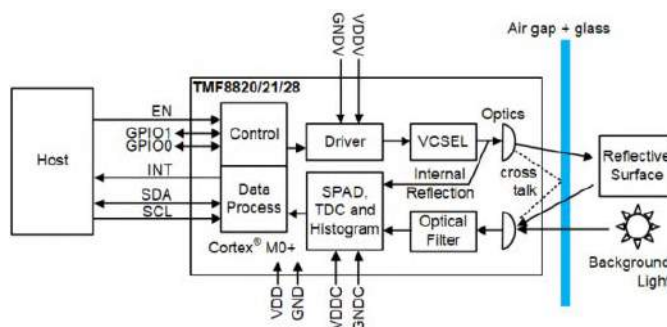


Rysunek 9. Kątowy rozkład mocy wiązki laserowej czujnika VL53L4CX; oś Y – kąt w płaszczyźnie horyzontalnej, oś X – kąt w płaszczyźnie pionowej(<https://t.ly/8DY3r>)

drastycznie spada do zaledwie 110...180 cm (w najlepszym wypadku). Nie należy jednak dziwić się tak dużemu rozrzutowi wartości, zważywszy na fakt, iż mamy do czynienia z miniaturowym czujnikiem (4,4×2,4×1,0 mm), w którym źródłem światła jest mały laser VCSEL klasy 1 oraz równie kompaktowe detektory, pracujące (dosłownie) z pojedynczymi fotonami. Drastycznie trudne warunki, narzucone przez zastosowaną technologię, wpływają ponadto na ograniczenie dokładności pomiarowej – ta znacząco maleje powyżej granicy 110 cm i w przypadku szarych obiektów o refleksyjności 17% wynosi zaledwie ±8% (w otwartym terenie), podczas gdy odległość przeszkód białych (88%) w pomieszczeniach jest wskazywana z dokładnością ±3%.



Rysunek 10. Charakterystyka dokładności pomiarowej czujnika VL53L4CX w funkcji odległości (<https://t.ly/8DY3r>)



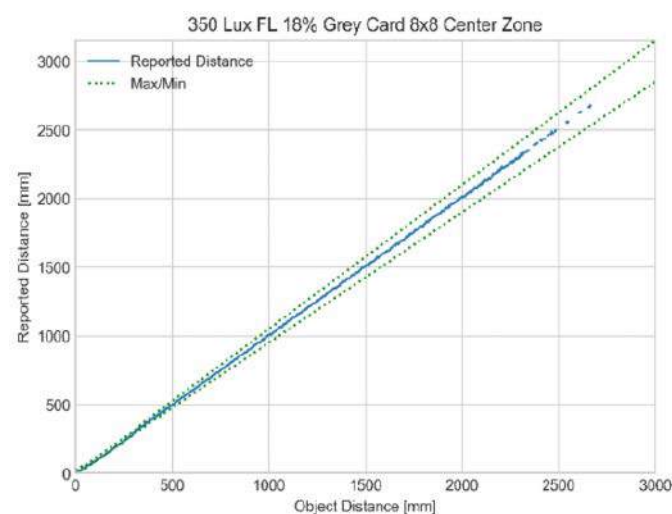
Rysunek 11. Schemat blokowy czujników z serii TMF8820/21/28 (<https://t.ly/TkP5t>)



Fotografia 2. Czujniki z serii TMF8820/21/28 marki ams OSRAM (<https://t.ly/M0D84>)

Warto zwrócić uwagę, że w sensorach dToF marki ST mamy do czynienia z charakterystycznym, dwuzakresowym sposobem określania dokładności – w przedziale bliskim (w tym przypadku 10...110 mm) granica błędu jest podawana w sposób bezwzględny (mm), zaś dopiero dalej – jako dokładność względna (% wskazania), patrz **rysunek 10**. W przypadku układów innych producentów – w tym ams-OSRAM – dokładność może być podawana w prostszy sposób, np. jako błąd względny dla danych warunków pracy. Przykład takiej charakterystyki dla wielostrefowego czujnika z serii TMF8820/21/28 (**rysunek 11**, **fotografia 2**) można zobaczyć na **rysunku 12**.

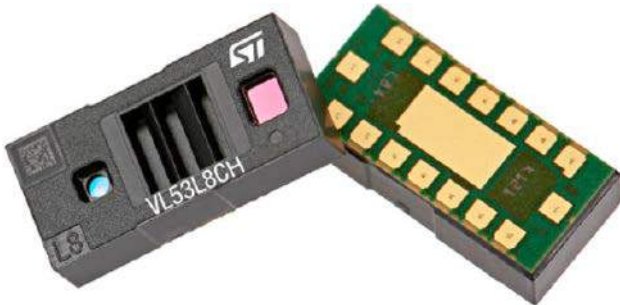
Co ciekawe, nawet w tak nowoczesnych technologicznie produktach, jak współczesne czujniki dToF, możemy już zauważyć swego rodzaju powrót do źródeł. Pierwsze tego typu sensory dokonywały obliczeń statystycznych w sposób całkowicie automatyczny, udostępniając użytkownikowi gotowy wynik pomiaru – i to było w zupełności wystarczające do większości aplikacji. Okazało się jednak, że wraz z wdrożeniem czujników wielostrefowych o coraz wyższej rozdzielczości przestrzennej układy dToF zaczęły być rozważane jako potencjalnie



Rysunek 12. Przykładowa charakterystyka dokładności czujników z serii TMF8820/21/28 w funkcji odległości (<https://t.ly/FEPNb>)

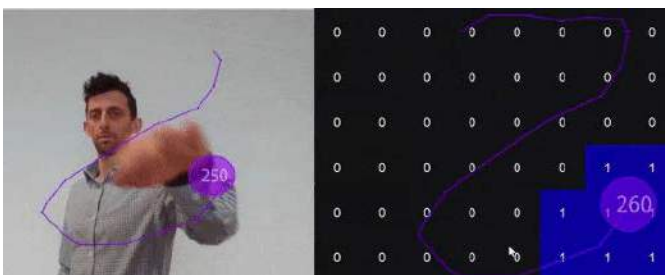


Rysunek 13. Przykładowy histogram, uzyskany za pomocą czujnika VL53L8CH (<https://t.ly/DyQB5>)



Fotografia 3. Zaawansowany czujnik ToF typu VL53L8CH marki STMicroelectronics (<https://t.ly/DyQB5>)

atrakcyjne rozwiązanie do wykrywania gestów bądź zliczania osób i innych obiektów w przestrzeni trójwymiarowej. Ekspansja zastosowań sztucznej inteligencji dotarła i tutaj, po czym szybko okazało się, że proste dane, będące wynikiem obliczeń dokonywanych przez procesor zintegrowany wewnątrz czujnika, mogą być niewystarczające dla sieci neuronowych, które z zasady „pożerają” spore ilości danych wejściowych. I tak, firma STMicroelectronics postanowiła zrobić swego rodzaju krok w tył – układ VL53L8CH (**fotografia 3**), dumnie określany mianem *Artificial intelligence enabler*, pozwala na pozyskiwanie z czujnika całego zestawu surowych danych pomiarowych w postaci pełnych histogramów (**rysunek 13**), zamiast pojedynczych wyników pomiaru dystansu dla poszczególnych pikseli 64-pikselowej matrycy, otrzymujemy więc pełne statystyki, na podstawie których algorytmy AI mogą wnioskować o obiektach znajdujących się w polu widzenia ze znacznie większą elastycznością. W praktyce może się to przełożyć na znaczną poprawę efektywności algorytmów i redukcję



Rysunek 14. Przykład zastosowania wielostrefowego czujnika ToF do rozpoznawania prostych gestów (<https://t.ly/FkD09>)



Rysunek 15. Technologia ToF przynosi znaczne korzyści robotyce mobilnej, m.in. automatycznym odkurzaczom – pozwala bowiem na detekcję oraz omijanie przeszkód z dużą precyzją i odpowiednim wyprzedzeniem (https://t.ly/H_KBd)

artefaktów, np. w przypadku rozpoznawania gestów (**rysunek 14**) czy też skanowania otoczenia, znajdującego się przed robotem sprzątającym (**rysunek 15**).

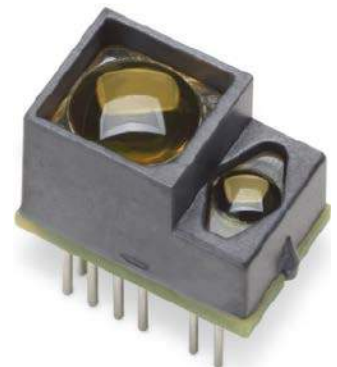
Moduły LIDAR

Podstawową wadą scalonych czujników odległości jest ich niewielki zasięg, przeważnie ograniczony do 4...6 metrów. W przypadku wielu aplikacji taki zakres pracy jest oczywiście wystarczający, jednak np. algorytmy omijania przeszkód bądź manewry automatycznego lądowania dronów wymagają już znacznie większego zasięgu detekcji. Kompaktowe dalmierze laserowe znajdują także zastosowanie w licznych aplikacjach robotycznych, przemysłowych czy też systemach bezpieczeństwa. Zapotrzebowanie na tego typu moduły jest niczym apetyt, który rośnie w miarę jedzenia – odpowiedzią na rosnący popyt w zakresie rozwiązań do pomiaru dużych odległości okazują się moduły LIDAR.

Do najbardziej kompaktowych LIDAR-ów o zasięgu przekraczającym 10 metrów należą wielostrefowe dalmierze z serii AFBR-S50 marki Broadcom (**fotografia 4**). W każdym modelu, w obudowie o rozmiarach 12,4×7,6×7,9 mm, umieszczono wydajny laser impulsowy VCSEL, macierz detektorów (w liczbie do 32) oraz układ ASIC, zasilany napięciem 5 V. Zasięg pomiaru dochodzi do 50 metrów, a prędkość odświeżania – w zależności od liczby obsługiwanych pikseli – osiąga nawet wartość 3 kHz. Co ważne, w przypadku pracy w trybie dwuczęstotliwościowym zasięg niektórych modeli wydłuża się nawet do 200 metrów, co przy tak kompaktowych wymiarach całości jest naprawdę nie lada osiągnięciem.

W podobnym do czujników AFBR-S50 zakresie mierzonych odległości pracują także liczne modele LIDAR-ów, oferowane na rynku elektroniki amatorskiej – choć w tym przypadku próżno szukać elementów o równie kompaktowych wymiarach. Jeżeli rozmiar możliwego do zamontowania modułu nie jest szczególnym ograniczeniem, można zdecydować się np. na LIDAR-Lite v3, którego dystrybutorem jest znana marka SparkFun, a producentem – firma Garmin (**fotografia 5**). Zasięg dalmierza wynosi 40 metrów, zaś moc szczytowa lasera (w tym przypadku zastosowano standardową diodę laserową o emisji bocznej zamiast źródła VCSEL) to 1,3 W. Wymiary obudowy sensora wynoszą 20×48×40 mm.

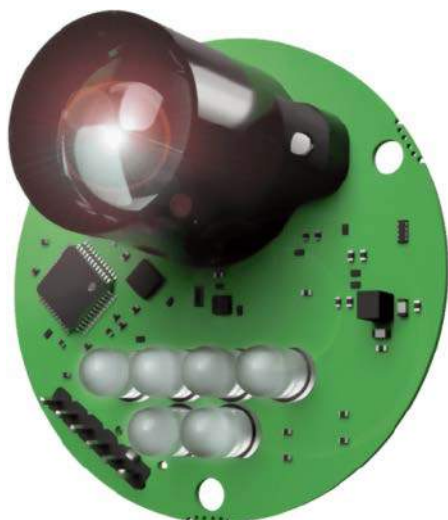
Co ciekawe, identyczny dystans można uzyskać także bez użycia lasera – takie rozwiązanie, z oświetlaczem wykonanym na bazie zestawu sześciu przewlekanych diod LED IR (850 nm) o kącie emisji 3°, zaproponowała marka



Fotografia 4. Dalmierz modułowy (LIDAR) z serii AFBR-S50 marki Broadcom (<https://t.ly/HbKr4>)



Fotografia 5. Moduł LIDAR-Lite v3 marki Garmin (<https://t.ly/rCoSy>)



Rysunek 16. Moduł LeddarOne marki LeddarTech (<https://t.ly/XKjTZ>)



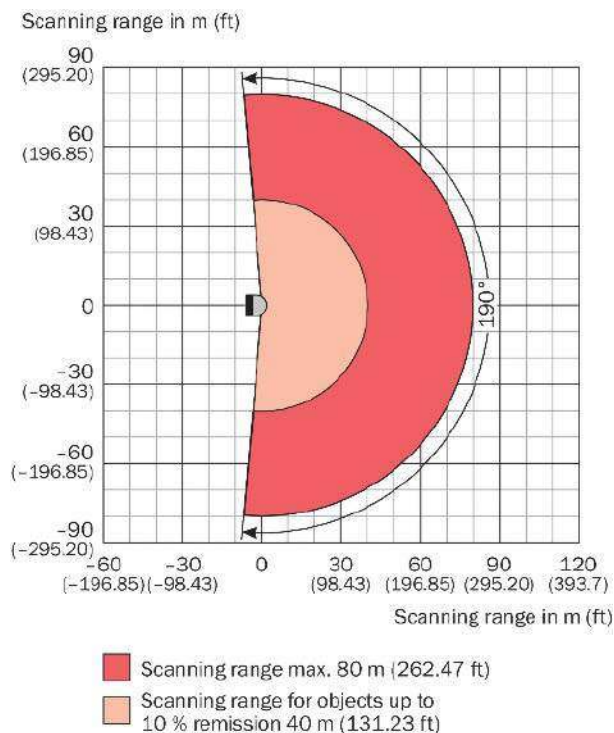
Fotografia 6. Moduły LIDAR marki Safran Vectronix AG (<https://t.ly/2xHvi>)

LeddarTech. Moduł LeddarOne (rysunek 16), zbudowany w oparciu na okrągłej płytce drukowanej o średnicy 51 mm, umożliwia akwizycję wyników pomiaru z częstotliwością dochodzącą do 140 Hz, a jego główną zaletą jest niewątpliwie... brak potencjalnego ryzyka, wynikającego z zastosowania doskonale skolimowanego źródła laserowego (w dodatku o niewidzialnym dla człowieka promieniowaniu). Większy kąt propagacji podczerwieni pozwala także niejako uśrednić wyniki pomiaru z większej powierzchni, co w niektórych sytuacjach będzie cechą wysoce pożądaną.

Na koniec tej części naszej prezentacji warto jeszcze zapoznać się z ofertą marki Safran Vectronix AG, która opracowała wysokiej klasy moduły LIDAR, przeznaczone m.in. dla branży militarnej. W tym przypadku mamy już do czynienia z naprawdę zaawansowanymi urządzeniami, zdolnymi do pracy w zakresie od 4,5 do nawet 25 km, zaś najwyższy model – LRF7047



Fotografia 7. Skaner LMS531-10100 Security marki SICK (<https://t.ly/K1MJG>)



Rysunek 17. Charakterystyka kątowna skanera LMS531-10100 Security marki SICK (<https://t.ly/K1MJG>)

(największy spośród pokazanych na fotografii 6) – pozwala na pomiar odległości aż do 30 km i to z doskonałą dokładnością 1 metra (!). Co szczególnie ważne w przypadku urządzeń wojskowych, LIDAR LRF7047 pracuje w paśmie 1550 nm, dzięki czemu pozostaje całkowicie niewidoczny dla konwencjonalnych kamer NIR.

Skanery laserowe OEM

Niektóre aplikacje – zwłaszcza w zakresie robotyki przemysłowej czy też autonomicznych pojazdów – wymagają dokładnej orientacji w otaczającej przestrzeni. W takich przypadkach pomiar odległości w jednym punkcie lub nawet na niewielkim, prostokątnym obszarze przestaje wystarczać – konieczne jest sięgnięcie po bardziej zaawansowane techniki. Skanery laserowe – bo o nich mowa – umożliwiają pełne mapowanie przestrzeni poprzez cykliczne przemiatanie badanego obszaru wiązką lasera i dokonywanie wielopunktowych pomiarów odległości. Tak działają m.in. szeroko stosowane w przemyśle skanery marki SICK – Czytelnicy zaznajomieni z automatyką przemysłową doskonale kojarzą urządzenia w charakterystycznych, półokrągłych obudowach, wyposażonych w czarne okno optyczne o zakresie pracy rzędu 190° (fotografia 7, rysunek 17).

Ten sam producent oferuje także skanery z oknami dookólnymi, o znacznie szerszym zakresie przemiatania (fotografia 8), a szczególnie podgrupę produktów stanowią skanery, przeznaczone do systemów bezpieczeństwa i certyfikowane na zgodność z normami ISO 13849-1 oraz IEC 62998 (fotografia 9). Tego typu urządzenia w wielu przypadkach mogą zastąpić lub przynajmniej efektywnie uzupełnić konwencjonalne zabezpieczenia, stosowane w halach produkcyjnych (np. bariery optyczne), choć równie dobrze spisują się w zastosowaniach mobilnych, np. wózkach AGV (fotografia 10).

Warto zwrócić uwagę, że opisane powyżej skanery LIDAR są przeważnie bardzo rozbudowanymi urządzeniami, bazującymi na zaawansowanych systemach elektroniczno-mechaniczno-optycznych (fotografia 11), co wpływa nie tylko na wysoką cenę, ale także na rozmiary i masę całości. W urządzeniach tego typu często można spotkać obrotowe lustro, przemiatające otoczenie wokół skanera (fotografia 12) – całość elektroniki może więc pozostać całkowicie nieruchoma, na stałe związana z obudową urządzenia. Diametralnie inne podejście zaproponowała firma LightWare Lidar LLC: skaner SF45/B ma



Fotografia 8. Portfolio standardowych skanerów LIDAR marki SICK (<https://t.ly/qrYqo>)



Fotografia 9. Portfolio skanerów LIDAR marki SICK, przeznaczonych do systemów zabezpieczeń (<https://t.ly/lZ9sm>)



Fotografia 10. Przykłady zastosowania certyfikowanych skanerów bezpieczeństwa: ochrona przed wejściem w strefę roboczą manipulatorów przemysłowych oraz zabezpieczenie robota mobilnego przed kolizją (<https://t.ly/WoFg8>)

bowiem... obrotową głowicę (fotografia 13). Kompaktowy LIDAR przeznacza otoczenie z częstotliwością 5 Hz, wykonując przy tym 5000 odczytów odległości w zakresie nawet do 50 m, podczas gdy nieruchoma pozostaje jedynie niskoprofilowa podstawa montażowa urządzenia.

Kamery ToF

Liczba dostępnych na rynku modułów kamer iToF jest rzecz jasna nieporównanie mniejsza niż w przypadku prostszych czujników wielostrefowych i jednostrefowych dToF. Nietrudno jednak zauważyć, że obiecujące możliwości, wynikające z połączenia obrazowania 3D

oraz sztucznej inteligencji, intensywnie napędzają koniunkturę w tym segmencie rynku, co przejawia się wprowadzaniem do sprzedaży kolejnych matryc iToF, front-endów ToF, modułów OEM oraz zestawów ewaluacyjnych, bazujących na ultranowoczesnych czujnikach głębi. I tak, znana z szerokiego portfolio rozmaitych czujników optycznych firma Melexis opracowała macierze MLX75026 (QVGA – fotografia 14) oraz MLX75027 (VGA), zaś Texas Instruments wprowadził do sprzedaży masowej matrycę OPT8241 (320×240 px – fotografia 15). Jeszcze dalej poszły firmy STMicroelectronics (VD55H1 – czujnik obrazu iToF o rozdzielczości 672×804 px – obecnie w fazie



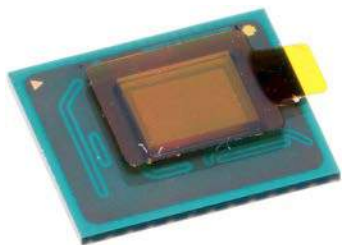
Fotografia 11. Wygląd wnętrza skanera laserowego SICK (<https://t.ly/d-G4P>)



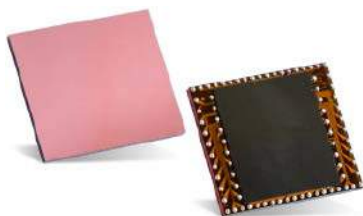
Fotografia 12. Lustro obrotowe, zastosowane w skanerze laserowym marki SICK (<https://t.ly/d10Tu>)



Fotografia 13. Miniaturowy skaner laserowy z obrotową głowicą – SF45/B marki LightWare Lidar LLC (<https://t.ly/qz67F>)



Fotografia 14. Matryca ToF typu MLX75026RTH-AAA-210-SP o rozdzielczości QVGA (<https://t.ly/YGCbd>)



Fotografia 15. Matryca ToF typu OPT8241 marki Texas Instruments (<https://t.ly/gbnO3>)



Fotografia 16. Matryca ToF typu ADSD3100 marki Analog Devices (https://t.ly/e_pVI)



Fotografia 17. Zestaw ewaluacyjny do kamery ToF na bazie matrycy MLX75026 (<https://t.ly/vCHFD>)

przedwdrożeńiowej) oraz Analog Devices (1-megapikselowa matryca ADSD3100 o rozdzielczości 1024×1024 px – **fotografia 16**).

Co ciekawe, matryce o mniejszej rozdzielczości da się już dziś kupić za 200...300 złotych, można zatem uznać, że ogromny potencjał, oferowany przez technologię iToF, otrzymujemy za półdarmo. Sprawa okazuje się jednak mniej optymistyczna, jeżeli weźmiemy pod uwagę stopień złożoności problemów natury mechanicznej i optycznej, związanych z budową modułu kamery od podstaw. Do tego dochodzi kwestia szybkiego oświetlacza, który pozwoli uzyskać w miarę możliwości jednorodne oświetlenie całej sceny podczerwienią...

Świadomi tych wszystkich aspektów inżynierowie opracowali zatem zestawy ewaluacyjne, pozwalające na szybkie rozpoczęcie prac z nowoczesnymi matrycami ToF. O ile jednak rozbudowane zestawy marki Melexis, wyposażone nawet w interfejs Ethernet (**fotografia 17**), są raczej niezbyt przyjazne zainteresowanym nimi użytkownikom pod względem finansowym (ceny zaczynają się od 9 tysięcy złotych i szybką – w zależności od wersji – nawet do ponad 20 tys. PLN), to już doskonały zestaw ADTF3175BMLZ (**fotografia 18**) marki Analog Devices jest w chwili pisania niniejszego artykułu dostępny za niewiele ponad 1000 zł (!). Warto też wspomnieć o gotowym do użycia module kamery ToF o nazwie flexx2 (**fotografia 19**), opracowanym przez firmę Pmdtechnologies ag – oferowana przez niego rozdzielczość jest wprawdzie bardzo przeciętna (zaledwie 224×172 px), ale za to wygodne



Fotografia 18. Zestaw ewaluacyjny do 1-megapikselowej kamery ToF marki Analog Devices (https://t.ly/Gz_7Z)



Fotografia 19. Widok rozstrzelony kamery flexx2 (<https://t.ly/AjM1w>)

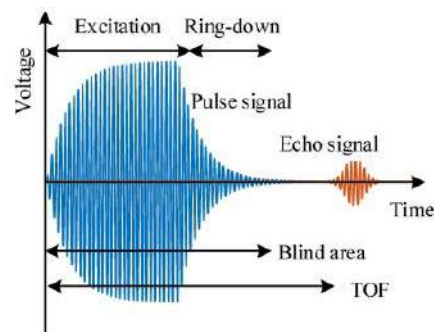
podłączenie (USB-C) i rozbudowane SDK ze wsparciem języków C/C++, Matlaba, czy też biblioteki OpenCV, zapewniają ułatwioną integrację oraz znacząco przyspieszają rozpoczęcie pracy z kamerą na różnych platformach sprzętowych i programowych.

Sonary ultradźwiękowe

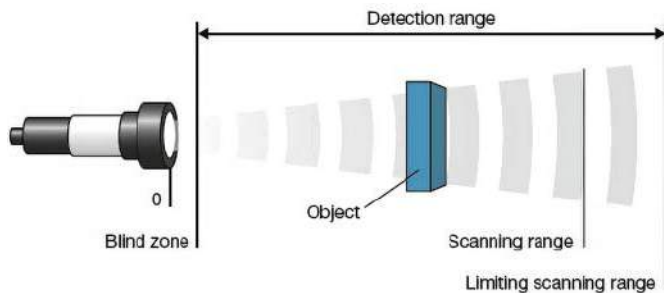
Pomiar odległości z użyciem sonarów (*SOund Navigation And Ranging*) ultradźwiękowych jest znany już od przeszło stu lat, a impulsem do rozpoczęcia prac nad nawigacją akustyczną była słynna katastrofa „Titanica”. Po upływie całego wieku, sonary nadal stanowią jedno z podstawowych narzędzi w branży morskiej, choć równie często można je spotkać w innych aplikacjach – rybołówstwie i wędkarstwie (echosondy do wykrywania pojedynczych ryb i całych ławic), robotyce mobilnej (pomiar odległości od przeszkód) czy motoryzacji (czujniki parkowania). Od strony konstrukcyjnej można wyróżnić dwie grupy rozwiązań, różniące się sposobem użytkowania przetworników. Systemy bazujące na dwóch przetwornikach – nadawczym i odbiorczym – są w aplikacjach komercyjnych spotykane nieporównanie rzadziej niż w elektronice amatorskiej, gdzie za sprawą furory, jaką zrobił kultowy moduł HC-SR04, na rynku wciąż pojawiają się rozmaite klony i ulepszone wersje słynnego dalmierza (**fotografia 20**). Z wielu względów znacznie lepsze okazują się jednak czujniki bazujące na pojedynczym przetworniku, pełniącym naprzemiennie funkcję nadajnika i odbiornika fal ultradźwiękowych.



Fotografia 20. Moduł taniego dalmierza ultradźwiękowego RCWL-1601, kompatybilny z HC-SR04 (<https://t.ly/vMx6u>)



Rysunek 18. Wyznaczenie czasu trwania paczki impulsów pobudzających przetwornik oraz czasu zaniku drgań rezonansowych pozwala na określenie strefy martwej (*blind area*), poniżej której dalmierz ultradźwiękowy nie jest w stanie wykrywać echa obiektów (<https://t.ly/ZP-AQ>)



Rysunek 19. Strefa martwa (blind zone) dalmierza ultradźwiękowego (<https://t.ly/LUZxZ>)

Rzecz jasna, takie rozwiązanie wiąże się z koniecznością bardziej efektywnego wytłumienia drgań rezonansowych po zakończeniu paczki impulsów pobudzających (rysunek 18), co wpływa bezpośrednio na wielkość strefy martwej (tj. obszaru znajdującego się tuż przed sensorem, w którym nie da się wykryć przeszkody z uwagi na zbyt silny udział drgań resztkowych z fazy nadawania – patrz rysunek 19). W zamian za tę niedogodność otrzymujemy jednak bardzo kompaktowe rozwiązanie, stosunkowo łatwe do uszczelnienia (co ma znaczenie np. w przypadku czujników parkowania) oraz nieporównanie prostsze w montażu. Czujniki jednoprzetwornikowe występują w różnorodnych wersjach – od niedrogich, prostych sonarów do podstawowych zastosowań (fotografia 21), aż po precyzyjne, przemysłowe czujniki z kompensacją termiczną (umożliwiające nawet synchroniczną współpracę wielu sensorów w jednym systemie – patrz fotografia 22).

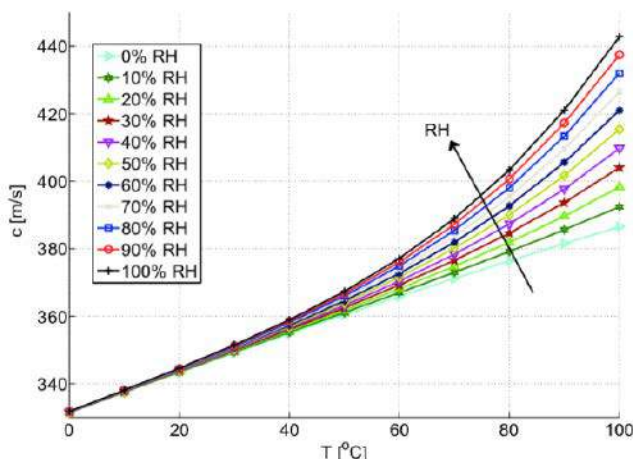


Fotografia 21. Prosty moduł dalmierza ultradźwiękowego – I²CXL-MaxSonar-EZ marki MaxBotix (<https://t.ly/h4u5m>)

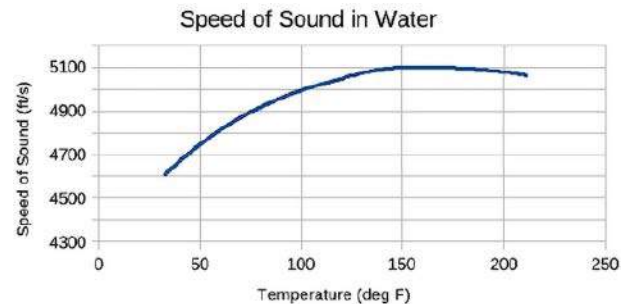


Fotografia 22. Przemysłowy czujnik odległości UMD123U035 (<https://t.ly/ejoP0>)

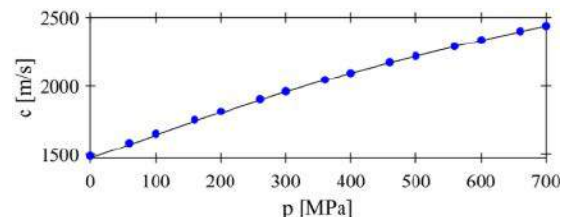
Istotną wadą dalmierzy bazujących na ultradźwiękach jest wzrost błęd pomiarowego podczas zmiany warunków otoczenia. Nie od dziś wiadomo, że prędkość dźwięku zależy – i to w dodatku nieliniowo – zarówno od temperatury, jak i wilgotności powietrza (rysunek 20) oraz (w mniejszym stopniu) od stężenia dwutlenku węgla, zaś w przypadku



Rysunek 20. Zależność prędkości propagacji dźwięku w powietrzu od temperatury i wilgotności w warunkach ciśnienia 101,3 kPa oraz przy stężeniu CO₂ równym 314 ppm (<https://t.ly/pfw9g>)



Rysunek 21. Zależność prędkości dźwięku w wodzie od temperatury (<https://t.ly/tN7Lm>)



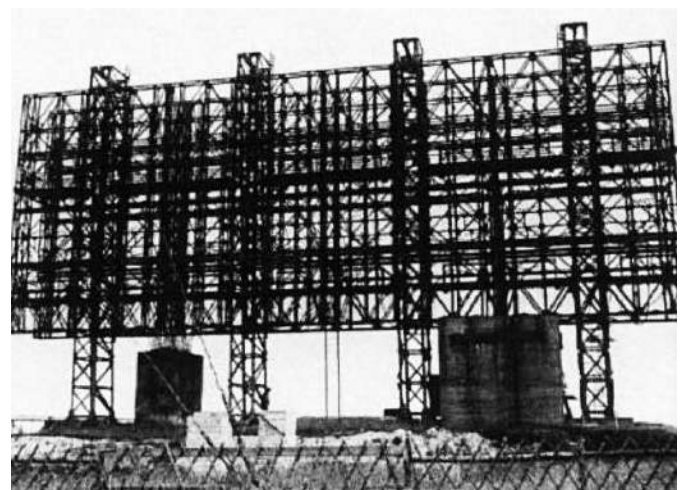
Rysunek 22. Zależność prędkości dźwięku od ciśnienia wody (<https://t.ly/IW0JX>)

sonarów wodnych trzeba wziąć pod uwagę nie tylko zasolenie ośrodka (wpływające na jego gęstość, a co za tym idzie – także prędkość propagacji fal akustycznych), ale także temperaturę (rysunek 21) oraz ciśnienie (rysunek 22).

Radary mikrofalowe ToF

Zastosowanie fal radiowych do zdalnej detekcji obiektów jest znane od początku ubiegłego stulecia, a największy udział w rozwoju technik radarowych miało – jak to zwykle bywa – wojsko. W 1904 roku niemiecki wynalazca Christian Hülsmeyer pokazał, że istnieje możliwość wykrywania obiektów metalicznych z użyciem fal elektromagnetycznych, choć jego pierwsze próby detekcji statku w gęstej mgle nie pozwoliły na określenie dokładnej odległości (a jedynie na stwierdzenie faktu obecności okrętu w danym kierunku). Technika radarowa była jednak na tyle obiecująca, że w ciągu kolejnych trzech dekad rozwinęła ją do poziomu, pozwalającego na efektywne zastosowanie w warunkach bojowych.

Mało tego – już pod koniec wojny Trzecia Rzesza dysponowała pierwszym na świecie radarem z szykiem fazowanym (FuMG 41/42 Mammut – fotografia 23), zdolnym do wykrywania celów lotniczych na odległość 300 km za pomocą fal radiowych w zakresie 120...150 MHz i to z dokładnością $\pm 0,5^\circ$. Dziś ta technologia jest szeroko stosowana w systemach obrony przeciwlotniczej, kontroli ruchu powietrznego,



Fotografia 23. Pierwszy w historii radar z szykiem fazowanym – niemiecki FuMG 41/42 Mammut (<https://t.ly/n7EOu>)

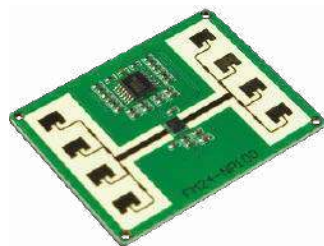


Fotografia 24. Współczesny zespół anten radarowych, wykorzystujący zysk fazowany (<https://t.ly/jglmk>)

aplikacjach kosmicznych, meteorologii, a nawet telekomunikacji (**fotografia 24**).

Zainteresowanie, jakie budzą kompaktowe radary mikrofalowe w zakresie aplikacji motoryzacyjnych, automatyki budynkowej i IoT, od dłuższego czasu napędza rozwój tego obszaru rynku, dzięki czemu niedrogo dostępne moduły do pomiaru odległości są coraz łatwiej dostępne w ofertach różnych producentów i dystrybutorów elektroniki. Dobrym przykładem może być moduł SEN0306 (**fotografia 25**) o rozmiarach 44×34 mm, pozwalający na pomiar odległości w przedziale od 50 cm do 20 metrów, przy użyciu mikrofal w paśmie K. Odczyt odbywa się z częstotliwością 10 Hz, zaś rozdzielczość pomiaru to 1 cm. Czujnik osiąga dokładność rzędu ± 10 cm.

Jeszcze lepsze osiągi oferuje miniaturowy moduł SMD o wymiarach 16×12 mm – V-LD1 marki RFbeam Microwave GmbH (**fotografia 26**). Urządzenie, bazujące na nowoczesnym układzie RF typu ASIC,



Fotografia 25. Moduł SEN0306 z oferty firmy DFRobot (<https://t.ly/qtfov>)



Fotografia 26. Moduł radarowy V-LD1 marki RFbeam Microwave GmbH (<https://t.ly/7W90n>)



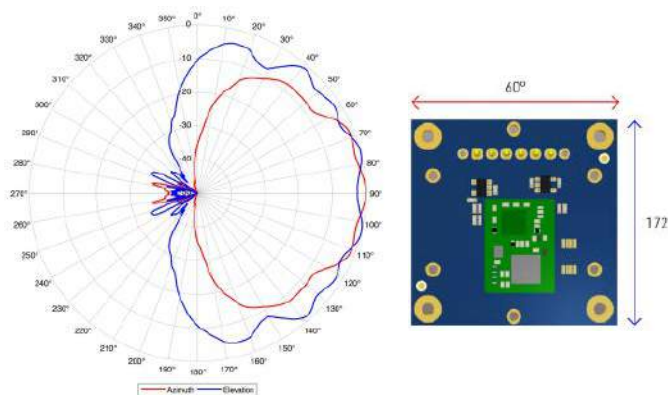
Fotografia 27. Zestaw ewaluacyjny dla modułu V-LD1 z zamontowaną soczewką (<https://t.ly/MSjUD>)



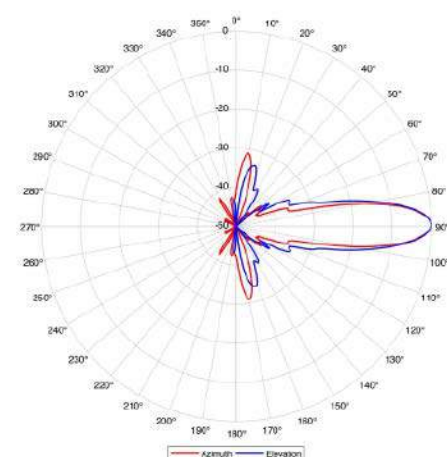
wspieranym przez procesor z serii STM32G4, jest w stanie osiągnąć zakres pomiarowy rzędu nawet 50 metrów, jednak do tego celu konieczne jest zastosowanie specjalnej, polimerowej soczewki (**fotografia 27**), znacząco poprawiającej charakterystykę kierunkową anteny (por. **rysunki 23 i 24**).

Warto zwrócić uwagę na fakt, że udział radarów impulsowych (opierających się tylko na pomiarze czasu przelotu fali) znacząco zmalał z uwagi na liczne problemy, jakie ta technologia stwarzała inżynierom, w szczególności w sektorze wojskowym (z uwagi na konieczność stosowania nadajników o ogromnej mocy chwilowej). Dziś szeroko stosowane są natomiast odmiany oryginalnej koncepcji, wykorzystujące zjawisko Dopplera: radary typu pulse-Doppler oraz MTI (Moving Target Indicator) są cennym narzędziem w aplikacjach militarnych czy też meteorologicznych. W przypadku czujników radarowych małej mocy stosowana jest natomiast fala ciągła z modulacją częstotliwościową (FMCW) – tak działają m.in. wspomniane wcześniej moduły SEN0306 oraz V-LD1.

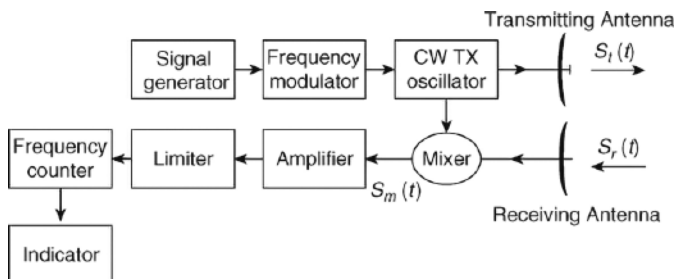
Zasada działania radarów FMCW bazuje na pomiarze przesunięcia częstotliwości pomiędzy sygnałem nadanym a odebrany (**rysunek 25**), przy czym owo przesunięcie zależy od odległości – im dalej znajduje się obiekt, tym większe opóźnienie, a zatem także – większe przesunięcie częstotliwości. Poszczególne implementacje różnią się m.in. kształtem przebiegu modulującego (sinusoida, fala trójkątna lub piłokształtna – **rysunek 26**), a – co ważne – zasięg pomiaru jest ogólnie ograniczony przez częstotliwość modulacji. Warto dodać, że na błąd pomiaru wpływa efekt Dopplera, przykładowo, jeżeli obiekt odbijający falę w kierunku radaru jest w ruchu względem anteny, to odbiornik pozyska sygnał nieco przesunięty w dziedzinie częstotliwości, zaś przesunięcie to będzie zależne od prędkości ruchu oraz częstotliwości nośnej.



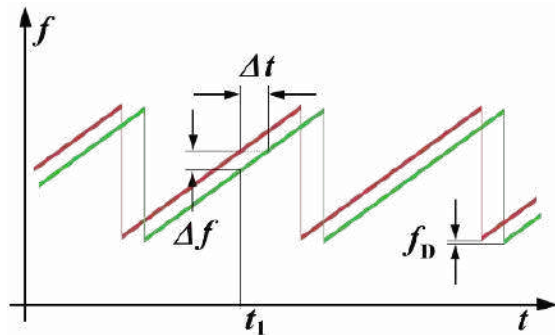
Rysunek 23. Oryginalna charakterystyka kierunkowa anteny modułu V-LD1 (<https://t.ly/MSjUD>)



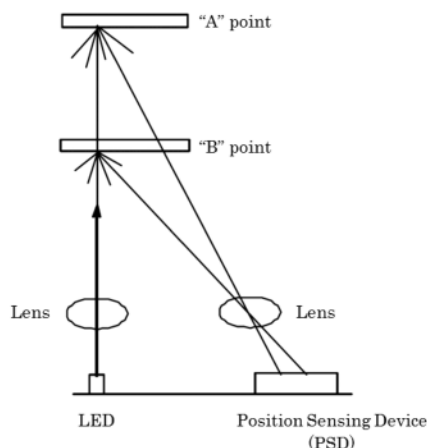
Rysunek 24. Charakterystyka kierunkowa anteny modułu V-LD1 po zastosowaniu soczewki, widocznej na fotografii 27 (<https://t.ly/MSjUD>)



Rysunek 25. Uproszczony schemat blokowy radaru FMCW (<https://t.ly/MIQkK>)



Rysunek 26. Przesunięcie częstotliwości sygnału powrotnego w radarze typu FMCW z modulacją przebiegiem piłokształtnym; Δf – różnica częstotliwości, Δt – opóźnienie pomiędzy sygnałem nadanym a odebranym; f_D – przesunięcie dopplerowskie, spowodowane ruchem obiektu (<https://t.ly/MJJ55>)



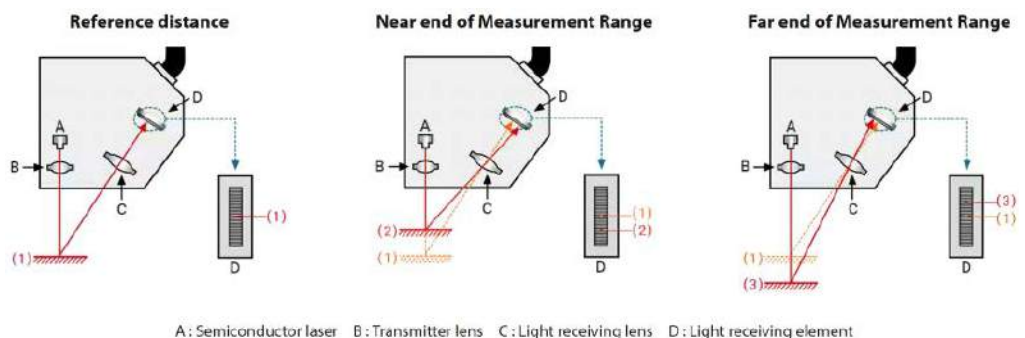
Rysunek 27. Triangulacja zastosowana w analogowych czujnikach odległości firmy Sharp (<https://t.ly/3Hwcs>)

Optyczne sensory triangulacyjne

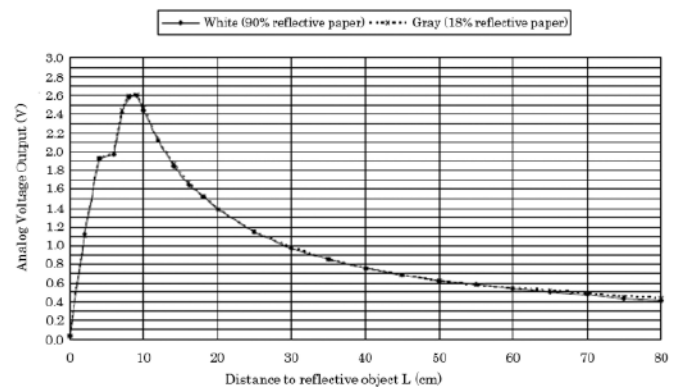
Drugą – obok ToF – grupą metod pomiaru odległości są wszelkiego rodzaju techniki bazujące na zależnościach geometrycznych. Triangulacja – bo o niej mowa – korzysta z pośredniego pomiaru kąta propagacji fal, odbitych od obiektu, a jedną z najprostszych jej



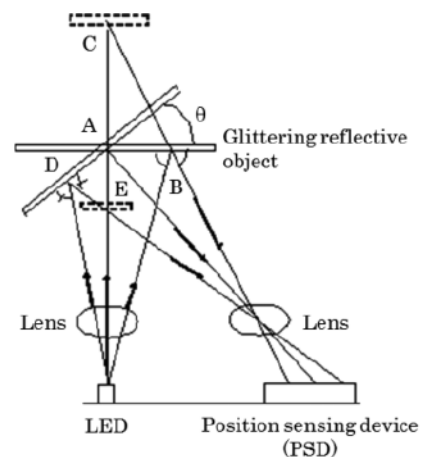
Fotografia 28. Analogowy czujnik triangulacyjny GP2Y0A02YK0F marki Sharp (<https://t.ly/z7P9h>)



Rysunek 30. Triangulacyjny czujnik przemieszczenia z cyfrowym linią pomiarowym (<https://t.ly/oRpQ3>)



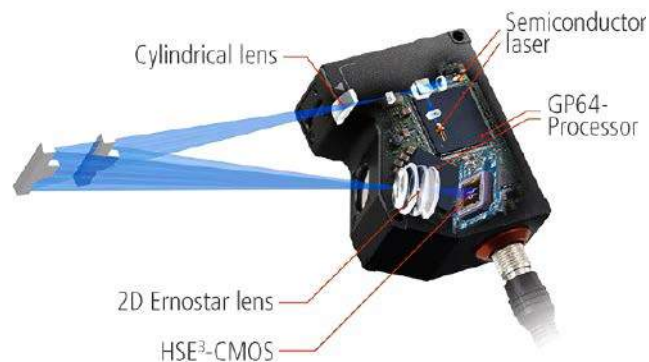
Rysunek 28. Charakterystyka wyjściowa analogowego czujnika odległości marki Sharp (<https://t.ly/8k5F4>)



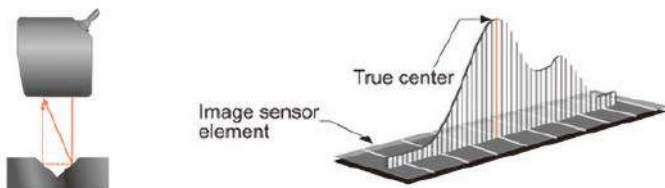
Rysunek 29. Błąd pomiaru odległości czujnikiem na bazie detektora PSD, wynikający z odbicia promieni od powierzchni, nachylonej względem osi optycznej detektora (https://t.ly/_8m2c)

realizacji są analogowe i cyfrowe czujniki odległości marki Sharp (rysunek 27, fotografia 28). Zasada ich działania bazuje na detektorach typu PSD, czyli specjalnych fotodiodach czułych na położenie plamki podczerwieni, padającej na strukturę światłoczułą – im dalej jest położony obiekt, odbijający promienie z dobrze skolimowanej diody LED IR, tym bliżej nadajnika pada punkt świetlny na powierzchnię detektora. Taka konstrukcja jest wprawdzie dość prosta pod względem elektronicznym, ale – niestety – generuje niejednoznaczność i silnie nieliniową odpowiedź (rysunek 28).

Niejednoznaczność wynika zarówno z artefaktów, powstających w bliskim zakresie odległości (poniżej pewnego progu, określającego dolny limit zasięgu danego modelu czujnika, sensor odpowiada napięciem takim samym, jak dla dalej położonych obiektów), jak i ze specyficznych warunków pracy czujnika. W szczególności jest to widoczne dla obiektów silnie odbijających światło (np. szkła) – przechylenie powierzchni względem osi optycznej sensora powoduje powstanie zafałszowanych wyników, co zobrazowano na rysunku 29.



Rysunek 31. Budowa dwuwymiarowego, laserowego czujnika przemieszczenia (https://t.ly/cCT_G)



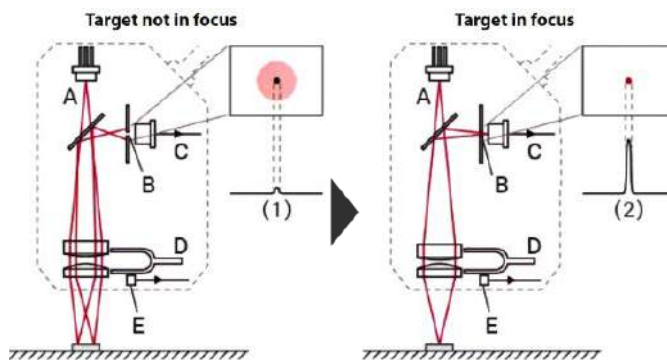
Rysunek 32. Analiza histogramu w zastosowaniu do obróbki danych obrazowych z matrycy laserowego czujnika przemieszczenia (<https://t.ly/U4h9H>)

Co ciekawe, ten sam problem może wystąpić także w znacznie dokładniejszych (i nieporównanie droższych), laserowych czujnikach przemieszczenia, stosowanych w precyzyjnej aparaturze przemysłowej (rysunek 30). Choć w tym przypadku stosowane są już nie detektory PSD, ale cyfrowe matryce jedno- lub dwuwymiarowe (rysunek 31), to efekt jest dokładnie taki sam – jeżeli czujnik współpracuje z powierzchnią rozpraszającą światło, błąd związany z kątem przechylenia nie wystąpi, ale w przypadku, gdy wykrywany obiekt ma błyszczące wykończenie, geometria układu bez trudu oszuka sensor. Na szczęście zastosowanie analizy histogramu pozwala poradzić sobie z innymi błędami, występującymi np. podczas pomiaru wąskich rowków, powodujących powstawanie dodatkowych odbić (echo w sygnale optycznym, pozyskiwanym przez matrycę – rysunek 32). Zaletą czujników triangulacyjnych jest szeroki zakres pomiarowy i duża odległość bazowa – przykładowo, niektóre sensory z serii HG-C marki Panasonic mogą mierzyć odległość w zakresie 200...600 mm z powtarzalnością na poziomie 300...800 μm .

Pomiary konfokalne i spektralne

Warto dodać, że zasygnalizowany wcześniej problem artefaktów pochodzenia geometrycznego skłonił inżynierów do opracowania innych topologii czujników laserowych. Metoda konfokalna (rysunek 33) bazuje na zastosowaniu przestrajanego układu optycznego, oscylacyjnie zmieniającego położenie soczewek obiektywu. Laser oświetla badaną powierzchnię za pośrednictwem obiektywu, przez który powraca też odbita od niej część wiązki. Za sprawą przesłony typu pinhole detektor rejestruje jednak nie pełny obraz plamki, a tylko jej centralną część. Jeżeli zatem oscylator znajdzie się w położeniu, w którym plamka zostanie idealnie zogniskowana na powierzchni detalu, detektor odbierze silny impuls światła (prawa część rysunku 33). W przeciwnym wypadku obraz rozmyje się, co spowoduje znaczący spadek jasności sygnału. Położenie soczewek w punkcie maksimum natężenia światła odbitego jest zatem pośrednią miarą odległości detalu od czujnika.

Zaletą systemów konfokalnych jest możliwość pracy ze znacznie lepszą

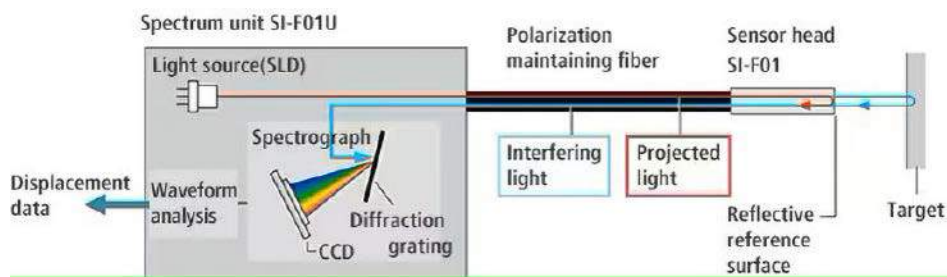


Rysunek 33. Zasada działania laserowego czujnika przemieszczenia z przemiataniem konfokalnym. A – laser, B – przesłona typu pinhole, C – fotodetektor, D – mechaniczny układ rezonansowy (oscylator), E – czujnik sprzężenia zwrotnego (<https://t.ly/Ma87s>)



Rysunek 34. Skanujący system pomiaru przemieszczeń o rozdzielczości pionowej równej 10 nm – seria LT-9000 marki Keyence (<https://t.ly/48E95>)

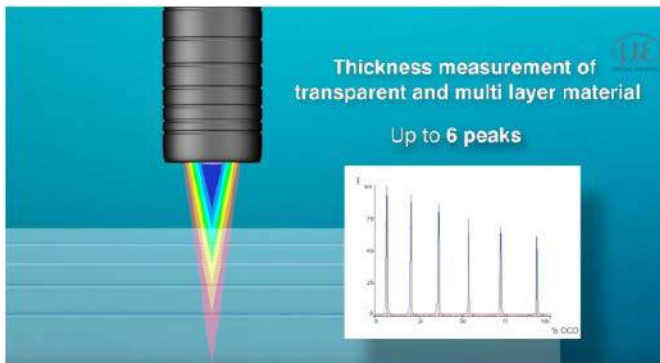
rozdzielczością i wyższą odpornością na artefakty, powstające w wyniku skomplikowanych zależności geometrycznych pomiędzy wiązką lasera a powierzchnią badanego detalu. Zastosowanie dodatkowego oscylatora w jednej z osi poziomych (X na rysunku 34) pozwala na dalsze zwiększenie możliwości aplikacyjnych w zakresie skanowania powierzchni (np. w systemach kontroli jakości, obrazowaniu profilu struktur mikroelektronicznych, itd.). Niestety, czujniki tego typu mają także znacznie bardziej ograniczony zasięg pracy, przez co nadają się one tylko do zastosowań precyzyjnych (przykładowo – modele LT-9010M oraz LT-9010 marki Keyence oferują obszar skanowania głębokości w zakresie $\pm 0,3$ mm przy odległości bazowej równej 6 mm).



Rysunek 35. Budowa spektralnego, konfokalnego czujnika przemieszczeń (<https://t.ly/pLh40>)



Rysunek 36. Zastosowanie konfokalnych czujników spektralnych pozwala na budowę kompaktowych systemów wielokanałowych, np. w aplikacjach analizy krzywizny powierzchni (<https://t.ly/rAUU8>)



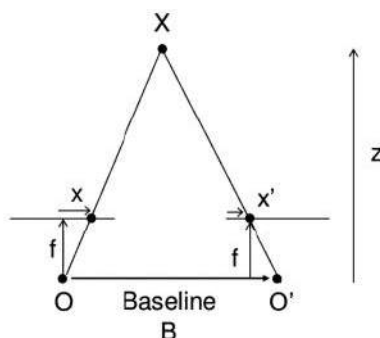
Rysunek 37. Zastosowanie systemu analizy konfokalnej marki Micro-Epsilon Messtechnik w celu pomiaru grubości warstw w materiałach wielowarstwowych (https://t.ly/6_ljg)

Inną grupą czujników konfokalnych są sensory spektralne (rysunek 35). W tym przypadku wiązka światła białego jest niejako rozszczepiana w taki sposób, że poszczególne długości fali ogniskują się na różnych głębokościach, a odczyt jest dokonywany za pomocą spektroskopu, analizującego piki widma w wiązce światła powracającego od badanego obiektu. Niezwykle istotną zaletą tego typu rozwiązań jest możliwość znacznego zmniejszenia rozmiarów głowicy pomiarowej dzięki zastosowaniu technik światłowodowych – w ten sposób integracja systemów mogą znacznie gęściej upakować czujniki w niewielkiej przestrzeni, co ma znaczenie w niektórych zaawansowanych aplikacjach przemysłowych oraz laboratoryjnych (rysunek 36). Mało tego – niektóre czujniki spektralne są w stanie wskazywać grubość kolejnych pięter materiałów wielowarstwowych (rysunek 37), używając do tego celu odbicia światła od granic kolejnych ośrodków.

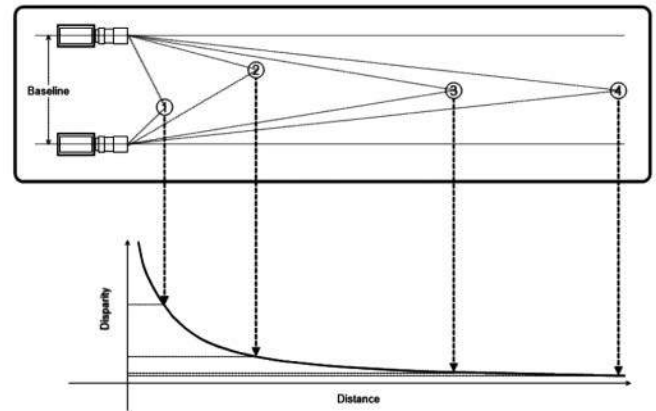
Stereowizja i skanowanie 3D ze światłem strukturalnym

Triangulacja może być z powodzeniem zrealizowana nie tylko na pojedynczych wiązkach światła, ale także na dwuwymiarowych obrazach – i w ten najprostszy sposób można chyba przejść do opisu stereowizji, będącej w istocie... bioniczną realizacją ludzkiego (i nie tylko) zmysłu wzroku.

Podstawowe założenia stereowizji można zwiualizować za pomocą trójkąta (rysunek 38), w którego dwóch wierzchołkach (O i O') znajdują się kamery, oddalone o pewną stałą odległość równą B. Obraz punktu X jest rzutowany na matryce w różnych ich punktach, a rozrzut pomiędzy współrzędnymi punktu w obu obrazach



Rysunek 38. Podstawowy układ geometryczny, rozpatrywany w zagadnieniach stereowizji (<https://t.ly/50QQS>)



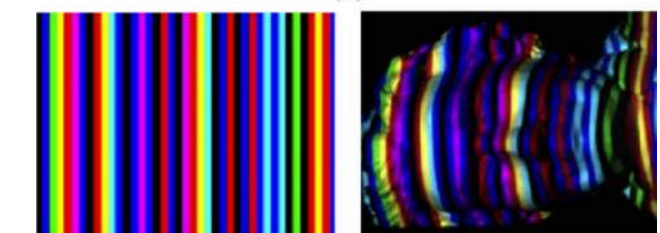
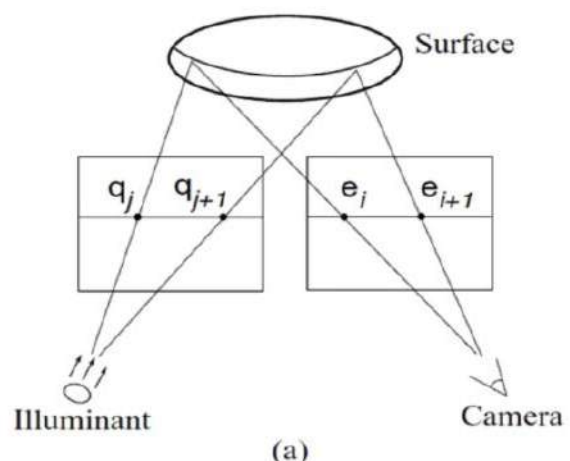
Rysunek 39. Zależność różnicy (disparity) współrzędnych odpowiadających sobie punktów z obu obrazów stereowizyjnych od rzeczywistej odległości obiektu od kamery – distance (https://t.ly/4L_5g)

(x oraz x') stanowi pośrednią miarę odległości Z (głębokości), będącej przedmiotem zainteresowania stereowizji.

Rzecz jasna, w praktyce sprawa nie jest aż tak prosta, jak mogłoby się wydawać – i to nie tylko z uwagi na fakt, iż wyznaczenie mapy odległości wymaga dość zaawansowanych metod rozpoznawania obrazu (trzeba przecież najpierw wykryć odpowiadające sobie punkty w obu widokach). Istnieje wszak jeszcze jeden szkopuł – owa różnica pomiędzy położeniem odpowiadających sobie logicznie punktów w obydwu obrazach drastycznie maleje wraz ze wzrostem rzeczywistej odległości przedmiotu od kamery (rysunek 39). Oznacza to, że w praktyce stereowizja najlepiej sprawdza się dla bliskiego otoczenia, podczas gdy obiekty w tle zlewają się w miarę oddalania od kamery – w większości praktycznych zastosowań tej techniki nie jest to jednak szczególnie istotny problem.



Fotografia 29. Moduł stereowizyjny IMX219-83 marki Waveshare (<https://t.ly/uf7Mb>)



Rysunek 40. Ilustracja tłumacząca ideę skanowania 3D z użyciem światła strukturalnego (<https://t.ly/QE8RI>)



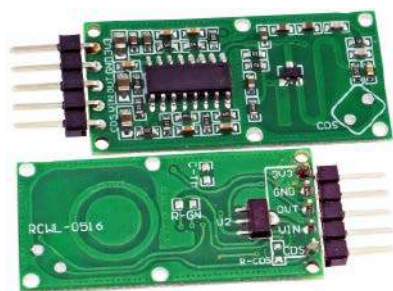
Fotografia 30. Macierz punktów świetlnych, rzutowana na twarz użytkownika przez projektor telefonu iPhone X (https://t.ly/d0_b2)

Stereowizja może być z powodzeniem zrealizowana na bazie niewielkich modułów kamer – dobrym punktem wyjścia może być platforma IMX219-83, oferowana przez markę Waveshare (**fotografia 29**) – dwa sensory obrazu Sony IMX219 o rozdzielczości 8 MPx każdy (3280×2464 px) są rozstawione na odległość 60 mm i zapewniają kompatybilność z minikomputerami jednopłytkowymi z serii Nvidia Jetson Nano oraz Raspberry Pi Compute Module.

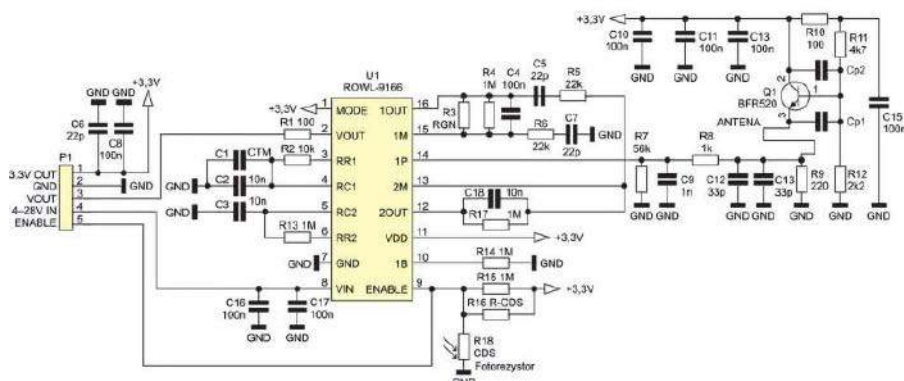
Swego rodzaju odmianą technik stereowizyjnych jest obrazowanie z użyciem światła strukturalnego. Idea metody jest zaskakująco prosta – zamiast rejestrować obraz z dwóch kamer, korzystamy tylko z jednego czujnika obrazu, podczas gdy rolę drugiego przejmuje projektor, oświetlający obiekt światłem o znanej strukturze barwnej lub geometrycznej (**rysunek 40**). Możemy więc przyjąć założenie, że obraz z drugiej, wirtualnej kamery to po prostu rzutowana przez projektor sekwencja pasków bądź punktów – reszta (obliczenia triangulacyjne)



Fotografia 31. Przykładowy chronograf balistyczny (<https://t.ly/lea75>)



Fotografia 32. Widok modułu radarowego RCWL-0516 (<https://t.ly/WVKix>)



Rysunek 41. Schemat ideowy sensora radarowego RCWL-0516 (<https://t.ly/WVKix>)

odbywa się na zasadzie analogicznej do poprzednio opisanej, klasycznej stereowizji. Systemy bazujące na świetle strukturalnym są szeroko stosowane w skanerach 3D, choć od pewnego czasu można je znaleźć nawet... w smartfonach – system zastosowany przez markę Apple oświetla twarz użytkownika tysiącami punktów podczerwieni (rzecz jasna, generowanych z użyciem laserów VCSEL), a wielkość plamek zarejestrowanych przez kamerę jest w istocie miarą odległości danego punktu od telefonu (**fotografia 30**).

Pomiary prędkości liniowej

Teoretycznie każda z metod pomiaru odległości może być też stosowana do pomiaru prędkości liniowej – wystarczy wsak dwukrotnie zmierzyć dystans, dzielący poruszający się obiekt od dalmierza, a następnie podzielić uzyskaną w ten sposób różnicę odległości przez różnicę czasu pomiędzy pomiarami. Mało tego – jeszcze prostsze rozwiązanie jest często stosowane np. w aplikacjach balistycznych, w których tzw. chronografy strzeleckie (**fotografia 31**) obliczają prędkość pocisku na podstawie czasu jego przelotu pomiędzy dwiema barierami optycznymi, umieszczonymi w stałej, znanej odległości.

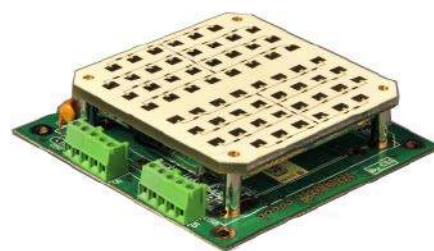
Znacznie większe możliwości i zdecydowanie większą elastyczność (w porównaniu do barier optycznych) daje zastosowanie technik dopplerowskich, wykorzystujących przesunięcie częstotliwości fali odbitej od poruszającego się obiektu. Podobnie jak w przypadku metodToF, także i tutaj te same zjawiska leżą u podwalin wielu różnych technologii – do wyboru mamy zarówno dopplery mikrofalowe, jak i ultradźwiękowe, a nawet optyczne (laserowe).

Radary dopplerowskie z falą ciągłą

Konstrukcja klasycznego radaru dopplerowskiego jest stosunkowo prosta i bazuje na emisji niemodulowanej fali ciągłej o częstotliwości nośnej (zwykle w przedziale od kilku do kilkudziesięciu gigaherców). Fala odbita od poruszającego się (w kierunku do lub od radaru) obiektu ma – za sprawą efektu Dopplera – nieco zmienioną częstotliwość. Wystarczy zatem zmieszać obydwa sygnały (nośną oraz sygnał powrotny), a wynikowy przebieg będzie zawierał informację o długości i zwrocie wektora prędkości w kierunku obiekt-radar.

Warto zauważyć, że tego typu radary dopplerowskie mogą być znacznie prostsze w konstrukcji niż radary pulse-Doppler, a nawet omówione wcześniej konstrukcje FMCW – tutaj mamy bowiem do czynienia z włączonym przez cały czas oscylatorem, którego częstotliwość nie musi być w żaden sposób kluczowana bądź przestrajana. Z tego właśnie względu topologia radaru dopplerowskiego fali ciągłej (CW) zyskała w ostatnich latach niebywałą popularność – na tej zasadzie działają m.in. wszechobecne już mikrofalowe czujniki ruchu (przykłady można zobaczyć na **fotografii 32** i **ryśunku 41**).

Wysokiej klasy moduły radarowe OEM są stosowane także w zastosowaniach drogowych, np. w celu monitorowania prędkości pojazdów – przedstawiony na **fotografii 33** moduł OEM umożliwia pomiar prędkości samochodów osobowych z odległości do 120 metrów lub pojazdów ciężarowych z 250 metrów, przy czym zakres mierzonych prędkości wynosi od 5 do 300 km/h, zaś kąt propagacji wiązki to 12×24°. Jako



Fotografia 33. Profesjonalny radar mikrofalowy OEM do zastosowań drogowych – IcomSpeed Doppler marki C&T Technology Ltd (<https://t.ly/DkLiG>)



Fotografia 34. Symulator celu ruchomego do testowania radarów dopplerowskich w paśmie 24 GHz – model K-DT1 marki RFbeam Microwave GmbH (<https://t.ly/zoey0>)

ciekawostkę warto dodać, że moduły tego typu mogą być walidowane za pomocą interesujących urządzeń, pełniących funkcję symulatorów efektu dopplerowskiego – przykład takiego testera, przystosowanego do pracy w paśmie 24 GHz, można zobaczyć na fotografii 34.

Dopplerowskie techniki ultradźwiękowe i laserowe

Co ciekawe, choć w starszej literaturze dla elektroników można znaleźć przykłady zastosowań ultradźwięków do wykrywania ruchu w pomieszczeniach, to w praktyce techniki takie są niemalże całkowicie zapomniane. Fale naddźwiękowe znalazły jednak swoje miejsce w pomiarach prędkości przepływu mediów (cieczy i gazów) w rurociągach, choć w tym przypadku stosowane są dwie techniki – dopplerowska (rysunek 42) oraz bazująca na pomiarze czasu propagacji wiązki pomiędzy nadajnikiem a odbiornikiem (czas ten jest modulowany przez kierunek oraz prędkość przepływu medium – rysunek 43).

Podobną technikę można zresztą znaleźć w konstrukcjach niektórych wiatromierzy, zawierających dwie pary



Fotografia 35. Dwuosioowy wiatromierz ultradźwiękowy (<https://is.gd/1CEVe3>)

przetworników ultradźwiękowych, ułożonych ortogonalnie względem siebie pod niewielkim zadaszeniem (fotografia 35). Rozwiązanie to ma szereg zalet w stosunku do wiatromierzy mechanicznych, a bodaj największą jest całkowicie statyczna konstrukcja oraz zerowa bezwładność, znacząco poprawiająca dynamikę pomiaru.

Dla ścisłości należy też wspomnieć o aplikacjach medycznych, w których efekt Dopplera jest używany do pomiaru i obrazowania prędkości przepływu krwi w naczyniach krwionośnych, a także do pomiaru tętna płodu w detektorach przenośnych oraz kardiokardiografach – zagadnienia te wykraczają jednak daleko poza ramy niniejszego artykułu.

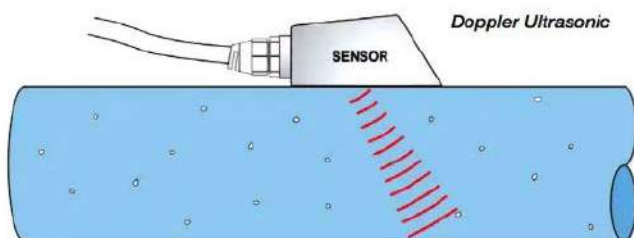
Dopplerowskie techniki laserowe także znalazły szereg zastosowań praktycznych w detekcji ruchu – choć zdecydowanie częściej wykorzystuje się je w wibrometrii laserowej, to istnieją także specjalistyczne urządzenia, pozwalające na pomiar prędkości przesuwających się (np. na taśmie produkcyjnej) obiektów. Moduły określane mianem LSV (*Laser Surface Velocimeter*) (fotografia 36) rzutują na badaną powierzchnię dwie wiązki światła laserowego, tworzące strukturę prążków interferencyjnych. Ruch powierzchni obiektu, pełniący funkcję „ekranu” dla obrazu interferencyjnego, powoduje, że intensywność światła powracającego do umieszczonego w urządzeniu detektora jest modulowana z częstotliwością proporcjonalną do prędkości przesuwu. W ten sposób można z dużą dokładnością, całkowicie bezkontaktowo, mierzyć szybkość poruszania się rozmaitych obiektów, w tym folii, blach, rur, arkuszy papieru i innych materiałów, stosowanych m.in. w produkcji przemysłowej. Mało tego – scałkowanie prędkości po czasie pozwala także na... pomiar długości tychże materiałów, a to otwiera zupełnie nowe możliwości w zakresie automatyzacji produkcji czy też kontroli jakości.

Podsumowanie

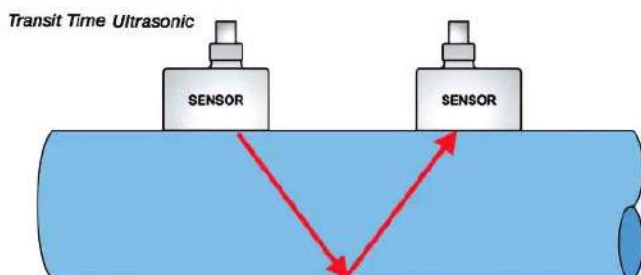
W artykule zaprezentowaliśmy szerokie spektrum zagadnień technologicznych, związanych z różnymi metodami pomiaru odległości i prędkości obiektów. Jak widać, nawet pomimo pozornie zbliżonych podstaw konstrukcyjnych poszczególnych urządzeń, wykorzystywane w nich zjawiska fizyczne są diametralnie różne, choć w niektórych przypadkach – np. w radarach typu pulse-Doppler – można zetknąć się z przemysłowym połączeniem dwóch odmiennych metod pomiarowych w ramach tego samego systemu. Światło, fale ultradźwiękowe oraz mikrofały dają niebywałą elastyczność w zakresie bezdotykowych pomiarów ruchu i odległości, sukcesywnie przekraczając kolejne granice technologiczne i podbijając nowe obszary aplikacyjne.

Oprócz opisanych technik istnieje także wiele innych – dość powiedzieć chociażby o odometrii, technologiach geolokalizacyjnych czy też sensorach inercyjnych. Umyślnie pominęliśmy również szeroki wachlarz rozwiązań przeznaczonych do pomiaru prędkości obrotowej – ich zróżnicowanie sprawia, że tematyka ta znacząco wykracza poza ramy niniejszego opracowania.

inż. Przemysław Musz, EP



Rysunek 42. Sposób zamocowania dopplerowskiego przepływomierza ultradźwiękowego na rurociągu (<https://t.ly/NFOfn>)



Rysunek 43. Przepływomierz mierzący czas przelotu wiązki ultradźwięków (<https://t.ly/8bKaB>)

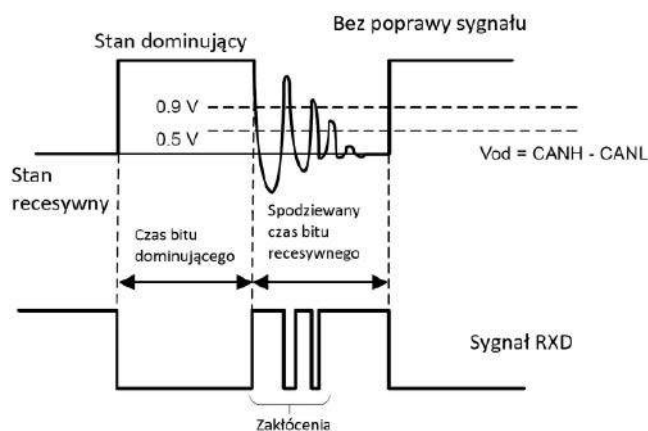


Poprawa jakości sygnału uwalnia prawdziwy potencjał transceiverów CAN-FD

Współczesne samochody mają mnóstwo funkcji poprawiających bezpieczeństwo, osiągi i komfort jazdy. Od układu napędowego po zaawansowane systemy wspomagające kierowcę. Od elektroniki nadwozia i oświetlenia po systemy informacyjno-rozrywkowe i bezpieczeństwa – w pojeździe zaszyta jest duża liczba tzw. elektronicznych jednostek sterujących (ECU). Moduły te nieustannie wymieniają pomiędzy sobą informacje za pośrednictwem magistrali sieciowych w pojeździe.

Typowymi magistralami komunikacyjnymi pomiędzy ECU w motoryzacji są obecnie Controller Area Network (CAN), Local Interconnect Network (LIN), FlexRay i Ethernet, gdzie magistrala CAN pozostaje najpopularniejszą ze względu na łatwość implementacji, odporność na zakłócenia, możliwość przesyłania wiadomości z określonym priorytetem, arbitraż bitowy do obsługi konfliktów na magistrali, wykrywanie błędów itd.

Łatwość, z jaką można skalować sieć w pojazdach poprzez dodawanie węzłów do istniejącej magistrali CAN, jest ogromną zaletą, jednak stać może się również problemem, gdy sieci te stają się złożone, na przykład połączone w topologii gwiazdy. Odbicia powodowane przez punkty w sieci bez terminacji, które są dołączane do tych sieci, mogą powodować wadliwą komunikację, szczególnie przy wyższych prędkościach. Z tego powodu transceivery standardu CAN-Flexible Data Rate (FD), choć klasyfikowane jako 5 Mb/s, muszą być używane z prędkością przesyłu danych poniżej 2 Mb/s w rzeczywistych warunkach. Zdolność do poprawy sygnału (tzw. SIC) umożliwia korzystanie



Rysunek 1. Przebieg w sieci CAN i na linii RXD w układzie bez SIC

z transceiverów CAN-FD z szybkością 5 Mb/s i wyższą w złożonych sieciach gwiazdowych bez konieczności większych zmian w projekcie.

Co to jest SIC?

SIC (Signal Improvement Capability) oznacza zdolność do poprawy sygnału. Jest to dodatkowa funkcja implementowana w transceiverach CAN-FD, która zwiększa maksymalną szybkość transmisji danych możliwą do osiągnięcia w złożonych topologiach gwiazdy poprzez minimalizację oscylacji sygnału. Transceivery CAN z SIC muszą spełniać lub przewyższać specyfikacje nakładane przez normę ISO 11898-2:2016 dla szybkiej warstwy fizycznej CAN oraz być zgodnym ze specyfikacją systemu poprawy sygnału organizacji CAN-in-Automation (CiA) numer 601-4.

Tabela 1. Porównanie specyfikacji zależności czasowych dla CiA 601-4 i ISO 11898-2

Parametr	Specyfikacja CiA 601-4		Specyfikacja ISO 11898-2:2016	
	Min. (ns)	Max. (ns)	Min. (ns)	Max. (ns)
Czas tłumienia oscylacji sygnału na TX	–	530	–	
Zmienność przesyłanej szerokości bitowej	–10	10	–65 dla 2 Mbps –45 dla 5 Mbps	30 dla 2 Mbps 10 dla 5 Mbps
Odebrana szerokość bitowa	–30	20	–100 dla 2 Mbps –80 dla 5 Mbps	50 dla 2 Mbps 20 dla 5 Mbps
Symetria czasowa odbiornika	–20	15	–65 dla 2 Mbps –45 dla 5 Mbps	40 dla 2 Mbps 15 dla 5 Mbps
Opóźnienie danych (czas propagacji) z nadajnika (TXD) do dominującej magistrali	–	80	Jedno opóźnienie dla pętli, od TXD do magistrali do RXD nieprzekraczające 255 ns	
Opóźnienie danych (czas propagacji) z nadajnika (TXD) do magistrali recesywnej	–	80		
Opóźnienie danych (czas propagacji) z magistrali do odbiornika (RXD) dominującego	–	110		
Opóźnienie danych (czas propagacji) z magistrali do odbiornika (RXD) recesywnego	–	110		

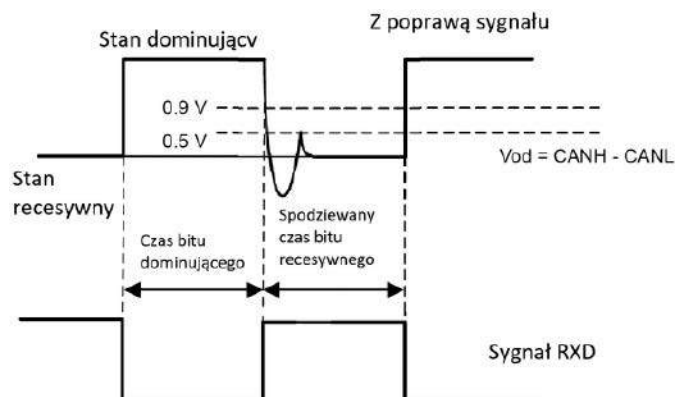
Na **rysunku 1** pokazano zwykły transceiver CAN-FD, w którym sygnał magistrali CAN oscyluje powyżej poziomu 900 mV (dominujący próg odbiornika CAN) i poniżej 500 mV (próg recesywny odbiornika CAN), co powoduje odbiór błędnych danych. Na **rysunku 2** pokazano, w jaki sposób transceiver obsługujący CAN SIC zgodne z CiA 601-4 tłumi oscylacje sygnału magistrali, dając w rezultacie prawidłowy sygnał RXD.

Jeśli chodzi o parametry elektryczne, transceiver CAN SIC zgodny z CiA 601-4 ma znacznie ściślejszą symetrię taktowania bitów i lepszą specyfikację opóźnienia pętli w porównaniu do zwykłego transceiwera CAN-FD, jak pokazano w **tabeli 1**. Dobranie opóźnień ścieżek nadawczych i odbiorczych może pomóc projektantom systemów w jasnym obliczeniu opóźnienia propagacji sieci w obecności innych elementów łańcucha sygnałowego. Należy zauważyć, że taktowanie określone w CiA 601-4 jest niezależne od szybkości transmisji danych i odnosi się zarówno do operacji z szybkością 2, jak i 5 Mb/s.

Ograniczenia klasycznej sieci CAN i CAN-FD

Protokół CAN pierwszej generacji, ISO 11898-2, znany również jako Classical CAN, został wprowadzony około 1993 roku. Protokół zezwalał na przesyłanie tylko 8 bajtów danych użytkowych z maksymalną szybkością transmisji danych określoną na 1 Mb/s. Ograniczenia te stały się szybko problemem w zastosowaniach motoryzacyjnych, gdzie pojazdy mają szereg elektronicznych węzłów komunikujących się ze sobą za pomocą magistrali CAN.

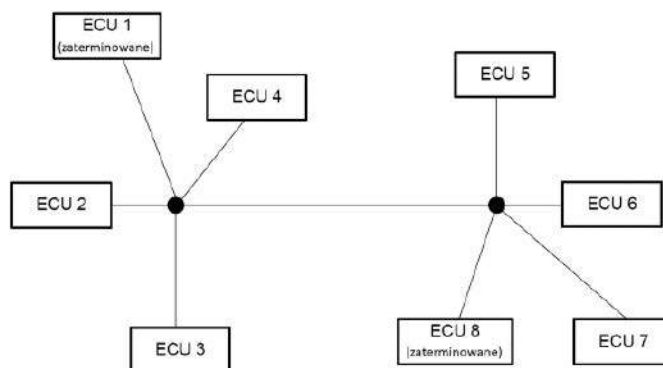
Około 2015 roku opublikowano specyfikację protokołu CAN-FD, która zwiększyła długość pakietu danych do 64 bajtów i maksymalną szybkość komunikacji w fazie danych do 5 Mb/s. Szybkość sygnalizacji w fazie arbitrażu była jednak nadal ograniczona do 1 Mb/s w celu zapewnienia wstecznej kompatybilności z Classical CAN. Podczas gdy CAN-FD zapewniał szybszą transmisję danych i większą pojemność sieci, to i tak nie był to postęp wystarczający, aby nadążyć za stale rosnącą liczbą modułów elektronicznych dodawanych do samochodowych sieci magistrali CAN. Projektanci zdali sobie sprawę, że nie mogą wykorzystać prawdziwego potencjału nadajników-odbiorników CAN-FD,



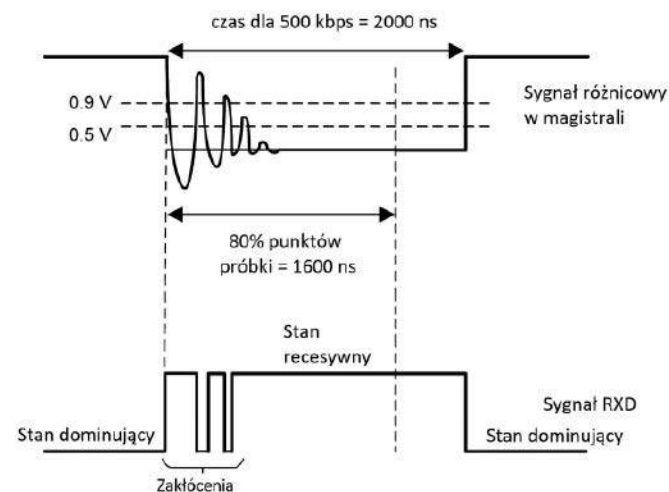
Rysunek 2. Przebieg w sieci CAN i na linii RXD w układzie z SIC

ponieważ oscylacje magistrali wynikające z powstania złożonych sieci o topologii gwiazdy wpływa na prawidłową komunikację sygnału.

Na **rysunku 3** pokazano przykładową sieć w topologii gwiazdy. W takich topologiach z wieloma odgażeniami, sygnał podróżujący po magistrali przechodzi wieloma ścieżkami o niedopasowanej impedancji, co może powodować m.in. odbicia sygnału. Zakłócenia



Rysunek 3. Węzły CAN połączone w sieć o topologii gwiazdy



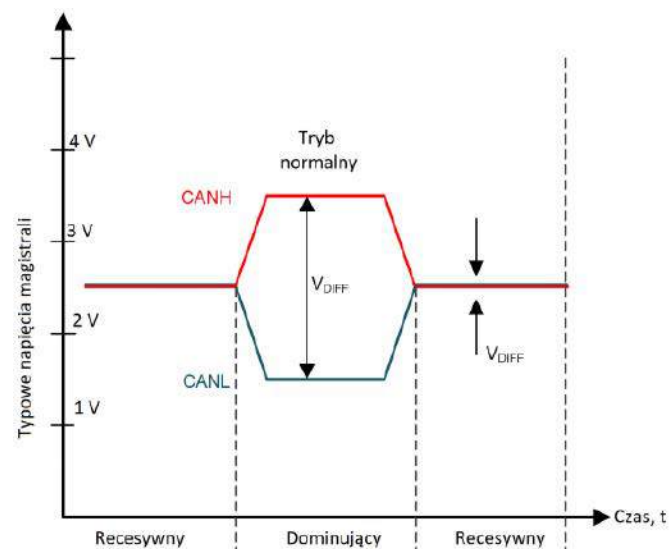
Rysunek 4. Oscylacje magistrali CAN i usterka RXD dla klasycznych prędkości CAN

te zniekształcają transmisję na magistrali CAN i powodują oscylacje, co skutkuje występowaniem nieprawidłowego poziomu na magistrali CAN i RXD w momencie próbkowania. Choć te efekty nie są specyficzne dla sieci CAN-FD, przy pracy z niższą szybkością w klasycznej implementacji CAN czas trwania bitu jest dłuższy, a oscylacje na magistrali na tyle mniejsze, że możliwe jest próbkowanie właściwego bitu, jak pokazano na **rysunku 4**, co skutkuje poprawną komunikacją, mimo występowania zakłóceń.

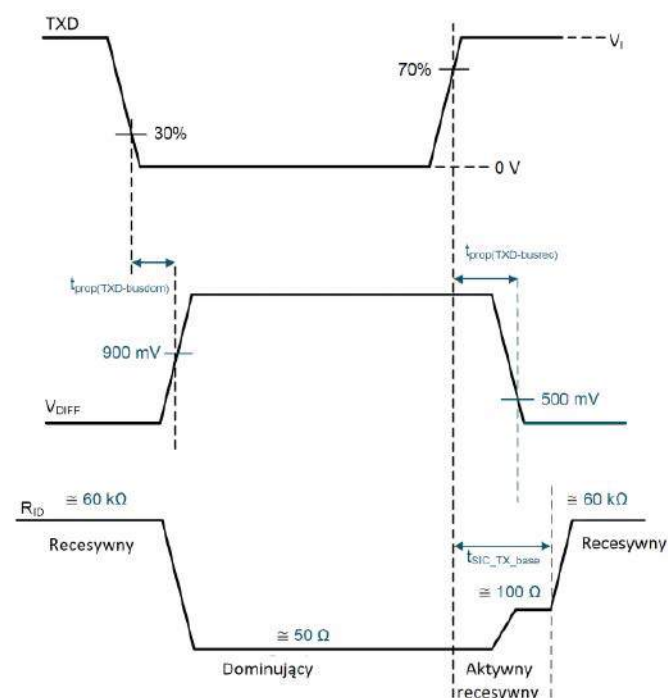
W przypadku działania CAN-FD z szybkością 5 Mb/s czas trwania bitu wynoszący 200 ns jest o wiele za krótki, aby zaniknęły oscylacje w złożonych topologiach gwiazdy, utrudniając niezawodną komunikację danych. To zniechęca projektantów systemów do korzystania z CAN-FD o pełnej prędkości 5 Mb/s. Wraz ze wzrostem ilości wymienianych danych w sieciach we współczesnych pojazdach i zwiększającym się zapotrzebowaniem na przepustowość, technologia CAN SIC toruje drogę dla nowej generacji magistrali komunikacyjnych w pojazdach. Mają one być szybsze i zapewniać większą elastyczność i skalowalność sieci w pojeździe.

Jak CAN SIC zmniejsza oscylacje w magistrali

Podczas normalnej pracy magistrala CAN ma dwa stany logiczne: recesywny i dominujący, jak pokazano na **rysunku 5**. Stan dominujący magistrali występuje podczas sterowania różnicowego magistrali i odpowiada niskiemu poziomowi logicznemu na pinach TXD lub RXD. Stan recesywny magistrali występuje, gdy magistrala jest spolaryzowana do połowy napięcia zasilania przez wewnętrzne rezystory wejściowe



Rysunek 5. Poziomy napięcia w magistrali CAN



Rysunek 6. Technologia CAN SIC: sekwencja zdarzeń

(RIN) odbiornika i odpowiada stanowi wysokiemu na pinach TXD i RXD. Stan dominujący zastępuje stan recesywny podczas arbitrażu. Zbocze sygnału od recesywnego do dominującego na szynie CAN jest zwykle czyste, ponieważ jest silnie sterowane przez nadajnik. Różnicowa impedancja wyjściowa transceivera CAN podczas fazy dominującej wynosi około 50 Ω i ściśle odpowiada impedancji charakterystycznej sieci. W przypadku zwykłego transceivera CAN-FD zbocze – przejście ze stanu dominującego do recesywnego – ma miejsce, gdy wyjściowa impedancja różnicowa sterownika nagle osiąga około 60 kΩ, a odbity sygnał natrafia na niedopasowanie impedancji, co powoduje oscylacje.

Układ SIC zaimplementowany w nadajniku wykrywa przejście ze stanu dominującego do recesywnego na linii TXD i aktywuje obwód tłumienia oscylacji na wyjściu drivera linii. Sterownik CAN kontynuuje silne recesywne sterowanie magistralą aż do osiągnięcia czasu $t_{SIC_TX_base}$, czyli czasu tłumienia oscylacji sygnału na TX (patrz tabela 1), tak że odbicie zmniejsza się, a bit recesywny jest czysty w punkcie próbkowania. W aktywnej fazie recesywnej impedancja wyjściowa nadajnika jest niska (wynosi około 100 Ω). Ponieważ odbity sygnał nie wykazuje dużego niedopasowania impedancji, oscylacje są znacznie wytłumione. Po zakończeniu tej fazy i wejściu urządzenia w pasywną fazę recesywną impedancja wyjściowa sterownika wzrasta do około 60 kΩ. **Rysunek 6** obrazuje całą sekwencję zdarzeń przy użyciu technologii CAN SIC.

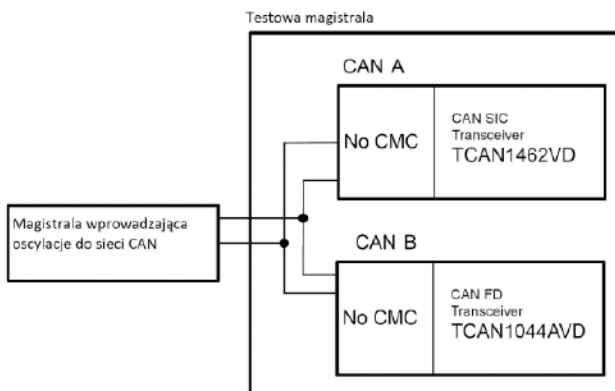
Ważnym czynnikiem w aktywnej fazie recesywnej jest to, że powinna ona trwać maksymalnie do 530 ns (jak podano w tabeli 1). Faza danych protokołu CAN-FD trwa maksymalnie do 200 ns (jeśli interfejs działa z szybkością 5 Mb/s), więc tłumienie oscylacji będzie aktywne przez cały czas trwania stanu recesywnego, co skutkuje występowaniem tylko prawidłowych sygnałów magistrali CAN i RXD. Jednak w fazie arbitrażu – gdzie najkrótszy czas trwania bitu wynosi 1 μs dla działania z prędkością 1 Mb/s, wiele nadajników może nadawać jednocześnie, a bit dominujący musi nadpisać bit recesywny – czas trwania włączenia układu tłumienia oscylacji może nakładać pewne ograniczenia na ogólną rozległość sieci i szybkość działania systemu arbitrażu. Więcej informacji na ten temat można znaleźć w specyfikacji CiA 601-4.

Wyniki eksperymentów na urządzeniu TCAN1462 firmy TI

Aby zaprezentować funkcję tłumienia oscylacji w transceiverze CAN TCAN1462, jego producent, firma Texas Instruments (TI), przeprowadził eksperyment z układem w następującej konfiguracji.

Tabela 2. TCAN1462 w porównaniu z konkurencyjnymi układami

Parametr	Konkurencyjne układy dostępne na rynku	TCAN1462	Znaczenie dla systemu
Napięcie logiki	od 3 V do 5,5 V	od 1,71 V do 5,5 V	układ może współpracować z systemami cyfrowymi 1,8 V
Zależności czasowe SIC	zgodne tylko dla $\pm 5\%$ VCC	w zakresie $\pm 10\%$ VCC	układ TI nie wymaga precyzyjnie stabilizowanego napięcia zasilania
Minimalne napięcie różnicowe CAN na poziomie 1,5 V	zgodne tylko dla $\pm 5\%$ VCC	w zakresie $\pm 10\%$ VCC	
Zakres zabezpieczeń sieci	od -36 V do 40 V	± 58 V	wyższe napięcie oznacza większą odporność na uszkodzenia
Odporność na ESD	6 kV	± 8 kV	wyższa ochrona przed ESD



Rysunek 7. Sieć z dwoma węzłami i obwodem oscylującym

Dwupunktowy układ komunikacji, gdzie węzeł 1 to TCAN1462, a węzeł 2 to TCAN1044 A, zwykły transceiver CAN-FD, bez funkcjonalności CAN SIC, pokazano na **rysunku 7**. Sieć oscylująca (znormalizowana przez CiA 601-4) emulująca złożoną topologię gwiazdy jest podłączona przez piny do magistrali CAN. Jak pokazują przebiegi na **rysunkach 8** oraz **9**, magistrala CAN i sygnały RXD wyglądają na czyste, gdy pracuje układ TCAN1462. Ale kiedy siecią steruje TCAN1044 A, w przebiegach magistrali widoczne są oscylacje i problemy z liniami RXD.

Schodzące bardzo nisko napięcie VOD nie stanowi problemu i nie występuje przeregulowanie na tej magistrali, co skutkuje czystymi liniami RXD.

Urządzenia CAN SIC firmy TI

Firma Texas Instruments wypuściła dwa urządzenia CAN SIC: ośmiopinowy TCAN1462 z obsługą trybu czuwania, który jest kompatybilny pin w pin z tradycyjnymi ośmiopinowymi transceiverami CAN oraz 14-pinowy układ TCAN1463, który ma tryb uśpienia i funkcję WAKE/INH i jest kompatybilny pin do pinu z klasycznymi układami transceiverów w obudowach z 14 pinami. Ta kompatybilność wsteczna umożliwia prostą wymianę układu w gotowych konstrukcjach, aby móc zwiększyć ich prędkość transmisji, dzięki użyciu CAN

SIC. Oprócz kompatybilnych wstecznie obudów, nowy układ jest dostarczany w miniaturowych obudowach z rodziny SOT23.

TCAN1462 jest dostępny w dwóch wariantach: TCAN1462 dla logiki 5 V oraz TCAN1462 V z obsługą poziomów logicznych od 1,8 V do 5 V. Urządzenia te mają znaczne zalety w porównaniu z konkurencyjnymi urządzeniami na rynku, jak pokazano w **tabeli 2**.

Podsumowanie

Transceivery wyposażone w technologię CAN SIC zapewniają znaczne korzyści systemowe w porównaniu ze zwykłymi transceiverami CAN-FD bez potrzeby zmian konstrukcyjnych na warstwie fizycznej czy w aplikacji. Układy te pozwalają na działanie z większą szybkością transmisji bitów, z większą swobodą w wyborze topologii sieci, przy jednoczesnym obniżeniu kosztów i wagi systemu komunikacyjnego.

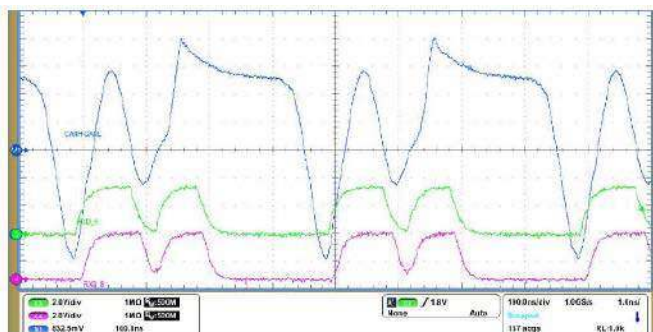
CAN SIC jest wstecznie kompatybilny z ISO 11898-2, więc może działać na tej samej magistrali, co CAN-FD. Omówione transceivery są wstecznie kompatybilne z innymi, standardowymi, dzięki czemu możliwe jest wprowadzenie zmian w istniejącym projekcie.

Jak pokazano w tabeli 1, transceivery CAN z SIC znacznie poprawiają symetrię taktowania bitów, co zapewnia większy margines dla efektów w sieci, które mogą pogarszać jakość sygnałów w magistrali CAN. Transceiver wprowadza znacznie mniejszą degradację przesyłanych i odbieranych danych, skracając czas trwania bitu, co pozwala mu działać niezawodnie nawet przy prędkości 8 Mb/s. Finalnie, opóźnienie w pętli dla transceiverów CAN SIC wynosi maksymalnie do 190 ns, w porównaniu z maksymalnym opóźnieniem równym 255 ns dla transceiverów CAN-FD. Dzięki temu możliwe jest zwiększenie maksymalnej długości sieci.

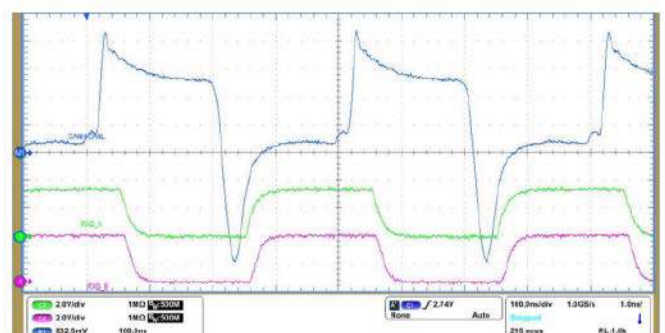
Nikodem Czechowski, EP

Bibliografia:

1. Vikas Kumar Thawani, „How Signal Improvement Capability Unlocks the Real Potential of CAN-FD Transceivers”, Texas Instruments Technical White Paper SLLA581, kwiecień 2022
2. Tony Adamson, „CiA 601-4: CAN signal improvement”, CAN Newsletter, kwiecień 2019

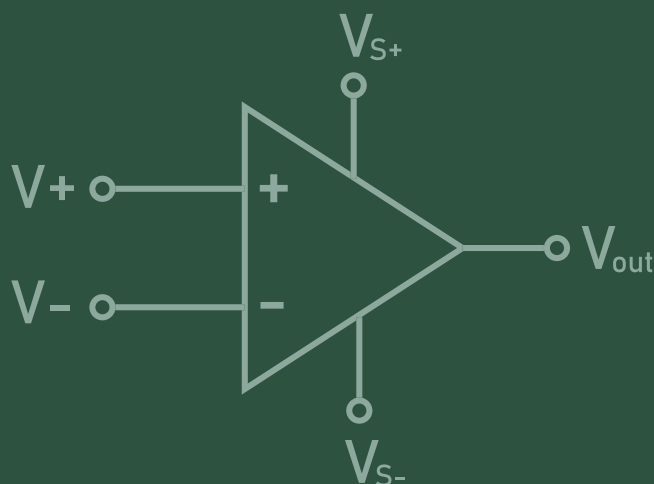


Rysunek 8. Przebiegi ze zwykłym układem CAN-FD sterującym magistralą



Rysunek 9. Przebiegi z układem z technologią CAN SIC sterującym magistralą

Operational Amplifier



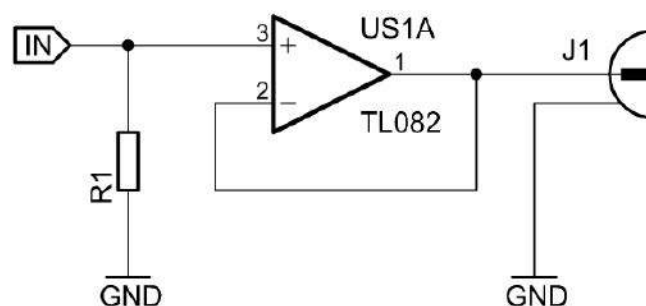
Wzmacniaczu operacyjny, nie wzbudzaj się!

W środowisku elektroników znane jest stwierdzenie, że szybkie wzmacniacze operacyjne często sprawiają problemy przez ich łatwe wzbudzenie. Równie często powtarzana jest sentencja, że te „wolne” układy działają bardziej stabilnie. O ile z pierwszą myślą ciężko jest się nie zgodzić, o tyle druga może być kontrowersyjna. W artykule opiszę moje doświadczenia w tej dziedzinie.

Jakie problemy może sprawiać układ typu TL082, pełniący funkcję wtórnika napięciowego na wyjściu RCA? Przecież wystarczy odpowiednio odsprzęgnąć jego zasilanie, poprowadzić ścieżki w miarę sensownie i zadbać o rozsądne wartości elementów. Włączamy zasilanie, przykładamy sondę oscyloskopu, wszystko działa – stabilny układ jest już gotowy, następny proszę! Schemat tego prostego rozwiązania pokazano na **rysunku 1**. Książkowy, bez dwóch zdań. Podłączam krótkim kablem wyjście tego układu z wejściem wzmacniacza. Jest pięknie, zniekształcenia są pomijalnie niskie, układ może jechać na testy do klienta.

Układ nie działa

Po kilku dniach odzywa się zbulwersowany klient, że dostał ode mnie generator ryku, skrzeku i pisku, a nie przedwzmacniacz. Przecież wszystko sprawdziłem, wzmacniacz operacyjny na wyjściu nie robi żadnych sztuczek – zresztą, nie ma nawet jak. Trudno, proszę odebrać. Oczywiście, jak to zazwyczaj w takich sytuacjach bywa, u mnie



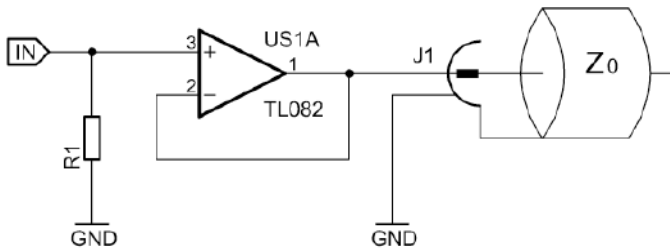
Rysunek 1. Schemat ideowy omawianego stopnia wyjściowego

działa. Jednak jak się okazuje – wszystko działa, dopóki obciążeniem tego układu jest kilka centymetrów cienkiej ścieżki i sonda oscyloskopu. Nawet krótki kabelek nie sprawia problemów.

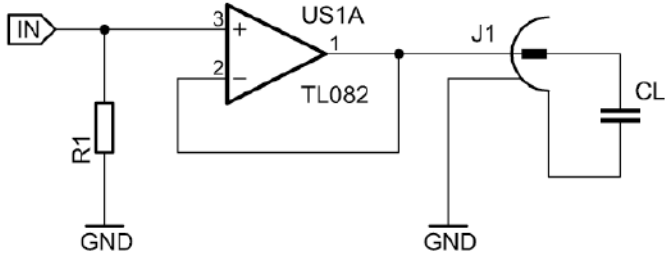
Konsultacje z klientem wykazały, że jego kabel sygnałowy nie ma półmetrowej długości, ale prawie 5 m! Zrobiłem podobne podłączenie u siebie i nie zawiodłem się – na wyjściu pokazał się przebieg sinusoidalno-trapezoidalny o bliżej nieokreślonej częstotliwości, na dodatek z jakąś modulacją amplitudową!

Czym jest kabel sygnałowy?

Teoretycznie – odcinkiem linii długiej o określonej impedancji charakterystycznej, który został rozarty na końcu (**rysunek 2**), bo impedancja wejściowa urządzeń audio z reguły sięga wielu kilohmów. Aż się prosi o wielokrotne odbicia od końca tej linii. Nie lubię myśleć „falowo”, czy zatem temat da się ugryźć od strony teorii obwodów?



Rysunek 2. „Falowe” ujęcie problemu

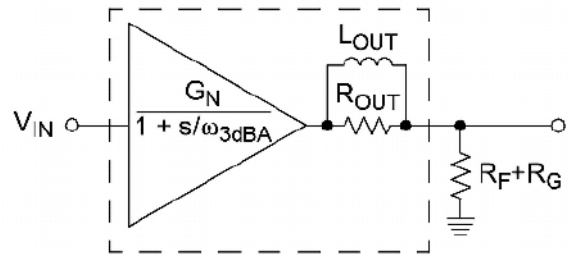


Rysunek 3. Zobrazowanie zgodnie z teorią obwodów

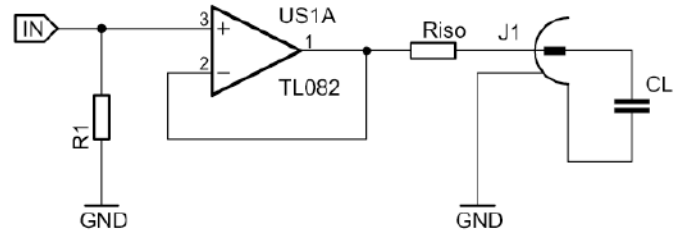
Owszem – przecież taki kabel to niemal idealny kondensator o pojemności kilku nanofaradów (**rysunek 3**). Z kolei wyjście wzmacniacza operacyjnego dla wysokich częstotliwości zaczyna zachowywać się jak indukcyjność z powodu przesunięcia fazowego, jakie wprowadza zewnętrzna pętla ujemnego sprzężenia zwrotnego. Po prostu sygnał zwrotny zaczyna się wydatnie opóźniać względem wejściowego, co tłumaczy wzrost impedancji wyjściowej, lecz można to również symulować jako pojawienie się indukcyjności na wyjściu (**rysunek 4**). A stąd już tylko krok do oscylacji tak powstałego obwodu LC.

Zapobieganie

W jaki sposób można do tego nie dopuścić? Najprościej, odpowiednio „psując” dobroć takiego obwodu. Jeżeli będzie ona niewystarczająca do powstania i podtrzymania oscylacji, to układ pozostanie stabilny. Microchip proponuje odpowiednie wzory, które umożliwiają dokładne obliczenie wartości takiego rezystora. Rodzi to jednak kilka problemów, zaś głównym z nich jest brak wiedzy o pojemności podłączonego kabla. W jednym przypadku kabel będzie krótki, w innym długi, jeszcze inny może być bardzo cienki (co zwiększa indukcyjność) i co wtedy? Zbyt wysoka wartość takiego rezystora izolującego (**rysunek 5**) zawęzi pasmo układu, wpłynie negatywnie na parametry



Rysunek 4. Zachowanie układu dla wysokich częstotliwości



Rysunek 5. Praktyczne i tanie rozwiązanie problemu

slew-rate oraz wprowadzi tłumienie przy niskiej impedancji wejściowej następnego stopnia. Zaś zbyt niska nie zda egzaminu.

Prowadziłem swoje doświadczenia i po kilku latach mogę stwierdzić, że uzyskanie na wyjściu układu stanu dopasowania (lub chociaż stanu zbliżonego do dopasowania) jest wystarczające do tego, by uniknąć opisanych na początku problemów. Typowe impedancje kabli stosowanych do prowadzenia sygnału analogowego audio wynoszą 50 Ω lub 75 Ω. Na wyjściu wzmacniacza operacyjnego stosuję rezystor izolujący o wartości 33 Ω i nie spotkałem jeszcze układu, który miałby ochotę na wzbudzenie się. Analogicznie robię z wyjściami układów sterujących kable do transmisji różnicowej, jak choćby w popularnym standardzie XLR – dwa rezystory 33 Ω włączone w szereg z każdym z wyjść wzmacniaczy operacyjnych. Takie połączenie nie zapewnia dopasowania impedancji, ale nie wprowadza też znaczącego tłumienia i co najważniejsze – spełnia swoją funkcję.

Michał Kurzela, EP

Bibliografia:

1. <https://tiny.pl/cc9m3>
2. <https://tiny.pl/cc9gh>

REKLAMA

Świat projektantów i programistów
dla elektroniki w nowej odsłonie.
Odwiedź wiecznie młody

ELPORTAL.pl

Solarna ładowarka akumulatorów z MPPT (2)

Prezentujemy drugą część artykułu o ładowarce, która pozwala ładować energią z ogniw PV akumulatory o napięciu do 50 V. Akumulatory te mogą być podstawą domowego magazynu energii, dzięki któremu efektywniej wykorzystamy energię słoneczną. Dodatkowo zaprezentowany projekt wyróżnia się tym, że realizuje algorytm śledzenia punktu mocy maksymalnej (MPPT), co gwarantuje uzyskiwanie maksymalnej ilości energii z ogniw PV.

Wybór mikrokontrolera i innych elementów

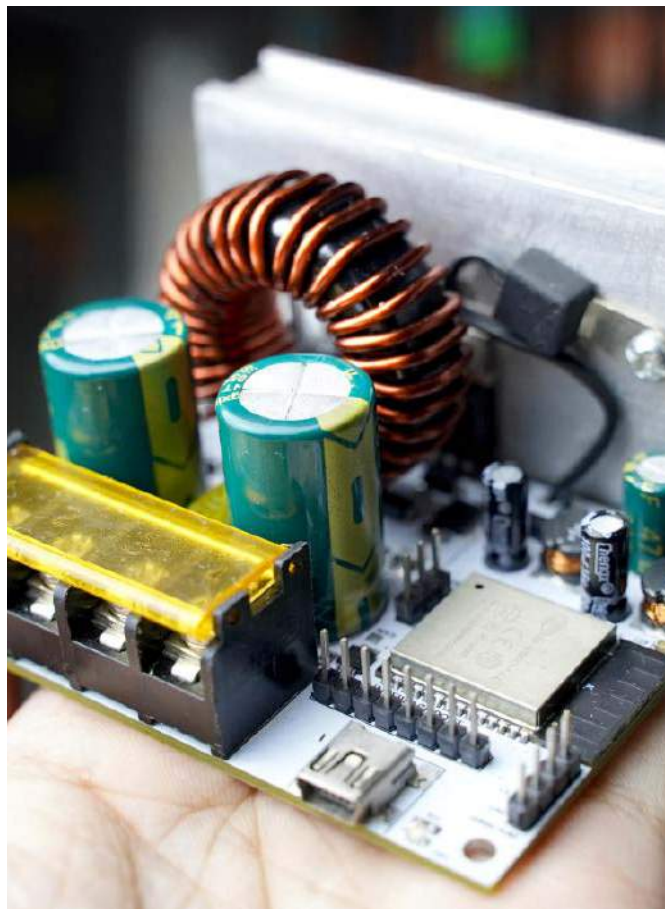
Układ ESP32 należy do dobrze znanej rodziny mikrokontrolerów. Istnieje szereg modułów z tymi komponentami, w tym moduł WROOM32 zawierający wszystko co niezbędne do działania układu. Moduł ten jest dostępny na wielu płytkach deweloperskich. Znajduje się na nich na ogół port USB do programowania układu. Sam moduł WROOM32 nie ma takiego portu USB. Aby można było go zaprogramować przez USB, potrzebny jest zewnętrzny interfejs USB-TTL UART, taki jak CH340C. Moduł WROOM32 kosztuje tylko około 1,80 dolara. Jest kompatybilny programowo z Arduino i można go używać jak każdego innego modułu Arduino.

Z drugiej strony ESP32 to bardzo szybki i wydajny 32-bitowy, dwurdzeniowy mikrokontroler taktowany zegarem 240 MHz. Ma wbudowane interfejsy Wi-Fi i Bluetooth BLE, więc nie trzeba kupować dodatkowych, drogich modułów. Układ ten oferuje maksymalnie 16-bitową rozdzielczość PWM, ma 12-bitowy wbudowany przetwornik analogowo-cyfrowy – ADC (nie będzie on używany w opisanym urządzeniu, zostanie zastosowany dokładniejszy, zewnętrzny układ) oraz przetwornik cyfrowo-analogowy – DAC. To tylko część z możliwości tych mikrokontrolerów, a cała lista jest długa. Dodatkowo, przy cenie 1,80 dolara, układ ten był tańszą alternatywą względem np. Arduino Nano, pomimo że oferuje możliwości, które przewyższają je pod każdym względem.

Najważniejsza w tym projekcie jest 16-bitowa rozdzielczość PWM. Wyższa rozdzielczość PWM jest znacznie bardziej preferowana w MPPT bazującym na przetwornicy Buck. Im wyższa rozdzielczość, tym dokładniejsze będą kroki zmiany napięcia i prądu. Z kolei, im niższa rozdzielczość PWM jest używana, tym wyższa częstotliwość PWM może zostać osiągnięta:

- dla 16-bitowej rozdzielczości PWM można osiągnąć maksymalnie 1,22 kHz PWM,
- dla 12-bitowej rozdzielczości PWM – 19,5 kHz PWM,
- dla 11-bitowej rozdzielczości PWM – 39,06 kHz PWM,
- dla 10-bitowej rozdzielczości PWM – 78,12 kHz PWM,
- dla 9-bitowej rozdzielczości PWM – 156,25 kHz PWM,
- dla 8-bitowej rozdzielczości PWM – 312,5 kHz PWM.

Im wyższa rozdzielczość PWM, tym dokładniejsze będą kroki sterowania napięcia i prądu na wyjściu stabilizatora Buck MPPT. Z kolei, im wyższa częstotliwość PWM, tym wyższa może być moc układu MPPT i tym mniejsze są napięcia tętnienia na wyjściu. Elektronika to gra kompromisów. Przetwornica może dostarczyć większą moc przy wyższej częstotliwości PWM, ale straty przełączania również będą wtedy, do pewnego momentu, rosły. Istnieje również granica między rozdzielczością i częstotliwością PWM a zegarem układu:



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

$$f_{PWM,MAX} = \frac{MCU_{clock}}{PWM_{resolution}}$$

gdzie:

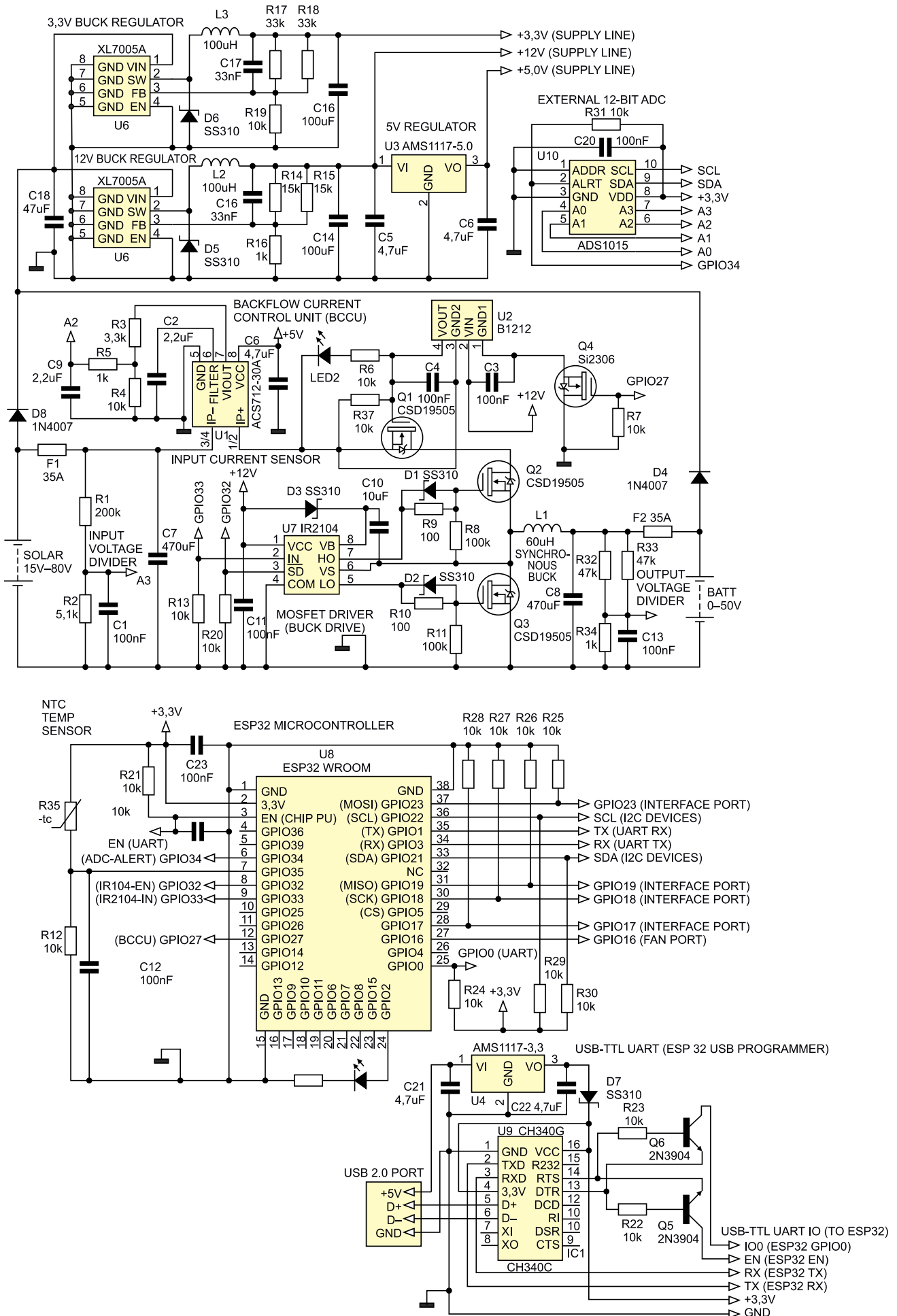
- f to częstotliwość,
- MCU_{clock} to zegar mikrokontrolera – dla tego układu wynosi 80 MHz,
- $PWM_{resolution}$ to rozdzielczość PWM, rozumiana jako liczba poziomów wypełnienia.

Dla n -bitowej rozdzielczości licznika $PWM_{resolution} = 2^n$. Dla rozdzielczości PWM równej 11 bitów maksymalna częstotliwość PWM wynosi:

$$\frac{80000000Hz}{2^{11}} = \frac{80000000Hz}{2048} = 39062,5Hz$$

Oznacza to, że dla 11-bitowego PWM możemy wybrać częstotliwość PWM wynoszącą 39,0625 kHz lub mniej. Po wypróbowaniu różnych konfiguracji PWM ostatecznie autor wybrał właśnie 11-bitową rozdzielczość PWM z częstotliwością 39 kHz jako domyślne ustawienia dla projektu i oprogramowania układowego MPPT. Mniejsze rozdzielczości nie zapewniały stabilnej pracy przetwornicy.

Oprogramowanie ESP32 w Arduino jest również proste, z uwagi na dostępność różnych bibliotek. Wystarczy, że wprowadzimy



Rysunek 7. Schemat modułu przetwornicy z śledzeniem punktu maksymalnej mocy i ładowarką akumulatorów

wymaganą rozdzielczość bitową PWM i częstotliwość, a resztę wykona oprogramowanie. Co więcej, w przypadku ESP32 można użyć niemal wszystkich wyprowadzeń jako wyjścia PWM.

Moduł ESP32 ma wbudowany 12-bitowy przetwornik ADC (4096 różnych wartości reprezentujących wartość analogową). To czterokrotnie więcej niż w przypadku Arduino Uno czy Nano. W przypadku obwodu z wyjściem 80 V, 30 A przekłada się to na rozdzielczość pomiaru napięcia równą 19,5 mV i rozdzielczość pomiaru prądu równą 7,32 mA. Jest to znacznie wyższa rozdzielczość pomiaru niż w Arduino. Przetwornik ten jest również szybszy niż w Arduino, co oznacza, że można częściej mierzyć wartości analogowe, a dzięki szybszemu procesorowi cały układ będzie bardziej responsywny. W rzeczywistości jednak nie jest tak prosto.

Niestety ESP32 znany jest z ADC o niezbyt dobrych parametrach – ma wysoce nieliniową odpowiedź, co oznacza, że ADC nie jest tak dokładny, jak można sądzić z wcześniejszych danych. Aby rozwiązać ten problem, autor zdecydował się na zewnętrzny, precyzyjny ADC. Do projektu wybrany został układ ADS1115 lub zamiennie ADS1015 firmy Texas Instruments. Przetworniki ADC z rodziny ADS1×15 mają wbudowane wewnętrzne źródło napięcia odniesienia. Oznacza to, że odczyty wartości analogowych są niezależne od VCC. Dzięki temu jest on mniej podatny na zakłócenia z linii zasilania. Układ ten ma też wbudowany programowalny wzmacniacz, a więc można programowo wybrać różne wzmocnienia. Co najistotniejsze, układ ADS1115 ma aż 16-bitową rozdzielczość przy pomiarze z próbkowaniem o częstotliwości 860 Hz. To 64 razy lepiej niż ADC w Arduino Uno. W przypadku konfiguracji MPPT z wyjściem 80 V, 30 A daje to rozdzielczość pomiaru napięcia na poziomie 1,22 mV i rozdzielczość pomiaru prądu na poziomie 0,457 mA.

Układ ADS1115 można w wielu aplikacjach zastąpić ADS1015. Pomimo nieco innych specyfikacji oba układy są kompatybilne z biblioteką Arduino ADS1×15 firmy Adafruit. ADS1115 ma 16-bitową rozdzielczość i częstotliwość próbkowania równą 860 Hz, a ADS1015 12-bitową rozdzielczość i częstotliwość próbkowania 3,3 kHz. Autor przetestował je oba i do testów i większości aplikacji rekomenduje tańszy i słabszy ADS1015, szczególnie jeśli nie chcemy korzystać z najwyższych poziomów napięć i prądów.

Pomimo takiej samej rozdzielczości, jak w modułach Arduino Nano czy ESP32, zewnętrzny ADC jest o wiele stabilniejszy i oferuje dużo wyższą efektywną dokładność pomiaru. Jest to istotne, ponieważ używając wyższej rozdzielczości PWM, konieczne jest teraz mierzenie coraz drobniejszych zmian napięć czy prądów, aby algorytm zaburzeń w MPPT działał bezbłędnie. Niewielki szum czy oscylacje pomiarów w odczytach czujników mają duży wpływ na śledzenie punktu maksymalnej mocy, w związku z czym zewnętrzny ADC daje o wiele lepsze parametry.

Na **rysunku 7** pokazano finalny schemat modułu przetwornicy MPPT. Autor zaprojektował PCB dla tego modułu (dokumentacja dostępna na stronie projektu oraz w jego repozytorium na GitHubie). Moduł po zmontowaniu pokazano na fotografii tytułowej. Pamiętajmy o zainstalowaniu wszystkich elementów mocy na radiatorze, aby zapewnić im chłodzenie. Jeśli zostaną one zainstalowane na jednym radiatorze, jak pokazano na zdjęciu, należy odizolować je elektrycznie od radiatora, inaczej może dojść do zwarcia – metalowa część obudowy TO-220 zazwyczaj połączona jest ze środkową nóżką elementu. Aby odizolować tranzystory od radiatora, stosuje się na ogół podkładki – silikonowe lub wykonane z miki – oraz izolujące podkładki pod śruby.

Oprogramowanie

Projekt MPPT bazuje na środowisku Arduino. Płytkę tę można traktować jak każdą inną płytkę rozwojową Arduino z ESP32. Kod oprogramowania układowego napisany został tak, aby był w miarę uniwersalny i nadawał się do stosowania również w przyszłych konstrukcjach MPPT. Program nazwano Fugu. Na projekt ten składają się tysiące

linii kodu w 9 różnych szkicach Arduino. Oprogramowanie to pobrać można z repozytorium autora na GitHubie (link na końcu artykułu). Nie trzeba w nim prawie nic zmieniać, aby móc od razu uruchomić przetwornicę. Jedyne, co trzeba uzupełnić, to identyfikator SSID i hasło naszej sieci Wi-Fi oraz token uwierzytelniający Blynk, jeśli chcemy, aby telemetria działała w aplikacji na smartfonie. Wszystko inne można skonfigurować w urządzeniu za pomocą interfejsu z LCD.

W dalszej części opisano, krok po kroku, jak przygotować i załadować oprogramowanie do pamięci ESP32.

1. Instalacja Arduino IDE. Wraz z nim należy zainstalować wymagane biblioteki wsparcia dla ESP32 itp.
2. Pobranie źródła Fugu MPPT z GitHuba. Składa się on z kilku plików `.ino`, które trzeba umieścić w jednym folderze. Folder musi nazywać się tak, jak główny plik `.ino` (na przykład `ARDUINO_MPPT_FIRMWARE_V1.1`). Bez tego Arduino IDE nie otworzy poprawnie źródła. Jeśli po kliknięciu któregoś z plików otworzy się Arduino IDE, w którym widoczne są wszystkie zakładki, oznacza to, że udało się poprawnie wczytać źródło. Dzięki podziałowi kodu na karty możliwe jest łatwiejsze uporządkowanie dużego kodu i podział go na osobne, wygodne w zarządzaniu sekcje.
3. Doinstalowanie bibliotek do Arduino IDE. Program korzysta z następujących bibliotek (poza tymi, które domyślnie instalowane są wraz z Arduino):
 - Blynk ESP32,
 - Liquid Crystal I²C LCD (Autor: Robojax),
 - Adafruit ADS1×15.
4. Dodanie poświadczeń dla systemu IoT do kodu. Jeśli system ma korzystać z aplikacji telemetrycznych w aplikacji na telefon, trzeba wykonać następujące kroki:
 - trzeba wprowadzić dane uwierzytelniające sieci Wi-Fi, aby ESP32 mógł połączyć się z siecią Wi-Fi:


```
ssid[] = " ";
pass[] = " ";
```
 - aplikacja Blynk na telefon ma funkcję prywatności i bezpieczeństwa. Blynk (Legacy) wysyła unikalny token uwierzytelniający do wszystkich użytkowników podczas rejestracji. Ten unikalny token zapewnia, że niepowołani użytkownicy nie mają dostępu do projektu. Trzeba skopiować token uwierzytelniania wysłany przez Blynk na e-maila:


```
auth[] = " ";
```
5. Wybór odpowiedniej biblioteki ADC dla ADS1015 lub ADS1115. W projekcie domyślnie stosowany jest przetwornik ADS1015. Jeśli taki stosujemy w naszym projekcie, nie trzeba nic zmieniać w kodzie. Jeśli zastosowano ADS1115, należy wykonać następujące kroki:
 - zakomentować wiersz `Adafruit_ADS1015 ads`,
 - odkomentować wiersz `Adafruit_ADS1115 ads`.
6. Wprowadzenie cen elektryczności. Aplikacja Blynk MPPT może śledzić, ile pieniędzy zaoszczędzono na zbieraniu energii. Aby ustawić walutę oszczędności energii, trzeba określić lokalny koszt elektryczności (np. 0,5 USD/kWh). Wartość ta znajduje się w głównym szkicu kodu, w sekcji `USER PARAMETERS`, jako zmienna `electricalPrice`.

Po wprowadzeniu wszystkich opisanych powyżej wymaganych zmian można przystąpić do kompilacji i załadowania firmware do mikrokontrolera. Aby to zrobić, należy wybrać odpowiednie ustawienia Arduino IDE (nazwy parametrów odpowiadają Arduino IDE 1.8.19 w polskiej wersji językowej):

- Płytki: ESP32 Dev Module,
- Upload Speed: 921600,
- CPU Frequency: 240 MHz (Wi-Fi/BT),
- Flash Frequency: 80 MHz,
- Flash Mode: QIO,
- Flash Size: 4 MB (32 Mb),

- Partition Scheme: Default 4 MB with spiffs (1,2 MB APP/1,5 MB SPIFFS),
- Core Debug Level: brak,
- PSRAM: Disabled,
- Port: Wybieramy port COM, pod jakim zgłasza się płytka po podłączeniu do komputera PC.

Po skonfigurowaniu Arduino IDE należy kliknąć ikonkę kompilacji, a następnie, gdy program skompiluje się bez błędów, przesłać program do Arduino ESP32 MPPT. Po pomyślnym przesłaniu wszystkich szkiców do mikrokontrolera wyświetlony zostanie komunikat o zakończeniu przesyłania programu.

Kalibracja

Kalibracja sensora prądu nie jest wymagana, ponieważ oprogramowanie układowe Fugu MPPT jest wyposażone w automatyczną kalibrację tego sensora. Kalibrację czujnika temperatury można przeprowadzić, zmieniając wartość zmiennej `ntcResistance = 9000,00`. Jest to wartość znamionowa rezystancji dla rezystora NTC – w projekcie znajduje się opornik NTC o rezystancji 9 kΩ, dlatego też parametr ten równa się 9000,00. Jeśli zastosujemy inny opornik, musimy zmienić tę wartość.

Należy wykonać kalibrację pomiaru napięcia. W tym celu wykonujemy następujące kroki:

1. Wchodzimy do menu ustawień i pozostawiamy je otwarte dla każdej wersji próbnej. Wejście do menu ustawień zatrzymuje wszystkie procesy ładowania. Trzeba to robić za każdym razem, gdy przesyłamy kod do MPPT, ponieważ MPPT resetuje się z powrotem do menu głównego za każdym razem, gdy przesyłany jest kod;
2. Podłączamy wejście panelu słonecznego MPPT do zasilacza ustawionego na 60 V;
3. Podłączamy wyjście baterii MPPT do zestawu baterii;
4. Podłączamy MPPT do portu USB;
5. Otwieramy Arduino IDE na komputerze;
6. Otwieramy monitor portu szeregowego;
7. Znajdujemy parametr VI w danych;
8. Mierzmy napięcie na wejściu za pomocą woltomierza;
9. Jeśli port szeregowy podaje napięcie wyższe lub niższe napięcie od pomiaru woltomierza, należy przeprowadzić kalibrację czujnika napięcia wyjściowego;
10. Znajdujemy parametr `inVoltageDivRatio = 40,2156`; w sekcji parametrów kalibracji głównego kodu. Zwiększamy lub zmniejszamy go, w zależności od odczytu. Przesyłamy zmodyfikowany kod do modułu;
11. Mierzmy ponownie napięcie wejściowe za pomocą woltomierza i porównujemy z parametrem VI z portu szeregowego. Powtarzamy kroki 7..10, aż woltomierz zmierzy takie samo napięcie jak VI raportowane na porcie szeregowym;
12. Analogicznie robimy z napięciem wyjściowym, którego dotyczy parametr `outVoltageDivRatio = 24,5000`;

Umieszczone w kodzie `inVoltageDivRatio` i `outVoltageDivRatio` to wstępnie obliczone wartości współczynnika dzielnika napięcia dla wejściowych i wyjściowych dzielników napięcia zastosowanych w schemacie. Zapisywanie wartości rezystorów indywidualnie mogłoby być bardziej zrozumiałe dla użytkownika, ale każdorazowe obliczanie współczynnika wymaga mocy obliczeniowej. W celu skrócenia czasu przetwarzania zapisano parametry jako już prze-liczone, zgodnie ze wzorem (jako odwrotność typowego przedstawienia dzielnika).

Listing 1. Fragmenty głównego szkicu programu

```
void setup() {
  // Inicjalizacja portu szeregowego
  // Konfiguracja pinów GPIO
  // Inicjalizacja generatora PWM
  // Ustawienie parametrów PWM
  ledcSetup(pwmChannel, pwmFrequency, pwmResolution);
  // Konfiguracja pinu wyjściowego jako PWM
  ledcAttachPin(buck_IN, pwmChannel);
  // Ustawienie wypełnienia PWM
  ledcWrite(pwmChannel, PWM);
  // Obliczenie maksymalnego wypełnienia
  pwmMax = pow(2, pwmResolution)-1;
  pwmMaxLimited = (PWM_MaxDC*pwmMax)/100.000;
  // Inicjalizacja ADC
  // Inicjalizacja GPIO
  // Uruchomienie multitaskingu na dwóch rdzeniach
  xTaskCreatePinnedToCore(coreTwo, „coreTwo”, 10000, NULL, 0, &Core2, 0);
  // Inicjalizacja i załadowanie danych z pamięci Flash
  // Inicjalizacja LCD
  // Komunikat powitalny
  Serial.println(„> MPPT HAS INITIALIZED”);
}

void loop() {
  //Zakładka #2 – Pomiar danych z sensorów i ich przeliczenia
  Read_Sensors();
  //Zakładka #3 – Algorytm wykrywania błędów
  Device_Protection();
  //Zakładka #4 – Proces systemowy
  System_Processes();
  //Zakładka #5 – Algorytm ładowania akumulatorów
  Charging_Algorithm();
  //Zakładka #6 – Wbudowana telemetria USB / UART
  Onboard_Telemetry();
  //Zakładka #8 – Algorytm wyświetlania danych na LCD
  LCD_Menu();
}
```

Listing 2. Funkcja `coreTwo`, która uruchamia subsystemy na rdzeniu 0

```
void coreTwo(void * pvParameters){
  setupWiFi(); // Konfiguracja Wi-Fi
  while(1){
    Wireless_Telemetry(); // Funkcja telemetryczna
  }
}
```

Listing 3. Fragment kodu opisującego algorytm sterowania MPPT i ładowarką

```
if(MPPT_Mode==0){
  if(currentOutput>currentCharging) {PWM--;}
  else if(voltageOutput>voltageBatteryMax){PWM--;}
  else if(voltageOutput<voltageBatteryMax){PWM++;}
  else{}
  PWM_Modulation();
}
else{
  if(currentOutput>currentCharging){PWM--;}
  else if(voltageOutput>voltageBatteryMax){PWM--;}
  else{
    if(powerInput>powerInputPrev && voltageInput>voltageInputPrev) {PWM--;}
    else if(powerInput>powerInputPrev && voltageInput<voltageInputPrev){PWM++;}
    else if(powerInput<powerInputPrev && voltageInput>voltageInputPrev){PWM++;}
    else if(powerInput<powerInputPrev && voltageInput<voltageInputPrev){PWM--;}
    else if(voltageOutput<voltageBatteryMax) {PWM++;}
    powerInputPrev = powerInput;
    voltageInputPrev = voltageInput;
  }
  PWM_Modulation();
}
```

Zasada działania

Jak łatwo się domyślić, oprogramowanie układowe jest zbyt obszerne, aby zaprezentować je tutaj w całości. Ogólna architektura programu warta jest w szkicu `ARDUINO_MPPT_FIRMWARE_V1.1.ino`, którego fragment pokazano na **listingu 1**. Program ten ma typową dla Arduino sekcję `void setup()` oraz `void loop()`. Program korzysta z obu rdzeni układu ESP32. Rdzeń 0, który normalnie odpowiada tylko za obsługę stosu Wi-Fi, wykorzystany został również do konfiguracji sieci bezprzewodowej oraz obsługi subsystemu telemetrii. Zastosowana do tego jest funkcja `coreTwo`, która uruchamia wymienione subsystemy – **listing 2**.

Kluczowa część programu realizowana jest w funkcji `loop`, jak pokazano na **listingu 1**. Funkcja ta, działając w pętli, uruchamia po kolei funkcje odpowiedzialne za odczyt danych z sensorów, wykrywanie błędów, sterowanie algorytmem ładowania itd. Poszczególne funkcje są definiowane w różnych modułach oprogramowania – osobnych szkicach.

Opis algorytmów układu

W pliku `4_Charging_Algorithm.ino` zawarto kluczową część systemu – dwa algorytmy sterujące systemem. Pierwszy to algorytm poszukujący punktu maksymalnej mocy ogniwa PV i utrzymujący

Tabela 1. Algorytm zmian wypełnienia PWM w funkcji zmieniającej się mocy i napięcia mierzonego na ogniwie fotowoltaicznym

Moc (P)	Napięcie (V)	Wypełnienie PWM
↑	↑	↓
↑	↓	↑
↓	↑	↑
↓	↓	↓

prąd i napięcie wyjściowe w zadanych granicach (są one ograniczone jedynie od góry, z wiadomych względów). Drugi algorytm jest prostszy – jest to tryb zasilacza, który stara się utrzymywać napięcie i prąd na wyjściu na zadanym poziomie, ale nie uwzględnia MPPT. Przeznaczony jest do pracy jako np. ładowarka ogniw zasilana z zasilacza sieciowego itp. To, jaki tryb jest aktywny, wybierane jest w menu konfiguracyjnym urządzenia, a na poziomie kodu decyduje o tym zmienna `MPPT_Mode`. Gdy jest ona równa 0, urządzenie działa jak prosta ładowarka, a gdy jest równa 1 – to jako ładowarka z MPPT.

Kod algorytmu pokazano na **listingu 3**. Instrukcje te znajdują się wewnątrz funkcji `Charging_Algorithm()`, która uruchamiana jest w głównej pętli programu (listing 1). Gdy `MPPT_Mode == 0`, aktywowana jest pierwsza część algorytmu, która sprawdza tylko, czy napięcie i prąd wyjściowe mają odpowiednią wartość. Gdy któreś z nich jest powyżej zadanego poziomu, układ zmniejsza współczynnik wypełnienia sygnału sterującego przetwornicą Buck, który przechowywany jest w zmiennej PWM. Jeśli natomiast napięcie wyjściowe jest poniżej zadanego napięcia ładowania, to PWM jest zwiększane. PWM jest zmieniany o jedną jednostkę, więc algorytm ten działa płynnie i bez żadnych nagłych skoków. Po zmianie parametru PWM uruchamiana jest funkcja `PWM_Modulation()`, która zmienia wypełnienie sygnału generowanego na pinie GPIO, sterującym przetwornicą.

W odmiennym przypadku aktywowana jest druga, bardziej złożona część algorytmu. W pierwszej kolejności program sprawdza, czy prąd i napięcie ładowania nie przekracza ustawionych wartości maksymalnych i jeśli tak jest, zmniejsza PWM o jedną jednostkę. Następnie uruchamiany jest algorytm MPPT, który poszukuje punktu maksymalnej mocy. Algorytm sprawdza wyjściową moc i napięcie ogniwa fotowoltaicznego i porównuje z tymi samymi parametrami w poprzednim cyklu programu. Algorytm zmian wypełnienia pokazany jest w **tabeli 1**.

Następnie algorytm sprawdza, czy napięcie na wyjściu przetwornicy nie przekracza maksymalnego napięcia ładowania ogniwa – jeśli tak jest, redukowany jest PWM.

Po wprowadzeniu tych zmian program zapisuje wartości napięcia i mocy na ogniwie, które będą używane w kolejnym cyklu działania pętli, a następnie uruchamia funkcję `PWM_Modulation()`, aby zapisać nową konfigurację do modułu generującego PWM na wyjściu mikrokontrolera.

Telemetria przez USB

System jest wyposażony w kanał telemetryczny na porcie USB, emulującym port UART. Funkcji tej można używać do diagnostyki, rozwiązywania problemów, poprawiania projektu, optymalizacji kodu lub eksperymentowania z tworzeniem własnego algorytmu MPPT. Dostęp do tych danych można uzyskać za pośrednictwem Arduino IDE, z pomocą monitora szeregowego lub dowolnego innego terminala dla portu szeregowego.

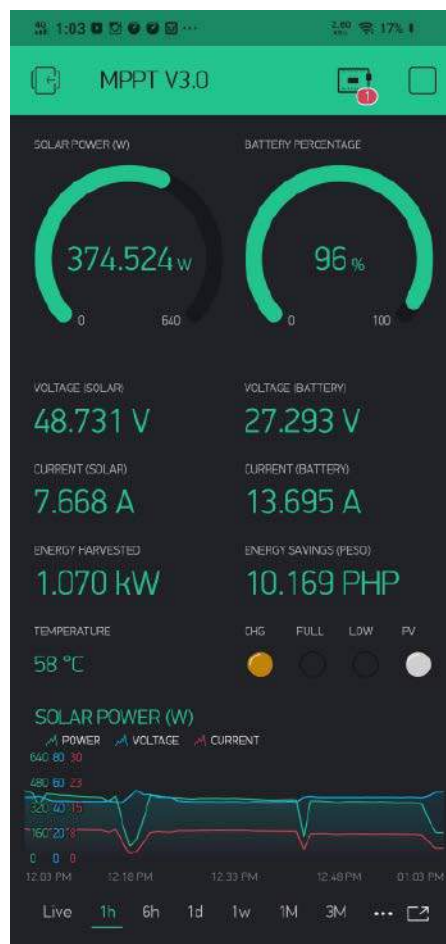
Zmienne w kanale telemetrycznym USB wysyłane są w postaci tekstu z następującymi oznaczeniami:

- ERR – liczba obecnych błędów,
- FLV – zbyt niskie napięcie systemu (nie można wznowić pracy),
- BNC – wskaźnik braku podłączenia akumulatora,
- IUW – wskaźnik zbyt niskiego napięcia wejściowego,
- IOV – wskaźnik zbyt dużego napięcia wejściowego,



Rysunek 8. Kod QR potrzebny do załadowania aplikacji MPPT w Blynk Legacy

- OOV – wskaźnik zbyt dużego napięcia wyjściowego,
- OOC – wskaźnik przetężenia wyjścia,
- OTE – wskaźnik przegrzania,
- REC – Powrót do sprawności po błędzie,
- MPPTA – wskaźnik włączenia algorytmu MPPT,
- CM – tryb wyjścia (1 = tryb ładowarki, 0 = tryb zasilacza),
- BYP – wskaźnik MOSFET jednostki sterującej prądem wstecznym,
- EN – wskaźnik włączenia Buck (1 = MPPT Buck działa, 0 = nie działa),
- FAN – wskaźnik pracy wentylatora,
- Wi-Fi – wskaźnik włączenia Wi-Fi,
- PI – moc wejściowa (moc pobierana z energii słonecznej w watach),
- PWM – PWM podawany do drivera IR2104 (wartość dziesiętna),
- PPWM – dopuszczalny limit poziomu PWM,
- VI – napięcie wejściowe (napięcie panelu słonecznego w woltach),
- VO – napięcie wyjściowe (napięcie akumulatora w woltach),



Rysunek 9. Ekran aplikacji telemetrycznej dla modułu MPPT w Blynk Legacy

- CI – prąd wejściowy (prąd panelu słonecznego w amperach),
- CO – prąd wyjściowy (prąd ładowania akumulatora w amperach),
- Wh – pobrana energia (w watogodzinach),
- Temp – temperatura radiatora (w stopniach Celsjusza),
- CSMPV – kalibrowany punkt zerowy czujnika prądu,
- CSV – wyjście analogowe (napięcie) z sensora prądu (chwilowe),
- VO%Dev – procentowe odchylenie VO od ustawionego maksymalnego wyjściowego napięcia akumulatora (w %),
- SOC – stan naładowania akumulatora (w %),
- LoopT – czas pętli, czas potrzebny na wykonanie jednego cyklu pętli kodu (w milisekundach).

Dla zmiennych typu bool, takich jak FLV na EN, 1 oznacza prawdę, a 0 fałsz.

Aplikacja telemetryczna

Moduł zapewnia również telemetrię przez Wi-Fi, jak opisano wcześniej. Zawiera do tego darmową aplikację do rejestrowania danych bazującą na serwerze. Aplikacja na telefon MPPT powstała na platformie Blynk Legacy. Jest dostępna zarówno na Androida, jak i iOS. Aby zainstalować i uruchomić aplikację, należy po kolei:

- pobrać aplikację ze sklepu i zarejestrować się jako nowy użytkownik,
- zeskanować kod QR z **rysunku 8** w aplikacji Blynk Legacy,
- aplikacja dla MPPT zostanie automatycznie załadowana,
- kliknięcie przycisku w aplikacji pozwoli uruchomić wyświetlanie danych,
- można dodać skrót do menu głównego, aby móc bezpośrednio otwierać aplikację MPPT bez konieczności ciągłego przeglądania menu.

Funkcje aplikacji MPPT Blynk:

- wyświetla na żywo transmisję energii zebranej z paneli słonecznych,
- wyświetla moc panelu słonecznego (waty),
- wyświetla napięcia i prądy panelu słonecznego,
- wyświetla poziom naładowania baterii w %,
- wyświetla napięcie i prąd ładowania akumulatora,

- wyświetla pobraną energię w kWh,
 - wyświetla wartość zaoszczędzonej energii,
 - wyświetla temperaturę,
 - prezentuje trzy różne wykresy w czasie rzeczywistym z wyświetlanych danych,
 - dane są rejestrowane w bezpłatnej bazie danych Blynk przez rok.
- Ekran aplikacji pokazano na **rysunku 9**.

Podsumowanie

Zaprezentowana konstrukcja może stać się podstawą wielu systemów, między innymi domowych magazynów energii itp. Autor podaje, że odnotował współczynnik sprawności konwersji na poziomie 98,6% przy mocy 270 W, napięciu wejściowym 61,4 V i napięciu wyjściowym 27,00 V. Nawet jeżeli sprawność ta jest nieco zawyżona z uwagi na błędy pomiaru, to tranzystory MOSFET, nawet bez aktywnego chłodzenia ledwo się nagrzewały.

Układ, jakkolwiek kompletny i gotowy do działania, pozostawia pewne pole do zmian. Autor wskazuje na potencjalne modyfikacje, jakie można wprowadzić do układu:

- wymiana tranzystorów na takie, które pozwalają na pracę z napięciem do 150 V DC. Oprócz tych elementów trzeba przeliczyć dzielniki napięcia, zmienić oprogramowanie itd.;
- zwiększenie napięcia akumulatorów do 80 V, co wymaga zmian w oprogramowaniu i modyfikacji dzielnika pomiarowego,
- zmiana wyświetlacza z LCD na AMOLED,
- zmiana indukcyjności na większą i zamontowanie na zewnątrz obudowy,
- połączenie równolegle czterech tranzystorów MOSFET w celu zmniejszenia rezystancji – pozwala to zwiększyć prąd przetwornicy.

Nikodem Czechowski, EP

Źródła:

1. <https://shorturl.at/ijmNZ>
2. <https://shorturl.at/iuyR8>

REKLAMA

Nie przegap wrześniowego wydania „Elektroniki dla Wszystkich”

przejrzyj i kupisz na
www.ulubionykiosk.pl





Kinetyczne wzory na piasku, sterowane przez mikrokontroler ESP32

Stałym elementem ogrodów Zen, inspirowanych azjatycką filozofią i japońskim designem ogrodów, są wzory na piasku. W szczególności w ogrodach Karesansui (suchego krajobrazu) starannie zgrabiony żwir tworzyć miał wrażenie fal, otaczających skały w ogrodzie. Taka forma ogrodnictwa miała sprzyjać wyciszeniu i medytacji. Pokazana konstrukcja w kreatywny sposób rozwija tę sztukę – linie już nie muszą być tylko równoległe i rysowane ręcznie, np. grabiami. Kinetyczne wzory tworzy robot, który samodzielnie rysuje motyw na piasku. Może być doskonałą ozdobą każdego wnętrza lub po prostu ciekawym projektem elektromechanicznym.

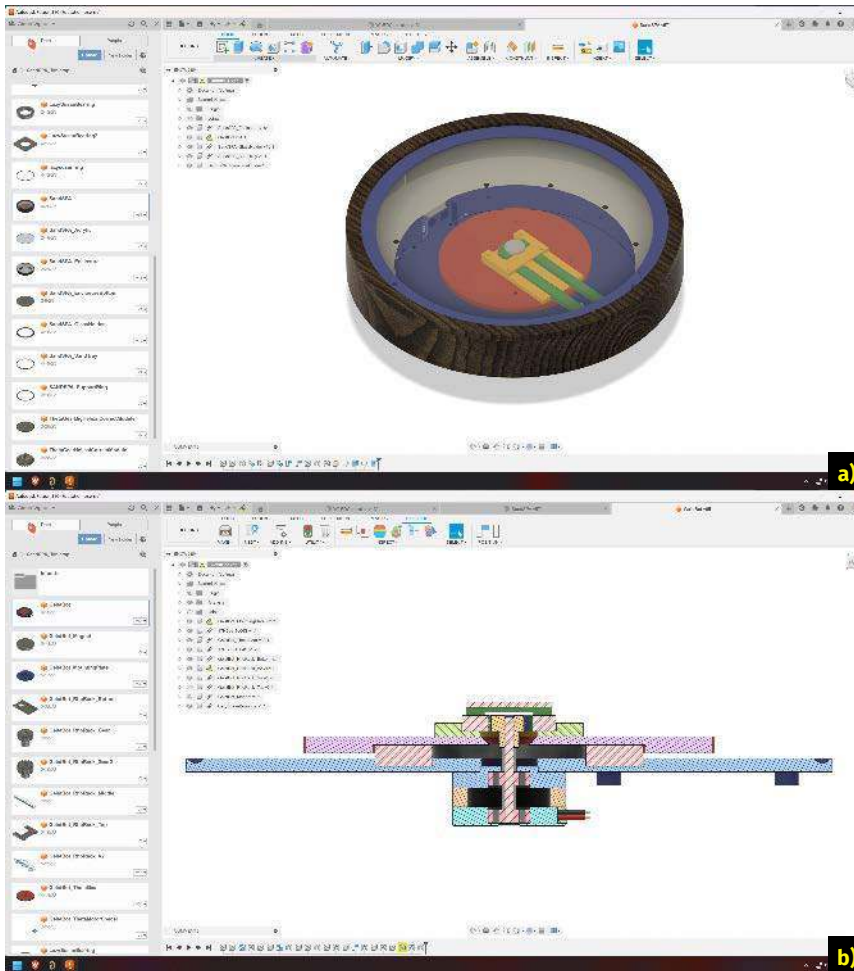
Początki tego projektu sięgają kilku lat wstecz – autor konstrukcji zbudował już w 2020 roku podobne urządzenie, jednak było ono bardzo proste i pozbawione wielu możliwości, jakie ma najnowsze opracowanie. Obecna konstrukcja jest kompaktową i unowocześnioną wersją urządzenia skonstruowanego w 2020 roku, które samo było zainspirowane reklamowanym w 2019 r. urządzeniem SANDSARA sfinansowanym w ramach kampanii crowdfundingowej w 2019 roku.

Przeprojektowanie trzyletniej konstrukcji zajęło autorowi sześć miesięcy. Po zakończeniu konstrukcji, którą dokumentował na swoim blogu, przygotował kompletny opis konstrukcji, który zawiera kompletny zestaw informacji dotyczących budowy tego układu.

Potrzebne elementy

Do zbudowania urządzenia potrzebna będzie drukarka 3D do wydrukowania szeregu elementów mechanicznych oraz narzędzia do lutowania, aby zmontować elektronikę do kontroli silników. Ponadto potrzebne będą:

- moduł ESP32 DevKit V1,
- dwa moduły TMC2208 lub TMC2209 – sterowniki silników krokowych,
- dwa silniki krokowe w rozmiarze NEMA 17,
- dwie krańcówki z magnesami (sensory Halla),
- sensor światła TSL2561,
- panelowe złącza USB-C oraz zasilania DC,
- płytką drukowaną dla układu (autor udostępnia projekt odpowiedniej płytki PCB),
- około 1 metra paska LED RGBW SK6812 – autor zastosował wersję z 144 diodami na metr,
- złącza JST z dwoma, trzema, czterema i pięcioma pinami oraz pasujące go nich gniazda SMD,



Rysunek 1. Projekt całego urządzenia po złożeniu: a) wygaszona tacka do wsypania piasku; b) przekrój układu; c) widok rozstrzelony urządzenia

Do wydrukowania elementów autor zużył niemalże 3 kg filamentu, który zakupił specjalnie do tego projektu. Dobrze jest zaopatrzyć się w odpowiednią ilość filamentu do druku, aby mieć zapas na wypadek jakichś awarii podczas drukowania itp. Autor wskazuje, że jeśli chodzi o druk, to PLA radziłby sobie lepiej jako materiał, szczególnie pod względem wypaczania się wydruku na skutek stygnięcia, jednak w tym projekcie zależało mu na właściwościach strukturalnych tworzywa ABS, stąd taka decyzja dotycząca materiału.

Montaż urządzenia

Podczas montażu należy zamontować 4-pinowe złącza JST do silników krokowych, 5-pinowe złącze do czujnika światła TSL2561 oraz 3-pinowe do obu krańcówek magnetycznych i paska LED. Następnie można zainstalować silnik krokowy theta za pomocą śrub z łbem stożkowym M3 do podstawy robota. Następnie na oś silnika zakładane jest koło pasowe GT2 i mocowane obrotowe łożysko zarówno do dużego centralnego koła zębatego osi theta, jak i do dolnej płyty podstawy. Wszystko montowane jest za pomocą śrub M3. Łożysko wskoczy na swoje miejsce i nie powinno się chybotać. W dokumentacji w Fusion dostępna jest opcjonalna przekładka, której można użyć, jeśli tak nie jest. Oś liniowa i szyny prowadzące wkręcają się bezpośrednio w górną część przekładni theta za pomocą śrub M3 i insertów osadzanych termicznie w wydruku 3D. Po złożeniu ten zespół powinien wyglądać, jak pokazano na **fotografii 1**. Widoczny jest zamocowany na końcu jednego z ramion magnes, który odpowiada za ruch kulki w piasku.



Fotografia 1. Ruchoma sekcja rzeźby kinetycznej po zmontowaniu

- 4-calowe (100 mm) łożysko z podstawki obrotowej,
- śruby M3 i inserty M3 do osadzania na ciepło,
- pasek GT2 o długości 400 mm oraz pasujące do niego koła zębate.

Projekt elementów mechanicznych

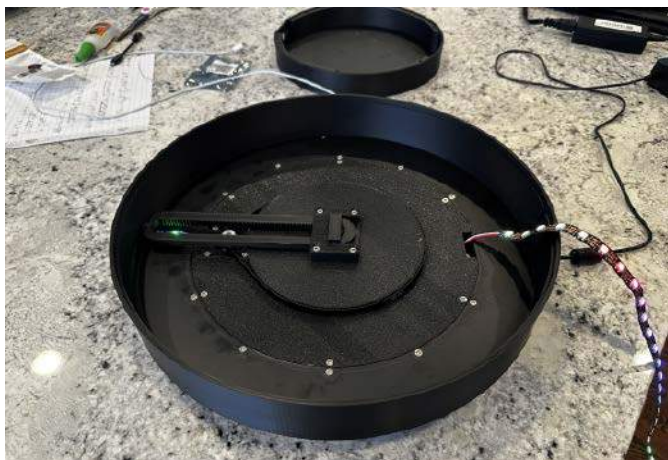
Zaprezentowany robot do kinetycznej rzeźby zaprojektowany został tak, aby był wyjątkowo smukły. Po złożeniu cały układ ma mieć nie więcej niż 75 mm wysokości. Z uwagi na to w systemie jest wiele nakładających się na siebie części, głównie w miejscach, w którym duża przekładnia theta mocuje się do płyty podstawy. Cały robot został zaprojektowany od podstaw w Fusion 360, a jego projekt jest udostępniony za darmo za pośrednictwem Fusion Team (link znaleźć można na stronie z projektem).

Po ściągnięciu projektu łatwo zauważyć, że jest to... 55 wersja tego projektu. Jak sam autor przyznaje: „Spędziłem w fazie projektowania więcej, niż chciałbym przyznać”. Projekt rozpoczął się od zgrubnego szkicu kształtu obudowy, a następnie została ona podzielona wzdłużnie, aby zapewnić możliwość łatwego drukowania. Następnie zaprojektowana została podstawa robota, wymodelowane zostało miejsce na duże łożysko do obrotu całości wokół własnej osi i dodany został centralny zestaw kół zębatach.

Projekt całego układu pokazano na **rysunku 1a**. Dodatkowo na **rysunku 1b** oraz **1c** widać, odpowiednio, przekrój przez kluczowe elementy oraz widok rozstrzelony układu.

Wydruk w 3D

Całość została wydrukowana z użyciem czarnego tworzywa ABS firmy Polymaker z wypełnieniem typu gyroid w wysokości od 30% do 50%, w zależności od części. Na ogół wyższe wypełnienie dotyczyło elementów wewnętrznych, przenoszących naprężenia i siły w układzie, a sama zewnętrzna obudowa wydrukowana została z najmniejszym wypełnieniem, jakie pozwalało zachować jej sztywność.



Fotografia 2. Zmontowany mechanizm wraz z górną częścią obudowy



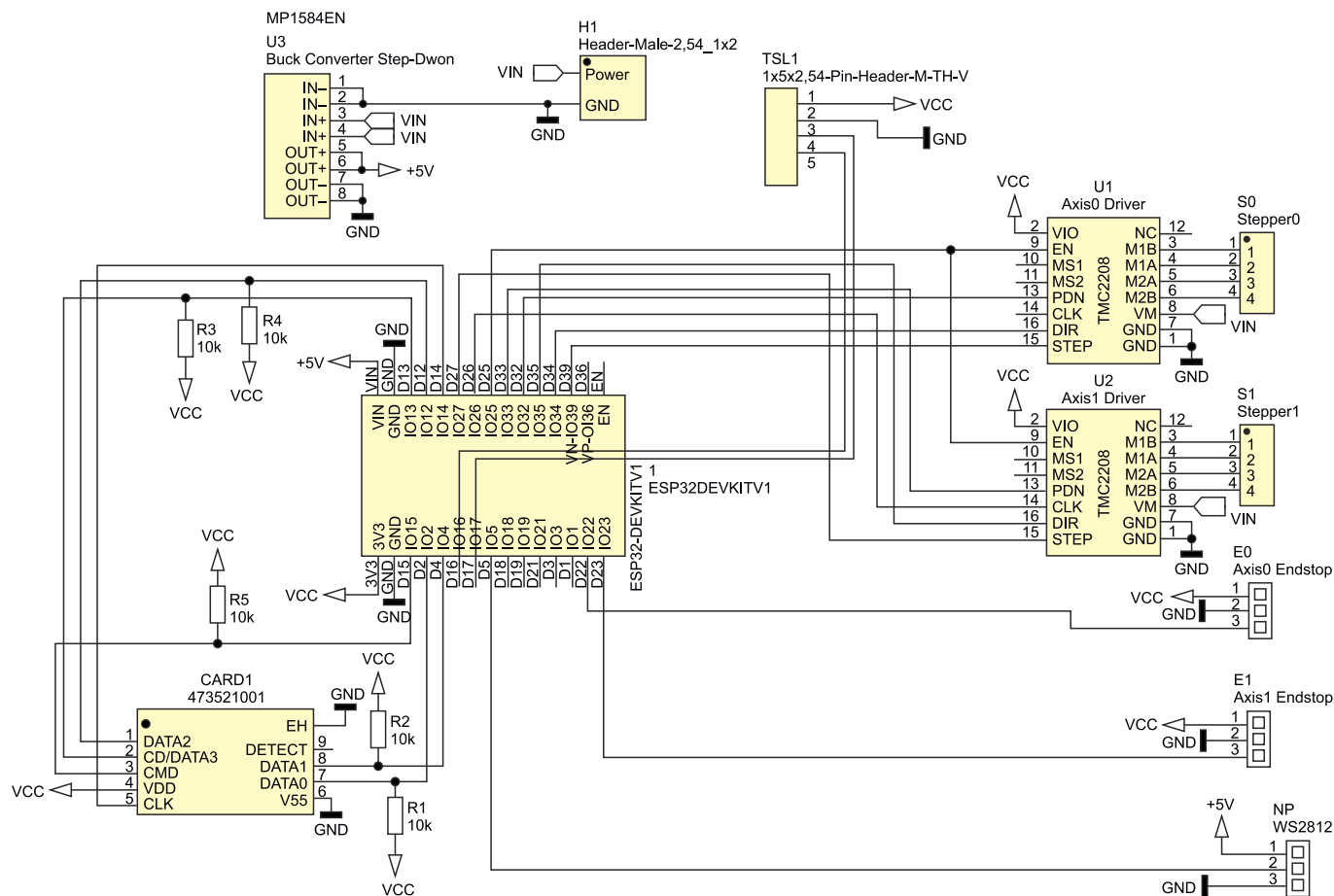
Fotografia 3. Gotowy układ z zainstalowanym podświetleniem i warstwą piasku na filcu

Po zamontowaniu elementów mechanicznych można obsadzić w podstawie krańcówki. Zawierają one czujniki Halla, które aktywują się w pobliżu pola magnetycznego magnesu. Jeden czujnik zainstalowany jest bezpośrednio w małym otworze przy środku, aby wyczuć, kiedy przekładnia theta dotarła do znanej pozycji.

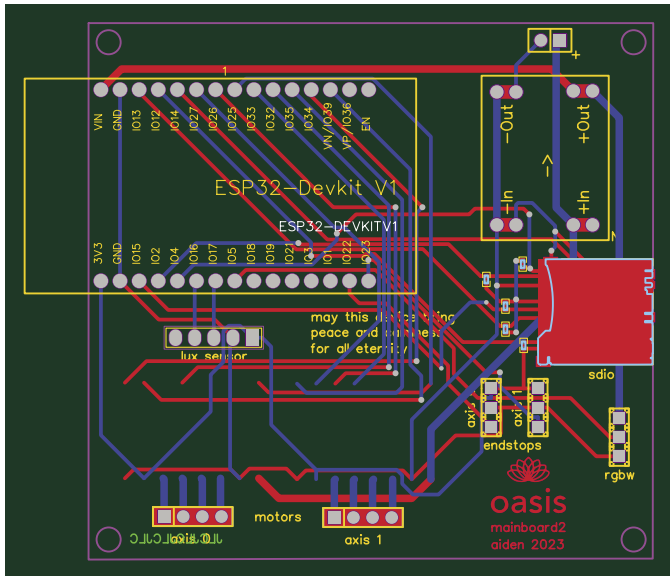
W pełni zmontowana płyta podstawy robota po prostu wsuwa się w środek dolnej połowy obudowy. Wystarczy umieścić górną połowę obudowy na dolnej połowie i można przykręcić ją za pomocą insertów osadzanych na ciepło w wydruku i 12 śrub M3 o długości 10 mm. Zmontowany moduł pokazano na **fotografii 2**. Na pokazany element zakładany jest jeszcze pojemnik na piasek. Aby kulka rzeźbiąca w piasku poruszała się bez oporów i hałasów, na dnie pojemnika autor wkleił cienki filc, który przykrył piaskiem. Aby konstrukcja wyglądała jeszcze ciekawiej, piasek został podświetlony łańcuchem LED, umieszczonym w pojemniku. Jego wklejenie powinno być ostatnim krokiem montażu części mechanicznej.



Fotografia 4. Gotowe urządzenie w trakcie pracy



Rysunek 2. Schemat układu sterującego kinetyczną rzeźbą



Rysunek 3. Projekt płytki drukowanej układu

Ostatnim etapem jest nasypianie drobnego piasku do przygotowanego pojemnika – autor zastosował zwykły piasek ogrodowy ze sklepu budowlanego, ale nic nie stoi na przeszkodzie, aby zastąpić go np. kolorowym piaskiem akwarystycznym, jasnym i bardzo drobnym piaskiem kwarcowym itp. Ostatecznie na piasku umieszczana jest stalowa kulka (koniecznie z normalnej stali, która jest ferromagnetyczna, a nie np. stali nierdzewnej). Poruszający się pod spodem magnes będzie wodził kulką po piasku, która będzie żłobiła w nim ślady. Finalny system pokazano na **fotografii 3** oraz **4**.

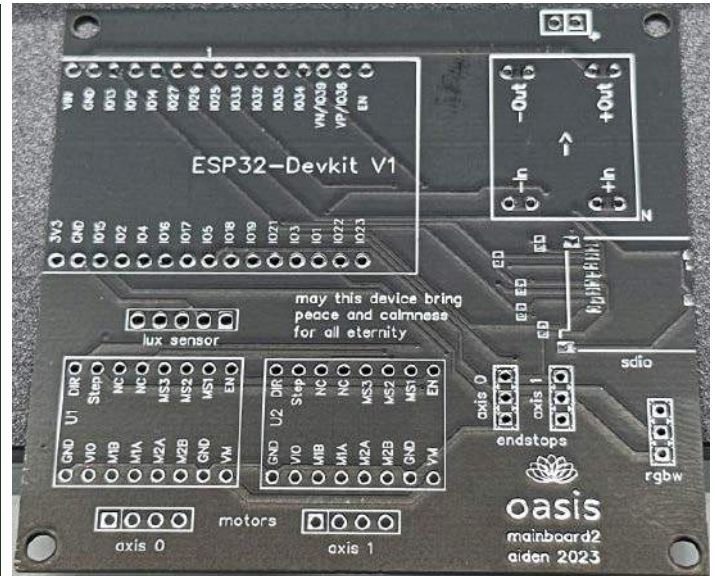
Sterowanie elektroniczne

Kinetyczna rzeźba jest sterowana prostym układem, który zawiera zaledwie kilka modułów. Centralnym elementem jest moduł ESP32DevKit V1, do którego podłączone są pozostałe elementy – dwa sterowniki silników krokowych (TMC2208), gniazda karty SD, krańcówki i diod LED ES2812. Schemat układu pokazany jest na **rysunku 2**.

System zasilany jest z zewnątrz, z zasilacza DC, dołączonego poprzez gniazdo H1. W układzie znajduje się przetwornica step-down typu Buck, która stabilizuje napięcie zasilania 5 V dla całego systemu.

Płytką drukowaną

Dla uproszczenia konstrukcji autor umieścił wszystkie moduły na jednej wspólnej płytce. Na płytce znajdują się też gniazda,

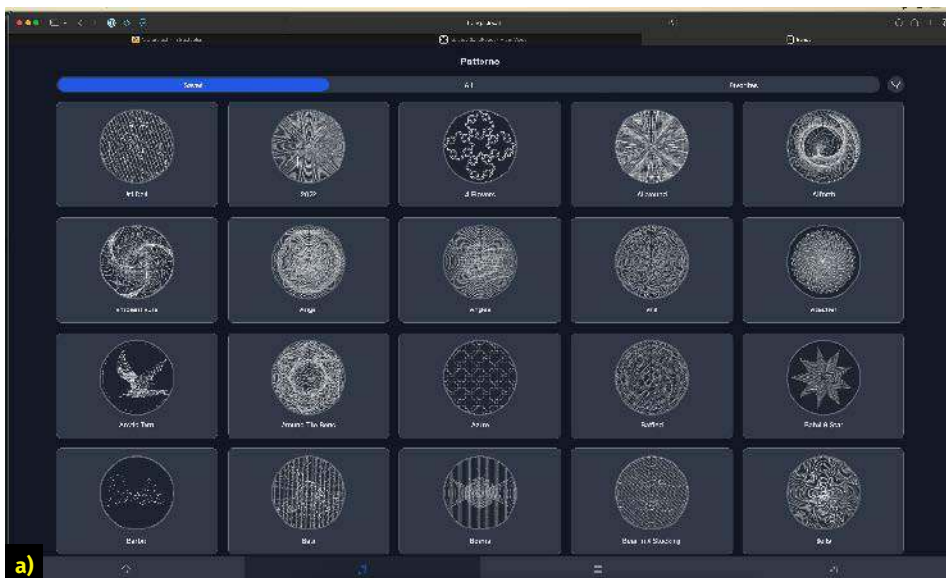


Fotografia 5. Płytką drukowaną systemu

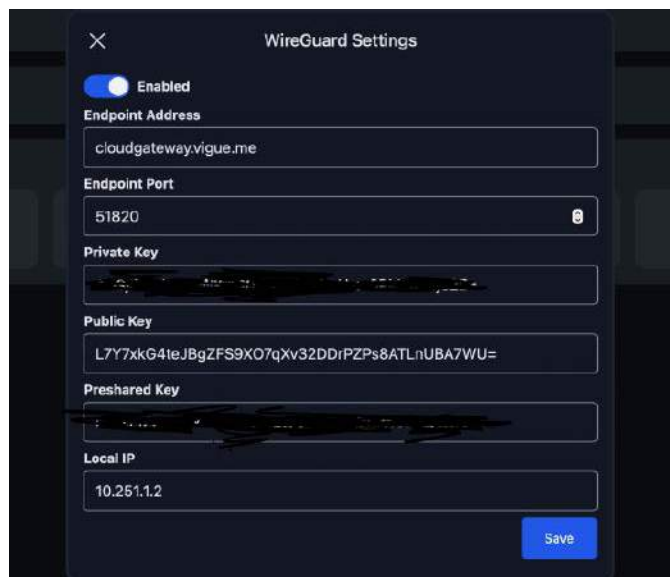
przeznaczone do podłączenia elementów zewnętrznych, takich jak dwa silniki krokowe, krańcówki, czujnik natężenia oświetlenia otoczenia (do dostosowywania natężenia światła z LED-ów) i diody LED. Projekt płytki pokazano na **rysunku 3**. Płytką została zamówiona w profesjonalnej fabryce i wykonana z czarną, matową soldermaską, dzięki czemu lepiej pasuje wizualnie do reszty elementów. Gotową płytkę drukowaną pokazano na **fotografii 5**.

Przy projektowaniu płytki do modułu ESP32 istotne jest, aby sekcja tego modułu z anteną wystawała poza obrys laminatu (najlepiej) lub przynajmniej część w otoczeniu anteny Wi-Fi była pozbawiona ścieżek, wylewki masy itd. Ma to zagwarantować poprawne działanie sekcji radiowej modułu. Podobnie jest z montażem gniazda karty SD – musi ono umożliwiać dostęp do karty SD, szczególnie, że taki sam dostęp zapewniać musi obudowa. Nazwa znajdująca się na PCB – „oasis”, to pierwotna nazwa projektu. Autor zmienił nazwę na „Tranquil”, ale jest ona obecna jedynie w oprogramowaniu.

Montaż PCB był bardzo prosty. Na płytce obsadzono drivery silników, moduł ESP32 i złącza JST. Dodano opcjonalne gniazdo kart SD. Złącze przylutowane zostało ręcznie za pomocą cienkiej końcówki lutownicy i topnika no-clean. Jest to dosyć delikatny element do montażu powierzchniowego, więc jego zalutowanie, jako jedynego elementu, może być pewnym wyzwaniem.



Rysunek 4. Interfejs urządzenia; a) widziany na komputerze; b) widziany na telefonie komórkowym



Rysunek 5. Konfiguracja VPN – WireGuard

Interfejs webowy

Interfejs sieciowy do kontroli aplikacji został napisany przez autora od podstaw przy użyciu TypeScript i Vue 3. Zarządzanie stanem odbywa się za pomocą Pinia, a pobieranie danych realizowane jest za pomocą zestawu przechwytywaczy zapytań (*request interceptor*) Axios.

Kod interfejsu webowego znajduje się w repozytorium tranquilvue autora (link do repozytorium znajduje się na końcu artykułu). Wystarczy sklonować repozytorium z GitHuba i zbudować ten projekt, aby kontynuować prace nad przygotowaniem oprogramowania dla układu dalej. Jeśli pracujemy pod Linuxem, aby zbudować ten pakiet, wystarczy w linii komend, znajdując się w folderze sklonowanym z repozytorium, wpisać następujące polecenie: `yarn && yarn build`.

Wygenerowane pliki kompilacji zapisywane są w katalogu `./dist`. Jeśli chcielibyśmy to samo zrobić na komputerze z Windowsem, musimy najpierw zainstalować menedżer pakietów Yarn, a następnie możemy przeprowadzić analogiczną operację.

Interfejs webowy pokazany jest na **rysunku 4**. Widoczne tam wzory nie są zawarte domyślnie w oprogramowaniu – trzeba je pobrać z Internetu. Interfejs tylko wyświetla wzory, które zainstalowane są na urządzeniu.

Oprogramowanie

Oprogramowanie modułu zawiera standardową instalację PlatformIO. Samo oprogramowanie zostało napisane od nowa, prawie od podstaw. Autor pracuje nad nim już od co najmniej roku.

Jest to normalny projekt dla PlatformIO i można go pobrać ze strony z projektem, z repozytorium TranquilFirmware autora na GitHubie.

Wystarczy teraz przeciągnąć zbudowane wcześniej zasoby interfejsu webowego do katalogu z projektem PlatformIO, a PlatformIO automatycznie przeniesie je na partycję SPIFFS w module ESP32.

Po zainstalowaniu oprogramowania układowe urządzenie skonfiguruje własny punkt dostępowy o nazwie Tranquil, w którym zostanie wykonana cała początkowa konfiguracja. Możesz następnie podłączyć urządzenie do naszego Wi-Fi, skonfigurować aktualizacje OTA, a nawet skonfigurować układ jako klienta WireGuarda (VPN), co opisano w kolejnym akapicie.

Konfiguracja WireGuard

WireGuard to szybki, prosty i nowoczesny tunel VPN (*Virtual Private Network*), który umożliwia bezpieczne i prywatne połączenie między różnymi urządzeniami w Internecie. Został stworzony jako alternatywa dla tradycyjnych narzędzi VPN, takich jak OpenVPN czy IPSec, i jest uważany za bardziej wydajny, bezpieczny i łatwiejszy w konfiguracji. Z tego powodu autor powyższej konstrukcji zdecydował się na implementację jego wsparcia w systemie.

Aby skonfigurować VPN w urządzeniu, należy przejść do karty ustawień urządzenia i wprowadzić odpowiednią parę kluczy oraz informacje o punkcie końcowym, aby robot mógł dołączyć do zdalnej podsieci WireGuard. Widok tych ustawień pokazano na **rysunku 5**. Autor wdrożył to rozwiązanie, ponieważ uczelnia, do której będzie niebawem uczęszczać, wdraża izolację klientów w swojej sieci, a bez tego nie miałby możliwości kontrolowania robota w akademiku. Ma to jednak o wiele więcej zastosowań – pozwala na opracowanie VPN, który umożliwi dostęp do konfiguracji i sterowania naszą kinetyczną rzeźbą z dowolnej sieci na świecie.

Gotowe urządzenie

Po zakończeniu budowy i programowania urządzenia można je skonfigurować do działania. Początkowo urządzenie uruchomi własny punkt dostępu o nazwie Tranquil, w którym zostanie wykonana cała początkowa konfiguracja. Moduł można podłączyć do Wi-Fi, skonfigurować aktualizacje OTA, a nawet skonfigurować opisanego powyżej klienta WireGuarda. To wszystko zapewnia bardzo dużą elastyczność układu i łatwe jego sterowanie.

Nikodem Czechowski, EP

Źródła:

1. <https://shorturl.at/aFI78>
2. <https://github.com/acvigue/tranquilvue>
3. <https://github.com/acvigue/TranquilFirmware>
4. <https://vigue.me/>

REKLAMA



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

Kurs FPGA Lattice (11)

Statyczna analiza czasowa i maksymalna częstotliwość zegara

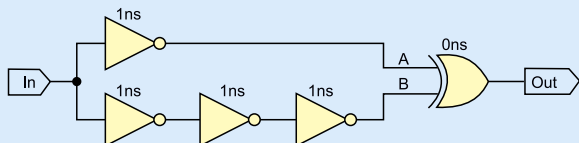
Producenci procesorów przyzwyczaili nas, że częstotliwość maksymalna zegara podawana jest na pierwszej stronie dokumentacji jako jeden z kluczowych parametrów, takich jak rozmiar pamięci Flash czy RAM. Z jakim zegarem może pracować układ FPGA? Niestety, jest to dość skomplikowane pytanie...

W tej części kursu będzie dużo teorii. Problemy, jakie będziemy dzisiaj poruszać, są niezwykle ważne, a wręcz fundamentalne. Brak znajomości tych zagadnień może prowadzić do bardzo irytujących błędów, które mogą być bardzo trudne do wykrycia i naprawienia. Jednak bądź dobrej myśli! Niedługo stwierdzisz, że to wcale nie jest takie trudne i nauczysz się zrecznie omijać pułapki.

Glitch

Zacznijmy od prostego eksperymentu, który uświadomi nam pewien problem, którym obarczony jest każdy układ kombinacyjny. Zastanówmy się nad banalnym schematem, który pokazano na **rysunku 1**. Na schemacie są tylko cztery bramki NOT i jedna bramka XOR. Ich tablice prawdy dla przypomnienia pokazano w **tabelach 1 i 2**. Układ ma tylko jedno wejście, oznaczone jako **In** oraz jedno wyjście **Out**.

Spróbujmy teraz stworzyć tablicę prawdy dla całego układu. Jest tylko jedno wejście, które może przyjmować stan 1 lub 0, więc musimy przeanalizować tylko dwie możliwości, jakie mogą wystąpić na wejściu **In**. W tablicy uwzględnimy także punkty pośrednie **A** i **B**,



Rysunek 1. Schemat z „glitchem”

Tabela 1. Tablica prawdy bramki NOT

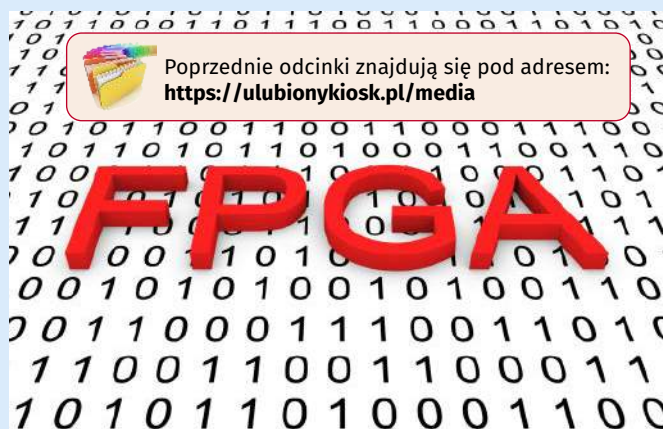
A	NOT
0	1
1	0

Tabela 2. Tablica prawdy bramki XOR

A	B	XOR
0	0	0
0	1	1
1	0	1
1	1	0

Tabela 3. Tablica prawdy dla układu ze schematu 1

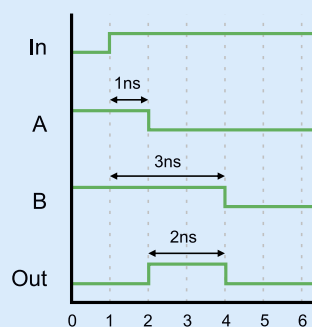
In	A	B	Out
0	1	1	0
1	0	0	0



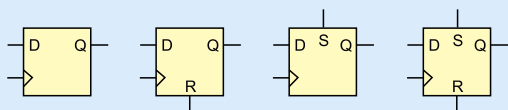
które są wejściami bramki XOR. Wynik tej symulacji zaprezentowano w **tabeli 3**. Okazuje się, że obojętnie jaki stan jest doprowadzony do wejścia **In**, to na wyjściu **Out** zawsze jest stan niski.

W naszej prostej symulacji nie uwzględniliśmy bardzo ważnego czynnika – czasu propagacji. Jest to czas, jaki mija pomiędzy ustaleniem się stanu na wejściu układu a ustaleniem się stanu jego wyjścia. Załóżmy, że bramki NOT mają czas propagacji równy 1 nanosekundzie, a bramka XOR niech ma zerowy czas propagacji (dla wygody) – czasy propagacji bramek podano na **rysunku 1** nad symbolami bramek. Aby przeprowadzić symulację czasową, już nie wystarczy prosta tabelka. Potrzebujemy wykresu, jaki znajduje się na **rysunku 2**. Naszą sekwencją testową, którą byśmy umieścili w test-benchu, pisząc kod w Verilogu, niech będzie jedynie przełączenie się sygnału **In** ze stanu 0 na 1, a następnie niech ten sygnał pozostaje bez zmian aż do końca symulacji.

Sprawdźmy, co będzie działo się w punkcie **A**. Sygnał przechodzi przez jedną bramkę NOT, więc zostanie zanegowany. Jednak ta bramka ma czas propagacji 1 ns, więc efekt negacji zobaczymy po upływie 1 ns. W punkcie **B** mamy podobną sytuację, lecz sygnał musi przejść przez trzy bramki NOT, więc stan tego punktu ustali się po upływie 3 ns. Co się dzieje wtedy z bramką XOR? Zwróć uwagę, że przez 2 nanosekundy na wejściu **A** jest stan niski, a na wejściu **B** jest stan wysoki. Zgodnie z tabelą 2, w takiej sytuacji wyjście bramki XOR powinno ustawić się w stan wysoki! Dopiero, kiedy łańcuch bramek NOT przetworzy zmianę sygnału, na obu wejściach **A** i **B** ustali się stan niski.



Rysunek 2. Przebiegi czasowe dla układu ze schematu z rysunku 1



Rysunek 3. Symbole różnych przerzutników D

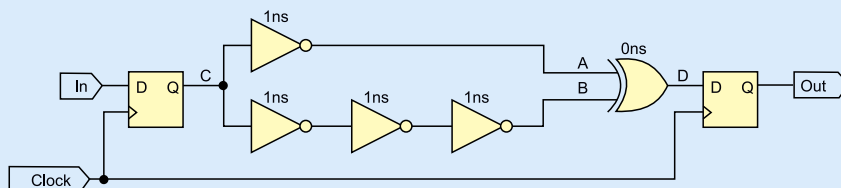
Tabela 4. Tablica prawdy przerzutnika D

D	Clock	Q_n
X	1	Q_{n-1}
X	\	Q_{n-1}
X	0	Q_{n-1}
0	/	0
1	/	1

Wyjście bramki XOR przełączy się ponownie w stan niski po upływie 3 ns od zmiany stanu wejścia **In** (gdybyśmy chcieli uwzględnić czas propagacji bramki XOR, wtedy wykres **Out** należałoby przesunąć w prawo o czas propagacji tej bramki).

Okazuje się, że przez dwie nanosekundy na wyjściu układu jest krótka szpilka, której nie powinno być. Jest to tzw. **glitch**. Ta krótka szpileczka może powodować bardzo poważne problemy. Gdybyśmy odczytali wyjście **Out**, zanim ustali się jego stan, to byśmy otrzymali nieprawidłową wartość tego sygnału. Ten nieprawidłowy stan przedostałby się do dalszych bramek, przerzutników, multiplexerów, przez co nasze urządzenie działałoby w sposób nieprawidłowy!

Glitch jest zjawiskiem występującym w układach kombinacyjnych. Im bardziej skomplikowany układ, tym więcej ścieżek jest do przeanalizowania i tym więcej różnych glitchy może powstać. W naszym prostym układzie mamy tylko dwie ścieżki. Jak ten problem rozwiązać? Z pomocą przychodzi przerzutnik D.



Rysunek 4. Schemat z przerzutnikami D

Przerzutnik D (flip flop)

Ten typ przerzutnika jest koniem pociągowym całej elektroniki cyfrowej. W zasadzie wszystkie przerzutniki, jakie mamy w FPGA, to w gruncie rzeczy przerzutniki D. Litera D oznacza *delay*, czyli opóźnienie. Przerzutnik jest najprostszą formą pamięci cyfrowej i ma pojemność jednego bitu. Może przechowywać stan niski lub wysoki przez dowolnie długi czas, tak długo, jak oczywiście zasilanie układu jest włączone. Kilka symboli przerzutników D pokazano na **rysunku 3**.

Przerzutnik D ma co najmniej trzy wyprowadzenia:

- **D** – wejście danych,
- **Q** – wyjście danych,
- **C** (oznaczane też znakiem trójkąta) – wejście zegara.

Tablicę prawdy przerzutnika D pokazano w **tabeli 4**. Przerzutnik jest układem sekwencyjnym. To znaczy, że stan jego wyjścia zależy nie tylko od obecnego stanu wejść, ale także od stanów, jakie były w przeszłości. Obecny stan wyjścia oznaczamy jako Q_n , a stan z poprzedniego cyklu zegarowego oznaczamy symbolem Q_{n-1} . Symbol / oznacza zbocze rosnące, czyli zmianę stanu z 0 na 1, a symbol \ zbocze opadające, czyli zmianę z 1 na 0. Znak X oznacza, że stan jest nieistotny, obojętnie niski czy wysoki.

Przerzutnik D reaguje jedynie na zbocze rosnące sygnału zegarowego. Wtedy stan wejścia D jest przepisywany na wyjście Q. Kiedy wejście zegarowe przyjmuje stan niski, wysoki lub jest zbocze opadające, to stan wejścia D może się dowolnie zmieniać, a wyjście Q pozostaje cały czas niezmienione. Oprócz tego można spotkać dodatkowe wyprowadzenia:

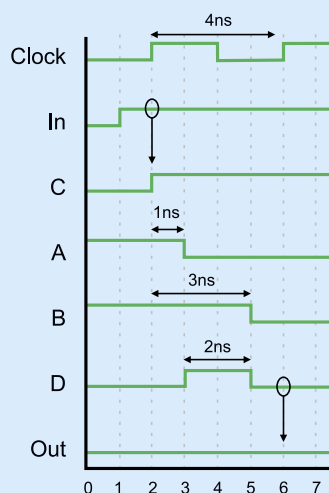
- **S** – wejście asynchroniczne set, ustawiające natychmiast wyjście Q w stan wysoki bez względu na stan pozostałych pinów,

- **R** – wejście asynchroniczne reset, ustawiające natychmiast wyjście Q w stan niski bez względu na stan pozostałych pinów,
- **!Q, Q z kreską na górze** – zanegowane wyjście.

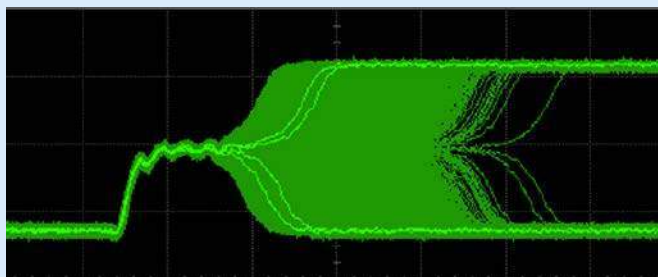
Jednoczesne ustawienie wejść S i R w stan wysoki jest niedozwolone. Wejścia S i R są najczęściej używane podczas inicjalizacji stanu przerzutnika po włączeniu zasilania lub zresetowaniu układu.

Dochodzimy do wniosku, że skoro przerzutnik D kopiuje stan wejścia na wyjście tylko w momencie zbocza zegara, to możemy go zastosować jako „izolator” pomiędzy blokami logiki kombinacyjnej. Usprawnijmy więc schemat, dodając do niego przerzutniki D, jak jak pokazano to na **rysunku 4**. Za wejściem oraz przed wyjściem zostały dodane dwa przerzutniki D, które taktowane są tym samym sygnałem zegarowym. Przerzutniki powinny mieć także wspólny sygnał resetujący, ale dla uproszczenia go pominięto. Przerzutnik przy wejściu **In** będziemy nazywać przerzutnikiem nadającym, a ten przy wyjściu **Out** będzie przerzutnikiem odbierającym.

Prześledźmy teraz przebiegi poprawionego układu, które pokazano na **rysunku 6**. Dla uproszczenia przyjmujemy, że przerzutniki mają zerowe opóźnienie. Podobnie jak wcześniej, sygnał wejściowy **In** zmienia swój stan z niskiego na wysoki, co widać w pierwszej nanosekundzie symulacji. Następnie, w drugiej nanosekundzie, mamy zbocze rosnące sygnału zegarowego (odstęp między zboczami sygnałów **In** oraz **Clock** jest celowy – będzie to wyjaśnione w dalszej części artykułu). Zbocze rosnące sygnału zegarowego powoduje, że przerzutnik nadający kopiuje stan swojego wejścia na wyjście. Z tego powodu stan w punkcie **C** zmienia się z niskiego na wysoki,



Rysunek 5. Przebiegi czasowe dla schematu z rysunku 4



Rysunek 6. Wyjście przerzutnika w stanie metastabilnym

a następnie bramki zaczynają przetwarzać sygnał z wyjścia przerzutnika nadającego. Przez kolejne kilka nanosekund wszystko dzieje się tak samo jak w poprzednim przykładzie.

Sygnał na wyjściu bramki XOR widzimy w punkcie **D**. Jest tam obecny ten sam glitch, co w pierwszym przykładzie. Jednak glitch nie przedostaje się do wyjścia **Out**, ponieważ między nim a wyjściem bramki XOR jest przerzutnik D. Stan bramki XOR trafia na wyjście układu dopiero w szóstej nanosekundzie, kiedy mamy kolejne zbocze rosnące sygnału zegarowego, co zaznaczono strzałką – jednak nie widzimy żadnej różnicy na linii **Out**, ponieważ stan niski jest zamieniany na stan niski. Wszak nasz układ ma zawsze dawać stan niski, niezależnie od tego, co znajduje się na jego wejściu.

Możemy dojść do wniosku, że zastosowanie przerzutników D w gruncie rzeczy nie rozwiązuje problemu, ale go ukrywa. Mimo to, umieszczanie przerzutnika D za blokiem logiki kombinacyjnej pozwala lepiej zapanować nad tym, co się dzieje w naszym projekcie. Dzięki temu rozwiązaniu wszystkie bloki kombinacyjne dostają „nowe dane” wraz ze zboczem zegara i muszą zdążyć je przetworzyć przed kolejnym zboczem sygnału zegarowego.

Zastanówmy się nad maksymalną częstotliwością zegara, jaką można zastosować w naszym układzie. Mamy dwie ścieżki o czasach propagacji 1 ns oraz 3 ns. Interesuje nas czas propagacji najdłuższej ścieżki. Musimy do niego dodać także czas setup time przerzutnika (o tym będzie za chwilę). Załóżmy, że setup time to 1 ns, czyli razem blok kombinacyjny z przerzutnikiem potrzebuje 4 ns na przetworzenie sygnału. Odwrotność z 4 ns to $1/4 \cdot 10^9$, czyli 250 MHz. To właśnie jest maksymalna częstotliwość sygnału zegarowego, jaką można zastosować w tym układzie. Gdybyśmy zastosowali zegar o częstotliwości większej niż 250 MHz, wtedy przerzutnik wyjściowy dokonałby skopiowania stanu wcześniej, zanim bramki ustaliły stany swoich wyjść. W takiej sytuacji przerzutnik skopiowałby glitch i na wyjściu **Out** byśmy zobaczyli stan wysoki, co byłoby niepoprawne.

Metastabilność

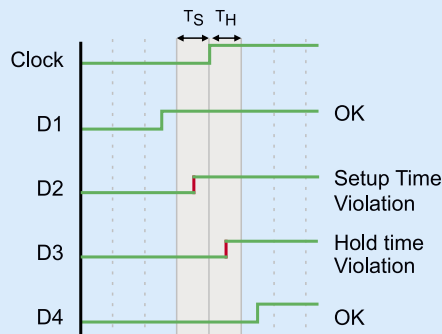
Wyjście przerzutnika może być w stanie niskim lub wysokim albo może wylać się poza zero-jedynkową konwencję. Wtedy mamy problem metastabilności. Stan metastabilny polega na tym, że przerzutnik „nie może się zdecydować”, czy chce być w stanie niskim, czy wysokim. Zobaczmy **rysunek 6** – pokazuje on wielokrotne, nałożone na siebie, przebiegi napięcia na wyjściu przerzutnika. Przerzutnik miał zmienić swój stan z 0 na 1, jednak wpada w stan metastabilny. Przez pewien krótki (ale zmienny!) czas napięcie na jego wyjściu jest zbliżone do połowy napięcia zasilania. Następnie przerzutnik przechodzi do stanu niskiego lub wysokiego. Problem pogarsza to, że stan, w jakim ustali się jego wyjście, to kwestia czysto losowa. Nie da się tego w żaden sposób przewidzieć.

Metastabilność powoduje dwa problemy:

1. Jeżeli stan wyjścia ustali się nieprawidłowo, to elementy odbierające sygnał z przerzutnika dostaną nieprawidłowe dane i prześlą je dalej, do kolejnych elementów układu;
2. Jeżeli stan wyjścia ustali się prawidłowo, to elementy odbierające sygnał z przerzutnika otrzymają go później, niż powinny. Może się okazać, że z tego powodu logika kombinacyjna nie zdąży przetworzyć sygnałów przed nadejściem zbocza zegara, a w rezultacie błędne dane zostaną przesłane do kolejnych elementów. Przerzutnik ma dwa kluczowe parametry czasowe:

- **Setup time T_s (czas ustalenia)** – jest to okres przed zboczem zegara, w którym sygnał na wejściu nie może się zmieniać;
- **Hold time T_h (czas utrzymania)** – jest to okres po zboczach zegara, w którym sygnał na wejściu nie może się zmieniać.

W ramach tych dwóch czasów zabroniona jest jakakolwiek zmiana na wejściu przerzutnika. To znaczy, że stan na wejściu musi się ustalić odpowiednio wcześniej przed zboczem zegara i musi pozostać stabilny jeszcze przez jakiś czas. Jeżeli w czasie setup lub



Rysunek 7. Przykłady prawidłowych i nieprawidłowych zmian na wejściu przerzutnika D

hold nastąpi zmiana sygnału wejściowego, to będziemy mieć błąd *setup time violation* lub *hold time violation*, a w rezultacie przerzutnik może wpaść w stan metastabilny. Na **rysunku 7** pokazano przykłady prawidłowych i nieprawidłowych zmian na wejściu danych przerzutnika.

W poprzednim przykładzie założyliśmy, że przerzutnik działa natychmiast i reaguje na zmiany sygnałów wejściowych w zerowym czasie. Przerzutnik ma także czas propagacji, podobnie jak elementy kombinacyjne. Czas propagacji przerzutnika liczy się od zbocza zegara do ustalenia się stanu na wyjściu przerzutnika i oznacza się TCQ. Możemy pokusić się teraz o podanie wzoru, pozwalającego obliczyć maksymalną częstotliwość sygnału zegarowego:

$$T_{MAX} = TCQ + T_{CLMAX} + T_S$$

$$f_{MAX} = \frac{1}{T_{MAX}}$$

gdzie:

- T_{MAX} – czas przetwarzania danych przez najdłuższą ścieżkę kombinacyjną,
- TCQ – czas propagacji przerzutnika nadającego, liczony od zbocza zegara do ustalenia się danych na wyjściu przerzutnika,
- T_{CLMAX} – czas propagacji najdłuższej ścieżki w logice kombinacyjnej pomiędzy przerzutnikiem nadającym i odbierającym,
- T_S – setup time przerzutnika odbierającego,
- f_{MAX} – maksymalna częstotliwość sygnału zegarowego.

Jednak dla bezpieczeństwa należy stosować zegar o częstotliwości trochę mniejszej niż maksymalna.

Dochodzimy ostatecznie do wniosku, że maksymalna częstotliwość zegara zależy od czasu propagacji w najbardziej rozbudowanym bloku logiki kombinacyjnej, czasu propagacji przerzutnika oraz jego setup time. Kiedy częstotliwość zegara jest zbyt duża, dostaniemy błąd *setup time violation*.

Kiedy zatem może wystąpić *hold time violation*? Ten problem wydaje się mniej oczywisty. Aby tego błędu nie było, musi zostać spełniony warunek:

$$TCQ + T_{CLMIN} > T_H$$

gdzie:

- TCQ – czas propagacji przerzutnika nadającego, liczony od zbocza zegara do ustalenia się danych na wyjściu przerzutnika,
- T_{CLMIN} – czas propagacji najkrótszej ścieżki w logice kombinacyjnej pomiędzy przerzutnikiem nadającym i odbierającym,
- T_H – hold time przerzutnika odbierającego.

TCQ oraz T_H to parametry przerzutnika i nie mamy na nie żadnego wpływu. Jedynie mamy wpływ na czas propagacji między przerzutnikami. Okazuje się, że nie może on być zbyt długi, bo będzie *setup time violation*, ale nie może być też zbyt krótki, bo wtedy dostaniemy *hold time violation*!

W jakiej sytuacji taki błąd może mieć miejsce? Najlepszym przykładem są rejestry przesuwające, czyli łańcuszek przerzutników, gdzie wyjście jednego jest bezpośrednio połączone z wejściem kolejnego.

W takiej sytuacji nie ma żadnych bramek między przerzutnikami, więc czas TCLMIN to jedynie czas ładowania pojemności pasożytniczej, która występuje między linią łączącą przerzutniki i masą. Ten czas jest bardzo mały.

Czy zatem w FPGA nie możemy robić rejestrów przesuwających? Możemy! Na szczęście Lattice Diamond w taki sposób tworzy rejestry przesuwające, żeby nie było tego problemu. Problem hold time violation jest bardziej istotny dla projektantów układów scalonych niż dla użytkowników FPGA.

Slack

Kolejnym terminem, jaki musimy poznać, jest *slack*. Jest to zapas czasu, jaki pozostaje przed kolejnym zboczem zegara. Mierzymy go od ustalenia się stanu na wyjściu ostatniej bramki logiki kombinacyjnej, które jest połączone z wejściem przerzutnika odbierającego, do wystąpienia zbocza sygnału zegarowego. Slack to czas, w którym nie dzieje się nic. Przerzutnik ma już na swoim wejściu gotowe dane, które się nie zmieniają, więc pozostaje już tylko czekać na zbocze sygnału zegarowego.

Im większy slack, tym lepiej, ale z drugiej strony – jeżeli slack jest bardzo duży, to znaczy, że moglibyśmy zastosować szybszy zegar lub kupić wolniejszy układ FPGA (czyli tańszy). Jeżeli slack jest ujemny, to znaczy, że zegar jest zbyt szybki. Układ kombinacyjny i/lub przerzutnik nie są w stanie przetworzyć danych przed nadejściem kolejnego zbocza zegara. Jeżeli slack jest zerowy, to jesteśmy na ostrzu noża. Teoretycznie układ powinien działać, ale rozsądnie byłoby mieć jakiś zapas czasu, ponieważ czas propagacji zmienia się w zależności od temperatury i napięcia zasilania.

Skew

Efekt skew polega na tym, że zbocze sygnału zegarowego nie występuje dokładnie w tej samej chwili we wszystkich przerzutnikach. Mogłoby się okazać, że przerzutnik nadający otrzymuje zbocze zegara szybciej niż przerzutnik odbierający lub odwrotnie. Przerzutniki leżące bliżej nadajnika otrzymują sygnał trochę wcześniej niż te, które leżą na drugim końcu struktury krzemowej. Z tego powodu sygnały zegarowe muszą być prowadzone z użyciem globalnych sieci zegarowych. W MachXO2 nazywają się one *Primary Clock* i mamy do dyspozycji osiem takich sieci. Są to specjalne linie zaprojektowane w taki sposób, aby zminimalizować skew. Syntezator automatycznie rozpoznaje sygnały zegarowe i umieszcza je w *Primary Clock Networks*.

Dokładnie ten sam problem dotyczy sygnału resetującego. Do prowadzenia sygnałów zerujących służy *High Fanout Network* i do dyspozycji mamy również osiem takich sieci.

W tym momencie musimy wspomnieć, że jeżeli sygnał zegarowy lub resetujący mają pochodzić spoza FPGA (np. z generatora kwarcowego), to trzeba je doprowadzić do ściśle określonych pinów. Są to piny PCLKxy, gdzie x oznacza bank GPIO, a y to numer wejścia zegarowego. Pinout każdego układu FPGA znajdziesz na stronie [1].

Dobłą praktyką jest, by we wszystkich blokach sekwencyjnych **always @(posedge Clock, negedge Reset)** stosować ten sam sygnał zegarowy i ten sam sygnał resetujący. Na liście czułości bloku **always** nie powinno się umieszczać sygnałów, które prowadzone są z użyciem uniwersalnych zasobów logicznych, ponieważ mają one zdecydowanie dłuższy czas propagacji niż sieci globalne, a co gorsza, efekt skew wtedy może być duży.

Jitter

Jitter to fluktuacja częstotliwości sygnału zegarowego. Częstotliwość zegara, jaką zwykle podajemy, to częstotliwość średnia. W rzeczywistości ta częstotliwość nieznacznie się zmienia i czasami jest trochę mniejsza lub trochę większa. Jest to jeden z powodów, dlaczego należy zachować pewien margines bezpieczeństwa pomiędzy maksymalną

częstotliwością, podawaną przez Diamond, a nominalną częstotliwością zastosowanego generatora sygnału zegarowego.

Duży jitter mają generatory tyłu RC, które są wbudowane w układy scalone – przykładem takiego jest generator OSCH wbudowany wewnątrz MachXO2. Jeżeli zależy nam na stabilnym sygnale zegarowym, to dobrym wyborem będzie zastosowanie scalonego generatora kwarcowego.

Synchronizator

Istnieje takie pojęcie jak domena zegarowa. Jest to zbiór przerzutników taktowanych tym samym sygnałem zegarowym. We wszystkich dotychczasowych przykładach (z wyjątkiem odcinka o dzielnikach częstotliwości) korzystaliśmy tylko z jednego zegara. Czyli mieliśmy jedną domenę zegarową.

Domena zegarowa wcale nie musi ograniczać się tylko do wnętrza układu FPGA. Wyświetlacz multipleksowany i klawiatura matrycowa, które omawialiśmy w poprzednich odcinkach, również działały w tej samej domenie zegarowej, co wszystkie pozostałe komponenty wewnątrz FPGA. Gdyby zaistniała potrzeba przesyłania sygnałów pomiędzy elementami w dwóch różnych domenach zegarowych, trzeba zastosować pomiędzy nimi synchronizator. W przeciwnym wypadku mogłoby dojść do naruszenia *setup time* i *hold time*.

Szczególnym przypadkiem są elementy takie jak przyciski, enkodery, sensory czy wejścia sterowane przez inne układy scalone. Są to wejścia, na których sygnały zmieniają się asynchronicznie i trzeba je najpierw zsynchronizować z zegarem. Jednym ze sposobów jest zastosowanie synchronizatora, którego kod pokazano na **listingu 1**. Możesz go zasymulować w symulatorze EDA Playground pod adresem [2]. Taki synchronizator może zostać zastosowany w celu synchronizowania sygnałów, pochodzących z domen taktowanych zegarem o mniejszej częstotliwości niż docelowa domena. Może być stosowany również do synchronizowania sygnałów, pochodzących z elementów całkowicie asynchronicznych.

Kluczowym elementem są dwa przerzutniki D, tworzące 2-bitową zmienną **Buffer** (linia #1). Zmienna ta jest wykorzystywana jako rejestr przesuwający. W linii #2 przesuwamy zerowy bit w miejsce pierwszego bitu, a w miejsce zerowego bitu umieszczany jest stan sygnału wejściowego **AsyncIn**. Bit pierwszy łączymy z wyjściem **SyncOut** za pomocą instrukcji **assign** (linia #3).

Układ powstały w wyniku syntezy tego kodu pokazano na **rysunku 8**. Składa się z zaledwie dwóch przerzutników D i nie ma żadnych elementów kombinacyjnych. Wejście przerzutnika **Buffer[0]** jest

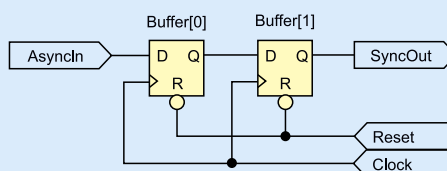
Listing 1. Prosty synchronizator

```
// Plik synchronizer.v
module Synchronizer(
    input Clock,
    input Reset,
    input AsyncIn,
    output SyncOut
);

    reg [1:0] Buffer; // #1

    always @(posedge Clock, negedge Reset) begin
        if (!Reset)
            Buffer <= 0;
        else
            Buffer[1:0] <= {Buffer[0], AsyncIn}; // #2
    end

    assign SyncOut = Buffer[1]; // #3
endmodule
```



Rysunek 8. Schemat powstały po syntezy kodu z listingu 1



Rysunek 9. Testowe przebiegi ilustrujące działanie synchronizatora

podłączone wprost do sygnału, który nie jest zsynchronizowany z zegarem **Clock**. W związku z tym na jego wyjściu może wystąpić stan metastabilny. Jednak nie stanowi to problemu, ponieważ jego wyjście jest połączone tylko z kolejnym przerzutnikiem **Buffer[1]**, który jest taktowany tym samym sygnałem zegarowym co przerzutnik zerowy. Nawet jeżeli wystąpi stan metastabilny, to nie zostanie przepuszczony dalej, bo odfiltruje go przerzutnik **Buffer[1]**.

Na **rysunku 9** pokazano przebiegi powstałe w wyniku symulacji synchronizatora. Na wejściu **AsyncIn** podawany jest sygnał, który nie jest zsynchronizowany z zegarem. Zbocza tego sygnału wypadają gdzieś w środku taktów zegara. Sygnał **SyncOut** odpowiada sygnałowi wejściowemu, ale przesuniętemu w taki sposób, by jego zbocza pokrywały się ze zboczami rosnącymi sygnału zegarowego.

Statyczna analiza czasowa

Wystarczy teorii – czas na ćwiczenia praktyczne! Utwórz nowy projekt i jako układ FPGA wybierz MachXO2-1200HC w obudowie TQFP100, a performance grade ustaw na 6. Dokładny part number tego FPGA to LCMXO2-1200HC-6TG100C. Kod z tego odcinka kursu możesz pobrać pod adresem: <https://tiny.pl/cchn6>.

Przeanalizujmy **listing 2**. Jest to prosty kod, który posłuży nam do zapoznania się z narzędziem do statycznej analizy czasowej. Układ ma 8-bitowe wejście **A** (linia #1) i 8-bitowe wyjście **Y** (linia #2). Porty nazwałem pojedynczymi literami, ponieważ omawiany układ nie pełni żadnej sensownej funkcji. W linii #3 tworzymy generator sygnału zegarowego, który już doskonale znamy, ale tym razem ustawiamy najwyższą możliwą częstotliwość, czyli 133 MHz. Następnie tworzymy dwa rejestry 8-bitowe, które będą służyć do synchronizacji wejścia **A** z domeną zegarową. W linii #6 przepisujemy stan wejścia **A** do rejestru **A_temp**, a w linii #7 kopiujemy wartość **A_temp** do **A_sync**. Operacje te są zawarte w jednym bloku **always**. Należy

Listing 2. Przykład kodu z dużym blokiem kombinacyjnym

```
// Plik top.v
module top(
    input Reset,
    input [7:0] A,
    output [7:0] Y
);

// Generator sygnału zegarowego
wire Clock;
OSCH #(
    .NOM_FREQ("133.00")
) OSCH_inst(
    .STDBY(1'b0),
    .OSC(Clock),
    .SESTDBY()
);

// Synchronizacja wejść
reg [7:0] A_temp;
reg [7:0] A_sync;
always @(posedge Clock, negedge Reset) begin
    if(!Reset) begin
        A_temp <= 0;
        A_sync <= 0;
    end else begin
        A_temp <= A;
        A_sync <= A_temp;
    end
end

// Czasochłonna operacja
reg [7:0] Temp;
always @(posedge Clock, negedge Reset) begin
    if(!Reset)
        Temp <= 0;
    else
        Temp <= (((A_sync ^ 8'b10101010) + 8'd123) * 8'd100) * A_sync + A_sync;
    end
end

assign Y = Temp;
endmodule
```

Listing 3. Przykład kodu z podziałem na mniejsze operacje

```
// Plik top.v
module top(
    input Reset,
    input [7:0] A,
    output [7:0] Y
);

// Generator sygnału zegarowego
wire Clock;
OSCH #(
    .NOM_FREQ("133.00")
) OSCH_inst(
    .STDBY(1'b0),
    .OSC(Clock),
    .SESTDBY()
);

// Synchronizacja wejść
reg [7:0] A_temp;
reg [7:0] A_sync;
always @(posedge Clock, negedge Reset) begin
    if(!Reset) begin
        A_temp <= 0;
        A_sync <= 0;
    end else begin
        A_temp <= A;
        A_sync <= A_temp;
    end
end

// Pipelinig
reg [7:0] Temp1;
reg [7:0] Temp2;
reg [7:0] Temp3;
reg [7:0] Temp4;
reg [7:0] Temp5;
always @(posedge Clock, negedge Reset) begin
    if(!Reset) begin
        Temp1 <= 0;
        Temp2 <= 0;
        Temp3 <= 0;
        Temp4 <= 0;
        Temp5 <= 0;
    end else begin
        Temp1 <= A_sync ^ 8'b10101010;
        Temp2 <= Temp1 + 8'd123;
        Temp3 <= Temp2 * 8'd100;
        Temp4 <= Temp3 * A_sync;
        Temp5 <= Temp4 + A_sync;
    end
end

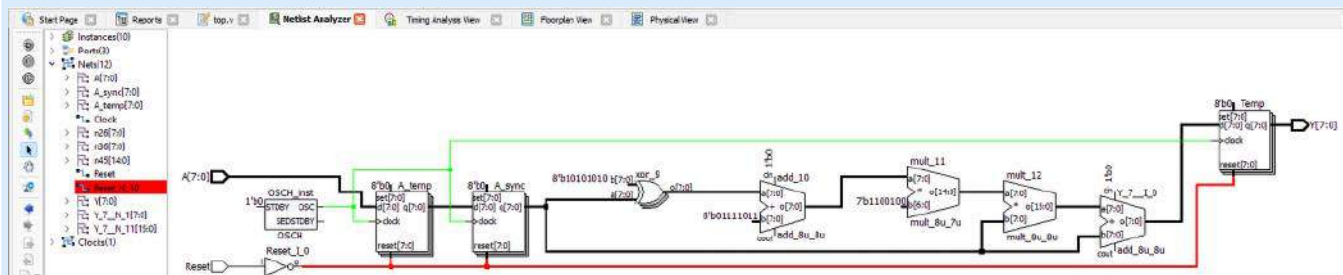
assign Y = Temp5;
endmodule
```

pamiętać, że linie #6 i #7 wykonują się jednocześnie w momencie wystąpienia zbocza rosnącego sygnału zegarowego. Oznacza to, że między zmianą stanu na wejściu **A** a pojawieniem się tej zmiany w **A_sync** mijają dwa takty zegarowe.

Rejestry **A_temp** i **A_sync** są syntezowane jako dwie grupy przerzutników D, a w każdej z nich jest 8 przerzutników. Zobacz **rysunek 10**, który przedstawia schemat powstały po syntezie tego kodu. Pomiędzy **A_temp** i **A_sync** może wystąpić stan metastabilny. Jednak nie stanowi to problemu, ponieważ zostanie odfiltrowany przez przerzutniki **A_sync** i do dalszych elementów układu przechodzą sygnały zsynchronizowane z sygnałem zegarowym.

W kolejnym bloku **always** mamy przykład czasochłonnej operacji. Operacja matematyczna, opisana w linii #9, skutkuje syntezą dużego układu kombinacyjnego. Składa się on z bramek XOR, sumatora dodającego stałą, multiplikatora (mnożarki) mnożącego przez stałą, multiplikatora mnożącego przez zmienną i finalnie sumatora dodającego zmienną. Wszystkie te elementy połączone są szeregowo. Widać je na **rysunku 10**.

W linii #8 tworzymy 8-bitowy rejestr **Temp**, który przechowuje wynik operacji, przeprowadzanej w linii #9. Wyjście tego rejestru podłączone jest bezpośrednio do wyjścia **Y** w linii #10 za pomocą instrukcji **assign**. Gdybyśmy pominęli ten



Rysunek 10. Schemat powstały po syntezy kodu z listingu 2

rejestr i wynik operacji z linii #9 skierowali bezpośrednio do wyjścia, wtedy mogłyby na nim wystąpić różne glitche. Zastosowanie dodatkowego rejestru, synchronizującego wynik bloku kombinacyjnego z zegarem, rozwiązuje ten problem.

Pozostaje nam jeszcze przygotować plik wymagań. W drzewku projektowym klikamy prawym przyciskiem myszy katalog **Synthesis Constraint Files**, po czym wybieramy **Add i New File**. Dodajemy plik **LDC** o dowolnej nazwie. W naszym przykładzie nazwałem go *timing.ldc*. Po utworzeniu pliku LDC pokaże się tabela, w której możemy zdefiniować różne wymagania odnośnie do zależności czasowych. Wypełniamy pierwszą linię tej tabeli tak, jak pokazano na **rysunku 11**. Niestety zamiast częstotliwości musimy podać okres zegara, więc na kalkulatorze musimy obliczyć odwrotność z $133 \cdot 10^6$, co daje wynik 7,518796 w nanosekundach. Zapisujemy plik i zamykamy go. Przejdź do okienka procesów i zaznacz **Place and Route Trace**. Następnie kliknij tę pozycję dwukrotnie, aby wykonać wszystkie niezbędne procesy. Nie ma potrzeby generować plików z bitstreamem.

Zwróć uwagę, że tym razem nie przypisujemy wejść i wyjść do fizycznych pinów FPGA w narzędziu **Spreadsheet**. Nie musimy tego robić. Syntezator sam przypisze piny w taki sposób, aby znaleźć jak najbardziej optymalne rozwiązanie.

Przejdźmy teraz do raportów. Zobacz **rysunek 12**. Przy raporcie **Place & Route** pojawił się znak ostrzeżenia. Otwórz raport **Place & Route Trace**. Czerwonym kolorem zaznaczone są wszystkie wymagania, jakich nie udało się spełnić. W naszym przypadku jest to żądanie, by zegar miał częstotliwość 133 MHz. Okazuje się, że w naszym układzie mamy 616 ścieżek kombinacyjnych, ale aż 50 z nich nie jest w stanie spełnić tego wymagania.

W dalszej części raportu mamy wyszczególnione wszystkie ścieżki, które nie spełniły wymogów. Raport jest bardzo rozbudowany i znajdziemy tam informacje nawet o tym, przez które slice'y w strukturze krzemowej będą poszczególnie sygnały. Niestety tak wygenerowany raport przytłacza mnogością szczegółów i jest nieczytelny. Z pomocą przychodzi narzędzie **Timing Analysis View**, które znajdziesz w menu **Tools**. Otwórz się okno, które pokazano na **rysunku 13**.

Największą częstotliwość zegara, jaką można zastosować, to 119,3 MHz. Wynika

to z faktu, że czas propagacji w najwolniejszej ścieżce to 8,328 ns (**rysunek 13**, kolumna **Arrival**, pierwszy wiersz). Odwrotność czasu propagacji najwolniejszej ścieżki pozwala obliczyć maksymalną częstotliwość sygnału zegarowego.

Enable	Source	Clock Name	reform High Edge	reform Low Edge	Period(ns)
1		Clock			7.518796
2		Clock			

Rysunek 11. Arkusz wymagań czasowych ze zdefiniowanym okresem sygnału zegarowego

Place & Route Trace Report

Loading design for application trace from file: kuzai1_imp1.msd.
 Design name: 500
 WCD Version: 3.3
 Vendor: LATTICE
 Device: L6000-1200HC
 Package: TQFP100
 Performance: 0
 Loading device for application trace from file: 'xilinx1000000' in environment: D:/Lattice/diagn
 Package status: Final Version 1.44.
 Performance Hardware Data Status: Final Version 36.4.
 Setup and Hold Report

Lattice TRACE Report - Backup, Version Diamond (64-bit) 2.12.1.454
 Sat Mar 11 12:12:09 2023

Copyright (c) 1991-1994 by Lattice Semiconductor Corporation. All rights reserved.
 Copyright (c) 1995-2001 Lattice Semiconductor Corporation. All rights reserved.
 Copyright (c) 2001-2002 Lattice Semiconductor Corporation. All rights reserved.
 Copyright (c) 2002-2003 Lattice Semiconductor Corporation. All rights reserved.

Report Information

Command line: tmap -v 10 -gn -astid m -sp 0 -astid m -m kuzai1_imp1.msd -xai -waguet D:/L
 Design file: kuzai1_imp1.msd
 File name: kuzai1_imp1.rpt
 Device speed: L6000-1200HC-0
 Report level: verbose report, limited to 10 items per preference

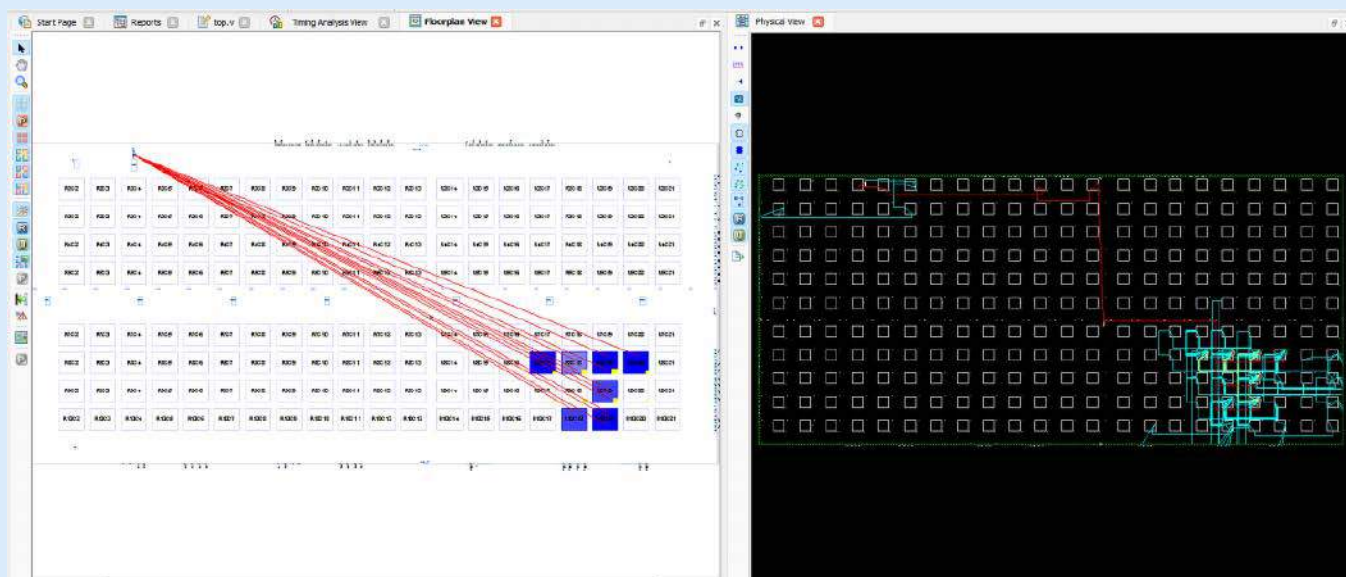
Performance Summary

* **FREQUENCY MET "Clock" 133.614099 MHz (90.427000ns)**
 If it items scored, 50 timing errors detected.
 Warning: 133.600000 is the maximum frequency for this performance.

Rysunek 12. Przykład raportu z błędem czasowym

Source	Destination	Weighted Slack	Arrival	Required	Data Delay	Route %	Level	Clk Skew	Setup/Hold	Other	Color
1	A_sync_0	0.000	0.000	0.000	0.000	0.000	0	0.000	0.000	0.000	Green
2	A_sync_0	0.000	0.000	0.000	0.000	0.000	0	0.000	0.000	0.000	Green
3	A_sync_0	0.000	0.000	0.000	0.000	0.000	0	0.000	0.000	0.000	Green
4	A_sync_0	0.000	0.000	0.000	0.000	0.000	0	0.000	0.000	0.000	Green
5	A_sync_0	0.000	0.000	0.000	0.000	0.000	0	0.000	0.000	0.000	Green
6	A_sync_0	0.000	0.000	0.000	0.000	0.000	0	0.000	0.000	0.000	Green
7	A_sync_0	0.000	0.000	0.000	0.000	0.000	0	0.000	0.000	0.000	Green
8	A_sync_0	0.000	0.000	0.000	0.000	0.000	0	0.000	0.000	0.000	Green
9	A_sync_0	0.000	0.000	0.000	0.000	0.000	0	0.000	0.000	0.000	Green
10	A_sync_0	0.000	0.000	0.000	0.000	0.000	0	0.000	0.000	0.000	Green

Rysunek 13. Timing Analysis View



Rysunek 14. Floorplan View i Physical View z sygnałem zegarowym zaznaczonym na czerwono

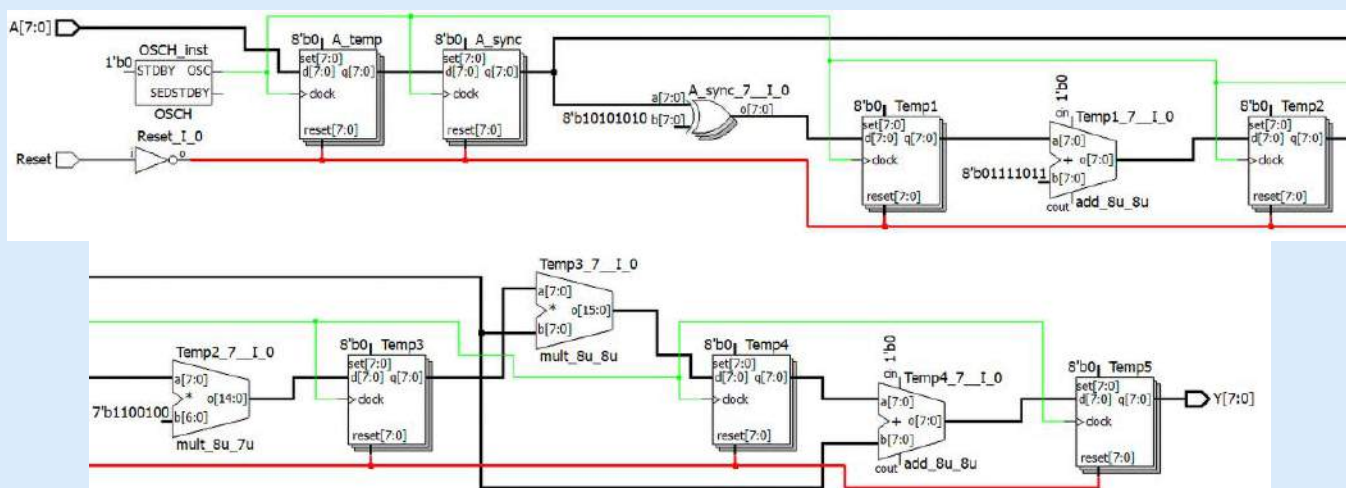
W pierwszej kolejności zmodyfikujemy parametr **Worst-Case Paths**, który domyślnie ustawiony jest na 10. Znajdziesz w lewej górnej części okna. Ten parametr określa liczbę najgorszych ścieżek, które są wyświetlane. Proponuję zmienić go na 200 – wtedy będziemy widzieć wszystkie 50 zbyt wolnych ścieżek, a także 150 ścieżek spełniających wymagania. W lewej dolnej części mamy raporty utworzone dla każdego z wymagań, jakie zostało określone w pliku LDC. Ustawiliśmy tam tylko wymaganie odnośnie do częstotliwości sygnału zegarowego, więc program obliczył tylko zależności czasowe dla zegara i sprawdził poprawność setup time i hold time. Raporty zawierające błędy zaznaczone są na czerwono. Kliknij raport **FREQUENCY NET „Clock” – setup**. W prawej górnej części okna pojawiły się wszystkie przeanalizowane ścieżki w blokach kombinacyjnych. Domyślnie posortowane są według **Weighted Slack** od najmniejszych do największych. Jeżeli slack jest ujemny, to znaczy, że czas propagacji analizowanej ścieżki jest dłuższy niż okres sygnału zegarowego. Taka sytuacja jest nieprawidłowa i wszystkie ujemne wyniki zaznaczone są kolorem czerwonym.

Interesować nas będą kolumny **Source** i **Destination**. Są w nich zawarte nazwy przerzutnika nadającego i odbierającego. Nazwy te pochodzą z naszego kodu w Verilogu, lecz program dokleja do nich różne cyferki, jeżeli są to rejestry wielobitowe. Widzimy, że wszystkie przekroczenia dotyczą sygnałów wychodzących z **A_sync** i wchodzących do **Temp**. Nic dziwnego – jest to jedyny blok kombinacyjny w naszym kodzie. Narzędzie **Timing Analysis View** pozwala nam

zlokalizować, gdzie są najdłuższe ścieżki kombinacyjne. To właśnie od nich powinniśmy rozpocząć optymalizację, aby możliwe było spełnienie wymagań czasowych.

Możemy zobaczyć też kilka ciekawych rzeczy. Kliknij prawym przyciskiem myszy na dowolną ścieżkę na liście, a następnie wybierz **Show in Physical View**. Zostanie pokazana struktura krzemowa z zaznaczonymi różnymi komponentami, jak slice, piny, ścieżki, itp. (nie wszystkie są domyślnie włączone – możesz włączać lub wyłączać różne obiekty, klikając przyciski w pionowym pasku narzędzi po lewej stronie okna).

Wróć do **Timing Analysis View** i kliknij tę samą ścieżkę prawym przyciskiem myszy i następnie wybierz **Show in Floorplan View**. Zobaczymy znów strukturę scaloną, ale narysowaną w sposób bardziej abstrakcyjny i nieco bardziej czytelny. Możemy wyświetlić oba widoki jednocześnie. Kliknij prawym przyciskiem myszy na zakładkę **Physical View** (na górze) i wybierz **Split Tab Group**. Okno zostanie podzielone w taki sposób, jak pokazano na **rysunku 14**. Każdy element, który klikniesz na jednym widoku, będzie automatycznie podświetlony też na drugim. Kliknij na tło w **Physical View**, a następnie naciśnij CTRL-F. Pokaże się okienko umożliwiające wyszukiwanie różnych obiektów. W polu **Find what** wpisz „Clock”, a w **Find type** ustaw „Net” i następnie kliknij **Select All**. Zostaną podświetlone wszystkie połączenia, które zostały wykorzystane do rozprowadzenia sygnału zegarowego **Clock**. Wychodzi on z generatora **OSCH** i jest rozprowadzany do wszystkich przerzutników. Na **rysunku 14** sygnał zegarowy



Rysunek 15. Schemat powstały po syntezie kodu z listingu 3

został zaznaczony na czerwono. Narzędzia te dają dość duże możliwości analizowania wewnętrznych zasobów na strukturze układu FPGA. Możemy tam znaleźć wszystkie sprzętowe peryferia, takie jak OSCH, PUR, TSALL, bloki pamięci EBR i całą masę innych rzeczy, które dopiero poznamy w przyszłości. Gorąco zachęcam, by pobrać się tymi narzędziami i poprzeglądać sobie wnętrza układu FPGA. Można zobaczyć, co siedzi wewnątrz każdego używanego elementu. W tym celu należy go kliknąć prawym przyciskiem myszy i wybrać **Logic Block View**.

Wróćmy jednak do optymalizacji czasowej. Musimy w jakiś sposób rozbić duży blok kombinacyjny, który znajduje się pomiędzy rejestrami **A_sync** i **Temp**. Jedną z możliwości jest pipelining.

Pipelining

Jest to metoda polegająca na dzieleniu dużego bloku kombinacyjnego na kilka mniejszych, rozdzielonych przerzutnikami D. Dzięki temu, że bloki kombinacyjne stają się mniejsze, ich czas propagacji jest krótszy. W rezultacie możemy zwiększyć częstotliwość sygnału zegarowego.

Spróbujmy przekształcić kod z **listingu 2**, aby jego funkcjonalność pozostała dokładnie taka sama, lecz z podziałem na mniejsze bloki kombinacyjne. Zobacz kod przedstawiony na **listingu 3**. Operacja z linii #9 listingu 2 składa się z pięciu operacji składowych. Z tego powodu na **listingu 3** w linii #1 i kolejnych przygotowujemy pięć rejestrów 8-bitowych. W linii #2 oraz kolejnych mamy wszystkie operacje, jakie dotychczas były wykonywane w jednym bloku kombinacyjnym. Jednak tym razem wynik każdej operacji zapisywany jest do kolejnego rejestru, który staje się argumentem kolejnej operacji.

Uruchom syntezę, a następnie otwórz **Netlist Analyzer**. Schemat powstały w wyniku syntezy pokazano na **rysunku 15**. Widzimy, że pomiędzy sumatorami i mnożarkami pojawiły się bloki przerzutników D, taktowane tym samym sygnałem zegarowym i wykorzystujące ten sam sygnał resetujący. Otwórz raport **Place & Route Trace**. Tym razem zobaczymy komunikat „167.870MHz is the maximum frequency for this preference”. Wymaganie co do częstotliwości zegara równej 133 MHz zostało spełnione.

Zobaczmy, co tym razem pokaże narzędzie Timing Analysis View. Na liście 200 ścieżek z najdłuższym czasem propagacji znajdują się głównie dwa rodzaje sygnałów. Są to:

- Biegające od rejestru **Temp3** do **Temp4**,
- Biegające od **A_sync** do **Temp4**.

Okazuje się, że odpowiedzialna za to jest operacja mnożenia dwóch zmiennych **Temp3** i **A_sync**. Gdybyśmy chcieli zoptymalizować tę operację, można by się pokusić o wykorzystanie gotowych bloków DSP, które przyspieszają operacje matematyczne. Możesz je znaleźć w programie **IP Express**, które omawialiśmy w 5 odcinku kursu.

Porównajmy zapotrzebowanie na zasoby w przypadku obu implementacji. Wyniki pokazano w **tabeli 5**. Zwiększyliśmy częstotliwość sygnału zegarowego, jednak ma to swoją cenę. Kod w wersji z pipeliningiem potrzebuje dwukrotnie więcej przerzutników. Jest też inna kwestia, o której nie wolno zapomnieć. W pierwszej wersji operacja matematyczna wykonywała się w jednym cyklu zegarowym przy maksymalnej częstotliwości 119 MHz. Druga wersja co prawda może działać z częstotliwością 168 MHz, ale potrzebuje pięciu cykli zegarowych. Jeżeli przeliczymy to na sumaryczny czas operacji, to operacja z pierwszego kodu zajmie 8,4 ns, a z zastosowaniem pipeliningu będziemy musieli poczekać 29,8 ns.

Tabela 5. Porównanie obu implementacji

Implementacja	Max. częstotliwość	Przerzutniki	Slice	LUT4
Combo	119 MHz	24	24	48
Pipelining	168 MHz	47	29	55

Czy to znaczy, że spowolniliśmy nasz projekt? Wszystko zależy od punktu widzenia, bowiem układ FPGA jest tak szybki, jak jego najwolniejszy blok kombinacyjny. Gdyby najważniejszym założeniem projektowym było jak najszybsze wykonywanie tej operacji matematycznej, to kod bez pipeliningu byłby właściwym rozwiązaniem. Jednak mogłoby się zdarzyć, że operacja matematyczna, omawiana w naszych przykładach, jest tylko mało istotną, poboczną operacją, która obniża nam częstotliwość zegara. W efekcie przez nią wszystkie inne układy muszą być taktowane takim zegarem, aby długi blok kombinacyjny dał radę przetworzyć sygnały w jednym takcie zegara. W takiej sytuacji rozbięcie dużego bloku kombinacyjnego na kilka mniejszych, ale wykonywanych w kilku taktach zegara, pozwoli przyspieszyć inne, bardziej istotne funkcjonalności.

Można na to spojrzeć jeszcze z innej strony. Pipelining ma świetne zastosowanie wszędzie tam, gdzie dane na wejściu zmieniają się z każdym taktem zegara. Przykładem takiej sytuacji jest lista instrukcji programu, które są odczytywane z pamięci i przetwarzane przez procesor. Jest to sprawa dość skomplikowana, więc często jest dzielona na kilka prostszych operacji. Pipelining w procesorach polega na tym, że instrukcja jest pobierana z pamięci w czasie, kiedy poprzednia instrukcja jest jeszcze w trakcie wykonywania. Można także dodać kilka etapów pośrednich. Więcej na ten temat możesz przeczytać pod adresem [3].

Reset synchroniczny czy asynchroniczny?

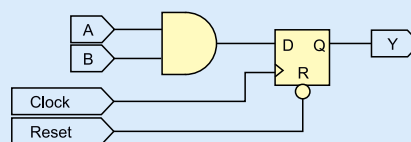
Wróćmy jeszcze do resetowania przerzutników. Sposób, w jaki ustawiamy wartość początkową przerzutników, ma również wpływ na czas propagacji i tym samym na maksymalną częstotliwość sygnału zegarowego.

Weźmy pod lupę **listing 4**. Jest to przykład banalnego kodu. Blok **always** reaguje na zbocze rosnące zegara i wtedy do przerzutnika D wpisywany jest wynik, podawany przez bramkę AND. Oprócz tego blok **always** reaguje także na zbocze opadające sygnału **Reset** i wtedy natychmiast ustawia wyjście w stan niski. Przerzutnik pozostaje w stanie niskim i ignoruje sygnał zegarowy tak długo, jak reset jest aktywny.

W wyniku syntezy otrzymujemy zaledwie dwa elementy – przerzutnik i bramkę AND. Zobacz **rysunek 16**. Wyjście bramki AND idzie prosto do wejścia D przerzutnika, a sygnał **Reset** połączony jest z dedykowanym wejściem zerującym w przerzutniku. Zobaczmy teraz kod z **listingu 5**. Jedyna różnica polega na tym, że z listy czułości bloku **always** zniknęło wyrażenie **negedge Reset**. W takiej sytuacji przerzutnik wyzeruje się tylko wtedy, kiedy sygnał **Reset** ma stan niski w momencie wystąpienia zbocza rosnącego sygnału **Clock**. Pomiędzy zboczami zegara **Reset** może się zmieniać zupełnie

Listing 4. Przykład bloku always z resetem asynchronicznym

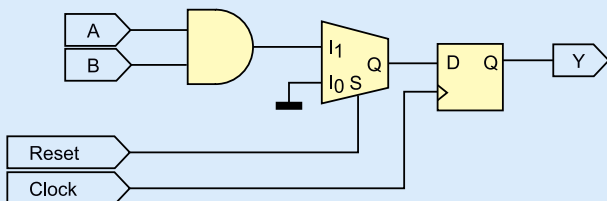
```
always @(posedge Clock, negedge Reset) begin
    if (!Reset)
        Y <= 0;
    else
        Y <= A & B;
end
```



Rysunek 16. Schemat powstały w wyniku syntezy kodu z listingu 4

Listing 5. Przykład bloku always z resetem synchronicznym

```
always @(posedge Clock) begin
    if (!Reset)
        Y <= 0;
    else
        Y <= A & B;
end
```

Rysunek 17. Schemat powstający w wyniku syntezy kodu z listingu 5

Tabela 6. Porównanie wydajności bloków always z resetem synchronicznym i asynchronicznym

Wersja kodu	Czas propagacji
z resetem asynchronicznym	1,498 ns
z resetem synchronicznym	1,669 ns

dowolnie i nie będzie to miało wpływu na wyjście układu (oczywiście z zachowaniem hold time i setup time).

Kod wydaje się krótszy i prostszy, ale w rzeczywistości prowadzi do syntezy bardziej złożonego bloku kombinacyjnego! Zobaczmy **rysunek 17**. Okazuje się, że pomiędzy bramką AND i przerzutnikiem pojawił się multiplexer! Jest on sterowany sygnałem **Reset**. Kiedy sygnał resetujący jest w stanie wysokim (czyli podczas normalnej pracy), multiplexer łączy wejście I1 z wyjściem Q. Natomiast kiedy **Reset** jest w stanie niskim, to multiplexer łączy wejście I0 z wyjściem Q, a z kolei wejście I0 jest na stałe połączone z masą, czyli stanem niskim. W taki sposób stan niski trafia na wejście przerzutnika i jest zapamiętywany w momencie wystąpienia zbocza rosnącego sygnału **Clock**. Dodatkowy multiplexer ma oczywiście wpływ na czas propagacji. Porównanie osiągnięć obu rozwiązań pokazano w **tabeli 6**.

Performance grade w MachXO2

Rodzina układów MachXO2 jest bardzo rozbudowana. Wybierając układ FPGA do naszego projektu, oczywiście musimy kierować się liczbą bloków LUT, jaka nam jest potrzebna, ale także szybkością. Zobacz **rysunek 18**, gdzie pokazany jest schemat nazewnictwa wszystkich układów z rodziny MachXO2. Pierwszy człon nazwy to oznaczenie LCMXO2, wspólne dla wszystkich układów rodziny MachXO2. Następnie, po myślniku, podaje się zakręgloną liczbę elementów LUT, dostępnych w układzie. Dalej mamy dwie literki określające wydajność i zasilanie. Możliwe są trzy kombinacje:

- **HC** – układy wysokiej wydajności, rdzeń jest zasilany napięciem 3,3 V lub 2,5 V,
- **HE** – układy wysokiej wydajności, rdzeń jest zasilany napięciem 1,2 V,
- **ZE** – układy energooszczędne, rdzeń jest zasilany napięciem 1,2 V.

Kolejna cyfra po oznaczeniu wydajności i napięcia zasilania określa szybkość układu. W Lattice Diamond nazywana jest **Performance Grade** w oknie, gdzie wybiera się układ FPGA stosowany w projekcie. Numer 1 oznacza układy najwolniejsze, a 6 to najszybsze. Dla układów serii ZE dostępne są Performance Grade od 1 do 3. Natomiast dla układów HC i HE możliwe są prędkości od 4 do 6.

Porównanie szybkości poszczególnych FPGA

Firma Lattice podaje w swoich datasheetach maksymalną częstotliwość, obliczoną dla kilku specyficznych aplikacji. W każdym Family Datasheet znajdziemy rozdział o nazwie Typical Building Block Function Performance. W **tabeli 7** porównano, z jakim zegarem mogą pracować różne układy FPGA firmy Lattice Semiconductor.

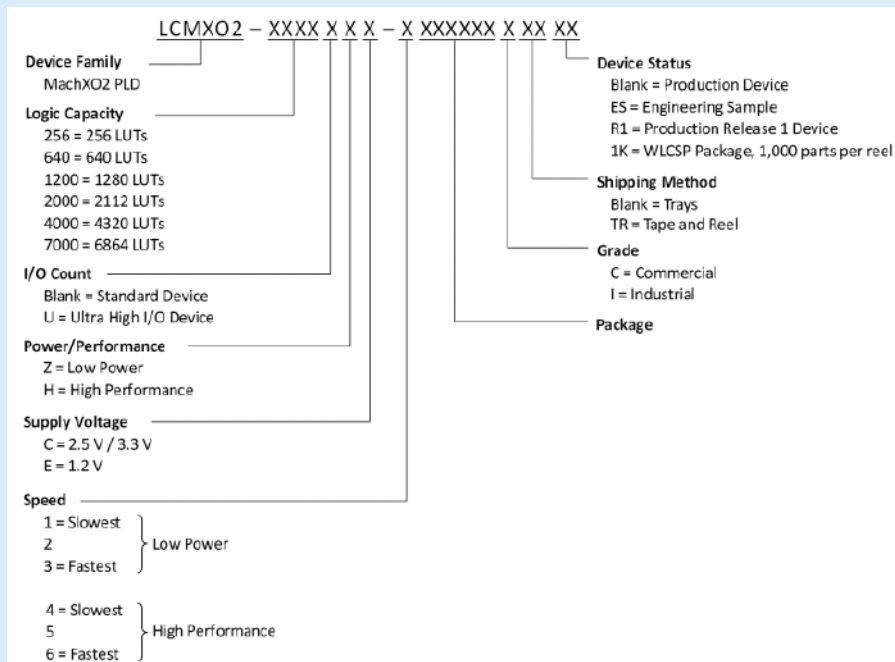
To wszystko na dziś! Pomimo że ten odcinek kursu jest najdłuższy ze wszystkich dotychczasowych, należy go traktować jedynie jako wstęp do tematyki statycznej analizy czasowej. Skupiliśmy się głównie na częstotliwości sygnału zegarowego, ale w pliku wymagań czasowych możemy zdefiniować jeszcze dużo różnych parametrów. Narzędzie do analizy czasowej ma jeszcze mnóstwo opcji, o których nawet nie wspominałem. W kolejnym odcinku kursu wrócimy do tematu symulacji i poznamy symulator Icarus Verilog.

Dominik Bieczyński
leonow32@gmail.com

Linki:

1. <https://tiny.pl/cchmc>
2. <https://www.edaplayground.com/x/PRaX>
3. https://en.wikipedia.org/wiki/Classic_RISC_pipeline

Repozytorium modułów wykorzystywanych w kursie:
<https://github.com/leonow32/verilog-fpga>



Rysunek 18. Nazewnictwo rodziny układów MachXO2

Tabela 7. Porównanie maksymalnej częstotliwości zegara różnych układów FPGA firmy Lattice

	iCE40 UltraPlus	iCE40 LP	iCE40 HX	MachXO2	ECP5U	ECP2	XP2
16:1 Mux	110	190	305	412	441	660	616
16-bitowy sumator	100	160	220	297	441	500	507
16-bitowy licznik	100	170	255	324	384	488	541
64-bitowy licznik	40	?	105	161	263	260	321

koktajl niusów



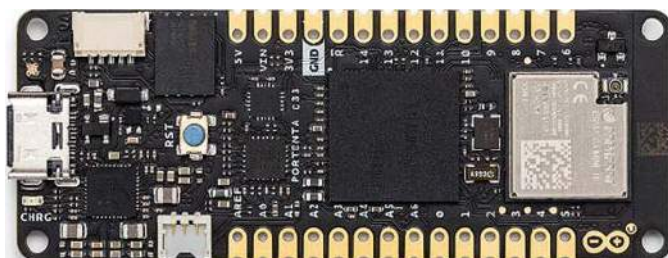
Wodorowe ogniwa paliwowe firmy Toyota będą napędzać bezemisyjny transport morski

Urządzenie Corvus Pelican Fuel Cell System (FCS) zawiera cztery moduły ogniw paliwowych firmy Toyota. To rozwiązanie ma zrewolucjonizować transport morski za sprawą bezpiecznej, czystej i wytwarzanej w wydajny sposób energii. Urządzenie waży 3100 kg i ma rozmiary: 2,3×1,4×2,1 m. Łączna moc całego systemu wynosi 340 kW. Średnie zużycie sprężonego wodoru zależy bezpośrednio od typu statku, w którym ma zostać zastosowane urządzenie.

Ogniwa paliwowe firmy Toyota potwierdziły już własną niezawodność. Teraz po drogach porusza się 20 tysięcy samochodów korzystających z tych ogniw. Przyszła kolej na transport morski. W trakcie liczącego 2 lata procesu rozwojowego pracowano nad obudową odporną na szeroki zakres ciśnień, która zapewnia najwyższy poziom bezpieczeństwa w czasie żegluga. Modułowa konstrukcja gwarantuje gazoszczelność w każdych warunkach. Nie wymaga stosowania systemu bezpieczeństwa i wentylacji, a integracja ogniw paliwowych z kadłubem statku jest całkowicie bezproblemowa.

Wodorowe ogniwa paliwowe firmy Toyota odgrywają znaczącą rolę w procesie dekarbonizacji branży transportowej. Tak twierdzi dyrektor jednostki biznesowej ds. ogniw paliwowych w firmie Toyota, Thiebault Paquet, który powiedział, że: „Nasza wielotechnologiczna strategia dotyczy nie tylko transportu drogowego, ale także z całą pewnością morskiego. Ważne jest holistyczne podejście, które pozwala połączyć podaż i popyt na wodór w ekonomicznie opłacalnych ekosystemach. Liczy się tutaj żegluga z portami i pozostałymi uwarunkowaniami”.

<https://tiny.pl/cc4r3>



Nowy komputer jednopłytkowy Arduino Portenta C33 dostępny w Farnell

Ta funkcjonalna platforma sprzętowa stanowi dobry wybór dla przemysłowych aplikacji czasu rzeczywistego. Jest to możliwe dzięki wydajnemu mikrokontrolerowi R7FA6M5BH2CBG od Renesas Electronics

z rdzeniem ARM Cortex-M33, taktowanego do 200 MHz, wyposażonego w pamięć SRAM o pojemności 512 kB i Flash o pojemności 2 MB. Dodatkowo zawiera pamięć zewnętrzną QSPI Flash o pojemności 16 MB oraz układ kryptograficzny SE050C2 od NXP Semiconductors.

Temperatura robocza Arduino Portenta C33 zawiera się w przedziale od -40 do 85°C. Do programowania oraz zasilania komputera służy gniazdo USB-C. Obsługiwane są interfejsy CAN, SD, SPI, GPIO, I²C, I²S i JTAG/SWD. Opisany komputer można również zasiląć z baterii litowo-polimerowej o napięciu 3,7 V. Oferuje 78 wyprowadzeń cyfrowych, 10 analogowych i 8 PWM. Wymiary płytki to 6,6×2,54 cm.

Dzięki komunikacji Wi-Fi oraz Bluetooth Low Energy (BLE) jednopłytkowy komputer Arduino Portenta C33 nadaje się do rozwiązań IoT, systemów zdalnego sterowania czy też sterowania flotami pojazdów. Może przeprowadzać aktualizacje oprogramowania układowego z pomocą Arduino Cloud czy innych usług hostingowych. Umożliwia szybkie prototypowanie dzięki obsłudze języka programowania MicroPython i dostępności wielu gotowych bibliotek, oprogramowania albo szkiców Arduino.

<https://tiny.pl/cc49t>



Firma Lapp Group zastosowała bioplastyczne tworzywo sztuczne firmy BASF

Oparcowany dla wielu zastosowań przemysłowych kabel ethernetowy ETHERLINE FD P Cat.5e sprawdza się m.in. w przewodnicach łańcuchowych. Teraz jest dostępny również w wersji zrównoważonej – termoplastyczny poliuretan (TPU) z surowców kopalnych, na wierzchu kabla, zastąpiono materiałem bazującym na surowcach odnawialnych. Nowa wersja kabla zawiera Elastollan N – biopolimer firmy BASF wytworzony z surowca na bazie kukurydzy. Udział surowca odnawialnego wynosi od 45% do 60%. Przytoczony biopolimer oferuje taką samą trwałość, elastyczność i właściwości mechaniczne oraz odporność na hydrolizę, chemikalia, a także promieniowanie UV, co zwykły elastollan. Nawet parametry przetwarzalności są zachowane. Ilość biopolimeru można precyzyjnie odmierzać według normy ASTM D 6866. Jak wyjaśnia menadżer ds. sprzedaży w firmie BASF, Oliver Mühren: „Nasze biopochodne TPU w niczym nie odstają od produktów wykonanych z konwencjonalnych paliw kopalnych. To słuszny krok, by zapewnić prawdziwą wartość dodaną, dzięki zrównoważonemu produktowi”.

<https://lapppoland.lappgroup.com/nawosci/prasa/prasa.html>

Rakieta AGH Space Systems czwarta na świecie

W czerwcu 2023 roku uczestnicy sekcji rakiet AGH Space Systems brali udział w największym na świecie międzuczelnianym konkursie inżynierii raketowej w USA. Zawody odbywały się na pustyni w Las Cruces w stanie Nowy Meksyk. Według pełnej klasyfikacji generalnej konkursu rakiet hybrydowa 3-TTK+ z AGH uplasowała się na czwartym miejscu, w kategorii lotu na 10 tysięcy stóp wysokości. Rakietę napędzającą paliwem hybrydowym to wyjątkowo skomplikowana konstrukcja. Dlatego pierwszy start i lot rakiet 3-TTK+ wzbudził w członkach AGH Space Systems spore emocje i nieco nerwów. Zespół ciężko pracował przez wiele miesięcy nad tym, żeby



wszystko poszło sprawnie. Sędziowie w Las Cruces wyrazili liczne słowa uznania pod adresem budowy rakiety i podzielili się cennymi wskazówkami. Teraz członkowie AGH Space Systems szykują się do następnych zawodów raketowych. Tym razem z Turbulencją, tzn. pierwszą w Polsce studencką raketą na paliwo ciekłe. Pojawia się ona na zawodach European Rocketry Challenge 2023. Będzie to w październiku 2023 roku. Przytoczone zawody odbędą się w Portugalii.

<https://tiny.pl/cc498>



Centrum Testowania Oprogramowania LG Electronics uzyskało najnowszą akredytację TÜV Rheinland

Centrum Testowania Oprogramowania LG Electronics zostało oficjalnie akredytowane przez TÜV Rheinland – światowego lidera w zakresie usług testowania, inspekcji i certyfikacji. Firma LG Electronics korzysta z międzynarodowych testów (IEC 60730-1 i IEC 60335-1), żeby ocenić oprogramowanie swoich inteligentnych urządzeń AGD. Dzięki nim firma może w pełni sprawdzić, czy spełniają one wymagania bezpieczeństwa. To pomaga znacznie skrócić czas dostarczenia klientom bezpiecznych oraz wysokiej jakości produktów. Firma podziwia się, że akredytacja w zakresie testowania oprogramowania ulepszy konkurencyjność urządzeń AGD – klienci będą mieli większe zaufanie do niezawodności i jakości zaawansowanych rozwiązań dla gospodarstw domowych.

LG Electronics stale rozwija swą ofertę, zwiększając możliwości oprogramowania – we wszystkich obszarach biznesowych. Już w czerwcu 2022 roku Centrum Oprogramowania zostało akredytowane przez TÜV Rheinland, stając się dosyć wiarygodnym, profesjonalnym ośrodkiem oceniającym bezpieczeństwo oprogramowania motoryzacyjnego i nie tylko. Wcześniej otrzymało akredytację Korea Laboratory Accreditation Scheme (KOLAS), dla oceny oprogramowania pojazdów autonomicznych i ustalania jakości produktów zarówno elektrycznych, jak i elektronicznych.

<https://tiny.pl/cc49b>

VCBL-HD – uchwyt uniwersalny firmy Schmalz do obróbki drewna

Niezależnie od tego, czy chodzi o cienkie i delikatne części składowe, czy o ciężkie, lite drewno o chropowatej powierzchni – nowy blok podciśnieniowy firmy Schmalz ma bezpiecznie utrzymywać



wszystkie elementy, podczas gdy maszyna CNC dokładnie frezuje ich kształty. VCBL-HD pozwala zatem usprawnić czasochłonne procesy i metody konfiguracji oraz przyspieszyć obróbkę drewna.

Firma Schmalz opracowała bezwężowy, podciśnieniowy system mocowania przewidziany dla maszyn CNC. System VCBL-HD zaciśnięta się na jedno- i dwuobwodowych stołach z konsolami i gwarantuje rewelacyjne efekty frezowania, dzięki dużej sile trzymania i absorpcji siły bocznej. Nowy blok podciśnieniowy łączy w sobie zalety dotychczasowych wersji VCBL – jego korpus wykonany z aluminium zapewnia wystarczającą stabilność głównie w przypadku ciężkich arkuszy oraz sporej siły cięcia. Dzięki adaptacyjnej uszczelce jego wymienna płyta ssąca bezpiecznie mocuje nawet powierzchnie o nie-standardowej konstrukcji. Mimo że przysawka oferuje wystarczającą siłę mocowania dla litego drewna, litej powierzchni czy też materiałów laminowanych, jej powierzchnia styku jest wystarczająco delikatna, aby zaciskanie cienkich, i co ważne, delikatnych powierzchni nie powodowało uszkodzeń. VCBL-HD pozwala zaoszczędzić sporo czasu przy ustawianiu. W celu zapewnienia dogodnej elastyczności firma Schmalz osiągnęła obrócone płyty ssące. Dzięki temu blok podciśnieniowy może adaptować się do zadanych geometrii przedmiotu obrabianego. Wbudowany tag NFC oraz związana z nim komunikacja ujawniają wszystkie kluczowe dane dotyczące produktu, np. numer części zamiennych, za pośrednictwem aplikacji Schmalz ControlRoom.

<https://tiny.pl/cc4wq>

Powerbank ZR-04 firmy RADWAG

Powerbank ZR-04 przeznaczony jest do akumulatorowego zasilania wag i akcesoriów do wag. Korzysta z technologii szybkiego ładowania i gwarantuje długi czas pracy, bez konieczności ładowania. Dzięki funkcji L-low Power zapewniane jest zasilanie urządzeń o niskim obciążeniu. Powerbank ZR-04 został wykonany z wysokiej jakości komponentów, które zapewniają niezawodną obsługę. Wyświetlacz, który prezentuje stan naładowania ZR-04, istotnie poprawia ergonomię pracy.

Urządzenie ma w sumie 7 zabezpieczeń: przed przeładowaniem, wyładowaniem albo też przepięciami, przegrzaniem, prądem przetężeniowym, krótkimi spięciami oraz polem elektromagnetycznym – dowodzą tego certyfikaty bezpieczeństwa. Powerbank ZR-04 spełnia standardy linii lotniczych – może być przewożony na pokładzie samolotów jako bagaż podręczny. Na czele listy atrybutów powerbanka ZR-04 znajduje się m.in. jego dwukierunkowość działania. Urządzenie może jednocześnie zasilać wagę, a także ładować się. Żeby naładować powerbank, wystarczy użyć dowolnej ładowarki do telefonu komórkowego. Dość optymalnym rozwiązaniem jest ładowarka z funkcją szybkiego ładowania, która przyspiesza proces.

<https://radwag.com/pl/powerbank-zr-04,w1,3KF,103-148-119>





PD20 – kompaktowy konwerter USB/RS-485 firmy Lumel

Jest to przenośne urządzenie przeznaczone do przetwarzania sygnałów z interfejsu USB na interfejs RS-485 i zawiera separację galwaniczną. Maksymalna przepustowość transmisji wynosi 115,2 kb/s. Obudowa konwertera PD20 spełnia wymogi normy szczelności IP40. Konwerter nie ingeruje w strukturę przesyłanych danych i jest zgodny z protokołami komunikacji przemysłowej.

Urządzenie jest zasilane napięciem stałym 5 V poprzez złącze USB-A, które jednocześnie umożliwia komunikację z komputerami zgodnie ze standardem USB 1.1 i USB 2.0. W systemie operacyjnym urządzenie działa jak wirtualny port szeregowy COM z transmisją danych w trybie half-duplex.

Temperatura robocza PD20 zawiera się w przedziale 0...40°C, dopuszczalna wilgotność względna wynosi 85% (bez kondensacji pary wodnej). Pobór mocy nie przekracza 1,5 W, a wymiary to 7,6×2,5×2 cm. Napięcie przebicia separacji galwanicznej wynosi 2,2 kV. Produkt spełnia trzy normy: PN-EN 61000 oraz PN-EN 61010 i PN-EN 61326.

<https://www.lumel.com.pl/katalog/produkt/konwerter-usb-rs-485>



Migracja automatyki WAGO do CODESYS 3.5

Przeważnie produkty firmy WAGO charakteryzują się dużym okresem dostępności na rynku. Postęp techniki wymusza jednak ciągłe udoskonalenia. Z tego właśnie powodu firma podjęła decyzję o rozpoczęciu procesu wycofywania ze swej oferty, środowiska e!COCKPIT, umożliwiając jednocześnie korzystanie z produktów automatyki WAGO w środowisku CODESYS 3.5. Środowisko CODESYS 3.5 uważane jest za wzorcową implementację normy IEC-61131. Jest to niepowtarzalne rozwiązanie, które cieszy się uznaniem za jakość i udostępnia istotne funkcje dla innowacyjnej automatyki.

Gwarantowana przez CODESYS 3.5 niezależność od dostawcy czy interoperacyjność sprzyjają zróżnicowanym możliwościom współpracy pomiędzy systemami i urządzeniami. Firma WAGO rozpoczęła proces dostosowywania programowalnych urządzeń automatyki (PFC200, CC100, BC100, EDGE oraz TP600) do pracy we wspomnianym środowisku. Dzięki zgodności rozwiązań systemowych istnieje możliwość, żeby w bezproblemowy sposób przenosić oprogramowanie aplikacyjne ze środowiska e!COCKPIT do CODESYS 3.5.

Od dnia 1 lipca 2023 roku wydawane będą tylko aktualizacje krytyczne dla środowiska e!COCKPIT. Z kolei do 31 grudnia 2023 roku będzie można uzyskiwać licencje dla tego środowiska. Całkowite zakończenie wsparcia e!COCKPIT nastąpi 31 marca 2025 roku. W późniejszym okresie będzie możliwa tylko instalacja ostatniej wersji e!COCKPIT (przy wsparciu technicznym WAGO).

<https://tiny.pl/cc4w4>



Kluczowe osiągnięcie firmy SK Hynix

Pamięć 4D NAND Flash została opracowana w sierpniu 2022 roku. Teraz pora na jej masową produkcję po zakończeniu testów, które były niezwykle wyczerpujące. Dzięki maksymalnej przepustowości 2,4 Gb/s zapewniony jest szybki odczyt, a także szybki zapis danych. Oznacza to rewelacyjną wydajność m.in. dla smartfonów i komputerów PC. Rozwiązanie jest przewidziane również dla dysków SSD PCIe 5.0. Mimo znacznej liczby warstw, nie odnotowano zwiększenia rozmiarów 4D NAND Flash lub spadku konkurencyjności cenowej tej pamięci. Jak wyjaśnia szef działu odpowiedzialnego za pamięci 4D NAND Flash w firmie SK Hynix, Jumsu Kim: „Firma SK Hynix opracowała pamięć 4D NAND Flash przeznaczoną m.in. do smartfonów czy dysków SSD. Biorąc pod uwagę konkurencyjność 4D NAND Flash pod względem ceny, wydajności i jakości, spodziewamy się przyrostu zysków w drugiej połowie 2023 roku. I nieustannie będziemy przewidywać ograniczenia technologii NAND”.

<https://tiny.pl/cc4dh>



Efektywny chłód ze słońca

Fotowoltaika pozostaje opłacalnym źródłem energii dla gospodarstw domowych. Na największe oszczędności mogą liczyć te gospodarstwa domowe, które zwiększą autokonsumpcję – aktualne zużycie produkowanej energii elektrycznej. Klimatyzacja może zwiększyć jej poziom nawet o 10% w skali roku. Gospodarstwo domowe, które korzysta z klimatyzatora o mocy 3,5 kW, który pracuje przez 4 godziny dziennie, w okresie 3 upalnych miesięcy, zużyje przez ten czas prawie 300 kWh. Przy założeniu, że produkcja z instalacji fotowoltaicznej i zużycie energii w budynku będzie rzędu 3 GWh, takie chłodzenie budynku w czasie, kiedy energii słonecznej jest najwięcej, umożliwia z całą pewnością przyrosty autokonsumpcji o przeszło 10% w skali roku. Jednocześnie aż 100% energii potrzebnej do zasilania urządzenia może pochodzić ze słońca. Wystarczy tylko odpowiednie sterowanie urządzeniem. Kiedy energii słonecznej jest dużo, za pomocą urządzeń elektrycznych można zamieniać ją na przyjemny chłód w pomieszczeniach.

Warto pamiętać o tym, że nowoczesne klimatyzatory to również rewersyjne pompy ciepła. Takie urządzenia, gdy jest okres przejściowy tj. wiosna-lato-jesień-zima, pozwalają na dogrzewanie budynku, co dodatkowo zwiększa komfort i pozwala zadbać o ekologię, jak również powiększyć wskaźniki autokonsumpcji.

<https://tiny.pl/cc4dm>

Jakub Tyburski
jakub.tyburski@elportal.pl

MultiHub dla komputerów SBC

Uniwersalny MultiHub rozszerza możliwości komputerów SBC takich jak Raspberry Pi o dodatkowy Ethernet, dwa porty USB2.0 oraz dwa porty szeregowy UART z wbudowanymi skonfigurowanymi driverami RS232. Moduł przydaje się podczas tworzenia aplikacji komunikacyjnych lub integracyjnych w systemach automatyki domowej lub IoT, eliminując płątaninę kabli i osobnych płytek konwerterów. Format nakładki jest zgodny mechanicznie z płytą HAT Raspberry Pi, co ułatwia jej pewne mocowanie, a wyfrezowane otwory nie utrudniają dostępu do złączy kamery i wyświetlacza Pi. Oczywiście nakładka, jako Hub USB z rozszerzoną funkcjonalnością, może zostać zastosowana z dowolnym komputerem wyposażonym w port USB, pracującym pod systemem Windows lub Linux.

Termostat do elektrozaworu z silnikiem DC

Zwykły termostat steruje grzałką lub chłodnicą na zasadzie włącz-wyłącz. W niektórych sytuacjach, na przykład podczas ogrzewania pomieszczeń kolektorami słonecznymi, układ zarządzający temperaturą musi sterować silnikiem obracającym zawór kulowy – właśnie do tego został zaprojektowany zaprezentowany układ. W odróżnieniu od znanych na rynku rozwiązań, jest w stanie regulować przepływ cieczy, która obiekt ochładza lub ogrzewa. Może sterować elektrozaworem, w którym znajduje się elektryczny silnik prądu stałego. Odpowiednie zarządzanie pracą takiego silnika pozwoli na jego bezawaryjną pracę przez długie lata. Co równie ważne, zadawanie odpowiednich nastaw jest banalnie proste, zatem nie trzeba być mistrzem elektroniki, aby móc takiego termostatu z powodzeniem używać.

Moduł 4 przełączników z interfejsem USB-C

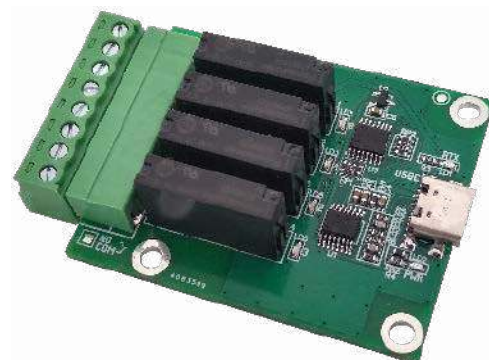
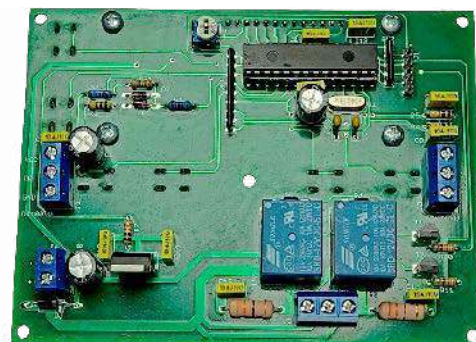
Moduł czterech wyjść przełącznikowych o obciążalności 30 V/5 A do łatwego sterowania urządzeń za pomocą komputera PC. Komunikacja odbywa się poprzez interfejs USB-C w trybie zgodności z USB, a zastosowanie mostka UART/GPIO/I²C ze zdefiniowanym protokołem komunikacyjnym zwalnia nas z tworzenia aplikacji dla mikrokontrolera, przenosząc oprogramowanie na komputer PC.

Sygnalizator napętnienia wanny

Kąpiel w wannie to wspaniała forma relaksu. Wystarczy zatkać odpływ, odkręcić kran i... przychodzić co chwilę, aby sprawdzać, czy jest już odpowiedni poziom wody? Na szczęście takie zadanie może wykonać za nas układ elektroniczny – zasygnalizuje nam osiągnięcie zadanego poziomu wody. Dużą zaletą tego układu jest zastosowanie gotowego modułu do bezkontaktowej detekcji poziomu cieczy. Nie trzeba wiercić w wannie jakichkolwiek otworów, nie będą nam przeszkadzały płytki, nie trzeba też martwić się o korozję metalowych płytek zanurzanych w wodzie. Plastikową puszczykę, szczelnie zaizolowaną przez jej producenta, przykleja się od zewnętrznej strony wanny (pod jej obudową) i gotowe! Warunek: wanna nie może być metalowa. Ceramika, akryl, szkło, tworzywa sztuczne i inne nieprzewodzące materiały są jak najbardziej dopuszczalne.

Tematy wiodące w EP 10/2023:

- Diagnostyka termowizyjna
- Obudowy szyte na miarę



Wykaz firm ogłaszających się w tym numerze „Elektroniki Praktycznej”.

AKSOTRONIK.....	29
ARMEL	13
BORNICO.....	9
COMPUTER CONTROLS.....	11
ELMAX.....	19
FERYSSTER.....	13
GAMMA	13
HAMMOND.....	7
MICROCHIP	5, 60, 108
MICRODIS.....	54, 55
PIEKARZ	13

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 Bryndza, 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”.





**Functional
Safety Ready**

Udoskonalony kontroler dotykowy PTC dla interfejsów użytkownika nowej generacji

Driven Shield Plus i najlepsze w swojej kategorii bezpieczeństwo funkcjonalne Tolerancja na wodę i niezawodne działanie

Branża motoryzacyjna nieustannie poszukuje sposobów na poprawę komfortu kierowców i pasażerów, z czego wynikają duże oczekiwania w stosunku do interfejsów dotykowych, aby działały dobrze w wilgotnym środowisku przy obecności zakłóceń i zawierały funkcje bezpieczeństwa funkcjonalnego. Sterowanie dotykowe zapewnia bardziej intuicyjny i nowoczesny sposób interakcji z pojazdem, ułatwiając dostęp do informacji i sterowanie funkcjami. Dzięki ulepszonej technologii kontrolera PTC sterowanie dotykowe jest jeszcze bardziej responsywne i precyzyjne, zapewniając płynniejsze i lepsze wrażenia użytkownika.

Nasza rodzina mikrokontrolerów PIC32CM JH ma zintegrowany kontroler PTC, który zapewnia tolerancję na wilgoć, odporność na zakłócenia, responsywne i precyzyjne wykrywanie dotyku dla interfejsów użytkownika nowej generacji. Dodatkowo funkcje FuSa dostarczane z tą rodziną mogą być używane do implementacji aplikacji ISO 26262, IEC 61508 i IEC 60730.

Kluczowe właściwości

- Driven Shield Plus zapewnia wiodącą w branży tolerancję na wilgoć
- Zasilanie 5V daje dużą odporność na zakłócenia
- Obsługa AUTOSAR ASIL B MCAL
- Obsługa powierzchni dotykowych, przycisków, suwaków i pokręteł



microchip.com/pic32cmjh



eprasa.pl 826d83c0a0

Nazwa i logo Microchip oraz logo Microchip są zastrzeżonymi znakami towarowymi firmy Microchip Technology Incorporated w Stanach Zjednoczonych i innych krajach. Wszystkie inne znaki towarowe są własnością ich zarejestrowanych właścicieli. © 2023 Microchip Technology Inc. Wszelkie prawa zastrzeżone. MEC2512A-POL-07-23